

RESEARCH ARTICLE OPEN ACCESS

Algorithms, Humans, and Racial Disparities in Child Protection Systems: Evidence From the Allegheny Family Screening Tool

Katherine Rittenhouse¹  | Emily Putnam-Hornstein² | Rhema Vaithianathan³

¹LBJ School of Public Affairs, University of Texas at Austin, Austin, Texas, USA | ²University of North Carolina at Chapel Hill, Chapel Hill, North Carolina, USA | ³Auckland University of Technology, Auckland, New Zealand

Correspondence: Katherine Rittenhouse (katherine.rittenhouse@austin.utexas.edu)

Received: 19 February 2025 | **Revised:** 19 May 2026 | **Accepted:** 28 May 2026

ABSTRACT

Equity concerns are a primary constraint in the adoption of algorithms across various settings. We ask how providing decision-makers with a predictive risk model affects racial disparities in the context of the child protection system. We exploit the implementation of the Allegheny Family Screening Tool (AFST), a tool that aims to help child protection workers decide which allegations of maltreatment to screen in for investigation. We find that the AFST reduced disparities in screening decisions and home removal rates for referrals involving Black versus White children. An analysis of false positive and false negative rates suggests that this reduction in disparities was achieved while improving welfare among both Black and White children.

JEL Classification: J13, J15, J18

1 | Introduction

Machine learning tools, actuarial tools, and predictive risk models (“algorithms”) can help human systems make better decisions (Kleinberg et al. 2018). As such, algorithms are increasingly promoted as useful complements to human decision-making in a wide variety of settings, including bail decisions (Chohlas-Wood 2020), resume screening (Raghavan et al. 2020), health care (Price 2019), and, more recently, the child protection system (Chouldechova et al. 2018). As their prevalence grows, so too do concerns that algorithms may entrench or exacerbate existing disparities, and in particular disparities across race. Equity concerns are a driving constraint in the adoption of these tools.¹ However, human bias is also well-established, and has been shown to cause racial disparities in many institutions.² In this paper, we ask whether the introduction of a predictive risk model increases or decreases disparities, relative to the relevant counterfactual of human decision-making.

We address this question in the context of the child protection system, an institution which interacts with approximately one third of US children by the time they reach 18 (Kim et al. 2017). The child protection system is comprised of several high-stakes decision points, including whether to investigate reports of alleged maltreatment, and whether to remove children from abusive or neglectful homes. This is also a setting with large racial disparities—Black children are 88% more likely than White children to be investigated for maltreatment, and over twice as likely to enter foster care by the time they reach 18 (Kim et al. 2017; Wildeman and Emanuel 2014).

Predictive risk models have been introduced as a tool to improve quality and consistency of decision-making in child protection systems, but the use of these tools has been stalled by concerns about the impact of algorithms on racial disparities.³ Improving racial equity is a top priority for child protection agencies.⁴ The United Nations has expressed “[concern] at the disproportionate

This is an open access article under the terms of the Creative Commons Attribution-NonCommercial-NoDerivs License, which permits use and distribution in any medium, provided the original work is properly cited, the use is non-commercial and no modifications or adaptations are made.

© 2026 The Author(s). *Journal of Policy Analysis and Management* published by Wiley Periodicals LLC on behalf of Association for Public Policy and Management.

number of children belonging to racial and ethnic minorities who are removed from their families and placed in foster care,” and the Biden administration noted the role of “systemic racism and economic barriers” in creating these disparities (UN2 2022; Biden 2021). However, there is no consensus around how to improve racial equity while maintaining standards of child safety and protection.

We study the implementation of the Allegheny Family Screening Tool (AFST), the first automated predictive risk model used to aid decision-makers in the child protection system. This algorithm aims to help workers decide whether to investigate (“screen in”) referrals alleging that a child is being maltreated.⁵ The AFST uses information about the referred families from linked administrative data to predict the risk that the child will be removed from their home if screened in, and shows the human decision-maker a risk score ranging from 1 to 20. For referrals with the highest risk scores, the AFST activates a high-risk protocol, defaulting the referral to a screen-in recommendation. In all cases, the call-screening supervisors remain the ultimate decision-maker, and may choose to override the default. Screened-in referrals proceed to an investigation, conducted by a different caseworker. In severe cases where the caseworker determines that a child’s safety is at risk, investigations may result in a child being removed from their home and placed in foster care. We ask how the implementation of the AFST affected: (1) the differential probability of screening in referrals involving Black versus White children, and (2) the differential probability of home removal for referrals involving Black versus White children.

