

Computing Care: Algorithmic Decision-Making and Inequality in Access to Medicaid
Long-Term Care

By

Ailish Burns

M.A., Brown University, 2022
B.A., University of Wisconsin-Madison, 2017

Dissertation

Submitted in partial fulfillment of the requirements for the degree of Doctor of
Philosophy in the Department of Sociology at Brown University

PROVIDENCE, RHODE ISLAND
May 2026

© Copyright 2026 by Ailish Burns

This dissertation by Ailish Burns is accepted in its present form by the Department of Sociology as satisfying the dissertation requirement for the degree of Doctor of Philosophy.

Date _____

Emily Rauscher, Advisor

Recommended to the Graduate Council

Date _____

Margot Jackson, Committee Member

Date _____

Josh Pacewicz, Committee Member

Date _____

Mark Suchman, Committee Member

Date _____

Pamela Herd, Committee Member

Approved by the Graduate School

Date _____

David Lindstrom, Dean of the Graduate School

CURRICULUM VITAE

Ailish Burns is a sociologist studying health inequality, technology, and policy with substantive expertise in Medicaid and maternal and infant health. She earned an A.M. in Sociology from Brown University in 2022 and a B.A. in Sociology with a certificate in Gender and Women's Studies from the University of Wisconsin-Madison in 2017. Her work, which has been funded by the National Science Foundation and the National Institute for Child Health and Human Development, uses both quantitative methods and in-depth interviews to understand how policies and organizational practices like algorithmic decision-making shape population health inequality. Another stream of research examines maternal and infant health inequality. She has published in *Social Science and Medicine*, *Demography*, *Sociological Perspectives*, and *RSF: The Russell Sage Foundation Journal of the Social Sciences*.

ACKNOWLEDGEMENTS

First, I would like to thank my brilliant committee: Emily Rauscher, Margot Jackson, Josh Pacewicz, Mark Suchman, and Pamela Herd. Emily Rauscher, my advisor since day 1 of graduate school, has been particularly important for this journey. I am thankful for her guidance and mentorship, both intellectual and personal. I am also grateful for other faculty who provided me with moral support, encouragement, and intellectual stimulation: Susan Short, Rachel Wetts, Scott Frickel, Nicole Gonzalez Van Cleve, John Eason, and Patrick Heller.

This dissertation would not have been possible without Shukria Sakhi and Gavin Schilling, who provided me with stellar research assistance during data collection. I am so appreciative of their hard work, patience, dedication, and equally importantly, for their kindness and sense of humor! I know they will go far in their academic and professional journeys.

The past seven years have been filled with long days reading, writing, and researching, but these are complemented by sunny afternoons on the green eating muffins from “the Faunce”, quick post-lunch cookie and coffee runs, weekend picnics at India Point Park, and evenings at GCB, the Point, Captain Seaweeds, and a rotating cast of graduate student apartments. All of this would have been much less enjoyable without my very special Providence community: Ester Fanelli, Bruno Cardinale Lagomarsino, Suvina Singal, Mathias Larsen, Ieva Zumbyte, Nallasivan V, Ike Uri, Maggie Eberts, Keenan Wilder, Alexandra Leonard, Jonathan Tollefson, Archana Ramanujam, Dan Kitson, Erika Dunn-Weiss, Liangdi Xu, Subadevan Boyd-Devan, Amari Boyd-Devan, Yasemin

Bavbek, Sertaç Sen, Holly Kirkman, Mritunjay Marol, Ale Cueto Piazza, Tomas Gold, Carilee Osborne, Juho Korhonen, and Danusha Selva Kumar. I also need to acknowledge Erin Ochocki and Olivia Cook, my Wisconsin gals who provided support from afar! Many of the aforementioned made valuable contributions to the 2025-2026 cohort of babies: Samuel Tachon, Andreas Leonard Wilder, Isla Korhonen, Amelia Mortveit, Vytas Zumbyte, and Lana Singal Larsen. You all give me hope for a better world!

Ester Fanelli and Suvina Singal deserve a special mention. We started this program together as nervous strangers and we ended it together as sisters. I am also especially thankful for Ike Uri, Maggie Eberts, Holly Kirkman, and Jay Marol who kept Dan and I fed during the most difficult 9 days of our lives.

Throughout my PhD (and the rest of my life) my family has been a steady source of laughs, TV debriefs and support for which I am immensely appreciative. Thank you to my parents, Joyce and Thomas, who raised me to be curious and care about the state of the world. Thank you to my brother Cooper and sister-in-law Johanna for your generosity, conversation, and ping pong lessons. Thank you to my in-laws – Carol, Gary, Emily, Kyle, Aaron, Sarah, Hannah, Finley, and Alaina – for the old fashioned, gator rides, and for introducing me to the Tachon way of life.

More than anything, I would like to thank Daniel Tachon, the person who started this journey with me as my boyfriend, then became my husband, and then the father of my child, our Sammy. You are the ultimate snack deliverer, cookie baker, DRM-free e-book reader, and chauffeur. In 2019, you agreed to move across the country with a person

you'd known for one year. Seven years later, we have a dog and a baby and a joint back account. I love our silly little songs and routines and WOW, I love you!

I would thank Francis and Sammy too, but honestly you two only made it more difficult to get this dissertation done on time. Still, there's no greater feeling than wrapping up a day of sociology-ing and cuddling on the couch with my boys. I love you both, you make everything worth it.

TABLE OF CONTENTS

INTRODUCTION	1
CHAPTER 1	6
Introduction.....	7
Policy Context.....	9
Theoretical Background.....	11
Algorithms and Technological Governance	11
Characteristics of Algorithmic Decision-Making	14
Methods.....	21
Data and Measures	21
Sample.....	24
Analysis.....	24
Results.....	26
Changes in Algorithmic Decision-Making from 2012 to 2020	26
Relationship Between Algorithm Characteristics and State Political Context	29
Discussion	32
Tables.....	35
Figures.....	40
CHAPTER 2	44
Introduction.....	45
Policy Context.....	47
Algorithm Typology	49
Theoretical Background.....	51
Algorithmic Decision-Making and Racial Inequality.....	51
Street-Level Bureaucrats and Resource Allocation	53
Methods.....	57
Data and Measures	57
Sample.....	61
Analysis.....	61
Sensitivity Analyses.....	63
Results.....	64
Descriptive Analysis	64
First Differences Regression Analysis.....	65
Differences by Medicaid Generosity	67
Sensitivity Analyses.....	68
Discussion	69
Limitations	72
Tables.....	74
Figures.....	80
Appendix.....	82
CHAPTER 3	86
Introduction.....	87
Policy Context.....	88
Theoretical Background.....	90
Algorithmic Decision-Making in Organizations	90

Deservingness and Cultural Inequality	93
Case: Medicaid Long-Term Care in Wisconsin.....	95
Methods.....	99
Interviews.....	99
Fair Hearing Appeals	100
Results.....	102
Medical Eligibility Screening: Objectivity and Deservingness	102
Judges and Adjudication: Legality and Credibility	107
Implications for Inequality.....	110
Discussion	113
Tables.....	116
Figures.....	117
CONCLUSION	118
REFERENCES.....	121

LIST OF TABLES AND ILLUSTRATIONS

TABLE 1.1.....	35
TABLE 1.2.....	36
TABLE 1.3.....	36
TABLE 1.4.....	37
TABLE 1.5a.....	38
TABLE 1.5b.....	38
TABLE 1.6.....	39
TABLE 2.1.	74
TABLE 2.2.	74
TABLE 2.3.	75
TABLE 2.4.	76
TABLE 2.5.	77
TABLE 2.6.	78
TABLE 3.1.	117
TABLE 3.2.	117
FIGURE 1.1.....	40
FIGURE 1.2.....	41
FIGURE 1.3.....	42
FIGURE 1.4.....	43
FIGURE 2.1.....	80
FIGURE 2.2a.....	80
FIGURE 2.2b.....	81
FIGURE 3.1.....	117
FIGURE 3.2.....	117

INTRODUCTION

This dissertation examines the contours of algorithmic eligibility determination for Medicaid Long-Term Services and Supports (LTSS). In doing so, I provide some of the first population-level evidence on how algorithmic decision-making has changed over the past decade and how different types of algorithmic decision-making are associated with racial inequality in access to long-term care. Prior research on algorithmic decision-making and inequality has been relatively idiosyncratic, examining the social dimensions of a single algorithm in a single context. By leveraging inter- and intra- state variation in Medicaid LTSS algorithmic eligibility determination, I show how individual characteristics of algorithms interact with the political and organizational context in which they are situated to shape inequality. As part of this work, I also propose a novel typology for characterizing algorithmic decision-making and have developed an original dataset of the algorithms used in Medicaid LTSS eligibility.

The first two chapters examine the predictors and consequences of different types of algorithmic decision-making from 2012 to 2020. This time period represents a general increase in algorithmic decision-making and marks the first time that the federal government pushed for increased standardization and digitization of medical eligibility processes in Medicaid LTSS.

Chapter 1: Balancing Legitimacy and Control: How Algorithmic Decision-Making in Medicaid Long-Term Care Changed from 2012 to 2020

While the literature is replete with claims that the use of automated or algorithmic decision-making has increased, there is little systematic evidence to support this, particularly in the public sector. Nor do we understand how particular characteristics of algorithmic decision-making have changed. To address this gap, Chapter 1 of this dissertation analyzes how the characteristics of algorithmic eligibility determination in Medicaid LTSS changed from 2012 to 2020. Specifically, I examine changes in automation, user expertise, organizational location, and the source of eligibility tools. I find that during this time period, the percentage of beneficiaries evaluated by automated processes which require a human to review eligibility decisions increased nationally. Conversely, fully automated assessment processes and processes where algorithms are just used as an aid to assist human decision-makers decreased.

Chapter 1 also examines how state political context is associated with the types of algorithms states use. I find that on average in both 2012 and 2020, states where a majority of beneficiaries were evaluated with automated processes which require human oversight had a lower level of Medicaid generosity and a higher percentage of votes for Republicans. It's possible that more conservative states prefer these types of algorithms because automation improves efficiency while the presence of a human "in-the-loop" enhances institutional legitimacy and provides the state with more oversight of eligibility decisions.

This chapter provides national evidence that in the context of Medicaid eligibility determination, fully automated decision-making has actually decreased in favor of processes which incorporate human decision-makers. I also show how the use of technical procedural mechanisms in public benefits allocation may reflect political goals.

Chapter 2: Algorithms in a Shifting Political Landscape: Algorithm Characteristics, Medicaid Generosity, and Racial Inequality in Long-Term Care

While research suggests that both automated decision-making and decision-making left entirely up to humans can exacerbate inequality, these are not mutually exclusive. As I show in Chapter 1, humans are deeply embedded in algorithmic decision-making and their role is actually increasing in the context of public benefits. Still, it is unclear how the extent of human involvement or the characteristics of these users is associated with racial inequality in resource allocation. To understand this, Chapter 2 analyzes the relationship between automation, user expertise, and racial inequality in access to long-term care from 2012 to 2020.

Using first differences regression analysis, I find that net of a host of political and demographic variables, shifts from fully automated algorithms implemented by medical professionals to other typologies of algorithmic decision-making are associated with decreases in the percentage of non-white nursing home residents. This suggests that fully automated algorithms implemented by medical professionals are the most beneficial for racial equality in the nursing home population. As I do not find differences in racial inequality according to algorithm automation or user expertise separately, these findings also demonstrate how the effects of algorithm characteristics can hinge on each other.

Furthermore, as Chapter 1 shows how political context is associated with the types of algorithms states use, Chapter 2 examines how this political context shapes the implications of algorithmic decision-making for racial inequality. I find that as Medicaid generosity increases, differences between algorithm types matter less for racial

inequality. This indicates that in more restrictive policy contexts, the characteristics of algorithmic decision-making may be more important mechanisms for determining access to resources. This is a particularly important finding as the federal government aggressively targets Medicaid for cutbacks.

Chapter 3: 100% Objective to Some Degree: Organizational Context and Inequality in a Medicaid Eligibility Algorithm

While the first two chapters examine algorithmic decision-making at the population level, Chapter 3 provides a case study of how a Medicaid LTSS algorithm operates in practice in a single state. Using in-depth key informant interviews and analysis of 235 court appeals in Wisconsin, I examine how two different groups of actors – eligibility screeners and appeal judges – implement and adjudicate the medical eligibility algorithm. I also assess how this process varies across different groups of applicants.

I find that while Medicaid applicants with greater health literacy and program knowledge are advantaged throughout the eligibility process, these resources are particularly important for gaining access to care during the initial screening process. Because of the institutional focus on objectivity and standardization, applicants must signal their limitations in very precise ways to be legible to the eligibility algorithm. Conversely, cultural resources are potentially less deterministic during the appeals process because judges can eschew the algorithm and make eligibility decisions that follow more lenient state regulations.

This chapter shows how institutional logics can lead to intra-state variation in the extent to which eligibility is automated. In the case of Wisconsin, the initial eligibility process,

which is fully automated, likely limits access to care for applicants with fewer cultural resources. On the other hand, the appeals process, in which judges use algorithms as decision-making aids, subjects applicants to fewer obstacles to care. By examining variation within the same eligibility process but across organizational contexts, this chapter highlights how individual bias and meso-level structures shape algorithmic resource allocation with implications for inequality.

In providing population-level evidence on the predictors and consequences of algorithmic decision-making in a public benefits setting, these three chapters demonstrate how organizational and political contexts shape patterns of algorithmic inequality, working in tandem with the specific characteristics of algorithms themselves.

CHAPTER 1

Balancing Legitimacy and Control: How Algorithmic Decision-Making in Medicaid

Long-Term Care Changed from 2012 to 2020

Introduction

Algorithms are often used to determine access to public benefits with varying consequences (Eubanks 2017). For instance, in 2016, Arkansas implemented a new benefit determination algorithm which resulted in immediate and drastic care reductions for many Medicaid long-term care beneficiaries (Lecher 2018). Like other tools embedded in the recesses of bureaucratic procedures, this Medicaid benefit algorithm was altered behind the scenes (ACLU New Jersey 2022), limiting public understanding of which characteristics of the algorithm actually changed or which factors motivated the state to make this change.

Though not new, algorithmic decision-making has become increasingly of interest to scholars and the general public. This is also evidenced by Google Books NGram results for the phrase “algorithmic decision - making” (Figure 1.1). The proportion of books that mention this phrase increased dramatically starting around 2014, with no sign of slowing as of 2022.

However, because most research on algorithms involves analyzing single, isolated cases, it is difficult to ascertain how the characteristics of algorithmic decision-making has changed over this time period or what political factors predict different types of decision-making. It is especially important to understand the landscape of algorithms in the context of public benefits. Governments frequently use these tools to allocate resources like healthcare, unemployment, disability insurance, and more (Eubanks 2017). Despite having meaningful implications for the life chances of the population, much of the features of algorithmic decision-making in governments are hidden from the public as

regulatory processes and rulemaking rather than formal legislation (Levy et al. 2021). This obscures how resources are actually allocated.

In this paper, I examine how algorithmic eligibility determination in Medicaid long-term services and supports (LTSS) changed from 2012 to 2020. U.S. states use algorithms to establish functional eligibility for these services. Each state uses a different eligibility process, which makes Medicaid LTSS a useful case to assess variation in algorithmic decision-making over time and across states.

I ask:

RQ1. How have the characteristics of algorithmic decision-making in Medicaid LTSS eligibility changed from 2012 to 2020?

RQ2. Are states' political context associated with the types of algorithms they use to determine functional eligibility in Medicaid LTSS?

I find that from 2012 to 2020, algorithmic decision-making in Medicaid LTSS moved away from fully automated processes to increasingly incorporate humans “in-the-loop.” In both years, states where the majority of beneficiaries were evaluated using automated programs that require human oversight had significantly less generous Medicaid programs and in 2012, had a significantly higher percentage of Republican voters.

Further, states tend to fall into two classes of algorithmic decision-making: those that rely more heavily on non-medicalized, fully automated processes that are outsourced to organizations external to the state, and those that rely more heavily on medicalized, in-house, human-in-the-loop processes.

This project leverages a context with widespread yet highly variable algorithmic decision-making to examine how different characteristics of decision-making have changed over time. In doing so, I show how state political context is associated with the technical, procedural mechanisms used to allocate resources. Specifically, I provide suggestive evidence that conservative states are more likely to use algorithmic eligibility processes that are conducted in-house and are automated but retain human oversight. In doing so, it's possible that conservative states benefit from the efficiency of automated decision-making while incorporating human oversight bolsters legitimacy and provides more control over who gets access to resources. Finally, I propose a novel typology of algorithmic decision-making that can be extended to other contexts.

Policy Context

Medicaid is the United States' publicly-funded health insurance program and the largest funder of LTSS in the country. This program provides care for the elderly and people with disabilities in-home, in the community, or in institutions like nursing facilities.

Applicants must meet financial thresholds to be eligible for Medicaid in general and must also be functionally or medically eligible to receive funding for LTSS through the program.

Compared to publicly-funded healthcare in other high-income countries, Medicaid is relatively restrictive, though the generosity of this program has changed over time. After the passage of the Affordable Care Act (ACA) in 2010, states were given the option to expand Medicaid to cover everyone under age 65 with incomes up to 133 percent of the Federal Poverty Level. Prior to this, adults without children or disabilities were ineligible

for Medicaid in most states. By January 1, 2020, 36 states, including Washington, D.C., had expanded their Medicaid programs. Non-expansion states largely included conservative southern and plains states¹ (Kaiser Family Foundation 2025).

The ACA also brought about changes to LTSS, including establishing the Balancing Incentive Program which, from 2011 to 2015, awarded grants to states² to increase access to home and community-based care. The grant stipulated that state functional eligibility tools must meet certain requirements, such as the inclusion of five core eligibility domains (ADLs, IADLs, medical conditions, cognition, and behavior). Further, grantees were encouraged to automate their functional eligibility processes as part of a larger push to store program data electronically. Thus, many states adjusted their eligibility protocols during this time period to meet grant requirements, presumably resulting in increased eligibility automation nationally starting after 2011. (Balancing Incentive Program n.d. ; Kako et al. 2013).

While the Balancing Incentive Program instituted several general requirements for functional eligibility determination, participating states were still given substantial latitude in shaping their assessment tools. Further, states that did not participate in the Balancing Incentive Program were not subject to any federal requirements regarding

¹ Non-expansion states as of January 1, 2020 include Alabama, Florida, Georgia, Kansas, Mississippi, Missouri, Nebraska, North Carolina, Oklahoma, South Carolina, South Dakota, Tennessee, Texas, Wisconsin, and Wyoming. Missouri, Nebraska, North Carolina, Oklahoma, and South Dakota have since expanded Medicaid (as of January 1, 2026).

² 18 states participated in the program: Arkansas, Connecticut, Georgia, Illinois, Iowa, Kentucky, Maine, Maryland, Massachusetts, Mississippi, Missouri, Nevada, New Hampshire, New Jersey, New York, Ohio, Pennsylvania, and Texas.

functional eligibility determination. Extensive variation in functional eligibility processes therefore persists.

Perhaps because of the state-level variation and lack of concrete policy guidelines, to my knowledge there has previously been only limited systematic data on the characteristics of each state's functional eligibility process and how these have changed over time (one exception is a 2016 report from the Medicaid and CHIP Payment Access Commission). This project uses an original data set that characterizes the functional eligibility determination process of each state in 2012 and 2020.

A typical functional eligibility screening involves a nurse or social worker evaluating the applicant through a combination of observations, interviews with family members and caregivers, and possibly a review of the applicant's medical records. The screener inputs this information into a standardized online form, and a computer program will return an eligibility determination. However, within this general framework, there are many differences across states. In this paper, I examine four characteristics of functional eligibility: the extent of automation, the expertise of eligibility screeners, the organizational location of eligibility determination, and the source of the eligibility tool.

Theoretical Background

Algorithms and Technological Governance

The word “algorithm” encompasses a wide variety of tools, from simple calculations based on several discrete variables to computationally intensive algorithms that use machine learning and big data to make predictions (Christin 2017). Though this term has become more widely recognized with the advent of machine learning, large language

models, and artificial intelligence, simpler algorithms and quantification have been used to categorize people and make decisions throughout history. Indeed, a key aspect of Weber's classic theory of bureaucracy involves rational, objective action according to "calculable rules" (1978, p. 975).

This quantification achieves several goals. First, it is a legitimating device: decisions made through "objective", quantified processes are perceived not as decisions but as inevitabilities (Porter 1995). When their decisions are challenged, bureaucrats can fall back on the results of algorithms. As Porter notes, "Objectivity lends authority to officials who have very little of their own" (1995; p. 8).

The use of procedural mechanisms to avoid accountability has been termed bureaucratic disentanglement and arises primarily as a way to contain costs (Lipsky 1984). Lipsky argues that because these mechanisms are often hidden in indirect "nondecisions", they are rarely subject to public scrutiny or opposition. Algorithms are key tools of bureaucratic disentanglement (Henman 2022) as they have been termed "black boxes" (Pasquale 2015) which obscure their inner workings.

Second, quantification renders populations legible (Scott 2020). In *Seeing like a State*, James Scott argues that governments engage in standardization to track, organize, and control their citizens. Practices like drawing maps and establishing naming conventions exemplify this standardization. This facilitates systematic data collection and makes people easier to compare through quantitative measures (Porter 1995) which is necessary for algorithms to operate.

Still, the use of quantification by governments has evolved, not least because algorithmic tools have become more pervasive and are now used to make decisions about individuals rather than primarily to shape policies and programs in the aggregate (Levy et al. 2021). Scholars have characterized the current era of algorithmic decision-making as one of “technological governance”, or “e-governance”, in which administrative processes and interactions between governments and citizens are increasingly digitized (Buffat 2015). Similarly, the phrase “algorithmic governance” has been used to capture the ways that computer-based tools are used to impose order on people and institutions (Katzenbach and Ulbricht 2019).

While technological governance serves political goals, it often does so at the expense of citizens. For one, algorithms are prone to producing unequal outcomes. While they are touted as objective and fair, the literature provides countless examples of these tools producing race- and gender-biased outcomes (Benjamin 2019; Berk et al. 2021; Zou and Schiebinger 2018). Their opacity and often proprietary nature also makes it more difficult for subjects to challenge decisions (Burrell 2016; Pasquale 2015).

States may rely on this opacity to alter or restrict the provision of resources without having to change legislation. In one such instance, in 2016 Arkansas implemented a new LTSS algorithm that resulted in dramatic reductions to care for many beneficiaries (Brown et al. 2020). In this case, the state used the eligibility determination tool to reduce Medicaid benefits with little warning.

Finally, algorithmic decision-making depoliticizes problems of resource allocation. For instance, in her book *Automating Inequality*, Virginia Eubanks (2017) argues that

governments use algorithmic decision-making to provide technical solutions to political and moral issues, thus sidestepping responsibility for the root causes of systemic societal ills like poverty.

While we understand many of the general motivations and consequences of algorithmic decision-making in the public sector, it is not clear how the characteristics of algorithms align with certain political motivations. Which types of algorithmic decision-making are preferred by states intending to restrict benefits vs. those intending to expand them? In answering this question, we can better uncover general, mutable pathways through which algorithmic decision-making regulates access to resources. To address this gap, this paper examines the relationship between the characteristics of algorithmic decision-making for Medicaid LTSS and the political context of states.

Characteristics of Algorithmic Decision-Making

I operationalize four dimensions of the process of algorithmic decision-making: automation, user expertise, organizational location, and tool source. These characteristics have been shown to matter for the fair and equitable allocation of resources. Further, they apply to algorithms used in a variety of settings which makes this typology more generalizable. See Table 1 for an overview of each dimension.

Automation

In the public imagination, algorithms are often thought of as being fully automated tools with limited to no human intervention. However, humans are involved in essentially all algorithmic decision-making processes, though to different degrees and levels of intentionality. Processes which explicitly incorporate human discretion or decision-

making are often described as “human-in-the-loop” systems, which are often proposed to keep algorithms in-check or to improve the reliability of results (Alqazlan et al. 2025; Jones 2017; Peeters 2020). For instance, Meg Leta Jones argues that in Europe, the right for citizens to request a human in the loop stems from concerns about dehumanization and threats to dignity brought about by fully computational decision-making (2017).

Human-in-the-loop systems are more humanized and may provide critical oversight to algorithmic decision-making (Alqazlan et al. 2025; Jones 2017). Yet, algorithmic output shapes human thought processes and understanding. For instance, users may misinterpret algorithmic output with negative consequences, such as interpreting risk scores as prescriptive and handing down stricter sentences (Hannah-Moffat 2013). Further, in some cases algorithms might produce more equitable outcomes than humans. A 2024 report found that when judges ignore the recommendations of pretrial risk assessments, racial inequality increases (Anwar et al. 2024).