Our setting is well-suited to studying the effects of implementing the algorithm, as we observe outcomes for referrals made both before and after the AFST was deployed. Crucially, this allows us to evaluate how humans use the algorithm in practice. We obtain data on referrals made to Allegheny County’s office of Children, Youth and Families (CYF) between 2010 and 2020. For each referral, we observe the AFST score (retroactively calculated for referrals made prior to implementation) and the screening decision. Referrals are also associated with one or more children, for whom we observe demographic information. We then link these individual children to the universe of foster care records between 2010 and 2020, in order to study home removals.

We first establish the extent of racial disparities in screening rates in our setting prior to the implementation of the AFST. Referrals involving Black children are 11 percentage points more likely than referrals involving White children to be screened in for an investigation. Of course, disparities may be warranted if the screen-in rate reflects differences in the riskiness of referred Black versus White children. Controlling for the (unobserved) underlying risk score, referrals involving Black children are still 6 percentage points more likely to be screened in for investigation. Finally, following the approach of Baron et al. (2026), we estimate unwarranted disparities in screening rates, defined as differences conditional on potential for future maltreatment.⁶ Referrals involving Black children are 8.8–10.5 percentage points more likely to be screened in than White children with the same potential for future maltreatment.

To study the effects of the AFST on disparities in screening rates we exploit the introduction of the AFST, comparing referrals

involving Black versus White children, made before versus after the AFST was implemented. While all children were potentially impacted by the introduction of the AFST, we test the hypothesis that Black and White children were differentially affected. We find that the AFST reduced unconditional disparities in screening rates by 3.3 percentage points, and disparities conditional on algorithm score by 3.1 percentage points. Effects are concentrated among high-risk referrals, which are likely to be defaulted to a screen in decision as per the high-risk protocol. Event studies are consistent with the necessary assumption, that screen-in rates for Black and White children followed similar trends prior to the introduction of the AFST.

Next, we study the effects of the AFST on disparities in foster care. For this analysis, we use two empirical strategies. First, we again compare referrals involving Black versus White children, before versus after the AFST was implemented. Second, we include a control group, comparing across referrals which were “treated” versus not “treated” by the AFST in a difference-in-differences design. Specifically, under Pennsylvania law, referrals which include certain allegations are automatically screened in for investigation, and as such were not affected by the AFST. This design enables us to control for any trends across time which might differentially affect Black versus White families. One drawback of this design is that the control-group referrals are expunged from the data prior to 2015, giving us a shorter timeframe to study. Both designs yield similar results, showing that the AFST reduced disparities in home removals within 3 months by approximately 1.7 percentage points (48% of the pre-existing gap). Again, effects are concentrated among high-risk referrals.

Next we discuss potential mechanisms, and provide evidence of welfare improvements among both Black and White children. We show that variation in screening disparities across workers is lower after the implementation of the AFST, suggesting that the algorithm may have improved consistency in decision-making. In theory, reductions in disparities have ambiguous welfare effects. We first show that *unwarranted* disparities were reduced after the AFST, in particular for high-risk referrals subject to the high-risk protocol. We investigate the potential welfare implications by studying the rates of false negatives and false positives, defined using downstream removals. False positives and false negative rates drop for both Black and White children after the implementation of the AFST. The reduction in false negatives is driven by high-risk referrals, while the reduction in false positives is driven by low- and medium-risk referrals.

We add to the growing literature which assesses the ways in which humans interact with algorithms within high-stakes decision-making, and the effects of that interaction on observed disparities. Several high-profile media reports have drawn attention to the potential for algorithms to discriminate.⁷ Academic research has in many cases validated these concerns, documenting and exploring the consequences of algorithmic bias in a variety of contexts, including healthcare (Obermeyer et al. 2019) and the criminal justice system (Arnold et al. 2021). For policymakers, a relevant question is whether algorithms *worsen* disparities, relative to the human-only systems they replace. Previous work addressing this question is limited, and has found mixed results (Stevenson and Doleac 2021; Albright 2019; Howell et al. 2021).