Additionally, incorporating humans into algorithmic decision-making may serve organizational goals. For one, human-in-the-loop systems may be perceived as more legitimate than fully automated decision-making particularly for high-impact decisions (Hillo et al. 2025; Waldman and Martin 2022). Integrating human discretion can also increase buy-in from bureaucrats. Tummers and Bekkers (2014) find that when bureaucrats perceive having discretion, they see their jobs as more meaningful and are consequently more willing to implement policies. Thus, incorporating these actors in algorithmic decision-making may make them more amenable to carrying out the goals of their organization.

On the other hand, fully automated algorithms are those which do not explicitly allow humans to determine eligibility or override algorithmic output. These represent the most depersonalized configuration, and much of the literature has highlighted pitfalls of these types of processes. They are prone to system-level errors which are more widespread than those that arise from discrepancies in bureaucratic discretion (Gules-Guctas 2025), though organizations may be better able to avoid accountability for resource allocation when decisions are fully automated (Eubanks 2017; Zalnieriute et al. 2019). While historically, more technical or quantified decision-making processes have boosted institutional legitimacy (Porter 1995), this may be eroding as problems with automation come to light (Grimmelikhuijsen and Meijer 2022; König and Wenzelburger 2021). Still, fully automated systems are especially efficient and can reduce labor costs (Administrative Conference of the United States 2021; Monarcha-Matlak 2021), which may be particularly attractive in settings with limited resources.

Although “human-in-the-loop” processes are often presented as diametrically opposed to fully automated processes, there are gradients within these categorizations (Citron and Pasquale 2014; Grønsund and Aanestad 2020). While all algorithmic decision-making processes involve humans in some capacity, even if primarily at the point of development or programming, not all processes involve people to the same extent. To understand this variation, I classify algorithmic decision-making processes as either “fully automated”, “automated but requires approval”, or “decision-making aids.”

Both automated/approved algorithms and decision-making aids represent different configurations of “human-in-the-loop” systems. In algorithmic decision-making processes which are automated but require approval, an algorithm produces an eligibility

determination but a human must review and/or approve the decision. Finally, in processes that use algorithms as decision-making aids, algorithms produce a recommendation or a standardized score that is then taken into account by a person or team of human decision-makers.

Because human-in-the-loop systems are increasingly recognized as beneficial for both governments and citizens, I expect to find that:

H1. The use of human-in-the-loop algorithms increased and the use of fully automated algorithms decreased from 2012 to 2020.

However, as fully automated algorithms are potentially more cost-saving and reduce government accountability, I expect to find that:

H2. States with lower Medicaid generosity indices and higher percentages of votes for Republicans are more likely to use fully automated algorithms than human-in-the-loop algorithms.

User expertise

The characteristics of the human who implements an algorithm likely shapes how they do so. Even in fully automated processes, the presence of technology does not eliminate discretion or personal judgment from decision-making. Indeed, scholars have proposed the enablement thesis to show how technology, including algorithms, shape the use of discretion in multifaceted ways (Buffat 2015). Notably, the way algorithms shape discretion depends on contextual factors, including the professional identities of users (Ibid.).

The use of healthcare professionals in eligibility determination may reflect a more medicalized approach to allocating care. Conversely, transitions away from requiring healthcare professionals may reflect professional deskilling which can accompany the implementation of algorithms (Ferdman 2025). Fears of deskilling may also provoke professionals to resist the implementation of algorithms in their work or ignore algorithmic output entirely (Brayne and Christin 2021). As a group with a strong professional identity, medical professionals may be especially resistant to algorithms in their work (Ackerhans et al. 2024), requiring states to balance medically appropriate care allocation with ease of implementation.

Accompanying the general move toward more automated decision-making in governments, I expect that:

H3. The percentage of beneficiaries evaluated by algorithmic eligibility processes that require the involvement of medical professionals decreased from 2012 to 2020.

As de-skilling often goes hand in hand with fully automated processes, I also expect that:

H4. States with lower Medicaid generosity indices and higher percentages of votes for Republicans are less likely to incorporate medical professionals in algorithmic eligibility determination.

Organizational location

Healthcare has become increasingly privatized (Himmelstein and Woolhandler 2008; Unruh and Rice 2025), with more services provided by health management organizations and private hospitals. The turn toward privatization has also affected Medicaid eligibility determination, as many states now contract with external organizations to perform this

function. One recent example is Maximus, which was contracted by several states to clear the Medicaid enrollment backlogs after the COVID Public Health Emergency ended (Watson 2022). As of March 2026, Maximus handles at least some of the functional eligibility screening in 20 states (Maximus n.d.).

Whether algorithmic decision-making is performed in-house or outsourced matters for the provision of services and may reflect state funding decisions and political priorities. For instance, there is evidence that when states hire private contractors to assist with Medicaid redetermination it often leads to disenrollment for administrative reasons rather than actual changes to eligibility (Watson 2022). Indeed, government functions that are contracted out are often not subject to the same oversight or regulations which in turn increases opacity and displaces responsibility for decision-making (Mulligan and Bamberger 2019).

I expect that:

H5. The percentage of beneficiaries evaluated by outsourced eligibility processes increased from 2012 to 2020.

Because outsourced functions are subject to less regulation and scrutiny, I expect that:

H6. States with lower Medicaid generosity indices and higher percentages of votes for Republicans are more likely to outsource eligibility determination.

Tool source

Finally, I examine the origin of each state's algorithmic decision-making tools. State agencies may use homegrown tools or adopt tools developed by outside organizations.

Whether states develop their own tool may reflect institutional isomorphism or a desire to boost legitimacy or streamline access to services.

As public critique of algorithmic decision-making grows, government organizations may be especially subject to coercive pressures that produce isomorphism, or similarities across organizations (DiMaggio and Powell 1991; Frumkin and Galaskiewicz 2004). As an organization becomes more similar to others in its field, its legitimacy increases (DiMaggio and Powell 1991). Thus, more states may turn to externally validated, standardized tools to resist scrutiny and boost legitimacy of their eligibility processes.

The use of pre-existing tools may also benefit a state's population, as this could streamline access to care by reducing state variation in functional eligibility standards and limiting learning costs (Herd and Moynihan 2019) for beneficiaries moving across states. Depending on the tool used, it could also reduce variation within programs. Several states have adapted the Minimum Data Set – which the federal government requires nursing homes to use to assess residents (U.S. Centers for Medicare and Medicaid Services 2024) – to determine initial eligibility for Medicaid LTSS. This reduces the overall number and type of assessments that Medicaid LTSS beneficiaries in nursing homes must undergo, again potentially reducing learning costs or other administrative burdens (Herd and Moynihan 2019).

Still, while external tools might provide more legitimacy to functional eligibility determination, these tools are likely not as targeted toward the specific needs of a state's population. For instance, some states, like Maine, have a larger elderly population who

might require a different type of eligibility screen to appropriately establish their eligibility than a state with a larger population of people with physical disabilities.

In light of the potential growing coercive pressures on states, I expect that:

H7. The percentage of beneficiaries evaluated using homegrown tools decreased from 2012 to 2020.

While there are potential competing political motivations for using external tools vs. homegrown tools, due to their capacity to increase access to services I expect that:

H8. States with higher Medicaid generosity indices and lower percentages of votes for Republicans are less likely to use homegrown algorithmic eligibility assessments.

Methods

Data and Measures

To examine how Medicaid LTSS eligibility algorithms have changed over time, I use original state-level data from 2012 and 2020. I collected these data using official state 1915c waiver applications, which I accessed through the Wayback Machine Internet Archive. 1915c waivers permit states to provide Medicaid LTSS in home- and community-based settings and applications provide detailed information on the eligibility criteria and processes for each program. Applications must be renewed at least every 5 years or whenever the state makes changes. I referenced the most recent application up through January of the analysis year, as many states update their waiver applications at

the beginning of the year. By ending data collection in January, I also avoid any issues with COVID-related changes to Medicaid eligibility in 2020.

These data classify algorithms according to their automation, whether a medical professional (physician or nurse) must either conduct or review the eligibility assessment, whether eligibility is conducted in-house or contracted out, and whether the eligibility tool is homegrown or developed by an outside organization. To examine how the characteristics of algorithms changed from 2012 to 2020 (RQ1), I measure automation as fully automated, automated and approved, or decision-making aid. However, for ease of analysis, I measure this as a binary variable (fully automated and human-in-the-loop) when conducting Latent Class Analysis.

Most states have several long-term care programs which may target different populations (often either the elderly, people with physical disabilities, or people with intellectual or developmental disabilities) or provide different types of services. Sometimes states use the same functional eligibility tool across programs (often termed “universal assessments”), but often different programs have different tools and eligibility requirements. Some programs use multiple tools to determine eligibility.

To account for this complexity, I measure each algorithm characteristic as the percentage of beneficiaries³ in each program whose eligibility is determined by an algorithm with the characteristic in question. For instance, take a hypothetical state with two Medicaid long-term care programs: Program A and Program B. Program A serves 20,000 people and

³ I use the number of beneficiaries collected by Katherine Miller at Johns Hopkins (<https://publichealth.jhu.edu/faculty/4587/katherine-miller>)

uses a fully automated tool to determine eligibility. Program B serves 5,000 people and uses a decision-making aid to determine eligibility. In this case, the state serves 80% of beneficiaries with fully automated tools and 20% with decision-making aids. Many states have values of either 0% or 100% on each of these variables, resulting in bimodal distributions. Therefore, I convert these to dummy variables that equal 1 if at least 50% of beneficiaries had their eligibility determined by an algorithm with the characteristic in question.

To explore how states with different “types” of algorithm usage differ, I use two time-varying indicators of state political context. First, I incorporate a measure of Medicaid generosity from the State Policy and Politics Database (SPPD; Montez 2025). This is measured as the average of scores on income eligibility, administrative burden, immigrant benefits, and Medicaid benefits. A higher value indicates a more generous Medicaid program. More generous states may be more likely to choose algorithms that are intended to increase equality or expand access, while more restrictive states may choose algorithms that restrict access, especially to less “deserving” groups (Schneider and Ingram 1993).

However, because LTSS eligibility determination is highly technical and largely decided at the organizational level, it may reflect more than formal Medicaid generosity.

Therefore, I include indices for the percentage of votes for Republican congresspeople (from the Federal Elections Committee). These variables capture political sentiment among the general public, likely also including the bureaucrats who implement these algorithms.

Sample

My final sample includes 42 states that are non-missing⁴ on my variables of interest in both 2012 and 2020. In 2020, these 42 states comprised about 84.3% of the U.S. population and 83.6% of the population on Medicaid.

Analysis

I use a descriptive, exploratory approach to analyze how algorithm characteristics changed from 2012 to 2020 (RQ1). I calculate the mean and standard deviation of each algorithm characteristic in both years and the mean difference from 2012 to 2020. Next, to assess whether there are different classes of algorithm types, I use Latent Class Analysis, which detects unobserved groups in the data and assigns each observation a probability of belonging to each class. For each year, I estimate state clusters using the poLCA package in R. I fit models for 2, 3, and 4 clusters and chose the model with the lowest Bayesian Information Criterion and Akaike Information Criterion. For both 2012 and 2020, the 2-class model was the most appropriate fit.

To assess intraclass homogeneity, I compute average posterior probabilities which are the average probabilities that each observation is a member of each class. Average posterior probabilities closer to 1 indicate greater separation and more distinct classes. In 2012, the average posterior probability for Class 1 was 0.963 and for Class 2 was 0.836. In 2020, the average posterior probability for Class 1 was 0.993 and for Class 2 was 0.776. These

⁴ Missing states include Arizona, Indiana, Massachusetts, Michigan, Ohio, Rhode Island, Tennessee, Vermont, and the District of Columbia.

indicate strong separation of observations in both years and support the suitability of the 2-class model.

I compute posterior-weighted means with 95% confidence intervals which indicate the estimated proportion of each characteristic in the latent class. This is more appropriate for small sample sizes because it avoids classification error by accounting for the uncertainty involved in assigning an observation to a particular class. I also estimate the effective sample size of each class which is the sum of the probabilities that each observation belongs to the class. This avoids assigning states to a single class, which may not be appropriate for states where the probabilities of belonging to each class are relatively equal.

To assess whether states' political context is associated with the types of algorithms they use (RQ2), I first examine the relationship between separate algorithm characteristics and the state political context. I calculate the mean and standard deviation of each political variable when each algorithm type equals 0 and 1. Algorithm variables equal 1 when at least 50% of the beneficiaries in the state were evaluated with the algorithm type and 0 otherwise. I use Welch's t-test to test for significant differences between 0 and 1 on this binary algorithm measure. Welch's t-test is robust to unequal variances and group sizes.

Finally, I examine differences in the political context of states in each algorithm classification group. To do so, I regress each context variable separately on latent class and weight the regression by posterior probabilities of class membership. I do not include any other control variables due to my small sample size. Because of limited statistical

power, this analysis should be interpreted as exploratory and descriptive rather than causal.

Results

Changes in Algorithmic Decision-Making from 2012 to 2020

Table 1 shows the characteristics of algorithmic decision-making in 2012 and 2020. Over this time period, there were small increases in both the mean number of waiver programs and the mean number of eligibility tools used by states. In 2020, on average states used more than 2 different tools to determine functional eligibility, totalling 94 tools up from 78 in 2012. This proliferation of eligibility processes could suggest either a more tailored approach to beneficiaries within states or an increase in complicated administrative procedures which act as barriers to care.

Next, I show how specific algorithm characteristics changed from 2012 to 2020 (see Figure 1.2). Supporting hypothesis 1, I find that the percentage of beneficiaries evaluated using fully automated tools decreased from 2012 to 2020 by almost 8 percentage points. While the use of human-in-the-loop algorithms increased, this was driven entirely by an increase in the use of automated/approved algorithms, which more than doubled over the time period while decision-making aids decreased. Still, decision-making aids remained the most prevalent type in both years. This suggests that on the whole, states are increasingly moving toward automated eligibility processes but only while retaining a human-in-the-loop. In fact, human-in-the-loop tools covered over 70% of beneficiaries in 2020. It's possible that automated/approved processes are perceived as beneficial for

states as automation increases efficiency while a human-in-the-loop legitimizes eligibility determination and gives bureaucrats ultimate control over decisions.

Turning to user expertise, there was a small decrease in the percentage of beneficiaries evaluated by tools which require a medical professional, as expected by hypothesis 3.

This shift toward a less medicalized approach to eligibility may reflect the de-skilling of professions that can accompany the implementation of algorithms (Ferdman 2025).

Unexpectedly, the percentage of beneficiaries whose eligibility was outsourced to a local non-state or private entity decreased over time in contrast to hypothesis 5, though by less than two percentage points. This suggests that while generally stable, states are increasingly relocating eligibility determination in-state. Similarly to the increase in automated/approved algorithms, this may reflect a desire for states to retain more control over functional eligibility.

Finally, patterns regarding tool source suggest a more standardized approach to eligibility. As expected by hypothesis 7, while a majority of beneficiaries in 2020 were still evaluated by tools developed within their state of residence, the percentage evaluated by homegrown tools decreased by over 8 percentage points from 2012 to 2020. This suggests that states are increasingly turning to externally validated tools such as the InterRAI suite of assessments, perhaps in response to the federal government's call for functional eligibility to be standardized as part of the Balancing Incentive Program. While states can still implement these tools differently, the decrease in homegrown tool use may lead to less variation across states in how eligibility is determined. This could limit lapses in care for LTSS users moving across state lines and decrease the

administrative burden brought on by learning to navigate a new eligibility process. However, in conjunction with the increase in average number of tools, administrative burden may persist if beneficiaries are having to navigate multiple eligibility screenings.

Latent class analysis. Next, I conduct Latent Class Analysis to examine whether the characteristics of different algorithms group together. In this analysis, I use a binary measure of algorithm use where each characteristic equals 1 if at least 50% of beneficiaries were evaluated by an algorithm with that characteristic. In both years, algorithms fall into two classes. Figures 1.3 and 1.4 show each state's predicted probability of being in each class as well as the dominant class in each state. In 2012, southern states were more likely to have Class 2 as their dominant class, but geographic patterns largely disappeared by 2020.

Table 3 shows the mean probability that observations in each class equal 1 on each algorithm characteristic. For instance, in 2012, Class 1 had about a 48% probability of having a majority of beneficiaries evaluated by a human-in-the-loop algorithm. That is, observations in this class are about equally as likely to rely on human-in-the-loop algorithms as fully automated algorithms. In contrast, observations in Class 2 have about an 87% chance of using mostly human-in-the-loop algorithms.

In both years, Class 1 is characterized by lower probabilities of human-in-the-loop algorithms and algorithms that require medical professionals, and higher probabilities of outsourced but homegrown algorithms. I call this class "Deskilled Outsourcers." Again, the association in this class between fully automated algorithms and less use of medical professionals reflects observations made in prior research that automated decision-

making can lead to the deskilling of professional fields (Ferdman 2025). Furthermore, the use of fully automated algorithms and the associated higher probabilities of outsourcing suggest that states in this class are potentially less concerned about retaining control over eligibility decisions.

Conversely, Class 2 in both years is characterized by greater probabilities of human-in-the-loop algorithms that require medical professionals. Class 2 also relies more heavily on homegrown tools, but has a lower probability of outsourcing eligibility algorithms, particularly in 2012. I call this class “In-House Professionals.” In sum, these states appear to prioritize relying on the expertise and discretion of medical professionals to make eligibility decisions, and keep eligibility primarily under state purview.

In 2012, the division between classes was more stark and algorithmic characteristics were more meaningful indicators of group membership, with probabilities closer to 0 and 1. For instance, in 2012, 22% of observations among Deskilled Outsourcers and 100% of observations among In-House Professionals had a majority of algorithms which required a medical professional, whereas in 2020 37% and 52% did, respectively. This characteristic more clearly divides observations in 2012. This indicates that in 2020 there was less grouping of algorithm characteristics with fewer systematic differences across states. In other words, automation, user expertise, organizational location, and tool source were more clearly associated with one another in 2012 than in 2020.

Relationship Between Algorithm Characteristics and State Political Context

From 2012 to 2020, the national political context shifted (Table 2). There was a small increase in Medicaid generosity brought on by the Affordable Care Act. This period also

marked a decrease in the percentage of votes for Republican congresspeople nationally.

But are these factors related to the algorithms states use to determine LTSS eligibility?

Tables 5a and 5b show the descriptive relationship between algorithm characteristics and state political context in 2012 and 2020 respectively.

Overall, I did not find political differences according to user expertise, but there are significant differences in political context for the other three characteristics. Though there is variation by year, it appears that states that mostly outsourced and used homegrown tools were *less* conservative, while states that mostly used automated/approved algorithms were *more* conservative.

I expected that states where the majority of beneficiaries were evaluated by fully automated algorithms would have significantly less Medicaid generosity and be more conservative (hypothesis 2). This hypothesis does not hold, however. Instead, states with more fully automated processes are slightly more generous in 2012.

There are differences in political context according to whether the state evaluated a majority of beneficiaries with automated/approved algorithms. States where at least 50% of beneficiaries were evaluated with automated/approved algorithms had significantly lower Medicaid generosity indices in both years and significantly higher percentages of Republican votes in 2012. It's possible that more conservative states are particularly likely to use automated/approved processes as a way to increase legitimacy and retain control over the outcomes of functional eligibility assessments.

I also expected that increased outsourcing is associated with less Medicaid generosity and a more conservative population (hypothesis 6). Again, the evidence does not support this

hypothesis. Rather, in 2012, states where at least 50% of beneficiaries were evaluated by outsourced algorithms had significantly higher Medicaid generosity indices and significantly lower percentages of votes for Republicans. There were similar patterns in 2020 though differences were not significant. This is surprising given that states have historically contracted with private companies to facilitate Medicaid disenrollment (Watson 2022). Again, conducting assessments in-house could provide conservative states with more eligibility oversight (Levy et al. 2021).

Finally, results do not support hypothesis 8: the use of homegrown tools is associated with greater Medicaid generosity and a lower percentage of Republican votes in 2020. More generous states may have more resources to devote to developing their own tools, or may prefer to use tools that can be more appropriately adapted to their populations.

Differences by latent class. Table 5 shows odds ratios from a logistic regression of political variables on latent class. Because of my small sample size, I run models for each political variable separately. Odds ratios represent the change in the odds of being in the class “In-House Professionals” relative to “Deskilled Outsourcers” with every one unit increase in the independent variable. Importantly, none of the odds ratios are significant for either year, which is unsurprising given a sample size of 42. Further, the odds ratios are all relatively close to 1, which suggests that changes to the independent variables are not associated with much change in the odds of being in the “In-House Professionals” class over “Deskilled Outsourcers”. It appears that the political context of states is more associated with individual characteristics of algorithms rather than the class of algorithms states use.

Despite this, it is notable that in 2012, “In-House Professionals” was the dominant class primarily in southeastern states (Figure 1.3), providing additional weak yet suggestive evidence that this combination of characteristics – particularly reliance on human-in-the-loop and in-house processes – may be favored by conservative states.

Discussion

While we know that governments use algorithms to allocate resources, we do not have a grasp on how the characteristics of algorithmic decision-making processes have changed in the last decade as interest in these tools have grown. Instead, to date, research on algorithmic decision-making has been relatively idiosyncratic, focusing on single algorithms in a single context. This project proposes a novel typology of algorithmic decision-making and examines how these characteristics change over time in the context of Medicaid LTSS.

I also examine whether there are differences in the types of algorithmic decision-making according to state political context – while “objective” algorithms are often explicitly depoliticized (Eubanks 2017), I show that the implementation of these tools varies according to Medicaid generosity and the political conservatism of the population.

Algorithms are commonly thought of as automated systems without human intervention, but I find that on average, from 2012 to 2020, states actually increased human involvement in Medicaid LTSS eligibility determination, particularly in more conservative states. While algorithms used in artificial intelligence and machine learning are becoming rapidly depersonalized, my findings suggest that actuarial algorithms in public benefits contexts often retain a human in the loop, perhaps to advance political

goals like garnering legitimacy or restricting eligibility without formal legislative or procedural changes.

Further, during this time period I find decreasing use of homegrown tools, outsourced eligibility, and eligibility processes that require medical professionals. Both the use of homegrown tools and eligibility outsourcing was associated with greater Medicaid generosity and lower percentages of Republican voters, though in different years. While the increased standardization of eligibility tools may reduce administrative burden associated with learning costs for beneficiaries (Herd and Moynihan 2019), it's possible that more generous states prefer tools that can be more easily adapted to the needs of the population.

The association of outsourcing with increased generosity in 2012 is surprising, as outsourcing can reduce transparency in decision-making and is often associated with increasing disenrollment (Watson 2022). However, this relationship is no longer significant in 2020, so it's possible that the political motivations for outsourcing eligibility are slowly changing. In general, more research is needed to understand *why* states choose particular types of algorithms in public benefit determination and what the consequences of these choices are for the population.

Finally, I find patterns in the types of eligibility algorithms states use. Using latent class analysis, I find two general classes of states in both 2012 and 2020: 1) “Deskilled Outsourcers” which is associated with greater use of fully automated algorithms that are outsourced and which do not require a medical professional, and 2) “In-House Professionals” which is associated with more human involvement in eligibility, more use

of medical professionals, and less outsourcing. Class membership was more salient in 2012, perhaps because states are becoming more similar over time in an isomorphic process. I do not find significant differences in political context across classes.

In short, states are increasing human oversight in algorithmic decision-making, but only when paired with automation. While “human-in-the-loop” systems are often endorsed as failsafes against the potential harm of algorithms, the predominant use of these systems by conservative states instead suggests that automated processes with human oversight may serve political goals of welfare retrenchment while maintaining legitimacy in decision-making. Additionally, the divergence in patterns between “automated/approved” and “decision-making aid” algorithms suggests that “human-in-the-loop” may not be a useful grouping term in the context of public benefits and may be hiding important variation within human-in-the-loop configurations.

Of course, the characteristics of algorithmic decision-making are likely not universally used for the same purpose. However, highlighting the characteristics that are more often associated with particular political contexts is a starting point for understanding the pathways through which algorithms regulate access to resources.

While this study is limited by a small sample size, it is an important starting point for categorizing algorithms and tracking changes in their use over time. By examining how the characteristics of algorithmic decision-making changed in the previous decade, we can hypothesize about where these processes may be headed in the future.

Tables

Table 1.1. Algorithm Characteristics

Characteristic	Description
<i>Automation</i>	
Fully automated	An algorithm produces an eligibility decision. Humans are not involved in making this decision.
Human-in-the-loop	
Automated/approved	An algorithm produces a decision and a human reviews and can overturn the decision.
Decision-making aid	An algorithm produces a recommendation and a human makes the eligibility decision.
<i>User Expertise</i>	
Requires medical professional	A medical professional (such as a physician or nurse) must conduct or sign off on eligibility assessment. This could mean either the eligibility screener is a medical professional or an independent physician or nurse must approve an applicant's request for services.
Does not require medical professional	A medical professional is not required to conduct or sign off on eligibility assessment. Eligibility screener is most likely social worker.
<i>Organizational Context</i>	
Outsourced	Screeners employed by a private company or local non-state agency conduct the eligibility assessment.
In-house	Screeners employed by the state conduct the eligibility assessment.
<i>Tool Source</i>	
Homegrown	The state uses an eligibility assessment tool that was developed by the state.
Pre-existing	The state uses an eligibility assessment tool that was developed by another state or a private company.