In related work, Grimon and Mills (2022) provide child protection system workers with randomized access to an algorithmic risk score. They find that providing access to the score reduced child injury hospitalizations. The pattern of their findings suggests that providing access to the algorithm score allows screeners to focus on other salient aspects of the allegations. Grimon and Mills (2022) also show that access to the tool reduced racial disparities in CPS contact, in particular among low-risk children. Our institutional context differs along several key dimensions. Their setting has a relatively small sample of Black families, with Black children making up 4% of CPS-involved children. We study disparities in a more diverse setting, where over 50% of referrals involve a Black child. Moreover, the high-risk protocol and associated default investigations in our setting allow us to speak to the effects of screening recommendations, as opposed to purely informational predicted risk scores.⁸

Prior studies in the computer science literature have studied the design and deployment of the AFST, as well as its possible implications for racial equity. Chouldechova et al. (2018) describe the development, validation, fairness auditing, and deployment of the AFST. De-Arteaga et al. (2020) study the effect of the AFST on call-screeners' decisions, finding that humans updated their behavior to align more closely with the risk score. They also study a period when a technical glitch led to some incorrect scores shown to call screeners, and find that screeners were less likely to adhere to the algorithm's recommendation in this case.⁹ Finally, Cheng et al. (2022) compare screening decisions under the AFST to a theoretical setting where the AFST is used without human input. However, this exercise is purely theoretical, as the AFST was not designed to be used without human supervision. In contrast, our paper studies the effects of the AFST as it was used in practice, relative to the prior decision-making process.

This paper also contributes to an interdisciplinary literature which attempts to understand causes of and solutions to racial disparities within child protection systems. While the existence of racial disparities is well-established, the causes for those disparities are not fully understood. Observed disparities may be driven by differences in the rates of maltreatment across race, and/or by biased decision-making. Several studies find that physicians may be more likely to report cases of potential abuse if the child is Black, and miss or overlook cases of abuse if the child is White (Lane et al. 2002; Jenny et al. 1999; Hampton and Newberger 1985). More recent work suggests that there are significant differences in risk of maltreatment across race, and that this is likely a primary driver of disparities in child protective service interactions. National differences in the rates of substantiated child abuse by race are largely consistent with racial differences in other public health outcomes, including infant mortality, low infant birth weight, and premature birth (Drake et al. 2011), suggesting that both sets of disparities may be driven by differences in exposure to the same underlying risk factors. Drake et al. (2021) provide an overview of the evidence linking poverty to child maltreatment. Not only is poverty strongly correlated with maltreatment rates, but recent evidence suggests that the link is causal.¹⁰ Several studies show that when income, or proxies for income, are controlled for, disparities in referrals, victimization rates and foster care placement rates are reduced and in some cases even reversed.¹¹ However, recent work in the economics literature suggests that gaps are at least in part a

product of discrimination: Baron et al. (2026) find in Michigan that "calls involving Black children are 55% more likely to result in foster care placement than calls involving white children with the same potential for future maltreatment in the home," with up to 19% of this unwarranted disparity attributable to call screeners. In subsequent work, Baron et al. (2024) extend this analysis to national data, finding similar patterns of unwarranted disparities across the United States.

The rest of the paper proceeds as follows. In Section 2, we describe the institutional context of the Allegheny County child protection system, as well as the AFST. Section 3 describes our data and Section 4 our empirical strategies. Section 5 presents and discusses results, Section 6 considers mechanisms and welfare implications, and Section 7 concludes.

2 | Allegheny County

Allegheny County, Pennsylvania, is home to 1.2 million people, and includes the city of Pittsburgh within its boundaries. While 13.5% of the population in Allegheny county is Black or African American, over 50% of referrals to the child protection system involve a Black child. The Office of CYF in Allegheny County is responsible for investigating allegations of child neglect and abuse. In Pennsylvania, the child protection system is state-supervised, and county-administered.