Table 1.2. Characteristics of algorithmic decision-making over time

	2012	2020	Change
	Mean (SD)	Mean (SD)	Mean
Number of waiver programs	2.98 (1.46)	3.50 (2.02)	0.52
Number of tools	1.88 (0.99)	2.05 (1.23)	0.17
<i>% of beneficiaries served by tools</i>			
<i>which:</i>			
are fully automated	35.94 (45.88)	28.23 (39.58)	-7.71
are human-in-the-loop	64.06 (45.88)	71.77 (39.58)	7.71
are automated/approved	11.75 (27.17)	27.26 (39.97)	15.51
are decision-making aids	52.32 (44.74)	44.51 (42.04)	-7.81
require medical professional	50.21 (44.88)	47.39 (44.70)	-2.82
are outsourced	59.61 (43.51)	58.02 (42.41)	-1.59
are homegrown	78.13 (36.78)	70.03 (39.90)	-8.09
n	42	42	42

Note: “human-in-the-loop” tools include both automated/approved and decision-making aids and are mutually exclusive with fully automated tools.

Table 1.3. State political context over time

	2012	2020	Change
	Mean (SD)	Mean (SD)	Mean
Medicaid generosity index	50.21 (9.75)	53.31 (10.02)	3.10
Percentage of votes for Republican congresspeople	53.68 (10.29)	46.75 (11.62)	-6.93
n	42	42	42

Table 1.4. Latent class of algorithm use in Medicaid LTSS eligibility determination

	Human-in-the-loop	Requires medical professional	Outsourced	Homegrown	Effective sample size
	Mean (95% CI)	Mean (95% CI)	Mean (95% CI)	Mean (95% CI)	
2012					
1: Deskilled Outsourcers	0.48 (0.30 - 0.66)	0.22 (0.07 - 0.37)	0.81 (0.67 - 0.95)	0.79 (0.64 - 0.94)	29.64
2: In-House Professionals	0.87 (0.72 - 1.03)	1.00 (1.00 - 1.00)	0.21 (0.02 - 0.40)	0.84 (0.68 - 1.01)	18.12
2020					
1: Deskilled Outsourcers	0.34 (0.02 - 0.66)	0.37 (0.05 - 0.70)	0.77 (0.49 - 1.06)	0.77 (0.49 - 1.06)	8.30
2: In-House Professionals	0.67 (0.52 - 0.82)	0.52 (0.36 - 0.68)	0.56 (0.40 - 0.72)	0.82 (0.69 - 0.94)	37.19

Note: means are posterior-weighted. The effective sample size is the sum of probabilities that each observation belongs to each class.

Table 1.5a. State political context by algorithm characteristic in 2012

Algorithm characteristic	Medicaid generosity		% of votes for Republicans		n when alg = 0	n when alg = 1
	Mean (SD) when alg = 0	Mean (SD) when alg = 1	Mean (SD) when alg = 0	Mean (SD) when alg = 1		
fully automated	48.81 (10.06)	52.50 (9.08)	53.79 (10.68)	53.52 (9.96)	26	16
human-in-the-loop	52.50 (9.08)	48.81 (10.06)	53.52 (9.96)	53.79 (10.68)	16	26
automated/approved	51.34 (9.22)	39.50 (9.04)*	52.77 (10.24)	62.35 (6.46)**	38	4
decision-making aid	49.90 (10.32)	50.50 (9.45)	55.28 (9.90)	52.23 (10.64)	20	22
requires medical professional	51.95 (8.99)	48.48 (10.39)	53.08 (9.74)	54.29 (11.02)	21	21
outsourced	46.41 (9.72)	52.80 (9.07)**	57.16 (10.33)	51.32 (9.76)*	17	25
homegrown	51.00 (11.87)	50.03 (9.39)	55.08 (8.68)	53.36 (10.72)	8	34

Table 1.5b. State political context by algorithm characteristic in 2020

Algorithm characteristic	Medicaid generosity		% of votes for Republicans		n when alg = 0	n when alg = 1
	Mean (SD) when alg = 0	Mean (SD) when alg = 1	Mean (SD) when alg = 0	Mean (SD) when alg = 1		
fully automated	53.27 (9.81)	53.42 (10.98)	45.51 (11.91)	49.84 (10.71)	30	12
automated/approved	55.47 (9.83)	47.92 (8.68)*	45.31 (11.82)	50.35 (10.75)	30	12
decision-making aid	51.64 (11.00)	55.76 (8.06)	48.78 (12.22)	43.77 (10.32)	25	17
human-in-the-loop	53.42 (10.98)	53.27 (9.81)	49.84 (10.71)	45.51 (11.91)	12	30
requires medical professional	54.64 (9.74)	51.85 (10.37)	46.22 (9.90)	47.33 (13.51)	22	20
outsourced	51.18 (9.75)	54.76 (10.13)	48.29 (11.73)	45.70 (11.67)	17	25
homegrown	48.36 (8.09)	55.06 (10.16)*	53.47 (7.69)	44.37 (11.94)**	11	31

*** p<.001, ** p<.01, * p<.05, + p<0.10

Note: algorithm characteristic variables = 1 when at least 50% of beneficiaries were evaluated by an algorithm with that characteristic and 0 otherwise. Welch's t-test used to evaluate whether differences between 0 and 1 are significant.

Table 1.6. Odds ratios from logistic regression of political factors on latent class

	2012		2020	
	Medicaid generosity	% Republican	Medicaid generosity	% Republican
Intercept	3.06 (1.74)	0.29 (1.76)	0.04 (2.47)	0.09 (1.92)
Medicaid generosity	0.97 (0.03)	-	1.03 (0.04)	-
% Republican	-	1.01 (0.03)	-	1.02 (0.04)

*** p<.001, ** p<.01, * p<.05, + p<0.10

Note: Odds ratios compare the odds of being in the “In-House Professionals” class relative to “Deskilled Outsourcers”. Robust standard errors are in parentheses.

Figures

Figure 1.1. Google NGram results for “algorithmic decision - making”

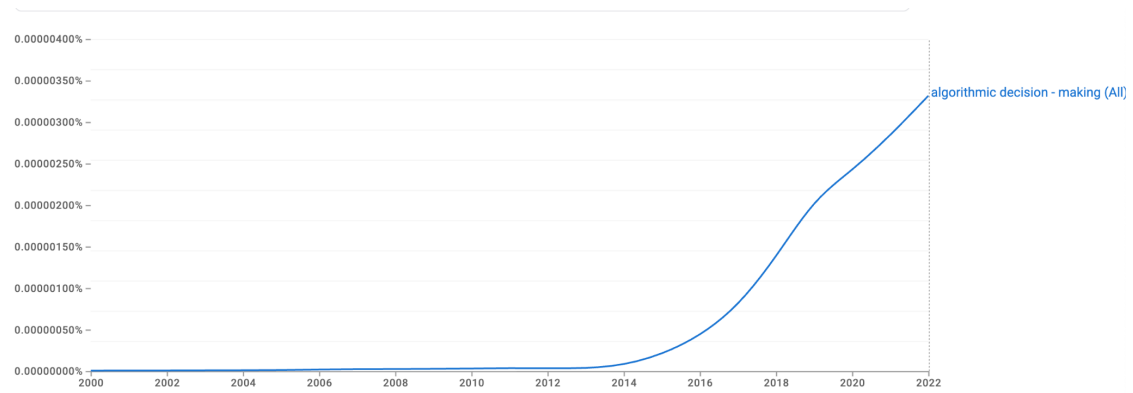


Figure 1.2. Percentage of Beneficiaries Evaluated by Algorithms with Each Characteristic in 2012 and 2020

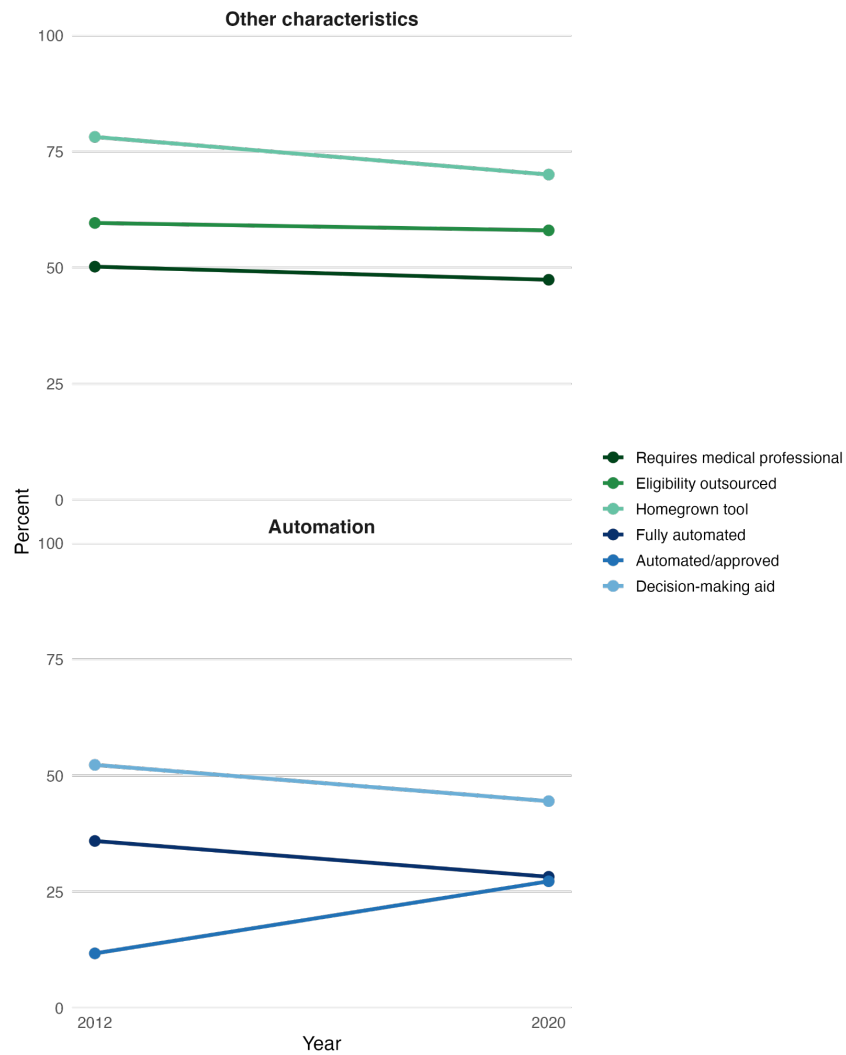


Figure 1.3. Posterior Probabilities and Dominant Class for States in 2012

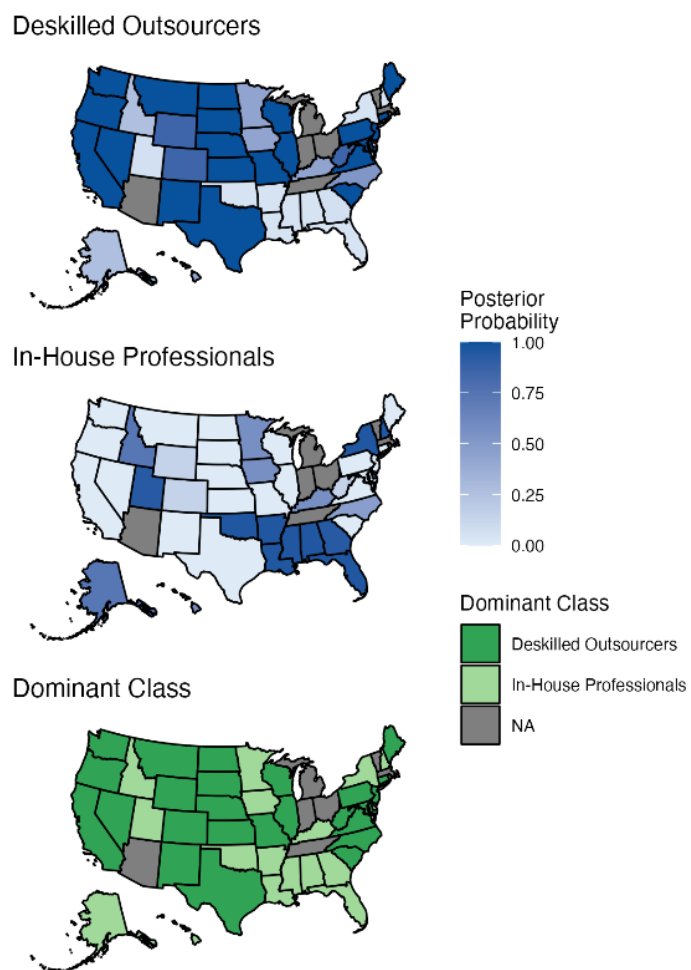
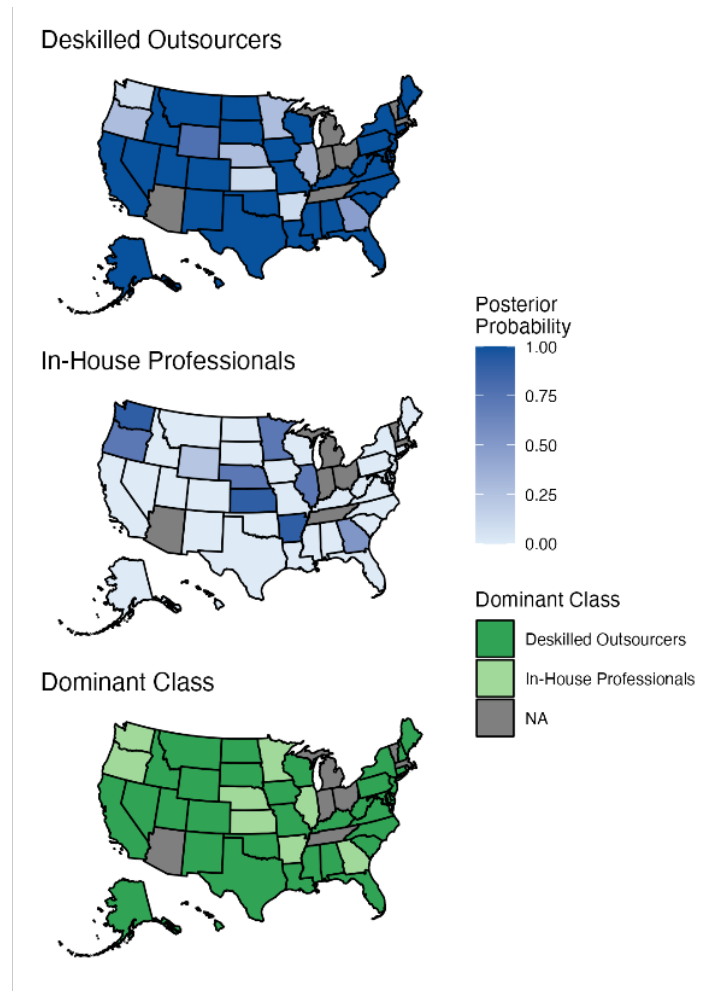


Figure 1.4. Posterior Probabilities and Dominant Class for States in 2020



CHAPTER 2

Algorithms in a Shifting Political Landscape: Algorithm Characteristics, Medicaid
Generosity, and Racial Inequality in Long-Term Care

Introduction

Current research on algorithmic inequality largely focuses on how a particular algorithm in a particular setting shapes who gets access to various resources. However, there has been limited attempt to systematically examine how different characteristics of algorithms are better or worse for racial inequality. One characteristic of algorithms that may have important implications for inequality is automation. Research suggests that both algorithmic decision-making and decision-making left entirely up to humans can exacerbate racial inequality (see Eubanks 2017; Webster et al. 2022). In reality, however, these are not mutually exclusive: human bias is deeply embedded even in processes that appear fully automated (Benjamin 2019). Still, it is unclear how the extent of human involvement – or the characteristics of the humans in question – in algorithmic decision-making are related to racial inequality in resource allocation. Finally, while scholars have noted the importance of the context in which algorithms are embedded, there is limited population-level evidence of how political context might shape the consequences of different types of algorithmic decision-making.

To address these gaps, this project uses the case of functional eligibility determination in Medicaid long-term services and supports (LTSS), which is typically established by some kind of algorithm with substantial variation across states. Specifically, I ask:

- RQ1. What is the relationship between the types of algorithms states use to determine functional eligibility for Medicaid long-term supports and services (LTSS) and racial inequality in access to institutional long-term care?

- 1a. What is the relationship between algorithm automation level and racial inequality in access to institutional long-term care?
- 1b. What is the relationship between eligibility screeners' professional backgrounds and racial inequality in access to institutional long-term care?
- 1c. What is the relationship between the intersection of algorithm characteristics and racial inequality in access to institutional long-term care?

RQ2. Are there differences in the relationship between algorithmic decision-making and racial inequality in long-term care access according to state Medicaid generosity?

I find that net of other social and political variables, the automation level of eligibility algorithms does not have a significant association with racial inequality in LTSS on its own. However, automation does matter depending on the expertise of an algorithm's users. Fully automated algorithms implemented by medical professionals are more beneficial for racial equality in the nursing home population than other types, suggesting that at least in the case of Medicaid LTSS, more "objective" processes are good for racial equality.

Importantly, though, differences between algorithm types matter less as Medicaid generosity increases. This finding underscores the crucial role political context plays in mediating the relationship between algorithmic decision-making and inequality. In generous policy contexts, algorithms are less deterministic for who gets access to care. Consequently, as Medicaid is increasingly threatened, algorithms may become more important tools for shaping inequality in access to care.

This project makes several empirical and theoretical contributions to the literature on algorithmic inequality. First, to my knowledge, this project provides some of the first population-level evidence on whether and how various characteristics of algorithmic decision-making shape inequality. Second, I operationalize two important attributes of algorithmic decision-making processes – automation and user expertise – and show how the effects of these elements hinge on each other and on the broader policy context.

Policy Context

Medicaid is the United States' publicly-funded health insurance program and the largest funder of LTSS in the country. This program provides care for the elderly and people with disabilities in-home, in the community, or in institutions like nursing facilities. Unlike home and community-based services, nursing home care is in entitlement, and Medicaid must provide this service for everyone deemed eligible. Further, nursing facilities are particularly important LTSS providers for racially marginalized groups: the percentage of Hispanic/Latinx, Black, and Asian nursing home residents has grown over time, partly because of unequal access to home and community-based services (Feng et al. 2011). From 2012 to 2020, the percentage of non-white nursing home residents increased from about 19% to 22% nationally (author's analysis of LTCFocus data). For these reasons, in this paper I focus on LTSS provided in nursing homes.

Applicants must meet financial thresholds to be eligible for Medicaid in general and must also be functionally or medically eligible to receive funding for LTSS through the program. Compared to publicly-funded healthcare in other high-income countries, Medicaid is relatively restrictive, though the generosity of this program has changed over

time. After the passage of the Affordable Care Act (ACA) in 2010, states were given the option to expand Medicaid to cover everyone under age 65 with incomes up to 133 percent of the Federal Poverty Level. Prior to this, adults without children or disabilities were ineligible for Medicaid in most states. By January 1, 2020, 36 states, including Washington, D.C., had expanded their Medicaid programs. Non-expansion states largely included conservative southern and plains states⁵ (Kaiser Family Foundation 2025).

The ACA also brought about changes to LTSS, including establishing the Balancing Incentive Program which, from 2011 to 2015, awarded grants to states⁶ to increase access to home and community-based care. The grant stipulated that state functional eligibility tools must meet certain requirements, such as the inclusion of five core eligibility domains (ADLs, IADLs, medical conditions, cognition, and behavior). Further, grantees were encouraged to automate their functional eligibility processes as part of a larger push to store program data electronically. Thus, many states adjusted their eligibility protocols during this time period to meet grant requirements, presumably resulting in increased eligibility automation nationally starting after 2011. (Balancing Incentive Program n.d. ; Kako et al. 2013).

While the Balancing Incentive Program instituted several general requirements for functional eligibility determination, participating states were still given substantial latitude in shaping their assessment tools. Further, states that did not participate in the

⁵ Non-expansion states as of January 1, 2020 include Alabama, Florida, Georgia, Kansas, Mississippi, Missouri, Nebraska, North Carolina, Oklahoma, South Carolina, South Dakota, Tennessee, Texas, Wisconsin, and Wyoming. Missouri, Nebraska, North Carolina, Oklahoma, and South Dakota have since expanded Medicaid (as of January 1, 2026).

⁶ 18 states participated in the program: Arkansas, Connecticut, Georgia, Illinois, Iowa, Kentucky, Maine, Maryland, Massachusetts, Mississippi, Missouri, Nevada, New Hampshire, New Jersey, New York, Ohio, Pennsylvania, and Texas.

Balancing Incentive Program were not subject to any federal requirements regarding functional eligibility determination. Extensive variation in functional eligibility processes therefore persists.

Perhaps because of the state-level variation and lack of concrete policy guidelines, to my knowledge there has previously been only limited systematic data on the characteristics of each state's functional eligibility process and how these have changed over time (one exception is a 2016 report from the Medicaid and CHIP Payment Access Commission). This project uses an original data set that characterizes the functional eligibility determination process of each state in 2012 and 2020.

A typical functional eligibility screening involves a nurse or social worker evaluating the applicant through a combination of observations, interviews with family members and caregivers, and possibly a review of the applicant's medical records. The screener inputs this information into a standardized online form, and a computer program will return an eligibility determination. However, within this general framework, there are many differences across states. In this paper, I examine differences according to the expertise of screeners and the balance between automation and human decision-making in eligibility determination.

Algorithm Typology

The national variation in functional eligibility determination makes Medicaid LTSS an excellent case to examine whether and how different characteristics of algorithmic decision-making are better or worse for inequality. I measure functional eligibility tools

according to their level of automation and the expertise of their users. See Table 2.1 for an overview of this typology.

In this project, I use Angèle Christin’s definition of algorithms as “computational procedures (which can be more or less complex) drawing on some type of digital data (“big” or not) that provide some kind of quantitative output (be it a single score or multiple metrics) through a software program” (2017, p. 2).

“Automation” captures whether eligibility determination is fully automated, automated but must be approved, or whether the results of the eligibility determination are merely used as decision-making aids by human evaluators. For example, in 2020 California used a fully-automated assessment tool for one program which “automatically tallies a score that identifies the participant’s tier level” and determines eligibility (Cantwell 2019). Colorado had an eligibility process that falls into the “automated but must be approved” category. They use a tool called the Uniform Long-Term Care Assessment Form, which assigns applicants a score upon which eligibility is determined. However, “the Department reviews [level of care] evaluations and determinations to ensure the assessments are completed for new waiver participants in accordance with State waiver policies” (Johnson 2019).

Finally, Iowa uses an assessment tool developed by the InterRAI company, but results from this assessment are used as decision-making aids, as the screening team “is responsible for determining the level of care based on the completed evaluation tool and supporting documentation from the medical professional” (Stier 2017). I also categorize programs that use multiple tools to determine functional eligibility as “decision-making

aids”, as the generation of multiple algorithmic scores or eligibility results necessarily requires a human to adjudicate between them.

The characteristic “expertise” is a binary variable that captures whether a medical professional (a physician or nurse) must either complete the eligibility screen or sign off on the assessment either before or after eligibility is determined. Often, this means an applicant’s personal physician must request a screening for the applicant or be consulted during the eligibility process. In others, this just means that the initial eligibility screener must have a background in nursing or another medical profession. In cases where screeners are not medical professionals they are most often social workers.

Theoretical Background

Algorithmic Decision-Making and Racial Inequality

Social scientists have repeatedly shown how algorithmic decision-making can reproduce existing inequality, particularly related to race. In contexts as varied as policing (Brayne 2021), healthcare (Obermeyer et al. 2019), housing (Rosen et al. 2021) and welfare (Eubanks 2017), scholars have identified several pathways through which algorithms exacerbate inequality. For instance, algorithms use data that reflect entrenched inequality (see Obermeyer et al. 2019), are built by biased humans (see Benjamin 2019), and are implemented and interpreted by humans with conscious and unconscious biases.

Regarding this third pathway, algorithms’ intended uses may not reflect how they are used in practice (Brayne 2021; Brayne and Christin 2021; Kellogg et al. 2020). The actors who are tasked with using algorithms may employ strategies to minimize their power or reject the algorithms entirely. Even when bureaucrats and other actors embrace

algorithms, they interpret algorithmic output via their own knowledge, expertise, and bias (Maiers 2017). For example, clinicians who use algorithms to predict infection interpret the algorithm's output in line with their pre-existing understanding of a patient's symptoms (Ibid). This decoupling could either reinforce or reduce racial bias depending on the user and organizational context.

While we know that humans are involved at all levels of algorithmic decision-making, it is unclear whether the *extent* to which humans have control over algorithmic decision-making matters for racial inequality in resource allocation. If humans have little opportunity to overturn decisions made by algorithms, individual bias may be less likely to permeate decisions, for better or for worse.

The evidence for whether algorithms or humans are more racially biased is mixed. Many scholars suggest that decisions made by humans are better for racial equality than algorithmic decision-making in part because algorithms prevent decisions made with empathy, legitimate racial and economic inequality, and prevent organizational accountability (Benjamin 2019; Bloch-Wehba 2020; Eubanks 2017). Despite these arguments, in some cases algorithms can reduce racial disparities relative to decisions made by humans. For example, research suggests that Black individuals are disproportionately penalized when judges ignore algorithmic pretrial risk assessments (Anwar et al. 2024).

This project clarifies this debate by measuring the level of human interference in algorithmic decision-making. Specifically, I classify Medicaid eligibility algorithms as either fully automated, automated but requiring human approval, or used just as decision-

making aids. Given widespread evidence of algorithmic bias, in response to RQ1a, I expect to find that:

H1. Relative to automated/approved and decision-making aid algorithms, increases in the percentage of programs using fully automated algorithms are associated with decreases in the percentage of non-White nursing home residents.

Importantly, research shows that organizational context matters for how algorithms are employed and their implications for inequality (Christin 2017; Hoffman 2017). In more liberal or generous policy environments where algorithms are explicitly implemented to increase equity in outcomes, fully automated algorithms may be better for racial equity. Conversely, in more unequal or conservative contexts where algorithms are used to restrict access to resources, humans may play a crucial role in reducing inequality. This paper examines the interaction between algorithmic decision-making and state policy context. In response to RQ2, I expect to find that:

H2. Increases in programs using fully automated algorithms (relative to automated/approved or decision-making aids) are associated with greater increases in the percentage of non-White nursing home residents when Medicaid generosity increases.

Street-Level Bureaucrats and Resource Allocation

In Medicaid LTSS eligibility and many other public benefits settings, the humans implementing and interpreting algorithms are street-level bureaucrats (SLBs), government workers who interact directly with the public (Lipsky 2010). These actors have significant discretion over resource allocation (Lipsky 2010; Portillo and Rudes

2014), though the nature of SLB discretion has changed in the wake of automated decision-making within governments. In Michael Lipsky's original formulation of SLBs, he argued that machines could not replace the human judgment required in public service (2010). However, in the past decade, increasing attention has been given to how technology shapes the work of SLBs (see Buffat 2015). While algorithms and other information technologies can give the appearance of objective, non-discretionary decision-making, evidence suggests that SLBs use their own judgment to apply these tools to their work (Angelova et al. 2025; Buffat 2015).

SLBs who conduct medical assessments for Medicaid LTSS may use discretion in ways that either disproportionately burden or benefit racially marginalized Medicaid applicants. For instance, these workers might categorize racially marginalized applicants as having "lower status" disabilities that provide access to fewer resources or are weighted less in the algorithm (see Skrtic et al. 2021). Alternatively, medical assessors who perceive issues of racial inequality could intentionally code Black or other racially marginalized applicants in ways that might facilitate their eligibility for services.

Even within the same organization or policy contexts, the characteristics of individual SLBs shape their discretionary behaviors (Tummers and Bekkers 2014). Professional status, such as educational background, expertise, and training, impacts how much discretion an SLB has, how they use it, and the values that underlie that it (Evans 2011) though some research suggests that the effect of professional field on SLB behavior is minimal (Scott 1997). In the case of Medicaid long-term care, eligibility screeners typically have backgrounds in either social work or healthcare. While social workers feature heavily in the SLB literature, nurses and physicians are not typically considered

SLBs. Because medical professionals operate as SLBs in Medicaid LTSS eligibility determination, this provides an interesting case to examine how different types of professionals use algorithms and the implications of this for racial inequality.

It's possible that compared to screening that only involves social workers, the presence of medical professionals in eligibility determination may increase racial inequality in access to care. The practice of medicine has long been embroiled in racial essentialism (Anderson et al. 2023; Phelan et al. 2013), and there is a legacy rooted in slavery and racial essentialism of physicians perceiving Black people as healthier and more resistant to pain (Bridges 2011; Hoffman et al. 2016). It's possible that when involved in functional eligibility determination, medical professionals may be less likely than social workers to perceive non-White applicants as disabled, potentially leading to fewer non-White applicants found eligible for care.

Additionally, requiring applicants' own medical providers to approve eligibility may disproportionately burden racially marginalized applicants who tend to have less access to healthcare (Hill et al. 2022; Yearby 2018; Smedley et al. 2003). These applicants may be less likely to have a medical provider with thorough knowledge of their needs who is willing to advocate for their Medicaid LTSS application.

However, because of their medical expertise, it's also possible that medical providers are better equipped to understand the nuances of applicants' limitations even if they don't manifest in ways visible to social workers. Further, if applicants' own physicians are involved, they may be more willing to advocate for their patients. Still, given potential

barriers to care and the literature on medical racism, in response to RQ1b I expect to find that:

H3. Increases in the percentage of programs that require a medical professional's involvement are associated with decreases in the percentage of non-White nursing home residents.

Further, it's possible that the expertise of eligibility screeners interacts with eligibility automation to affect inequality. For example, the background of eligibility screeners might matter less in fully automated processes than when the algorithm operates as a decision-making aid. Social workers using decision-making aids might be better equipped to consider an applicant's social circumstances, and thus facilitate eligibility for racially marginalized applicants. On the other hand, medical professionals using decision-making aids might be likely to defer to their medical training and make "race blind" eligibility decisions that may actually increase inequality (Hirschman and Bosk 2020; Penner and Dovidio 2016). In response to RQ1c I expect that:

H4. Social workers using decision-making aids are better for racial equality relative to other types of algorithms, and medical professionals using fully automated tools are the worst for racial equality.

State-level Medicaid generosity may affect how eligibility screeners use algorithms to facilitate or deny eligibility for Medicaid LTSS. If restrictive states have more exacting eligibility protocols or feature stronger political messaging about limiting Medicaid benefits, both social workers and medical professionals may have less discretion (real or perceived) in implementing these tools even when the processes explicitly incorporate

humans (Evans 2020; May and Winter 2009). On the other hand, it's possible that if states become more restrictive applicants may face greater barriers in accessing a physician to sign off on their eligibility (Shen and Zuckerman 2005; Sommers et al. 2017) while social workers may intentionally use their discretion to facilitate access to care (Kelly 1994; Maynard-Moody and Musheno 2000).

Because of potential countervailing effects of generosity on medical professionals and social workers, in response to RQ2 I expect to find that:

H5. Differences between medical professionals and social workers are less pronounced when Medicaid generosity increases.

Methods

Data and Measures

This project combines multiple sources of data. First, I use original state-level data from 2012 and 2020 on the types of algorithms states use to determine functional eligibility for Medicaid LTSS. I collected these data using official state 1915c waiver applications, which I accessed through the Wayback Machine Internet Archive. Applications must be renewed at least every 5 years or whenever the state makes changes. 2012 is the earliest year where waiver applications were consistently available through the internet archive, and by using 2020 as the comparison year I avoid changes in both my independent and dependent variables related to the COVID-19 pandemic. I referenced the most recent application up through January of the analysis year, as many states update their waiver applications at the beginning of the year. By ending data collection in January, I also avoid any issues with COVID-related changes to Medicaid eligibility in 2020. As my

primary dependent variable is captured in April of each year, this allows three months for states to implement potential changes to the eligibility tools.

1915c waivers allow states to “waive” certain federal Medicaid program requirements and target programs directly to specific groups, such as the elderly or people with autism or HIV/AIDS. While 1915c waivers are often used to expand home-and community-based services, these applications also include information on the tools used to determine functional eligibility for “state plan” LTSS, which primarily include institutional care. The majority of waiver programs use the same eligibility process for state plan services (70.29%⁷), or are only slightly different in ways that do not change their categorization according to the typologies I use.

These data provide my independent variables, which are the characteristics of algorithms each state uses according to the typology in Table 2.1. This classifies algorithms according to their automation (fully automated, automated but must be approved, or results used as decision-making aids) and whether a medical professional (physician or nurse) must either conduct or review the eligibility assessment.

Most states use multiple functional eligibility tools and have several long-term care programs. Therefore, my independent variables capture the percentage of beneficiaries⁸ in each program whose eligibility is determined by an algorithm with the characteristic in question. For instance, take a hypothetical state with two Medicaid long-term care programs: Program A and Program B. Program A serves 20,000 people and uses a fully

⁷ Author’s calculation. 265 out of 377 waiver programs in 2012 and 2020 use the same eligibility process for state plan services.

⁸ I use the number of beneficiaries collected by Katherine Miller at Johns Hopkins (<https://publichealth.jhu.edu/faculty/4587/katherine-miller>)

automated tool to determine eligibility. Program B serves 5,000 people and uses a decision-making aid to determine eligibility. In this case, the state serves 80% of beneficiaries with fully automated tools and 20% with decision-making aids. I do not include tools that are only used to assess eligibility for applicants with intellectual or developmental disabilities, as this population is usually served either in home- and community-based settings or in Intermediate Care Facilities.

Second, to assess changes to the institutional long-term care population, I use county-level data on the demographics of all nursing homes in the United States from LTCFocus.org. These data are compiled from administrative and claims data by researchers at Brown University and cover all U.S. states except for Alaska due to sample size limitations. Washington, D.C. is also excluded for the same reason. These data include measures of the percentage of non-Hispanic Black (from here on, “Black”), non-Hispanic White (from here on, “White”), and Hispanic (of any race) nursing home residents present in all facilities in the county on the first Thursday in April. This timing avoids issues stemming from seasonal changes in the nursing home population or churning of the same residents moving in and out of facilities. Further, by April 2020 nursing homes had not yet seen drastic declines in the population due to facility closures or residents moving out or dying from COVID.

My primary dependent variable is the county-level change in the percentage of all non-White nursing home residents. Variables indicating the percentage of Black or Hispanic nursing home residents are missing for many counties (n=916 [38.44%] and n=1168 [49.01%] respectively), primarily due to small numbers of these residents. There is much less missingness for the percentage of White nursing home residents (n=31, [1.30%]).

Therefore, my primary analyses examine the change in the percentage of non-White nursing home residents (calculated as 100 minus the percentage of White residents). This way, I reduce the chance that my results are biased by the exclusion of counties where the majority of residents are White.

Though many people receive Medicaid LTC outside of facilities, nursing home or institutional care is an entitlement and not subject to long waitlists or other barriers that are common for home and community-based care. Assessing access to institutional care therefore provides a baseline estimate of inequality that is likely only more pronounced in home and community settings.

Importantly, not all residents of nursing homes in the LTCFocus data are enrolled in Medicaid, though the Kaiser Family Foundation estimates that about 63% of residents have Medicaid as their primary payer (Chidambaram and Burns 2025). For this reason, my results likely underestimate the effects of algorithmic eligibility determination on the Medicaid long-term care population: patterns in the general nursing home population associated with changes to Medicaid eligibility are likely stronger among residents on Medicaid. Still, to adjust for differences across counties, all analyses are weighted by the percentage of nursing home residents on Medicaid.

Third, I include both county- and state-level controls from several administrative sources. See Table 2.2 for an overview of all variables used in the analysis. Controlling for the proportion of Medicaid LTSS recipients in home- and community-based care in each state minimizes the possibility that changes to the population receiving nursing home care is driven by beneficiaries switching to in-home or community-based care. Controlling for

the proportion of White Medicaid enrollees adjusts for the possibility that changes in the nursing home population reflect overall changes to the Medicaid population. I also control for several measures of the social and political context of states, including median income, percent in poverty, total population, and political alignment.

Lastly, I use a measure of Medicaid generosity from the State Policy and Politics Database (Karas Montez 2025). This index ranges from 0-100 (100 is most generous) and represents the average generosity value of income eligibility, administrative burden, immigrant benefits, and Medicaid benefits.

Sample

My final analytic sample includes all nursing facilities in states with no missing data on any variables in my analysis for either 2012 or 2020 (see Table 2.2). This includes 2,238 counties in 35 states⁹ which are non-missing on all variables in my analysis. This is equal to 78% of counties in the 2012 data and 80% in the 2020 data. Combined, these 35 states represent about 81.2% of the total population of Medicaid long-term care beneficiaries (Kaiser Family Foundation 2020).

Analysis

To analyze the relationship between my algorithm typology and racial inequality in access to institutional long-term care, I conduct a first differences Ordinary Least Squares

⁹ Excluded states include Alaska, Arizona, Delaware, Washington, D.C., Hawaii, Indiana, Kansas, Maine, Michigan, New Jersey, Ohio, Rhode Island, South Carolina Tennessee, and Vermont

regression analysis, which regresses the change in the percentage of non-White nursing home residents on the change in algorithms from 2012 to 2020:

$$\Delta NonWhite_{ct} = \beta_1 \Delta Alg_{st} + \beta_2 X_{ct} + \beta_3 \Delta Z_{st} + \epsilon_{it} \quad (1)$$

For each county c , Equation 1 predicts the change over time t in the percentage of non-White nursing home residents ($\Delta NonWhite$) as a function of the change in the percentage of beneficiaries subject to each category of eligibility algorithms (ΔAlg) in each state s . I model each algorithm characteristic (automation and expertise) separately and then as the whole typology. Each type is included as separate continuous variables with the reference group excluded. In models where automation is the independent variable, fully automated is the reference group. In models where expertise is the independent variable, non-medical is the reference group. Finally, in the whole typology, fully automated/medical is the reference group, as this is theoretically the most “objective” category that is least subject to human bias. Each coefficient on Alg can thus be interpreted as the effect of a positive one percentage point change from the reference category to the category in question from 2012 to 2020.

I run each model with and without a set of controls at the county and state level. ΔX represents the change in a set of county-level controls, and ΔZ represents the change in a set of state-level controls. See Table 2.2 for a complete list of these controls. Standard errors are robust and clustered within states. All models are weighted by percentage of nursing home residents on Medicaid in each county at the midpoint of the values in 2012 and 2020. By examining changes within states from 2012 to 2020, I control for any time-invariant characteristics of states.

To assess potential variation according to Medicaid generosity, I run separate models interacting algorithm typology with a continuous measure of Medicaid generosity:

$$\Delta NonWhite_{ct} = \beta_1 \Delta Alg_{st} + \beta_2 \Delta MedGen_{st} + \beta_3 \Delta Alg_{st} * \Delta MedGen_{st} + \beta_4 X_{ct} + \beta_5 \Delta Z_{st} + \epsilon_{it} \quad (2)$$

Equation 2 estimates how the effect of a change in algorithm typology on county-level changes in the non-White nursing home population depends on changes in Medicaid generosity. β_3 represents the percentage point change in the effect of algorithm typology on the non-White nursing home residents for every one unit increase in Medicaid generosity.

Sensitivity Analyses

Many of the counties in my sample have very small non-White populations. Therefore, results may be skewed by counties where a change of only a few residents resulted in large percent changes. To address this, I re-run my analysis limited just to counties where the percentage of non-White nursing home residents is greater than the 50th percentile (8.42%).

As an additional robustness check testing whether results vary with Medicaid generosity, I divide my sample into two parts: 1) counties in states with less than or equal to the median change in the Medicaid generosity index (median = 2.5) and 2) counties in states with greater than the median change in the Medicaid generosity index. I use Equation 1 to assess the relationship between algorithmic decision-making and racial inequality in each of these samples.

Results

Descriptive Analysis

Table 2.3 shows descriptive statistics for all the eligibility tools used in states in my analytic sample. From 2012 to 2020, there is an overall trend toward more automated eligibility determination processes involving medical professionals. In 2012, decision-making aids covered the largest number of beneficiaries, with about 49% of all LTSS beneficiaries in programs using these tools. However, the use of decision-making aids covered only less than 32% of beneficiaries by 2020. In this year, fully automated algorithms were most common, covering almost 37% of beneficiaries. Across this time period, the use of a medical professional was increased from covering around 57% to around 65% of LTSS beneficiaries. Considering both characteristics of algorithms, the most prevalent type in both years was decision-making aids used by a medical professional. Interestingly, there was a sharp decrease, from 14.58 to 1.60, in the percentage of beneficiaries covered by decision-making aids used by non-medical staff. It appears that states are increasingly likely to give more discretionary power to medical professionals than to social workers in eligibility determination.

Table 2.4 shows descriptive statistics for my dependent variable and state and county-level controls. Overall Medicaid generosity increased slightly from 2012 to 2020, reflecting the implementation of the Affordable Care Act. During this time period, the percentage of non-White nursing facility residents increased from 14.86% to 17.20%. The percentage of non-Hispanic White Medicaid similarly decreased over the time period, but by less than 1 percentage point. This suggests that the non-White population

requiring long-term care is outpacing the non-White population on Medicaid. Finally, the percentage of LTSS beneficiaries in home and community-based services increased from 2012 to 2020, reflecting a national effort to “rebalance” Medicaid LTSS away from institutional settings.

First Differences Regression Analysis

See Table 2.5 for results from the first differences regression analysis for the entire sample. In fully adjusted models, there is no statistically significant relationship between changes to either the level of automation or the expertise of the algorithm’s users and changes in the percentage of non-White nursing home residents. Therefore, I do not find support for hypotheses 1 or 3. It is worth noting, however, that the coefficients for automated/approved and decision-making aids are in the expected direction despite being insignificant.

The combination of automation and expertise does have a statistically significant relationship with the racial composition of nursing homes (Figure 2.1). Specifically, relative to the “medical/automated” category, the categories “medical/approved”, “medical/aid”, “non-medical/automated”, and “non-medical/approved” are significantly associated with a decrease in the percentage of non-White nursing home residents. Each one percentage point change from medical/automated eligibility processes to medical/approved eligibility processes from 2012 to 2020 is associated with a 0.023 percentage point decrease in the percentage of non-White nursing home residents. Similarly, one point shifts from “medical/automated” to the “medical/aid”, “non-medical/automated, and “non-medical/aid” categories are associated with 0.017, 0.061,

and 0.023 percentage point decreases in the percentage of non-White nursing home residents. Again, this is unexpected and does not support hypothesis 4.

For context, I estimate the predicted change in the percentage of non-White nursing home residents if all algorithms were changed to each type from 2012-2020. If all algorithms were changed to medical/automated during this time period, the percentage of non-White nursing home residents would increase by about 2.60 percentage points. In contrast, if all algorithms were changed to non-medical/automated processes, the percentage of non-White nursing home residents would *decrease* by about 3.54 percentage points.

While these effect sizes are small on their own, they have substantial implications for racial inequality in the nursing home population. In 2020, there were about 1.3 million residents in nursing homes (Lendon et al. 2024). A 2.60 percentage point increase in the population totals to nearly 34,000 additional non-White residents.

In general, this indicates that among algorithms that require a medical professional to be involved in eligibility determination, fully automated processes are best for non-White long-term care applicants. However, for algorithms that do *not* require a medical professional to be involved, only those involving the most human intervention – decision-making aids – are not associated with declining access to institutional long-term care for non-White applicants. It's possible that when given more discretionary power, social workers are more likely to capture the social context of long-term care applicants and thus facilitate eligibility for non-White applicants.

Differences by Medicaid Generosity

Table 2.6 shows results from models interacting algorithm type with state Medicaid generosity index. In contrast to hypothesis 2, there is no statistically significant relationship when interacting Medicaid generosity with automation level. However, the relationship between expertise and racial inequality in the nursing home population is significantly different based on state Medicaid generosity. When there is no change in Medicaid generosity, each percentage point increase in the percent of consumers served by algorithms requiring medical professionals is associated with a 0.036 point increase in the percent of non-White nursing home residents. With each 1 unit increase in Medicaid generosity, this association decreases by 0.005 percentage points. That is, Medicaid generosity appears to make the expertise of eligibility screeners less consequential for racial inequality in long-term care. This supports hypothesis 5.

Turning to the full typology, Medicaid generosity appears to make the type of algorithm less consequential for racial inequality. When there is no change to Medicaid generosity, changes from medical/automated to other types of algorithms (except non-medical/aid) are associated with decreases in the percentage of non-White nursing home residents. Changes from medical/automated to non-medical/automated algorithms are less strongly associated with decreases in non-White nursing home residents as Medicaid generosity increases (Figures 2.2a and 2.2b). This indicates that for fully automated eligibility processes, the expertise of the algorithm's user matters less in generous states. The findings in support of hypothesis 5 are therefore mostly driven by fully automated processes, as Medicaid generosity does not seem to change the relationship between

expertise and racial inequality for decision-making aids or algorithms that are automated/approved.

Conversely, in the absence of changes to Medicaid generosity, changes from medical/automated to non-medical/aid are associated with an increase in the percentage of non-White nursing home residents. However, this positive association decreases as Medicaid generosity increases, suggesting that in more generous states, shifts to more “subjective” forms of eligibility screening are not as important for racial equality in access to care.

Sensitivity Analyses

Results are largely the same when limited to counties where the non-White nursing home population is greater than the 50th percentile (Appendix Tables A and B). I do not find significant differences based on automation or the expertise of algorithm users alone.

When examining the full typology, changes from medical/automated to medical/approved, non-medical/automated, and non-medical/approved are associated with significant decreases in the non-White nursing home population. Changes from medical/automated to non-medical/aid are associated with significant increases in the non-White nursing home population. The coefficients for these are larger than in the full sample, suggesting that the effect of changes to a state’s algorithm type is stronger in counties with more non-White nursing home residents.

When interacting with Medicaid generosity, coefficients are all in the same direction as for the full sample. For the full typology, only the interaction coefficient for non-medical/aid is significant, and the coefficient is larger than in the main analysis. This

suggests that the effect of changes to Medicaid generosity on the relationship between changes to non-medical decision-making aids and the percentage of non-White nursing home residents is stronger in counties with more racial diversity.

In samples stratified according to changes in Medicaid generosity, results are comparable to interaction models and overall findings support the idea that changes in algorithmic decision-making matter less when Medicaid generosity expands more (see Appendix Tables C1 and C2). The relationship between automation and expertise and racial inequality are only significant in states that improved Medicaid generosity less. Further, there are more significant differences in the whole algorithm typology in these less improved states.

Interestingly, however, in states that expanded generosity more, shifts toward automated/approved processes that incorporate humans (medical professionals or not) are worse for racial equality in access to long-term care. This suggests that fully automated processes used by medical professionals may be better in more expansive states, perhaps because the discretion of medical professionals does not keep pace with the structural benefits of improvements to Medicaid generosity.

Discussion

I find that on their own, the extent to which algorithmic functional eligibility processes are automated or incorporate human decision-making matters little for racial inequality in access to Medicaid LTSS. However, automation *is* an important determinant of access to care when we factor in the characteristics of the algorithms' users and the state political context. In short, fully automated algorithms employed by medical professionals appear

to be best for racial equality in nursing home care. Shifts away from these types of algorithms to either less-automated algorithms used by medical professionals or to algorithms with some extent of automation but not used by medical professionals are associated with decreases in the percentage of non-White nursing home residents. This is surprising, as a growing segment of the literature shows how automation and more “objective” forms of resource allocation and decision-making can perpetuate harm especially for Black people and other non-White groups (Benjamin 2019; Eubanks 2017; Siddique et al. 2024). My divergent findings point to the important role of an algorithm’s users in mitigating inequality – I find that fully automated processes are only beneficial for equality in the hands of medical professionals. This result warns against the deskilling that often accompanies the implementation of fully automated algorithms (Ferdman 2025) – the knowledge and experience of humans still matters in algorithmic decision-making even when they do not have control over the output of algorithms. More work should be done to examine the mechanisms underlying this and whether the expertise of users is important in other settings.

Notably, as Medicaid generosity increases, shifts from medical/automated to non-medical/automated or non-medical/aid processes matter significantly less for racial inequality. These results cannot be explained by changes in LTSS uptake, poverty levels, population political leanings, or the racial demographics of the population. Lipsky’s theory of bureaucratic disentanglement argues that states use procedural mechanisms to restrict access to benefits (Lipsky 1984). It’s possible that more restrictive states manipulate the characteristics of their functional eligibility algorithms to ration resources, thus making these characteristics matter more for inequality. Conversely, in more

generous states an overall culture of generosity may reduce the role that algorithmic eligibility processes play in granting access to care. As Medicaid benefits are slashed at both the federal and state levels, eligibility algorithms then may become increasingly important mechanisms of inequality.

Prior research suggests that the data used by algorithms or other aspects of how an algorithm is programmed matters for inequality (Benjamin 2019; Noble 2018). However, my findings underscore the importance of studying algorithms as embedded in a particular context. In a public benefits setting, the exact same algorithm may have vastly different consequences for racial inequality depending on the person implementing it or the political environment in which it is used.

Further, this project advances a typology of algorithmic decision-making that categorizes algorithms according to their level of automation and the expertise of their users. To my knowledge, this is one of the first attempts to systematically examine how different categories of algorithms matter for racial inequality in access to healthcare. Future research could build on this work by examining other characteristics of algorithms. For instance, there is mixed evidence on whether race-blindness in decision-making is good or bad for racial equality in access to resources (Hirschman and Bosk 2019; Obermeyer et al. 2019). Whether algorithms incorporate the race of individuals or communities could matter in cases like school admissions or funding allocations.

While automation and user expertise are less consequential for inequality on their own, these characteristics matter more in particular combinations, and especially when Medicaid is more restrictive. Future work should seek to measure additional

characteristics of algorithmic decision-making and how they intersect to shape outcomes. Researchers should also continue to situate algorithms in practice, examining how their consequences are shaped both by their organizational context and the broader political landscape.

Limitations

This project is limited by several factors. First, this analysis is not causal. However, I examine changes within states at two time points which reduces the likelihood that time-invariant differences across states underlie my findings. Further, my results are robust to a host of social and political factors like changes in the racial composition of the population, the proportion of the population in home- and community-based services, and state conservatism. Still, there may be other factors that changed from 2012 to 2020 and drive differences in both algorithm use and the racial demographics of nursing homes.

Second, I only examine changes in the nursing home population and cannot assess how changes in algorithmic decision-making are associated with racial inequality in home- and community-based services. Many states have gradually undertaken a rebalancing of LTSS which attempts to move beneficiaries out of institutions and into their homes and communities (Musumeci 2015). However, the fact remains that as an entitlement, nursing home care is more accessible and potentially more responsive to functional eligibility changes. Future research could repeat a similar analysis in home and community settings using claims data or LTCFocus Home Health Data which currently only covers 2016 through 2019 (LTCFocus n.d.).

Third, due to data limitations, I measure racial inequality as changes in the non-White population and cannot further disaggregate by race. Barriers to care and treatment by SLBs vary across racial groups (Einstein and Glick 2017; Harrits 2019; Hernandez and Sparks 2020). More work is needed to understand how the effects of algorithmic decision-making vary across and within populations by race and ethnicity.

Despite these limitations, this analysis provides important and robust population-level evidence that algorithmic decision-making is associated with racial inequality in access to Medicaid long-term care.

Tables

Table 2.1. Typology of algorithms according to automation and expertise of users

	Automated	Automated & approved	Decision-making aid
Medical professional involved	Medical/automated	Medical/approved	Medical/aid
No medical professional involved	Non-medical/automated	Non-medical/approved	Non-medical/aid

Table 2.2. Variables used in analysis

Variable	Level	Source	Timeframe
<i>Independent variables</i>			
Algorithm typology	State	Original data	2012 and 2020 (through January)
Medicaid generosity	State	State Policy and Politics Database	2012 and 2020
<i>Dependent variables</i>			
Percentage of nursing home residents who are non-White	County	LTCFocus	2012 and 2020 (measured on the first Thursday in April)
<i>Controls</i>			
Median household income	County	Small Area Income and Poverty Estimates Program (SAIPE)	2012 and 2020
Percentage of residents in poverty	County	SAIPE	2012 and 2020
Total population (including those living in group quarters like nursing homes)	County	Census	2012 and 2020
Percentage of Medicaid LTSS beneficiaries in home and community-based services	State	Kaiser Family Foundation analysis of T-MSIS	2012 and 2020
Percentage of Medicaid beneficiaries who are White	State	American Community Survey	2012 and average of 2019/2021
Percentage of total noninstitutionalized population who are White	State	KFF analysis of the American Community Survey	2012 and average of 2019/2021
Percentage of votes for Republican candidates in Senate	State	Federal Elections Committee	2010 and 2018

and House races

Note: all variables are measured as the absolute change from 2012 to 2020.

Table 2.3. Descriptive statistics of functional eligibility tools

	Beneficiaries covered in 2012		Beneficiaries covered in 2020	
	n	%	n	%
<i>Automation</i>				
Fully automated	491,201	35.06	931,340	36.55
Automated but results must be approved	217,404	15.52	809,874	31.78
Decision-making aid	692,563	49.43	806,847	31.67
Medical professional must complete or review assessment	802,521	57.28	1,668,462	65.48
<i>Typology: Automation and Expertise</i>				
Medical / fully automated	109,269	7.80	193,069	7.58
Medical / automated but approved	205,033	14.63	709,310	27.84
Medical / decision-making aid	488,219	34.84	766,083	30.07
Non-medical / fully automated	381,032	27.26	738,271	28.97
Non-medical / automated but approved	12,371	0.88	100,564	3.95
Non-medical / decision-making aid	204,344	14.58	40,764	1.60
Total	1,401,168		2,548,061	

Table 2.4. Descriptive statistics of state and county-level variables

	2012		2020	
	Mean	SD	Mean	SD
<i>State-level variables</i>				
Medicaid generosity index	49.89	9.98	53.23	9.91
Proportion of LTSS beneficiaries in home- and community-based services	63.40	10.77	71.31	12.90
Proportion of Medicaid enrollees who are non-Hispanic White	51.28	18.31	50.31	17.77
Percentage of votes for Republican Senate and House candidates	54.70	9.69	47.49	11.77
Number of functional eligibility tools	1.60	0.91	1.51	1.01
<i>County-level variables</i>				
County population	105,128.47	348,487.78	112,882.67	365,352.93
Median household income	59,305.64	15,176.40	67,374.72	17,249.29
Percentage of residents in poverty	17.42	6.61	13.95	5.48
Percentage of nursing facility residents who are non-White (dependent variable)	14.86	17.59	17.20	18.22

Table 2.5. First Differences Regression Results

	Automation	Med professional involved	Combined typology
% of nursing home residents who are non-White			
<i>Automation</i>			
Fully automated	Reference	-	-
Automated/approved	0.003 (-0.013)	-	-
Decision-making aid	0.01 (-0.015)	-	-
Medical professional involved in eligibility determination	-	0.007 (-0.012)	-
<i>Typology</i>			
Medical automated	-	-	Reference
Medical approved	-	-	-0.023*** (-0.006)
Medical aid	-	-	-0.017** (-0.005)
Non-medical automated	-	-	-0.061** (-0.019)
Non-medical approved	-	-	-0.023*** (-0.006)
Non-medical aid	-	-	0.008 (-0.019)
<i>Controls</i>			
% of population that is White	-0.38 (-0.396)	-0.36 (-0.426)	-0.541+ (-0.31)
% of Medicaid beneficiaries who are White	0.022 (-0.203)	0.036 (-0.2)	0.213 (-0.139)
% of Medicaid LTSS beneficiaries in home and community-based services	0.073* (-0.034)	0.066* (-0.031)	0.061* (-0.03)
County population	0.000*** (0)	0.000*** (0)	0.000*** (0)
% of votes for Republican candidates	-0.08 (-0.059)	-0.086 (-0.058)	-0.035 (-0.04)
% of residents in poverty	0.02 (-0.113)	0.04 (-0.12)	0.054 (-0.095)
Median household income	0 (0)	0 (0)	0 (0)
N	2238		

Standard errors are clustered by state and are in parentheses. Coefficients represent change from 2012 to 2020.

*** p<.001, ** p<.01, * p<.05, + p<0.10

Table 2.6. First Differences Regression Results Interacting with Medicaid Generosity

	Automation	Med professional involved	Combined typology
	% of nursing home residents who are non-White		
<i>Automation</i>			
Fully automated	Reference	-	-
Automated/approved	0.008 (-0.019)	-	-
Decision-making aid	0.014 (-0.018)	-	-
Medical professional involved in eligibility determination	-	0.036* (-0.017)	-
<i>Typology</i>			
Medical automated	-	-	Reference
Medical approved	-	-	-0.044* (-0.019)
Medical aid	-	-	-0.024*** (-0.005)
Non-medical automated	-	-	-0.072*** (-0.006)
Non-medical approved	-	-	-0.069*** (-0.013)
Non-medical aid	-	-	0.080+ (-0.043)
Medicaid generosity	-0.007 (-0.031)	-0.023 (-0.028)	0.004 (-0.027)
<i>Interactions</i>			
Fully automated x generosity	Reference	-	-
Automated/approved x generosity	-0.003 (-0.003)	-	-
Decision-making aid x generosity	-0.004 (-0.003)	-	-
Medical professional x generosity	-	-0.005** (-0.002)	-
Medical automated x generosity	-	-	Reference
Medical approved x generosity	-	-	0.002 (-0.002)
Medical aid x generosity	-	-	-0.003 (-0.002)
Non-medical automated x generosity	-	-	0.007*** (-0.002)
Non-medical approved x generosity	-	-	0.001 (-0.003)
Non-medical aid x generosity	-	-	-0.014* (-0.006)
<i>Controls</i>			
% of population that is White	-0.600+ (-0.341)	-0.814** (-0.293)	-0.766** (-0.249)
% of Medicaid beneficiaries who are White	0.149 (-0.175)	0.195 (-0.155)	0.483** (-0.144)
% of Medicaid LTSS beneficiaries in home and community-based services	0.064+ (-0.032)	0.047 (-0.028)	0.036 (-0.022)
County population	0.000** (0)	0.000** (0)	0.000** (0)
% of votes for Republican candidates	-0.085 (-0.051)	-0.074* (-0.032)	-0.033 (-0.026)
% of residents in poverty	-0.008 (-0.103)	0.022 (-0.102)	0.028 (-0.089)
Median household income	0 (0)	0 (0)	0 (0)

N

2238

Standard errors are clustered by state and are in parentheses. Coefficients represent change from 2012 to 2020.

*** $p < .001$, ** $p < .01$, * $p < .05$, + $p < 0.10$

Figures

Figure 2.1. Coefficient plot of the change in the percentage of non-White nursing home residents with each one percentage point change from fully automated algorithms requiring medical professionals to other algorithm types

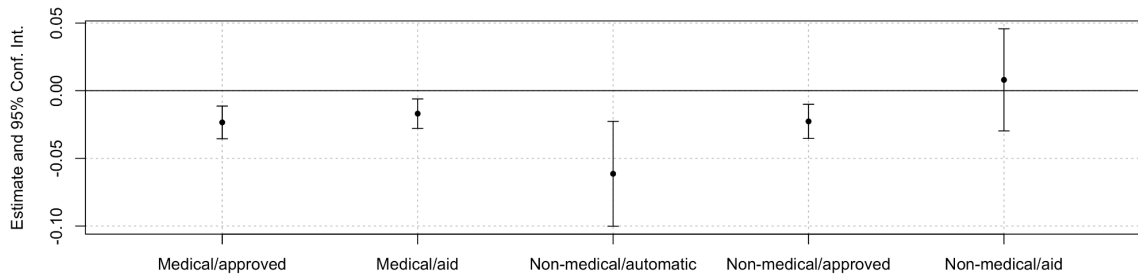


Figure 2.2a. Effects of changes from medical/automated algorithms to medical/approved and medical/aid on the percentage of non-White nursing home residents across levels of Medicaid generosity

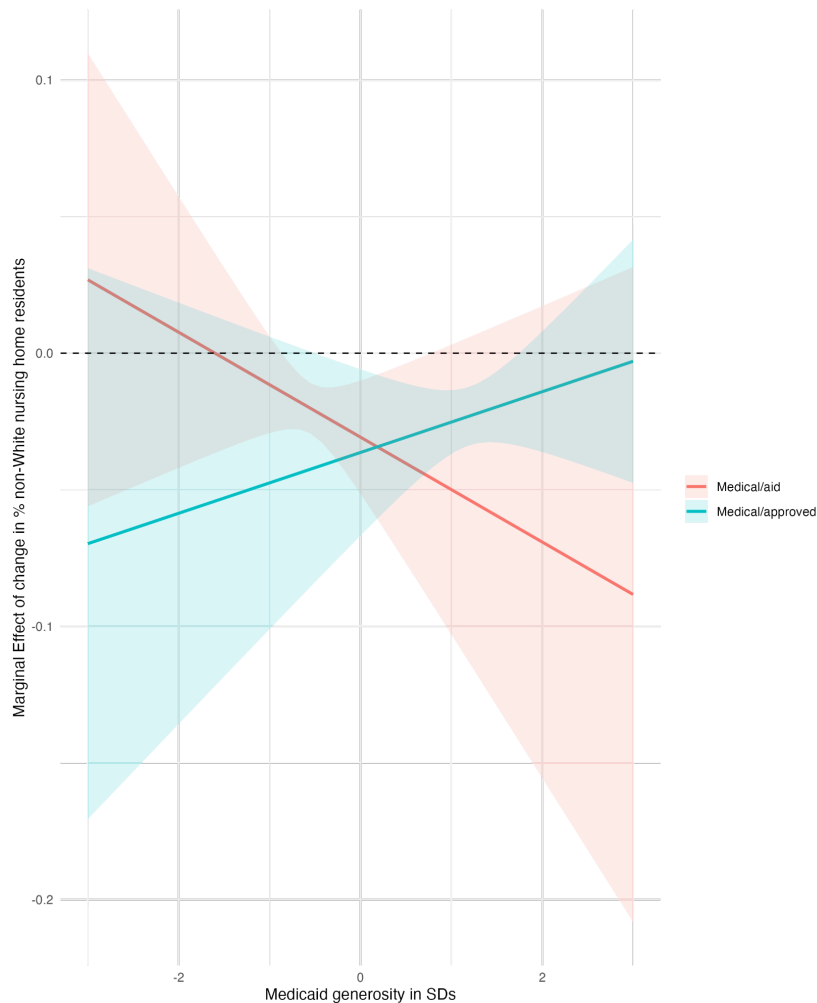
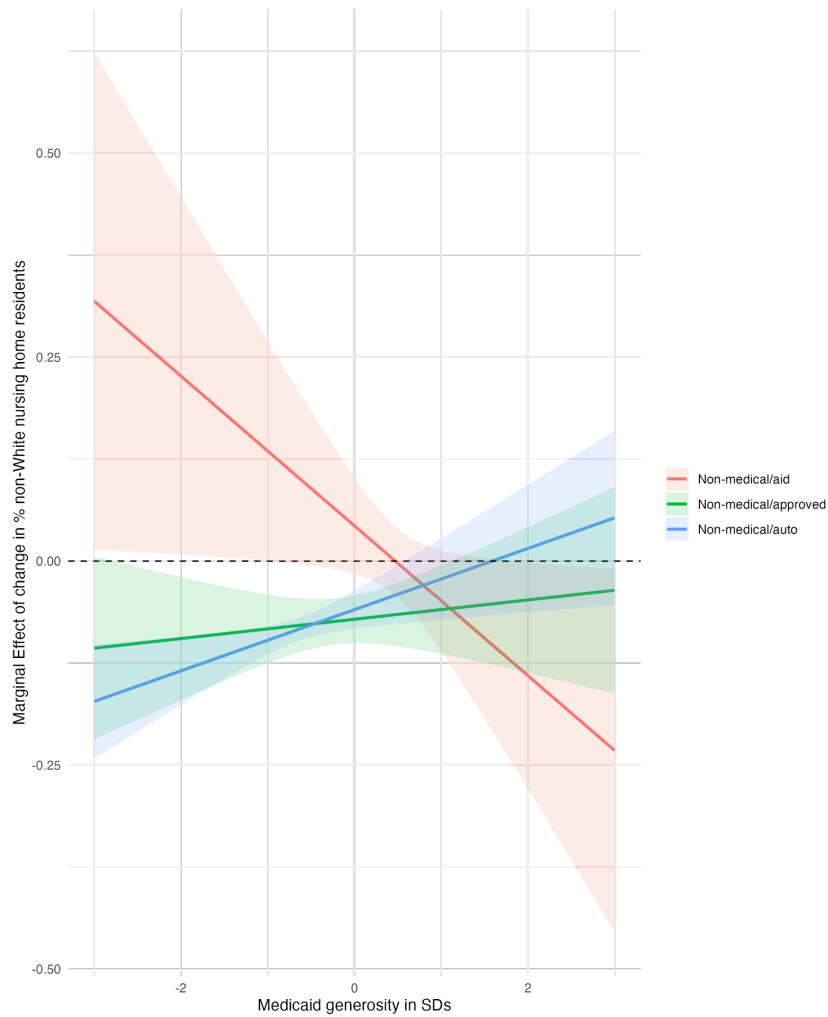


Figure 2.2b. Effects of changes from medical/automated algorithms to non-medical types on the percentage of non-White nursing home residents across levels of Medicaid generosity



Appendix

Table A. First Differences Regression Results for Counties with a High Percentage of Non-White Residents

	Automation	Med professional involved	Combined typology
% of nursing home residents who are non-White			
<i>Automation</i>			
Fully automated	Reference	-	-
Automated/approved	-0.01 (-0.017)	-	-
Decision-making aid	0.034 (-0.024)	-	-
Medical professional involved in eligibility determination	-	0.022 (-0.03)	-
<i>Typology</i>			
Medical automated	-	-	Reference
Medical approved	-	-	-0.053*** (-0.009)
Medical aid	-	-	0 (-0.01)
Non-medical automated	-	-	-0.112*** (-0.02)
Non-medical approved	-	-	-0.052** (-0.015)
Non-medical aid	-	-	0.065* (-0.031)
<i>Controls</i>			
% of population that is White	0.122 (-0.601)	-0.207 (-0.637)	0.189 (-0.487)
% of Medicaid beneficiaries who are White	-0.457 (-0.452)	-0.248 (-0.43)	0.184 (-0.231)
% of Medicaid LTSS beneficiaries in home and community-based services	0.153** (-0.052)	0.094+ (-0.047)	0.165*** (-0.044)
County population	0.000** (0)	0.000** (0)	0.000** (0)
% of votes for Republican candidates	0.003 (-0.082)	-0.064 (-0.09)	0.123+ (-0.064)
% of residents in poverty	0.021 (-0.137)	0.1 (-0.151)	-0.015 (-0.11)
Median household income	0 (0)	0 (0)	0 (0)
N	1154		

Standard errors are clustered by state and are in parentheses. Coefficients represent change from 2012 to 2020.

*** p<.001, ** p<.01, * p<.05, + p<.10

Table B. First Differences Regression Results Interacting with Medicaid Generosity for Counties with a High Percentage of Non-White Residents

	Automation	Med professional involved	Combined typology
	% of nursing home residents who are non-White		
<i>Automation</i>			
Fully automated	Reference	-	-
Automated/approved	-0.008 (-0.025)	-	-
Decision-making aid	0.032 (-0.024)	-	-
Medical professional involved in eligibility determination	-	0.058+ (-0.032)	-
<i>Typology</i>			
Medical automated	-	-	Reference
Medical approved	-	-	-0.077+ (-0.041)
Medical aid	-	-	-0.018 (-0.011)
Non-medical automated	-	-	-0.128*** (-0.015)
Non-medical approved	-	-	-0.134*** (-0.016)
Non-medical aid	-	-	0.201*** (-0.044)
Medicaid generosity	-0.003 (-0.117)	-0.027 (-0.092)	0.051 (-0.096)
<i>Interactions</i>			
Fully automated x generosity	Reference	-	-
Automated/approved x generosity	-0.002 (-0.003)	-	-
Decision-making aid x generosity	-0.004 (-0.003)	-	-
Medical professional x generosity	-	-0.008** (-0.002)	-
Medical automated x generosity	-	-	Reference
Medical approved x generosity	-	-	0.002 (-0.004)
Medical aid x generosity	-	-	-0.006 (-0.005)
Non-medical automated x generosity	-	-	0.002 (-0.004)
Non-medical approved x generosity	-	-	0.002 (-0.005)
Non-medical aid x generosity	-	-	-0.030** (-0.011)
<i>Controls</i>			
% of population that is White	-0.19 (-0.529)	-0.818+ (-0.447)	-0.402 (-0.436)
% of Medicaid beneficiaries who are White	-0.272 (-0.435)	0.138 (-0.337)	0.705* (-0.282)
% of Medicaid LTSS beneficiaries in home and community-based services	0.134* (-0.054)	0.068 (-0.045)	0.104* (-0.04)
County population	0.000** (0)	0.000** (0)	0.000** (0)
% of votes for Republican candidates	-0.024 (-0.093)	-0.065 (-0.069)	0.09 (-0.076)
% of residents in poverty	0.003 (-0.126)	0.029 (-0.125)	0.006 (-0.111)

Median household income	0 (0)	0 (0)	0 (0)
N	1154		

Standard errors are clustered by state and are in parentheses. Coefficients represent change from 2012 to 2020.
 *** p<.001, ** p<.01, * p<.05, + p<0.10

Table C1. First Differences Regression Results for States with Changes to Medicaid Generosity Greater than the Median

	Automation	Med professional involved	Combined typology
% of nursing home residents who are non-White			
<i>Automation</i>			
Fully automated	Reference	-	-
Automated/approved	-0.013 (-0.012)	-	-
Decision-making aid	-0.003 (-0.016)	-	-
Medical professional involved in eligibility determination	-	-0.015 (-0.009)	-
<i>Typology</i>			
Medical automated	-	-	Reference
Medical approved	-	-	-0.022** (-0.008)
Medical aid	-	-	-0.006 (-0.012)
Non-medical automated	-	-	-0.012 (-0.015)
Non-medical approved	-	-	-0.020+ (-0.011)
Non-medical aid	-	-	0.029 (-0.019)
<i>Controls</i>			
% of population that is White	-0.926** (-0.241)	-1.188*** (-0.278)	-0.994*** (-0.192)
% of Medicaid beneficiaries who are White	0.24 (-0.271)	0.241 (-0.164)	0.224 (-0.164)
% of Medicaid LTSS beneficiaries in home and community-based services	0.113* (-0.04)	0.096*** (-0.024)	0.138*** (-0.026)
County population	0.000** (0)	0.000** (0)	0.000** (0)
% of votes for Republican candidates	-0.027 (-0.025)	-0.071* (-0.03)	-0.03 (-0.025)
% of residents in poverty	0.018 (-0.112)	0.102 (-0.094)	0.059 (-0.101)
Median household income	0 (0)	0 (0)	0 (0)
N	1119		

Standard errors are clustered by state and are in parentheses. Coefficients represent change from 2012 to 2020. The median change in the Medicaid generosity index from 2012 to 2020 is 2.5.
 *** p<.001, ** p<.01, * p<.05, + p<0.10

Table C2. First Differences Regression Results for States with Changes to Medicaid Generosity Less than or Equal to the Median

	Automation	Med professional involved	Combined typology
% of nursing home residents who are non-White			
<i>Automation</i>			
Fully automated	Reference	-	-
Automated/approved	0.089* (-0.039)	-	-
Decision-making aid	-0.013 (-0.015)	-	-
Medical professional involved in eligibility determination	-	0.040* (-0.017)	-
<i>Typology</i>			
Medical automated	-	-	Reference
Medical approved	-	-	0.167* (-0.068)
Medical aid	-	-	-0.033** (-0.009)
Non-medical automated	-	-	-0.052** (-0.015)
Non-medical approved	-	-	-0.009 (-0.028)
Non-medical aid	-	-	-0.077*** (-0.016)
<i>Controls</i>			
% of population that is White	0.887 (-0.704)	-0.388 (-0.464)	0.333 (-0.43)
% of Medicaid beneficiaries who are White	0.035 (-0.239)	0.312 (-0.204)	0.116 (-0.155)
% of Medicaid LTSS beneficiaries in home and community-based services	0.017 (-0.031)	0.01 (-0.02)	-0.008 (-0.016)
County population	0.000** (0)	0.000** (0)	0.000** (0)
% of votes for Republican candidates	-0.458** (-0.144)	-0.284* (-0.122)	-0.354** (-0.101)
% of residents in poverty	-0.023 (-0.181)	-0.039 (-0.195)	0.058 (-0.163)
Median household income	0 (0)	0 (0)	0 (0)
N	1119		

Standard errors are clustered by state and are in parentheses. Coefficients represent change from 2012 to 2020. The median change in the Medicaid generosity index from 2012 to 2020 is 2.5.

*** p<.001, ** p<.01, * p<.05, + p<.10

CHAPTER 3

100% Objective to Some Degree: Organizational Context and Inequality in a Medicaid Eligibility Algorithm

Introduction

Scholars have argued that standardized, algorithmic tools do not accurately or fully capture the contours of human experience (Cruz 2022; Star and Bowker 2007). This can have important consequences when algorithms are used to allocate resources such as public benefits. However, algorithms do not operate independently. Rather, human actors engage with decision-making algorithms during their development, implementation, and during debates about their results and could shape how these tools measure their subjects.

Current research on algorithmic decision-making “in practice” (Christin 2017) has shown how humans, influenced by their professional identities and organizational contexts, implement algorithms in ways that are decoupled from their original purpose. Another largely separate strand of work shows how algorithms exacerbate inequality because they use data that reflects entrenched inequality and because the actors who implement or otherwise interact with them are biased (Eubanks 2017; Benjamin 2019; Obermeyer et al. 2019).

Despite this, there is limited understanding of how inequality-producing mechanisms vary across contexts for the same algorithm. Current work primarily examines how algorithms are used in a single context, and largely at the stage of output interpretation. To my knowledge, there is no research examining how different actors engage with the same algorithm at different points of the resource allocation process.

I use the case of Medicaid Long Term Services and Supports (LTSS) in Wisconsin to examine how a functional eligibility algorithm, which is used to classify applicants as disabled or not, is implemented in practice at two distinct points: eligibility screening and

appeal hearings. I ask: how do different groups of actors implement and adjudicate a functional eligibility algorithm? How does this vary across different groups of applicants? Based on 19 key informant interviews with eligibility screeners and other state personnel and analysis of 235 court appeals documents, I find that Medicaid applicants with particular sources of cultural capital are advantaged throughout the functional eligibility process. However, resources like health literacy and program knowledge are more important during the initial screening process, during which applicants must signal their limitations in very precise ways to be legible to the eligibility algorithm. During appeals, cultural resources are potentially less important because judges' organizational position and institutional logics largely allow them to disregard the algorithm in favor of broader state regulations.

By examining variation *within* the same algorithmic decision-making process but *across* organizational contexts, this project begins to untangle the separate but linked roles of individual bias and meso-level structures in reproducing inequality in algorithmic resource allocation.

Policy Context

Medicaid LTSS is a critical source of long-term healthcare for the elderly and people with disabilities in the United States. In 2022, Medicaid paid for over 60% of all spending (over \$400 billion) on LTSS in the United States (Chidambaram and Burns 2024). In 2019, over 10% of all Medicaid beneficiaries received some type of LTSS through the program, totaling almost 8.8 million users (Kim et al. 2022).

Long-term care is expensive - often prohibitively so. For those who must pay out-of-pocket for long-term care, the average annual cost for a nursing home is over \$100,000. For those with intensive healthcare needs who prefer to stay in their homes, an around-the-clock home health aide averages almost \$300,000 a year (Chidambaram and Burns 2024). As the population ages (Vespa et al. 2020), an increasing number of U.S. residents will need long-term care in the coming decades.

In all states, applicants to Medicaid LTSS must be deemed functionally eligible before they receive services. In other words, they must have qualifying long-term health conditions or disabilities. While previously states might have relied solely on physicians to decide whether a LTSS applicant is functionally eligible, it is now commonplace for states to use algorithms to determine functional eligibility and how much care and what kind of care they receive (Medicaid and CHIP Payment Access Commission 2016).

There are no federal requirements regarding how functional eligibility is determined, and a 2016 report estimated that across U.S. states, over 120 different tools are used to measure functional eligibility (Ibid.).

In short, Medicaid is a critical source of LTSS in the United States, and given astronomical costs of both institutional and in-home care, the only funding option for many people. However, functional eligibility requirements add an additional layer of complexity to the program, requiring applicants to prove their medical need according to whichever metric their state of residence has implemented – or risk losing care altogether.

Theoretical Background

Algorithmic Decision-Making in Organizations

Scholar Angèle Christin describes algorithms as “computational procedures (which can be more or less complex) drawing on some type of digital data (“big” or not) that provide some kind of quantitative output (be it a single score or multiple metrics) through a software program” (2017, p. 2). By this definition, algorithms include anything from simple decision trees to computationally complex tools that use machine learning to generate predictions.

Many algorithms are classification tools: they automatically and efficiently sort data into categories. When applied to humans, classification can determine one’s life chances, access to resources, and well-being. For instance, Cecilia Menjívar shows how being categorized as an undocumented immigrant in the United States (as opposed to an ‘asylum seeker’) restricts access to jobs and other social benefits (2002; 2023). Credit scores provide another example of a fine-grained classification system that grants or prohibits access to various economic and social opportunities (Fourcade and Healy 2013). These scores determine whether people can get loans, purchase real estate, and generally cement class stratification (Ibid.). Once made official, the categories assigned to people become “real” and are given legal power (Porter 1995; Scott 2020).

Algorithmic resource allocation. Organizations and governments increasingly use algorithms to make decisions and allocate resources in part because, as proponents argue, algorithms are more objective than humans; they don’t have personal stakes in decisions and subject all data to the same code. In many cases, organizations explicitly state that

they use algorithms to increase objectivity or reduce bias in decision-making. For example, the Wisconsin state health department describes their eligibility algorithm as “an automated and objective way to determine the long-term care needs of elders and people with physical or intellectual/developmental disabilities” (Wisconsin Department of Health Services 2024).

The use of algorithms to determine public benefits eligibility reflects a general trend of “bureaucratic disentanglement” where bureaucrats defer to procedural mechanisms to avoid responsibility for resource allocation decisions (Lipsky 1984). The use of algorithms as more ‘objective’ means of decision-making can legitimize government organizations, algorithm users, and the decisions made by algorithms (Porter 1995). As public concerns about implicit bias and inequality become more prevalent, the perceived objectivity of algorithms lends authority to classification decisions.

Bureaucrats and the implementation of algorithms. Despite appearing objective, however, bureaucrats still have a hand in shaping and implementing eligibility algorithms. Though higher-level public officials and policymakers typically define eligibility categories, *how* people are classified in practice occurs almost entirely among local, on-the-ground workers (Lipsky 2010). Street-level bureaucrats have substantial discretion in enacting policy. That is, they have flexibility in determining how to implement policy and what types of resources to provide citizens (Ibid.; Maynard-Moody and Musheno 2000; Sandfort 2000; Tummers and Bekkers 2014).

Lipsky (2010) argued that managers regulate the use of discretion among street-level bureaucrats. Today, this responsibility is often delegated to or supplemented by

algorithms (Buffat 2015). However, current scholarship highlights how algorithms are often decoupled from their stated purpose (Christin 2017). Algorithms rarely operate on their own but are shaped within work contexts by the people interacting with them. For example, Brayne (2017) shows how police officers defer to their own knowledge about crime patterns even when it conflicts with algorithmic predictions. Strategies of algorithmic resistance (or adoption) also vary across work contexts; Brayne and Christin (2021) show that legal professionals in criminal courts are better able to resist predictive algorithms than police officers who are often subject to stricter hierarchical oversight. To that end, Christin (2017) calls for work on “algorithms in practice” which documents the specific contexts and practices that govern how algorithms are used.

One potential reason for algorithmic decoupling is that the humans who implement and interpret algorithms don’t necessarily share the motivations of their organization. Rather, professionals are guided both by organizational priorities and their own motivations (Maynard-Moody and Musheno 2000). For instance, Donatz-Fest (2024) shows how police officers in the Netherlands sometimes avoid entering interactions with citizens into a crime database if the citizens are cooperative or have no prior offenses. In this case, police officers shape how crime data is collected in ways that are likely antithetical to their organization’s purpose.

Institutional logics. The organizational priorities guiding individual work have been theorized as institutional logics, which Thornton et al. (2012) define as “frames of reference that condition actors’ choices for sensemaking” (p. 2). The theory of institutional logics argues that individual action within organizations is in part shaped by (and shapes) norms and practices at meso and macro levels. In the context of algorithmic

decision-making, Esthappen (2024) finds that criminal judges use algorithmic risk scores strategically to navigate competing bureaucratic, legal, and political logics. For example, judges often followed risk scores to legitimate their decisions and avoid public criticism, but also occasionally ignored algorithmic recommendations when they conflicted with their understanding of legally-defined risk.

While some research examines how different algorithms are used within related contexts (such as how various predictive algorithms are used within the criminal justice system; Brayne and Christin 2021), to my knowledge, there is no research examining how different actors engage with the *same* algorithm at different points. This is especially important when algorithmically-produced outcomes made at one stage can later be overturned or otherwise altered by actors in different organizational contexts. This project adds to this growing body of knowledge by examining algorithmic decision-making as an unfolding process that involves multiple actors within multiple organizational contexts.

This paper examines how institutional logics shape how actors engage with a functional eligibility algorithm in Wisconsin's Medicaid LTSS program at two different points – data collection (by eligibility screeners) and adjudication (by judges). First, I explore how eligibility screeners use the algorithmic screening tool to classify applicants as disabled or not. Second, I examine how administrative judges adjudicate the results of the functional eligibility algorithm in appeal cases.

Deservingness and Cultural Inequality

Classification, especially when used to allocate public benefits, is deeply intertwined with conceptions of morality and deservingness (Hasenfeld 2000; Schneider and Ingram

1993). The social and political construction of various groups as deserving or not facilitates program eligibility and resource allocation (Schneider and Ingram 1993). Though the social construction of populations varies over time and place, the sick, elderly, and disabled are typically considered the most deserving groups, while the poor (especially those who are unemployed) are often considered undeserving of government assistance (Watkins-Hayes and Kovalsky 2017). There are differences in perceptions of deservingness among the sick and disabled too, as research suggests that mental illnesses are among the most stigmatized disabilities (Best & Arseniev-Koehler 2023).

The role of deservingness rhetoric on algorithmic classification is not well understood, perhaps because algorithms often disguise the presence of human discretion and morality. In theory, algorithms should reduce the effect that individual perceptions of an applicant's deservingness have on eligibility. However, the theory of decoupling and evidence that human bias permeates algorithmic decision-making suggests that *how* humans implement algorithms may reflect their views of deservingness.

The perceived deservingness of an applicant may also relate to applicants' ability to marshal certain resources to signal their disabilities to screeners and judges. For example, applicants with health literacy and the ability to advocate for themselves may be better able to provide evidence of their disabilities in ways that align with the categories measured by the eligibility algorithm.

Resources associated with cultural capital can shape whether screeners and judges view an applicant as deserving of long-term care. As theorized by Pierre Bourdieu, cultural capital in its embodied form refers to a person's disposition, way of communicating and

presenting themselves, and knowledge of social practices (1986). Janet Shim highlights the importance of cultural capital specifically in the context of healthcare, defining what she terms “cultural health capital” (2010). Cultural health capital as a concept captures “certain socially-transmitted and differentially distributed skills and resources [that are] critical to the ability to effectively engage and communicate with clinical providers” (Shim 2010, p. 1). Cultural health capital can include skills like health literacy, self-advocacy in healthcare settings, and the ability to “communicate health-related information to providers in a medically intelligible and efficient manner” (Ibid., p. 4).

This paper investigates how screeners and judges use the eligibility algorithm to characterize applicants as deserving of publicly-funded LTSS or not. In doing so, I show how the effects of algorithmic decision-making depend on an applicant’s resources. I find that Wisconsin’s algorithmic functional eligibility process disproportionately advantages applicants who are better able to wield cultural health capital to signal their disabilities. However, I also find that the value of an applicant’s cultural health capital varies between screening and appeals, particularly according to the institutional logics that guide decision-making.

Case: Medicaid Long-Term Care in Wisconsin

Wisconsin’s Medicaid LTSS program is moderately generous. There is no wait list for services, and the state spends about \$705 a year per resident on LTSS (as compared to \$678 nationwide; Murray et al. 2023). The state primarily provides Medicaid LTSS

services through a program called Family Care¹⁰, which is administered by the state Department of Health Services (DHS). To establish functional eligibility for Family Care, the state contracts with county-level screening agencies, who employ social workers (“screeners”) to conduct assessments using an 18-page questionnaire called the Long-Term Care Functional Screen (hereafter the “screen”). This tool collects information about the applicants’ medical diagnoses, functional limitations (Activities of Daily Living [ADLs] and Instrumental Activities of Daily Living [IADLs]), and health service requirements. The screen is accompanied by a 149-page instructions document which details how to complete the screen and attempts to account for various scenarios a screener might face. For example, the instructions list the following illustrative examples for the scenario when an applicant is considered to need help with the bathing ADL:

- *Due to a physical, cognitive, or memory loss impairment requires supervision, cueing, and/or partial or complete hands-on assistance with ALL the above components of Bathing and another person needs to be present throughout the task.*
- *Bathes independently but doing so results in a significant, negative health outcome and another person should be present to help with ALL the components of the task.*
- *Uses an improvised or homemade item and without it, they would need assistance from another person to complete ALL the components of Bathing.*
- *Requires assistance with ALL the components of Bathing but he or she can be left alone to soak in the bathtub (without negative health and/or safety concerns). Soaking in the bathtub is not a component of the Bathing ADL.*

¹⁰ Wisconsin also has long-term care programs called IRIS (“Include, Respect, I Self-direct”), “Family Care Partnership”, and PACE (“Program of All-Inclusive Care for the Elderly”). IRIS participants manage their own care within a predetermined budget. Partnership and PACE are administered under the umbrella of Family Care and include a more expansive suite of care, including coordination across Medicaid and Medicare services. Functional eligibility is determined the same way for all of these programs.

- *Requires supervision due to having an uncontrolled seizure disorder, evidenced by one or more seizures in the last three months AND requires standby assistance during the entire task of Bathing.*

(Division of Medicaid Services 2024, p. 5-8).

After completing the in-person assessment, the screener reviews the applicant's medical records and may follow up with the applicant's caregivers, physicians, or family members if there are remaining questions. The screener then inputs this information into an algorithm, which determines eligibility.

There are three functional eligibility thresholds in Wisconsin: the "nursing home" level of care, the "non-nursing home" level of care, and ineligible. Applicants who meet the nursing home level of care are eligible for the full suite of long-term care benefits. Those who meet the non-nursing home level of care are only eligible for case coordination services. Ineligible applicants cannot receive Medicaid funding for any long-term care services. This process is outlined in Figure 3.1.

Applicants who meet either level of care (in addition to financial eligibility requirements) can enroll in Family Care. Enrollees join a managed care organization (MCO) which coordinates their care. All LTSS enrollees must undergo another screen at least once a year, or more if their care teams document a change in condition. A subsequent screen can result in level of care changes or disenrollment from the program.

If an applicant's care is denied, reduced, or ended, they have the right to appeal through a Fair Hearing. During this process, a judge employed by the state Department of Administration hears the case and makes a ruling. The judge will either dismiss the case or require the organization to provide whatever benefit was under debate.

Figure 3.2 provides a simplified overview of the organizations and key actors involved in Medicaid LTSS functional eligibility determination. The key departments at the state level are DHS which administers the state Medicaid program, and the Department of Administration, which employs the judges who hear appeal cases. Within DHS, administrators oversee initial and ongoing programming and maintenance of the functional eligibility algorithm, and oversee the work of the eligibility screeners at local screening agencies. The work of both departments theoretically follow the eligibility requirements outlined in state administrative code (Appendix A). However, as I will discuss in more detail in the results section, while appeal judges follow the administrative code very closely in their decisions, the current functional eligibility algorithm uses different standards.

Wisconsin is an empirically useful case for several reasons. First, compared to other states, Wisconsin's functional screen is quite long with extensive and nuanced instructions for screeners. The incredibly exacting standardization of Wisconsin's functional eligibility process positions the state as a best case scenario in which we might expect the least amount of subjectivity in screening.

Second, the overall functional eligibility process is fairly straightforward. Applicants only undergo one assessment and the same tool is used for all Medicaid LTSS programs. This is not the case in all states. Many have different eligibility thresholds and assessment tools for different long-term care programs (Medicaid and CHIP Payment Access Commission 2016). Some states with waitlists for services may have one functional assessment to get on the waitlist, and another to actually receive care (like Florida). The relative simplicity of Wisconsin's eligibility determination allows me to better isolate the

effects of organizational context as I am examining the same algorithm in the same program.

Methods

This paper combines insights from 19 in-depth key informant interviews with an in-depth analysis of 235 Fair Hearing appeals. These sources allow me to examine how key actors engage with the eligibility algorithm in two distinct organizational contexts.

Interviews

From February 2023 to February 2024, I conducted semi-structured in-depth Zoom interviews with 18 Medicaid LTSS functional eligibility screeners (n=7), their supervisors (all of whom had previously been screeners; n=4), state administrators (n=4), and LTSS ombudsmen (n=3). The interviews averaged 1 hour. Most of the screeners and supervisors I interviewed were employed by county agencies across the state, where functional eligibility is initially determined. These interviewees represent about 12% of the 65 total county and tribal screening agencies. Two of those interviewees were employed by MCOs, where LTSS beneficiaries are re-screened at least once every 12 months to determine ongoing eligibility.

The state administrators I interviewed are employed by DHS, and oversee everything related to the screen and the eligibility algorithm, including maintaining and updating the screen instructions and eligibility logic. Finally, the ombudsmen are state-employed

advocates for LTSS beneficiaries ages 60 and over. They accompany applicants¹¹ to appeal hearings, mediate problems between beneficiaries and their care teams, and otherwise help applicants navigate the program. See Table 3.1 for a summary of interviewees.

Finally, In May 2025, I interviewed a Wisconsin appeal judge, David. This interview lasted for 33 minutes. Interviewees were not compensated.

I contacted potential interviewees either directly if their email was publicly available, or indirectly by emailing their organization's general inbox. I asked for referrals to other interviewees at the end of interviews but most interviews resulted from direct contact. I purposively sampled to get interviewees from both rural and urban counties, as applicants' access to resources may vary across geographic settings. I also interviewed one screener who works for a Tribal Nation in Wisconsin and exclusively screens enrolled tribe members in a very rural part of the state. To maintain confidentiality, I do not specify which of my interviewees is a tribal screener.

Fair Hearing Appeals

To gain insight into Fair Hearing appeals, I analyze 235 publicly available records of Medicaid LTSS appeals from 2020-2023 related to functional eligibility or level of care.

Wisconsin publishes records of all appeals online. These records include a detailed description of the issue under debate, a timeline of events, the arguments of both parties,

¹¹ This applies to both people already receiving services (beneficiaries) and those applying for services for the first time (applicants). I use the term "applicant" throughout for simplicity, except where findings relate specifically to those already receiving care.

and the judge’s reasoning for whatever ruling they make. The records also occasionally include information about the beneficiary’s age and diagnoses if relevant. Potentially identifiable information about the beneficiary is redacted. In appeal records, the applicant who makes the appeal is referred to as the “petitioner”, and the county agency, MCO, or state health department is referred to as the “respondent.”

I downloaded all appeals related to Medicaid LTSS, and hand-coded each appeal for the result (dismissed or awarded), whether it was related to functional status, and whether it directly challenged a screen result. I separate out appeals related to functional status because these show how judges adjudicate issues of “medical necessity”. Functional status appeals include debates about functional eligibility, level of care, and the medical necessity of various services requested by the petitioner. I create an additional variable indicating whether the appeal directly challenges a functional screen result. Appeals in the “other” category include those debating financial eligibility, those where the petitioner requested that their start date be backdated, those discussing whether a certain service was legally covered by the program, or anything else not related to functional status. I do not do an in-depth analysis of appeals in the “other group” because they do not provide insight into functional eligibility determination.

Table 3.2 shows a summary of the appeals. During the COVID public health emergency which lasted from March 2020 to May 2023, Medicaid LTSS recipients were not allowed to be disenrolled from the program due to loss of functional eligibility. The lower numbers of appeals in 2021 and 2022 in part reflect this. However, the state could still find first-time applicants ineligible for the program, could reduce the amount of care a

recipient received, and could deny requests for additional services. A majority of appeals across the 4 years were dismissed, with only 27% awarded in favor of the petitioner.

When broken down by appeal reason, a higher percentage (37% vs. 20%) of appeals related to functional status were awarded than those appealed for other reasons. This might reflect the more subjective nature of functional status as compared to financial status, for instance. An even higher number of appeals that challenged a screen result were awarded in favor of the petitioner (44%), partially driven by cases where the applicant was eligible based on the state administrative code but not based on the algorithm. All of these cases were awarded.

Results

In the following sections, I describe the logics that guide how eligibility screeners and administrative law judges use the functional eligibility algorithm. Next, I describe how these different logics may disproportionately benefit applicants with particular forms of cultural health capital.

Medical Eligibility Screening: Objectivity and Deservingness

Within screening agencies, screeners are subject to the logic of objectivity. In this case, “objectivity” primarily concerns the ability of the functional screen algorithm to remove individual screener bias from the eligibility determination process. Materials and guidance from DHS, which oversees all eligibility screening (Figure 3.2), establish this logic within screening agencies. For example, the functional screen instructions state that “The [screen] is designed to be as *objective* as possible to achieve the highest inter-rater reliability (two screeners would answer the same way for an individual), to ensure fair

and proper functional eligibility determinations, and provide statewide consistency” (emphasis added; Division of Medicaid Services 2024, p. 1-1).

Members of the four-person “screen team” at DHS – the administrators responsible for maintaining and updating the screen and algorithm – often invoked objectivity when describing their work. One member of this team, Stacey, told me the following:

Sometimes we look at our instructions, and we're like 'ooooo we see why the screener did it that way.' You know, there's some flexibility, or there's some grayness in how our instructions are written, and we need to clarify that. So a lot of what we do is through those desk reviews. We have a lot of discussion with each other, and then we add it to a list to attempt to clarify it with our instruction updates.

This illustrates one of the main roles of DHS administrators – to stamp out any opportunity for discretion in the screening tool and make the process as objective as possible. The organization’s focus on inter-rater reliability also reflects this logic. Inter-rater reliability, or how similarly two screeners rate the same applicant, is a key measure DHS uses to evaluate the screen. As Shannon, another member of the screen team told me, one of her team’s main responsibilities is to “make sure that all of our screeners are interpreting our instructions the same way and the way that we intend them to be interpreted.” Objectivity, established as a priority within DHS, is enforced upon screeners as the measure by which their work is evaluated.

At the same time, screeners’ personal beliefs about who deserves state-funded long-term care shaped eligibility screening. Indeed, many of my interviewees acknowledged the role of their own discretion in determining eligibility outcomes. For example, screener Jean told me:

I think that [the screen] is intended to be a hundred percent objective. And you know it is, I guess, to some degree. [...] I know the intent is that it's objective. and I know in practice... for me, anyways, there's some areas that are not a hundred percent. I mean you just gotta decide where you're going to put a person.

Here, Jean seemed torn between the ideal of the screen as “a hundred percent objective” and acknowledging her own role in producing the eligibility outcome. To her, discretion is unavoidable even though this seemingly conflicts with the stated goal of objectivity.

Screeners Kylie also pointed to the role of discretion in eligibility screens, telling me, “We have like total discretion. [...] It definitely sounds easier than it is because basically this job is entirely in the gray area and my perception might be different than somebody else’s.”

Importantly, how screeners use their discretion varies, generally favoring applicants who screeners deem as more deserving of care. Across my interviews, screeners tended to view the elderly and physically disabled as the most deserving, and those with stigmatized or invisible disabilities, especially mental health conditions, as the least deserving.

For an example of how screeners differentially construct disabilities as stigmatized and undeserving of care, see Julia’s description of elderly applicants:

There's the very typical older person, like 78 and above. 78 to 100 that lived their whole normal life, had a job, but they've used their nest egg. You know, they're sort of a typical person who now is in the frail elder category. They need help with lots of things, but it could be from COPD or dementia, or peripheral neuropathy, or you know, whatever it is.

Contrast this with Julia’s description of mentally ill applicants:

There's a huge subset of people that have mental illness that want to be screened because they need help with chores. And they're not eligible. [...] They are 30 to 50 and they have depression or anxiety or personality disorder, or all 3 of those. And they are managing, but they're struggling in life. Life kinda sucks for them because they don't have family. Maybe they have alcohol issues. They burned every bridge. [...] This long term care program is not designed for people with [mental illness]. [...] The world would go bankrupt if every single person was enrolled because they didn't feel like doing their laundry anymore.

Julia frames elders as “normal” and “typical” people who contributed to society until they eventually needed LTSS, through no fault of their own. This starkly contrasts with how Julia describes mentally ill applicants to the program. These people are responsible for their bad situation, who don’t want to do chores, who aren’t deserving of publicly funded LTSS.

Perceptions of applicants as “deserving” or not matter for eligibility outcomes, as screeners find ways to influence assessments according to their personal motivations. For instance, several screeners reported “coaching” applicants to answer questions in particular ways to increase their likelihood of eligibility. The following quotes from Lily and Lindsey exemplify this:

Lily: I'm probably a little more biased in favor of the person. If they need the help, then we're gonna make sure we have conversations in a way to make sure that they get the help. [...] That's how I am. I want to provide everybody with whatever opportunities that they can have.

Lindsey: If I think somebody should qualify... I try to be very aware of what my screen instructions are and what the manual says I can and can't mark. So I might ask a question specifically knowing if they say yes, I can mark it. I'm not like trying to lead them or coerce them into lying. [...] I'm just making sure that I am asking enough questions so that if they do qualify, I have captured that.

I find that the organizational emphasis on objectivity did not hinder screeners' ability to facilitate access to care for deserving groups of applicants. Rather, for some screeners, the functional eligibility algorithm both absolved them of responsibility and even provided them with more autonomy. For instance, Lindsey noted that she is in favor of the algorithm because it removes the pressure of decision-making. She said:

To be honest with you, I would prefer it that way [have the algorithm determine eligibility]. I feel like it's less pressure as a screener, like, I'm gonna trust the system. If I do my job like I'm supposed to do my job, that system is going to do its job.

For Lindsey, the objectivity of the algorithm allowed her to feel removed from denying care to applicants, even as she shapes the screening process to facilitate eligibility for applicants who “should” qualify.

Screening supervisor Julia also lauded the algorithm as a legitimizing device. She told me:

Sometimes people really aren't eligible, and that is why it's nice, then, to have the algorithm behind us. So it's not all bad, by any means. [...] I'm happy to have the algorithm have our back and be able to fall back on that. So certainly not all bad.”

Similarly to Lindsey, Julia appreciates that the algorithm provides an ‘objective’ way to deny care. As screeners generally tend to understand how the algorithm determines eligibility, this focus on objectivity allows them to expand access to care by coaching more “deserving” applicants while avoiding responsibility for the denial of care for “undeserving” applicants.

Ultimately, however, screeners are constrained by their capacity to fit applicants into the algorithm's eligibility framework. For example, Julia shared the case of an applicant, an elderly woman with "moderate dementia", who relied on an automated medication dispenser to remember to take her medicine. Typically, needing adaptive equipment carries substantial weight in the eligibility algorithm. However, the functional screen had no place to list adaptive equipment for medication use. In Julia's words: "there is no selection on the screen for adaptive equipment [under medication administration]. So we had to say she was independent with medicine. She is not independent with medication administration in any way, shape, or form."

Despite Julia's perception of this applicant as particularly deserving and even after raising the issue with DHS, she was found ineligible for the program. In sum, although screeners can guide eligibility determination according to their personal beliefs, applicants still must fit into particular categories in the algorithm, driven by the focus on objectivity.

Judges and Adjudication: Legality and Credibility

I find that legality, as defined by adherence to the state administrative code, is the primary guiding institutional logic experienced by appeal judges. While judges use screen results as supplementary evidence of applicants' abilities, their organizational position as separate from DHS (Figure 3.2) gives them latitude to disregard the eligibility algorithm if they see fit. Ultimately, judges based appeal decisions on the administrative code (Appendix A). The difference between the algorithm and administrative code came up repeatedly in appeals, as evidenced by the following two appeal excerpts:

The petitioner was found to be ineligible going forward, consistent with the DHS-directed result. However, the computer program infrequently yields a result that is not consistent with state code. [...] The Department determination regarding level of care is inconsistent with the Administrative Code criteria and petitioner remains at the nursing home level of care. (198590.pdf)

There are times when the logic or algorithm used by the computer program produces results that are at odds with the state regulations that govern the Family Care program. When such conflict is present, the regulations, not the computer program, control the outcome (201983.pdf)

The judge I interviewed, David, confirmed that the functional screen algorithm has limited power in appeals, stating:

I should let you know, even though I have a functional screen in front of [me], that's not how I make the decision. Recently, [a petitioner's] attorney attacked three findings in the functional screen and wanted to re-do them, but I don't do that. I use the functional screen as a guide rather than the do-all and end-all.

Importantly, the algorithm has stricter eligibility standards than the administrative code, which defines eligibility almost entirely according to the applicant's ability to perform ADLs and IADLs. This has implications for inequality in access to care, which I elaborate on later.

Because the administrative code is more lenient, judges are also potentially better able to affect the outcome according to their own motivations. As judge David told me, "I find that if it's close, I'm probably gonna swing to the person's behalf. I guess I'm liberal in wanting these people to get help who need it. That's one advantage of these being so subjective." While screeners Lindsey and Lily also reported wanting to help applicants, they did so by coaching applicants to fit into the discrete categories in the screen. Judges are not bound by those same parameters.

As screeners' personal views of deservingness affected how they screened applicants, judges assessed the "credibility" of applicants and witnesses when determining eligibility. Indeed, whether a judge awards or dismisses a case often hinges on the perceived credibility of either party's claims. For instance, in one dismissed case, a judge writes that they "give petitioner's mother's testimony no weight based on a lack of credibility and questionable motive" (202429.pdf). Conversely, in one successful appeal, the judge noted that "Petitioner's son credibly testified that Petitioner requires just as much care as she has always needed" (197327.pdf).

Judges' power to ignore the eligibility algorithm occasionally frustrated screeners and their supervisors. When I asked Suzanne, a supervisor at an MCO, if appeals are usually upheld or awarded, she said:

I'd say it's about 50/50. I was involved in 4 when I was the lead and it was 50/50, and it was just... sometimes it was the [judge] just completely ignoring the [functional screen] instructions and saying, 'I think this is the way it should be.' And then we just go, 'okay, fine.'

Julia echoed this sentiment, telling me, "If the [judge] says 'yep, they should be eligible' regardless of what the screen algorithm says, the team tells us how to manipulate it to make them eligible. So we're putting something inaccurate on the screen because an administrative law judge told [us] to." This frustration is emblematic of a conflict between the logics guiding DHS and the screening agencies (primarily objectivity) and the judges at the Department of Administration (primarily legality).

Implications for Inequality

Cruz (2022) argues that algorithms can both capture and obscure the social world they aim to represent. In this case, whether applicants' needs are sufficiently captured by the eligibility algorithm varies according to the logics under which screeners and judges employ the algorithm.

In both screening and adjudication, I find that eligibility determination disproportionately advantages applicants with greater cultural health capital, who are potentially better able to signal their limitations. For instance, eligibility determination requires applicants to present their limitations in medically intelligible ways, which may benefit applicants with greater health literacy. Indeed, when I asked appeal judge David what makes an appeal particularly challenging, he told me: "The people involved, how they describe what their functional limitations are. You know, you're very much shackled by their ability to describe what they can do."

Applicants must also physically demonstrate their limitations. For instance, if an applicant cannot prepare meals or bathe independently, a screener must observe this, or at least not observe behavior to the contrary. In fact, if applicants don't behave according to expectations during the screen, screeners may suspect overstating or dramatizing their limitations to become eligible. Screener Nicole told me, "I'm sure people are overstating sometimes. [...] if they're saying they can't get around their house without something, but then you're seeing them in the visit get up and grab a glass of water or something like that..."

While these sources of capital are advantageous in both screening and appeals, my findings suggest that they are especially necessary during screening because screeners are bound by the structure of the algorithm, which requires applicants to signal their limitations in very particular ways. The quantification of an applicant's disability, necessitated by the logic of objectivity and concerns about inter-rater reliability, means that there is no room for nuance or qualitative descriptions. Screeners must fit applicants into discrete categories, which are often arbitrarily defined.

The screen also only captures a moment in time, so whether a screener assesses an applicant on a "good" day or "bad" day can and does affect eligibility. Screener Lily describes how other screeners in her county will sometimes judge an applicant's eligibility based on one meeting with them:

There'll be times that [other screeners] are talking about cases where they're like, 'Oh, well, this person got on transit, and they got to the [screening agency], and they sat down in my office and they were able to have this conversation. They're able to get up and leave and do all this stuff themselves. So they're not gonna qualify.' And I'm like, 'Okay, but what if this is a good day?' I don't understand that... like, maybe the day before was one that they were just really struggling, and it was really cold that day. And so the osteoarthritis was really just making it difficult for them. Maybe the day that they came and asked was a good one.

In contrast, appeal judges are positioned to view applicants more holistically. In one example from appeals, the judge incorporated additional evidence about a petitioner's cognitive abilities not captured by the algorithm, which resulted in the petitioner being found eligible for care:

Petitioner has a diagnosis of major neurocognitive disorder. I find that neurological condition meets the above state definition for developmental disability. There was no dispute at hearing that her neurological condition was expected to continue indefinitely and that it constitutes a substantial handicap as

it results in the need for assistance with 1 ADL and 6 IADLs. Based on the foregoing, I find that petitioner meets the state definition for the developmental disability target group for Family Care and that she requires care at a comprehensive/nursing home level. (206000.pdf)

Further, judges often consider testimony about petitioners' abilities on good days and bad days, not just on the day of the initial eligibility screen.

Still, appeal judges use the screen results as evidence of applicants' ability to perform I/ADLs. Therefore, even though judges rarely use the algorithm's standards to determine eligibility, applicants' ability to fit their needs into the categories in the screen can still affect their eligibility during appeals. The following appeal excerpt illustrates how judges can use the screen to make appeal decisions:

Unfortunately, the testimony of petitioner and his wife is simply not corroborated by the medical documentation in the record. The petitioner's testimony regarding his specific needs for assistance is rebutted by the [screen] findings based upon observations of petitioner's capabilities. The two recent [screens] found that the petitioner requires assistance with two IADLs: laundry/chores, and transportation. He does not meet the criteria for the nursing home level of care which is required to demonstrate functional eligibility for the [long-term care] program. (209455.pdf)

Finally, as discussed previously, both screeners and judges use their discretion in ways that primarily benefit applicants perceived as more "deserving" or "credible". These perceptions are based in part on whether the applicant has the cultural health capital to signal the legitimacy of their disability or need. For instance, when I asked screener Wanda how much discretion she has while screening, she told me:

We have quite a bit. We get to decide if the person's credible or not. If someone's telling me that they can't do this and can't do that, but then I see them, like,

walking without assistance, with no problem, but they're telling me that they can't do that... Then I'll mark it from what I could see, not what they told me.

In this case, and others like it, the applicant may just have lacked the cultural health capital to signal their mobility limitations to the screener.

Discussion

Organizations within both healthcare and the government increasingly rely on algorithmic decision-making which privileges quantifiable information about patients and citizens (Burrell and Fourcade 2021; Timmermans 2010). Medicaid LTSS operates at the intersection of these institutions, as U.S. states assess functional eligibility to allocate healthcare resources. However, far from being a straightforward process, the implementation of these algorithms is complicated by varying organizational priorities and the personal motivations of the actors involved.

Through in-depth analysis of key informant interviews and appeal records, I find that Wisconsin's Medicaid functional eligibility algorithm is employed according to different institutional logics: functional eligibility screeners, contracted by the Department of Health Services, use the algorithm according to the logic of objectivity, while appeal judges, as a separate entity, primarily employ the logic of legality to make eligibility decisions. These logics have differing implications for the algorithmic decision-making process. While screeners must fit applicants into inflexible categories in order to ensure inter-rater reliability, appeal judges have the ability to disregard algorithmic output, which allows them to incorporate a more nuanced understanding of applicants' situations.

Further, both screeners and judges enact their personal beliefs about who deserves access to care within these logics. Screeners tend to view the elderly and people with physical disabilities as more deserving, and will occasionally coach these applicants to fit their needs into the discrete categories in the algorithm. Similarly, judges evaluate the credibility of applicants and their witnesses and use their discretion to make eligibility determinations according to these perceptions.

Interestingly, I do not find that the emphasis on objectivity within screening agencies conflicts with screeners' personal discretion in facilitating access to care for certain groups of applicants. Rather, objectivity provides screeners with distance from decision-making that allows them to shape eligibility for deserving applicants while denying care for the undeserving. Meanwhile, judges are directly accountable for appeal decisions as they have more explicit latitude within the logic of legality.

Finally, I find that these institutional logics shape the role of the algorithm in producing inequality. Applicants with greater health literacy are likely particularly advantaged during the screening process, as they are better able to signal their limitations in ways that are legible to the algorithm. Applicants with greater cultural resources are also likely better positioned to appeal an ineligible decision, which provides them access to a context in which their health needs are considered more fully without requiring them to fit the discrete categories in the eligibility algorithm.

Importantly, these findings do not suggest that applicants are consciously manipulating screeners or attempting to overstate their limitations to get LTSS. Embodied cultural capital, which I argue shapes applicants' ability and willingness to signal their

disabilities, is informed by the *habitus* (Bourdieu 1986) and therefore largely unconscious. Further, because LTSS beneficiaries are rescreened at least annually, it's highly unlikely that beneficiaries without substantial care needs will continuously be found eligible, regardless of their levels of capital. My results *do* suggest that many applicants who need LTSS are likely found ineligible if they lack the cultural capital to signal their disabilities and navigate the complex functional eligibility process. These may be people with low educational attainment, prior experience navigating bureaucracy, or limited capacity to advocate for themselves.

In places with significant demographic homogeneity, like rural counties in Wisconsin, cultural capital may operate as an invisible axis of inequality, where Medicaid LTSS applicants with particular cultural resources are better positioned to receive care. To the extent that cultural capital is racialized (Richards et al. 2023) and associated with financial resources (Bourdieu 1986), these sources of capital may also reinscribe more commonly documented forms of social inequality by race and SES – even in the context of a program that intentionally targets the most marginalized.

My findings suggest that algorithmic inequality is shaped in large part by the organizational context in which algorithms are used. As algorithms are increasingly integrated into the workplace, it is crucial to understand variations in the ways algorithms are implemented and adjudicated throughout the entire process of algorithmic decision-making.

Tables

Table 3.1. Summary of Interviewees

Role	Pseudonyms
State bureaucrats	Stacey, Leslie, Mariah, Sue
Ombudsmen	Patty, Carly, Meredith
Functional screeners	Carson, Nicole, Kylie, Jean, Wanda, Brandi, Lily
Functional screening supervisors	Suzanne, Julia, Lindsey, Ed
Appeals judge	David

Table 3.2. Summary of Appeals

	Appeal Reason			Total
	Functional status	Other	Challenged screen*	
Year				
2020	97 (41%)	105 (34%)	42 (42%)	202 (37%)
2021	49 (21%)	58 (19%)	15 (15%)	107 (20%)
2022	26 (11%)	64 (21%)	12 (12%)	90 (16%)
2023	63 (27%)	84 (27%)	32 (32%)	147 (27%)
Result				
Dismissed	148 (63%)	248 (80%)	57 (56%)	396 (73%)
Awarded	87 (37%)	63 (20%)	44 (44%)	150 (27%)
Total	235 (43%)	311 (57%)	101 (18%)	546

**Appeals with the reason “challenged screen” are a subset of appeals with the reason “functional status”. “Other” and “functional status” are mutually exclusive.*

Note: Percentage totals might not equal 100% due to rounding.

Figures

Figure 3.1. Functional Eligibility Process in Wisconsin

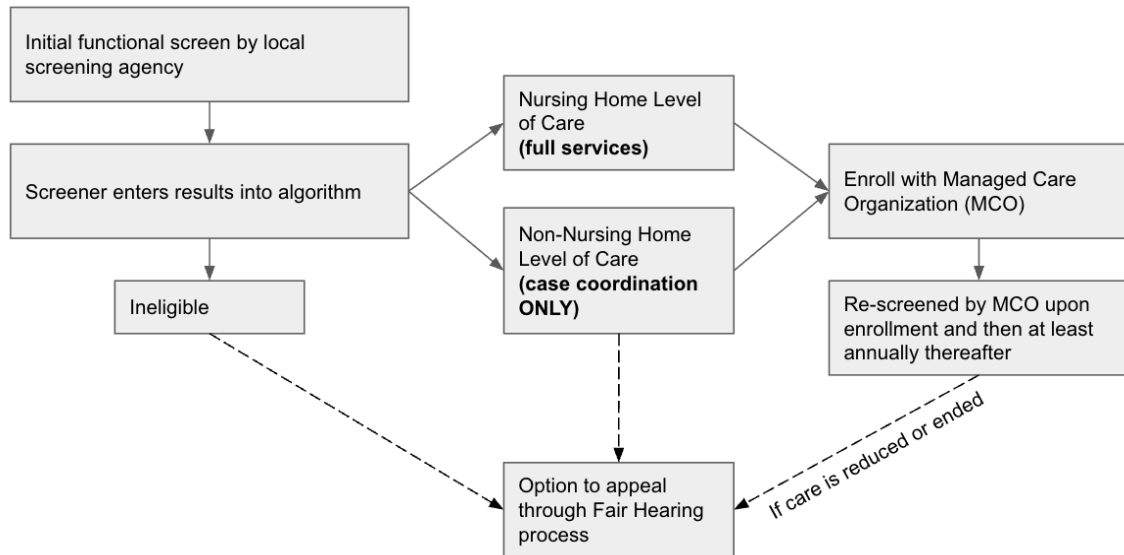
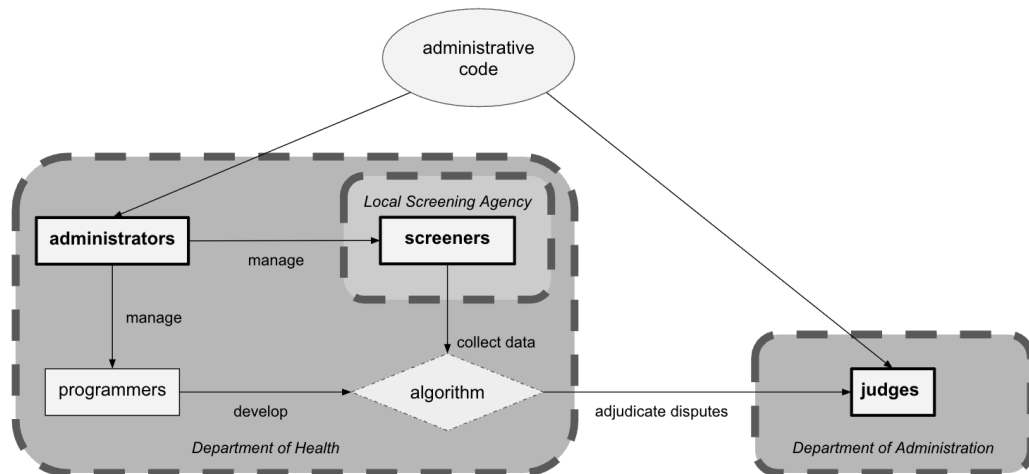


Figure 3.2. Organizational Arrangement of Algorithmic Decision-Making in Wisconsin



CONCLUSION

Using the case of functional eligibility determination in Medicaid LTSS, this dissertation provides some of the first population-level evidence of the features, predictors, and consequences of algorithmic decision-making. In Chapter 1, I show that fully automated algorithmic eligibility determination decreased from 2012 to 2020, while automated processes which require human oversight (“automated/approved” algorithms) increased. Interestingly, states which primarily used these automated/approved processes had lower Medicaid generosity and a higher percentage of Republican voters. These political differences may reflect strategies in more conservative states of increasing efficiency and cutting costs through automation while retaining control and bolstering legitimacy by giving humans oversight of eligibility decisions.

Chapter 2 builds on this to show how two dimensions of algorithms – automation and user expertise – jointly predict racial inequality in access to nursing home care. I find that fully automated algorithms implemented by medical professionals were associated with the most racial equality in the composition of nursing homes. While much of the current literature suggests that depersonalization in public benefits allocation can have harmful effects, this finding suggests the opposite: functional eligibility algorithms which do not explicitly incorporate humans at the point of output interpretation may be better for racial inequality, but only if medical professionals are involved in eligibility determination in some way. This finding underscores the important role of an algorithm’s users in

mitigating inequality and warns against the deskilling that often accompanies the implementation of fully automated algorithms. However, this chapter also shows that differences between algorithm types are less consequential for inequality as Medicaid generosity increases, suggesting that algorithms are especially important for regulating access to care in restrictive contexts.

Finally, Chapter 3 uses the case of functional eligibility determination in Wisconsin to show that even within the same state, the implementation of a single algorithm can vary across organizational contexts and institutional logics – eligibility screeners are bound to the highly standardized structure of the algorithm while appeal judges have more flexibility to disregard the algorithm. I argue that the emphasis on objectivity and standardization in screening requires applicants to signal their limitations in particular ways to be legible to the algorithm, which may disadvantage applicants with limited health literacy or program knowledge. During appeals, these resources are potentially less important as judges prioritize state regulations over algorithmic output.

Above all, this project shows that despite concerns about automation overtaking decision-making in government settings, bureaucrats still play a critical role in shaping these processes. In fact, their role may actually be increasing over time as the public questions the legitimacy of fully automated decision-making. However, my results show that increasing human oversight is not always beneficial for the public and may instead advance state organizational and political goals. While much of the literature centers on debating whether fully automated processes or “human-in-the-loop” processes are better for equality, there is variation within these categories. To fully understand the mechanisms underlying algorithmic inequality, we should turn toward examining how

automation, user expertise, and other characteristics of algorithms interact with each other and the broader political landscape.

REFERENCES

- Ackerhans, Sophia, Thomas Huynh, Carsten Kaiser, and Carsten Schultz. 2024. "Exploring the Role of Professional Identity in the Implementation of Clinical Decision Support Systems—a Narrative Review." *Implementation Science* 19(1):11. doi:10.1186/s13012-024-01339-x.
- Alqazlan, Lama, Zheng Fang, Michael Castelle, and Rob Procter. 2025. "A Novel, Human-in-the-Loop Computational Grounded Theory Framework for Big Social Data." *Big Data & Society* 12(2):20539517251347598. doi:10.1177/20539517251347598.
- Anderson, Nientara, Mytien Nguyen, Kayla Marcotte, Marco Ramos, Larry D. Gruppen, and Dowin Boatright. 2023. "The Long Shadow: A Historical Perspective on Racism in Medical Education." *Academic Medicine* 98(8S):S28–36. doi:[10.1097/ACM.0000000000005253](https://doi.org/10.1097/ACM.0000000000005253).
- Angelova, Victoria, Will Dobbie, and Crystal S. Yang. 2025. "Algorithmic Recommendations and Human Discretion." *Review of Economic Studies* rdaf084. doi:[10.1093/restud/rdaf084](https://doi.org/10.1093/restud/rdaf084).
- Anwar, Shamena, John Engberg, Isaac M. Oppen, and Leah Dion. 2024. *What Happens When Judges Follow the Recommendations of Pretrial Detention Risk Assessment Instruments More Often?* RAND Corporation. doi:[10.7249/RR3299-1](https://doi.org/10.7249/RR3299-1).

- Balancing Incentive Program. n.d. Baltimore, Maryland, USA: Centers for Medicare and Medicaid Services. <https://www.medicaid.gov/medicaid/long-term-services-supports/balancing-long-term-services-supports/balancing-incentive-program>.
- Benjamin, Ruha. 2019. *Race after Technology: Abolitionist Tools for the New Jim Code*. Medford, MA: Polity.
- Berk, Richard, Hoda Heidari, Shahin Jabbari, Michael Kearns, and Aaron Roth. 2021. "Fairness in Criminal Justice Risk Assessments: The State of the Art." *Sociological Methods & Research* 50(1):3–44. doi:10.1177/0049124118782533.
- Best, Rachel Kahn, and Alina Arseniev-Koehler. 2023. "The Stigma of Diseases: Unequal Burden, Uneven Decline." *American Sociological Review* 88(5):938–69.
- Bloch-Wehba, Hannah. 2020. "Access to Algorithms." *Fordham Law Review* 88(4):1265–1314.
- Bourdieu, Pierre. 1986. "The Forms of Capital." Pp. 15–29 in *Handbook of Theory and Research for the Sociology of Education*, edited by J. Richardson.
- Brayne, Sarah. 2017. "Big Data Surveillance: The Case of Policing." *American Sociological Review* 82(5):977–1008. doi:[10.1177/0003122417725865](https://doi.org/10.1177/0003122417725865).
- Brayne, Sarah. 2021. *Predict and Surveil: Data, Discretion, and the Future of Policing*. New York, NY: Oxford University Press.

- Brayne, Sarah, and Angèle Christin. 2021. "Technologies of Crime Prediction: The Reception of Algorithms in Policing and Criminal Courts." *Social Problems* 68(3):608–24. doi:[10.1093/socpro/spaa004](https://doi.org/10.1093/socpro/spaa004).
- Bridges, Khiara M. 2011. *Reproducing Race: An Ethnography of Pregnancy as a Site of Racialization*. Berkeley: University of California Press.
- Brown, Lydia X. Z., Michelle Richardson, Ridhi Shetty, Andrew Crawford, and Timothy Hoagland. 2020. Challenging the Use of Algorithm-Driven Decision-Making in Benefits Determinations Affecting People with Disabilities. Center for Democracy and Technology.
- Buffat, Aurélien. 2015. "Street-Level Bureaucracy and E-Government." *Public Management Review* 17(1):149–61. doi:[10.1080/14719037.2013.771699](https://doi.org/10.1080/14719037.2013.771699).
- Burrell, Jenna. 2016. "How the Machine 'Thinks': Understanding Opacity in Machine Learning Algorithms." *Big Data & Society* 3(1):2053951715622512. doi:10.1177/2053951715622512.
- Burrell, Jenna, and Marion Fourcade. 2021. "The Society of Algorithms." *Annual Review of Sociology* 47(1):213–37. doi: [10.1146/annurev-soc-090820-020800](https://doi.org/10.1146/annurev-soc-090820-020800).
- Cantwell, Mari. 2019. "Application for a 1915(c) Home and Community-Based Services Waiver."

- Chidambaram, Priya, and Alice Burns. 2025. "A Look at Nursing Facility Characteristics in 2025." Kaiser Family Foundation. <https://www.kff.org/medicaid/a-look-at-nursing-facility-characteristics/>.
- Christin, Angèle. 2017. "Algorithms in Practice: Comparing Web Journalism and Criminal Justice." *Big Data & Society* 4(2):205395171771885. doi:[10.1177/2053951717718855](https://doi.org/10.1177/2053951717718855).
- Citron, Danielle Keats, and Frank Pasquale. 2014. "The Scored Society: Due Process for Automated Predictions." *Washington Law Review* 89(1).
- Cruz, Taylor Marion. 2022. "The Social Life of Biomedical Data: Capturing, Obscuring, and Envisioning Care in the Digital Safety-Net." *Social Science & Medicine* 294:114670. doi:[10.1016/j.socscimed.2021.114670](https://doi.org/10.1016/j.socscimed.2021.114670).
- DiMaggio, Paul J., and Walter W. Powell. 1991. "The Iron Cage Revisited: Institutional Isomorphism and Collective Rationality in Organizational Fields." in *The new institutionalism in organizational analysis*, edited by P. J. DiMaggio and W. W. Powell. Chicago: University of Chicago Press.
- Division of Medicaid Services. 2024. "Wisconsin Long Term Care Functional Screen Instructions."
- Donatz-Fest, Isabelle. 2024. "The 'Doings' behind Data: An Ethnography of Police Data Construction." *Big Data & Society* 11(3):1–12. doi:[10.1177/20539517241270695](https://doi.org/10.1177/20539517241270695).
- Einstein, Katherine Levine, and David M. Glick. 2017. "Does Race Affect Access to Government Services? An Experiment Exploring Street-Level Bureaucrats and

- Access to Public Housing.” *American Journal of Political Science* 61(1):100–116.
doi:[10.1111/ajps.12252](https://doi.org/10.1111/ajps.12252).
- Esthappan, Sino. 2024. “Assessing the Risks of Risk Assessments: Institutional Tensions and Data Driven Judicial Decision-Making in U.S. Pretrial Hearings.” *Social Problems* spae060. doi:[10.1093/socpro/spae060](https://doi.org/10.1093/socpro/spae060).
- Eubanks, Virginia. 2017. *Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor*. First Edition. New York, NY: St. Martin’s Press.
- Evans, T. 2011. “Professionals, Managers and Discretion: Critiquing Street-Level Bureaucracy.” *British Journal of Social Work* 41(2):368–86.
doi:[10.1093/bjsw/bcq074](https://doi.org/10.1093/bjsw/bcq074).
- Evans, Tony. 2020. “Street-Level Bureaucrats: Discretion and Compliance in Policy Implementation.” in *Oxford Research Encyclopedia of Politics*. Oxford University Press.
- Feng, Zhanlian, Mary L. Fennell, Denise A. Tyler, Melissa Clark, and Vincent Mor. 2011. “Growth Of Racial And Ethnic Minorities In US Nursing Homes Driven By Demographics And Possible Disparities In Options.” *Health Affairs* 30(7):1358–65.
doi:[10.1377/hlthaff.2011.0126](https://doi.org/10.1377/hlthaff.2011.0126).
- Ferdman, Avigail. 2025. “AI Deskillling Is a Structural Problem.” AI & SOCIETY.
doi:[10.1007/s00146-025-02686-z](https://doi.org/10.1007/s00146-025-02686-z).

- Fourcade, Marion, and Kieran Healy. 2013. "Classification Situations: Life-Chances in the Neoliberal Era." *Accounting, Organizations and Society* 38(8):559–72. doi: [10.1016/j.aos.2013.11.002](https://doi.org/10.1016/j.aos.2013.11.002).
- Frumkin, Peter, and Joseph Galaskiewicz. 2004. "Institutional Isomorphism and Public Sector Organizations." *Journal of Public Administration Research and Theory* 14(3):283–307. doi:10.1093/jopart/muh028.
- Grimmelikhuijsen, Stephan, and Albert Meijer. 2022. "Legitimacy of Algorithmic Decision-Making: Six Threats and the Need for a Calibrated Institutional Response." *Perspectives on Public Management and Governance* 5(3):232–42. doi:10.1093/ppmgov/gvac008.
- Grønsund, Tor, and Margunn Aanestad. 2020. "Augmenting the Algorithm: Emerging Human-in-the-Loop Work Configurations." *The Journal of Strategic Information Systems* 29(2):101614. doi:10.1016/j.jsis.2020.101614.
- Gules-Guctas, Esra. 2025. "How Do Algorithmic Decision-Making Systems Used in Public Benefits Determinations Fail? Insights From Legal Challenges." *Public Administration Review* puar.70043. doi:10.1111/puar.70043.
- Hannah-Moffat, Kelly. 2013. "Actuarial Sentencing: An 'Unsettled' Proposition." *Justice Quarterly* 30(2):270–96. doi:10.1080/07418825.2012.682603.
- Harrits, Gitte Sommer. 2019. "Stereotypes in Context: How and When Do Street-Level Bureaucrats Use Class Stereotypes?" *Public Administration Review* 79(1):93–103. doi:[10.1111/puar.12952](https://doi.org/10.1111/puar.12952).

- Hasenfeld, Yeheskel. 2000. "Organizational Forms as Moral Practices: The Case of Welfare Departments." *Social Service Review* 74(3):329–51. doi: [10.1086/516408](https://doi.org/10.1086/516408).
- Henman, Paul W. Fay. 2022. "Digital Social Policy: Past, Present, Future." *Journal of Social Policy* 51(3):535–50. doi:10.1017/S0047279422000162.
- Herd, Pamela, and Donald P. Moynihan. 2019. *Administrative Burden: Policymaking by Other Means*. Russell Sage Foundation.
- Hernandez, Stephanie M., and P. Johnelle Sparks. 2020. "Barriers to Health Care Among Adults With Minoritized Identities in the United States, 2013–2017." *American Journal of Public Health* 110(6):857–62. doi:[10.2105/AJPH.2020.305598](https://doi.org/10.2105/AJPH.2020.305598).
- Hill, Latoya, Samantha Artiga, and Sweta Haldar. 2022. *Key Facts on Health and Health Care by Race and Ethnicity*. Kaiser Family Foundation. <https://www.kff.org/racial-equity-and-health-policy/report/key-facts-on-health-and-health-care-by-race-and-ethnicity/>.
- Hillo, Jaakko, Isak Vento, and Tero Erkkilä. 2025. "Algorithmic Governance: Experimental Evidence on Citizens' and Public Administrators' Legitimacy Perceptions of Automated Decision-Making." *Public Administration* padm.70028. doi:10.1111/padm.70028.
- Himmelstein, David U., and Steffie Woolhandler. 2008. "Privatization in a Publicly Funded Health Care System: The U.S. Experience." *International Journal of Health Services* 38(3):407–19. doi:10.2190/HS.38.3.a.

- Hirschman, Daniel, and Emily Adlin Bosk. 2020. "Standardizing Biases: Selection Devices and the Quantification of Race." *Sociology of Race and Ethnicity* 6(3):348–64.
- Hoffman, Kelly M., Sophie Trawalter, Jordan R. Axt, and M. Norman Oliver. 2016. "Racial Bias in Pain Assessment and Treatment Recommendations, and False Beliefs about Biological Differences between Blacks and Whites." *Proceedings of the National Academy of Sciences* 113(16):4296–4301.
doi:[10.1073/pnas.1516047113](https://doi.org/10.1073/pnas.1516047113).
- Hoffman, Steve G. 2017. "Managing Ambiguities at the Edge of Knowledge: Research Strategy and Artificial Intelligence Labs in an Era of Academic Capitalism." *Science, Technology, & Human Values* 42(4):703–40.
doi:[10.1177/0162243916687038](https://doi.org/10.1177/0162243916687038).
- Johnson, Tracy. 2019. "Application for a 1915(c) Home and Community-Based Services Waiver."
- Jones, Meg Leta. 2017. "The Right to a Human in the Loop: Political Constructions of Computer Automation and Personhood." *Social Studies of Science* 47(2):216–39.
doi:[10.1177/0306312717699716](https://doi.org/10.1177/0306312717699716).
- Kaiser Family Foundation. 2020. "Medicaid Enrollees Using Long-Term Care (LTC) as a Percent of Full-Benefit Medicaid Enrollees." <https://www.kff.org/medicaid/state-indicator/medicaid-enrollees-using-ltc-as-a-percent-of-full-benefit-medicaid->

[enrollees/?currentTimeframe=3&sortModel=%7B%22colId%22:%22Location%22,%22sort%22:%22asc%22%7D.](#)

Kaiser Family Foundation. 2025. “Status of State Action on the Medicaid Expansion Decision.” <https://www.kff.org/affordable-care-act/state-indicator/state-activity-around-expanding-medicaid-under-the-affordable-care-act/?currentTimeframe=0&sortModel=%7B%22colId%22:%22Expansion%20Implementation%20Date%22,%22sort%22:%22asc%22%7D>.

Kako, Edward, Rebecca Sweetland, Kerri Melda, Elizabeth Coombs, Mason Smith, and John Agosta. 2013. *THE BALANCING INCENTIVE PROGRAM: IMPLEMENTATION MANUAL*. San Francisco, CA: Mission Analytics Group, Incorporated.

Karas Montez, Jennifer. 2025. “State Policy and Politics Database.” Syracuse University: Center for Aging and Policy Studies.

Katzenbach, Christian, and Lena Ulbricht. 2019. “Algorithmic Governance.” *Internet Policy Review* 8(4). doi:10.14763/2019.4.1424.

Kellogg, Katherine C., Melissa A. Valentine, and Angèle Christin. 2020. “Algorithms at Work: The New Contested Terrain of Control.” *Academy of Management Annals* 14(1):366–410. doi:[10.5465/annals.2018.0174](https://doi.org/10.5465/annals.2018.0174).

Kelly, Marisa. 1994. “Theories of Justice and Street-Level Discretion¹.” *Journal of Public Administration Research and Theory*. doi:[10.1093/oxfordjournals.jpart.a037201](https://doi.org/10.1093/oxfordjournals.jpart.a037201).

- Kim, Min-Young, Edward Weizenegger, and Andrea Wysocki. 2022. *Medicaid Beneficiaries Who Use Long-Term Services and Supports: 2019*. Chicago, IL: Mathematica.
- König, Pascal D., and Georg Wenzelburger. 2021. “The Legitimacy Gap of Algorithmic Decision-Making in the Public Sector: Why It Arises and How to Address It.” *Technology in Society* 67:101688. doi:10.1016/j.techsoc.2021.101688.
- Lecher, Colin. 2018. “What Happens When an Algorithm Cuts Your Health Care.” *The Verge*, March 21.
- Lendon, Jessica P., Christine Caffrey, Amanuel Melekin, Priyanka Singh, Zhaohui Lu, and Manisha Sengupta. 2024. *Overview of Post-Acute and Long-Term Care Providers and Service Users in the United States, 2020*. *National Health Statistics Reports*. 208. Centers for Disease Control and Prevention.
<https://www.cdc.gov/nchs/fastats/nursing-home-care.htm>.
- Levy, Karen, Kyla E. Chasalow, and Sarah Riley. 2021. “Algorithms and Decision-Making in the Public Sector.” *Annual Review of Law and Social Science* 17(1):309–34. doi:10.1146/annurev-lawsocsci-041221-023808.
- Linzer, Drew, and Jeffrey Lewis. 2026. “Polytomous Variable Latent Class Analysis.”
- Lipsky, Michael. 1984. “Bureaucratic Disentitlement in Social Welfare Programs.” *Social Service Review* 58(1):3–27. doi:[10.1086/644161](https://doi.org/10.1086/644161).

Lipsky, Michael. 2010. *Street-Level Bureaucracy: Dilemmas of the Individual in Public Services*. Updated edition. Publications of Russell Sage Foundation. New York: Russell Sage Foundation.

LTCFocus Public Use Data sponsored by the National Institute on Aging (P01 AG027296) through a cooperative agreement with the Brown University School of Public Health. Available at www.ltcfocus.org. <https://doi.org/10.26300/h9a2-2c26>

LTCFocus. n.d. Brown University Center for Gerontology and Healthcare Research: Shaping Long-Term Care in America Project. ltcfocus.org.

Maiers, Claire. 2017. “Analytics in Action: Users and Predictive Data in the Neonatal Intensive Care Unit.” *Information, Communication & Society* 20(6):915–29. doi:[10.1080/1369118X.2017.1291701](https://doi.org/10.1080/1369118X.2017.1291701).

Maximus. Clinical Services. n.d. <https://maximusclinicalservices.com/home>.

May, P. J., and S. C. Winter. 2009. “Politicians, Managers, and Street-Level Bureaucrats: Influences on Policy Implementation.” *Journal of Public Administration Research and Theory* 19(3):453–76. doi:[10.1093/jopart/mum030](https://doi.org/10.1093/jopart/mum030).

Maynard-Moody, S., and M. Musheno. 2000. “State Agent or Citizen Agent: Two Narratives of Discretion.” *Journal of Public Administration Research and Theory* 10(2):329–58. doi:[10.1093/oxfordjournals.jpart.a024272](https://doi.org/10.1093/oxfordjournals.jpart.a024272).

- Medicaid and CHIP Payment and Access Commission. 2016. *Chapter 4: Functional Assessments for Long-Term Services and Supports*. Report to Congress on Medicaid and CHIP.
- Menjívar, Cecilia. 2002. “The Ties That Heal: Guatemalan Immigrant Women’s Networks and Medical Treatment.” *International Migration Review* 36(2):437–66. doi: [10.1111/j.1747-7379.2002.tb00088.x](https://doi.org/10.1111/j.1747-7379.2002.tb00088.x).
- Menjívar, Cecilia. 2023. “State Categories, Bureaucracies of Displacement, and Possibilities from the Margins.” *American Sociological Review* 88(1):1–23. doi: [10.1177/00031224221145727](https://doi.org/10.1177/00031224221145727).
- Monarcha-Matlak, Aleksandra. 2021. “Automated Decision-Making in Public Administration.” *Procedia Computer Science* 192:2077–84. doi:10.1016/j.procs.2021.08.215.
- Mulligan, Deirdre K., and Kenneth A. Bamberger. 2019. “Procurement As Policy: Administrative Process for Machine Learning.” *SSRN Electronic Journal*. doi:10.2139/ssrn.3464203.
- Murray, Caitlin, Michelle Eckstein, Debra Lipson, and Andrea Wysocki. 2023. *Medicaid Long Term Services and Supports Annual Expenditures Report: Federal Fiscal Year 2020*. Chicago, IL: Mathematica.
- Musumeci, MaryBeth. 2015. *Measuring Long-Term Services and Supports Rebalancing*. Kaiser Family Foundation. <https://www.kff.org/medicaid/measuring-long-term-services-and-supports-rebalancing/>.

- Noble, Safiya Umoja. 2018. *Algorithms of Oppression: How Search Engines Reinforce Racism*. New York: New York University Press.
- Obermeyer, Ziad, Brian Powers, Christine Vogeli, and Sendhil Mullainathan. 2019. “Dissecting Racial Bias in an Algorithm Used to Manage the Health of Populations.” *Science* 366:447–53.
- Pasquale, Frank. 2015. *The Black Box Society: The Secret Algorithms That Control Money and Information*. Cambridge, MA: Harvard University Press.
- Peeters, Rik. 2020. “The Agency of Algorithms: Understanding Human-Algorithm Interaction in Administrative Decision-Making” edited by S. Giest and S. Grimmelikhuijsen. *Information Polity* 25(4):507–22. doi:10.3233/IP-200253.
- Penner, Louis A., and John F. Dovidio. 2016. “Racial Color Blindness and Black-White Health Care Disparities.” Pp. 275–93 in *The myth of racial color blindness: Manifestations, dynamics, and impact.*, edited by H. A. Neville, M. E. Gallardo, and D. W. Sue. Washington: American Psychological Association.
- Phelan, Jo C., Bruce G. Link, and Naumi M. Feldman. 2013. “The Genomic Revolution and Beliefs about Essential Racial Differences: A Backdoor to Eugenics?” *American Sociological Review* 78(2):167–91. doi:[10.1177/0003122413476034](https://doi.org/10.1177/0003122413476034).
- Portillo, Shannon, and Danielle S. Rudes. 2014. “Construction of Justice at the Street Level.” *Annual Review of Law and Social Science* 10(1):321–34. doi:[10.1146/annurev-lawsocsci-102612-134046](https://doi.org/10.1146/annurev-lawsocsci-102612-134046).

- Porter, Theodore M. 1995. *Trust in Numbers: The Pursuit of Objectivity in Science and Public Life*. Princeton, N.J: Princeton University Press.
- Richards, Bedelia Nicola, Hugo Ceron-Anaya, Susan A. Dumais, Jennifer C. Mueller, Patricia Sánchez-Connally, and Derron Wallace. 2023. “What’s Race Got to Do With It? Disrupting Whiteness in Cultural Capital Research.” *Sociology of Race and Ethnicity* 9(3):279–94. doi: [10.1177/23326492231160535](https://doi.org/10.1177/23326492231160535).
- Rosen, Eva, Philip M. E. Garboden, and Jennifer E. Cossyleon. 2021. “Racial Discrimination in Housing: How Landlords Use Algorithms and Home Visits to Screen Tenants.” *American Sociological Review* 86(5):787–822. doi: [10.1177/00031224211029618](https://doi.org/10.1177/00031224211029618).
- Sandfort, J. R. 2000. “Moving Beyond Discretion and Outcomes: Examining Public Management from the Front Lines of the Welfare System.” *Journal of Public Administration Research and Theory* 10(4):729–56. doi: [10.1093/oxfordjournals.jpart.a024289](https://doi.org/10.1093/oxfordjournals.jpart.a024289).
- Schneider, Anne, and Helen Ingram. 1993. “Social Construction of Target Populations: Implications for Politics and Policy.” *American Political Science Review* 87(2):334–47. doi: [10.2307/2939044](https://doi.org/10.2307/2939044).
- Scott, James C. 2020. *Seeing like a State: How Certain Schemes to Improve the Human Condition Have Failed*. Veritas paperbacks edition. New Haven: Yale University Press.

- Scott, P. G. 1997. "Assessing Determinants of Bureaucratic Discretion: An Experiment in Street-Level Decision Making." *Journal of Public Administration Research and Theory* 7(1):35–58. doi:[10.1093/oxfordjournals.jpart.a024341](https://doi.org/10.1093/oxfordjournals.jpart.a024341).
- Shen, Yu-Chu, and Stephen Zuckerman. 2005. "The Effect of Medicaid Payment Generosity on Access and Use among Beneficiaries." *Health Services Research* 40(3):723–44. doi:[10.1111/j.1475-6773.2005.00382.x](https://doi.org/10.1111/j.1475-6773.2005.00382.x).
- Shim, Janet K. 2010. "Cultural Health Capital: A Theoretical Approach to Understanding Health Care Interactions and the Dynamics of Unequal Treatment." *Journal of Health and Social Behavior* 51(1):1–15. doi: [10.1177/0022146509361185](https://doi.org/10.1177/0022146509361185).
- Siddique, Shazia Mehmood, Kelley Tipton, Brian Leas, Christopher Jepson, Jaya Aysola, Jordana B. Cohen, Emilia Flores, Michael O. Harhay, Harald Schmidt, Gary E. Weissman, Julie Fricke, Jonathan R. Treadwell, and Nikhil K. Mull. 2024. "The Impact of Health Care Algorithms on Racial and Ethnic Disparities: A Systematic Review." *Annals of Internal Medicine* 177(4):484–96. doi:[10.7326/M23-2960](https://doi.org/10.7326/M23-2960).
- Skrtic, Thomas M., Argun Saatcioglu, and Austin Nichols. 2021. "Disability as Status Competition: The Role of Race in Classifying Children." *Socius: Sociological Research for a Dynamic World* 7:237802312110243. doi:[10.1177/23780231211024398](https://doi.org/10.1177/23780231211024398).
- Smedley, Brian D., Adrienne Y. Stith, Alan R. Nelson, and Institute of Medicine (U.S.), eds. 2003. *Unequal Treatment: Confronting Racial and Ethnic Disparities in Health Care*. Washington, D.C: National Academies Press.

- Sommers, Benjamin D., Bethany Maylone, Robert J. Blendon, E. John Orav, and Arnold M. Epstein. 2017. “Three-Year Impacts Of The Affordable Care Act: Improved Medical Care And Health Among Low-Income Adults.” *Health Affairs* 36(6):1119–28. doi:[10.1377/hlthaff.2017.0293](https://doi.org/10.1377/hlthaff.2017.0293).
- Star, Susan Leigh, and Geoffrey C. Bowker. 2007. “Enacting Silence: Residual Categories as a Challenge for Ethics, Information Systems, and Communication.” *Ethics and Information Technology* 9(4):273–80. doi:10.1007/s10676-007-9141-7.
- Stier, Mikki. 2017. “Application for a 1915(c) Home and Community-Based Services Waiver.”
- Thornton, Patricia H., William Ocasio, and Michael Lounsbury. 2012. *The Institutional Logics Perspective: A New Approach to Culture, Structure, and Process*. Oxford: Oxford University Press.
- Timmermans, Stefan. 2010. “Evidence-Based Medicine: Sociological Explorations.” in *Handbook of medical sociology*, edited by C. E. Bird, P. Conrad, A. M. Fremont, and S. Timmermans. Nashville: Vanderbilt University Press.
- Tummers, Lars, and Victor Bekkers. 2014. “Policy Implementation, Street-Level Bureaucracy, and the Importance of Discretion.” *Public Management Review* 16(4):527–47. doi:[10.1080/14719037.2013.841978](https://doi.org/10.1080/14719037.2013.841978).
- Unruh, Lynn, and Thomas Rice. 2025. “Private Equity Expansion and Impacts in United States Healthcare.” *Health Policy* 155:105266. doi:10.1016/j.healthpol.2025.105266.

U.S. Centers for Medicare and Medicaid Services. 2024. “Minimum Data Set (MDS) 3.0 for Nursing Homes and Swing Bed Providers.”

<https://www.cms.gov/medicare/quality/nursing-home-improvement/minimum-data-sets-swing-bed-providers>.

Vespa, Jonathan, Lauren Medina, and David M. Armstrong. 2020. *Demographic Turning Points for the United States: Population Projections for 2020 to 2060*. P25-1144. United States Census Bureau.

Waldman, Ari, and Kirsten Martin. 2022. “Governing Algorithmic Decisions: The Role of Decision Importance and Governance on Perceived Legitimacy of Algorithmic Decisions.” *Big Data & Society* 9(1):20539517221100449. doi:10.1177/20539517221100449.

Watkins-Hayes, Celeste, and Elyse Kovalsky. 2017. “The Discourse of Deservingness: Morality and the Dilemmas of Poverty Relief in Debate and Practice.” Pp. 193–220 in *The Oxford Handbook of the Social Science of Poverty*, edited by D. Brady and L. M. Burton. Oxford University Press.

Watson, Libby. 2022. “The Coming Medicaid Purge.” in *Profits Over Patients: The Corporatization of American Health Care and What You Can Do to Protect Yourself*. The Lever.

Weber, Max. 1978. “Chapter XI: Bureaucracy.” in *Economy and Society: An Outline of Interpretive Sociology*. Vol. 1. University of California Press.

- Webster, Craig S., Saana Taylor, Courtney Thomas, and Jennifer M. Weller. 2022. “Social Bias, Discrimination and Inequity in Healthcare: Mechanisms, Implications and Recommendations.” *BJA Education* 22(4):131–37. doi:[10.1016/j.bjae.2021.11.011](https://doi.org/10.1016/j.bjae.2021.11.011).
- Wisconsin Department of Health Services. 2024. “Overview of the Long-Term Care Functional Screen.” Retrieved November 14, 2024 (<https://www.dhs.wisconsin.gov/functionalscreen/lcfs-overview.htm>)
- Yearby, Ruqaiijah. 2018. “Racial Disparities in Health Status and Access to Healthcare: The Continuation of Inequality in the United States Due to Structural Racism.” *The American Journal of Economics and Sociology* 77(3–4):1113–52. doi:[10.1111/ajes.12230](https://doi.org/10.1111/ajes.12230)
- Zalnieriute, Monika, Lyria Bennett Moses, and George Williams. 2019. “The Rule of Law and Automation of Government Decision-Making.” *The Modern Law Review* 82(3):425–55. doi:10.1111/1468-2230.12412.
- Zou, James, and Londa Schiebinger. 2018 “Design AI so That It’s Fair.” *Nature*.