2.1 | Referral Process

Allegations of maltreatment are brought to the attention of the county by mandated reporters and community members. The State of Pennsylvania categorizes each referral under either Child Protective Services (CPS) or General Protective Services (GPS), based on the type of allegation. CPS referrals include an allegation of abuse as defined in state statute, whereas GPS referrals primarily allege neglect, implying a child may be at risk due to inadequate parental care.¹² After this classification is made by state staff, all referrals are sent to the county for further review and possible investigation. Figure 1 presents a simplified depiction of how maltreatment referrals move through the child protection system in Allegheny County. By state law, CPS referrals must always be investigated. For the remainder of referrals (GPS), staff (or "screeners") must decide whether or not to screen in the referral for investigation.

Screened-in referrals are assigned to an investigator operating out of a regional office.¹³ The investigator visits the home of the alleged victim, speaks to collateral contacts (e.g., teachers, other family members), and may gather medical and other information to evaluate the allegations of maltreatment and determine whether the child or family is in need of additional monitoring or services. State law requires that the investigation is concluded within 60 days of receipt of the report.

At any time during or after the investigation, the investigator and their supervisor decide whether or not to open a case for services. In general, opening a case indicates that the family requires ongoing services or involvement from social workers to ensure the safety of the child. An opened case might result in

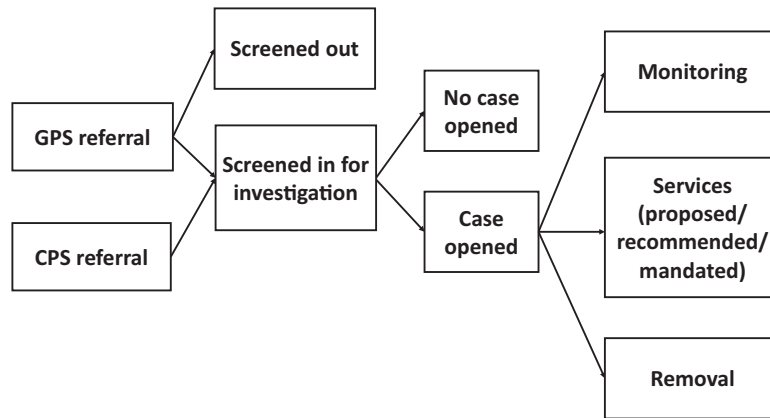


FIGURE 1 | Referral process. *Note:* This figure presents the steps by which referrals to CYF move through the child protection system in Allegheny County, PA.

continued monitoring by a social worker, voluntary or mandated participation in services, or, often as a last resort, a court-ordered removal of the child from the home. A court will order removal if there are imminent and unresolved concerns for a child's safety and well-being. Removals can occur at any time during an investigation, or after a case has been opened.

Other than the subset of referrals which are automatically screened in for investigation, a family's involvement with the child protection system is determined by the screener's decision. Prior to August 2016, screeners relied solely on professional judgment to recommend which of the GPS referrals to screen in. In making this recommendation, screeners could use both information from the referral itself (e.g., reporter, allegations, age of children), as well as data on each child and adult included in the referral from the linked Allegheny Data Warehouse (e.g., a child's history of foster care placements, adult arrest records).¹⁴ The Data Warehouse provides individual-level information on previous child protection system involvement, as well as involvement in a range of other county systems.¹⁵ However, while these data were available and the county expected information to be systematically reviewed, there was little guidance for screeners on exactly how those data should be incorporated into their screening decision.¹⁶ There was also no way for the county to confirm whether or not a screener had reviewed data to inform their decision. Call screeners' recommendations are reviewed and approved by a supervisor. Note, after making their recommendation, call screeners would not learn about any outcomes for investigated or screened-out families.¹⁷ As such, there was little opportunity for improvement in decision-making.

2.2 | Referral Process and the AFST

In August 2016, Allegheny County implemented the AFST, a predictive risk model to help screeners decide which referrals to screen in for investigation. Note, the AFST score is generated for both CPS and GPS referrals, but is only included in the decision-making process for GPS referrals. CPS referrals are screened in 100% of the time both before and after AFST implementation. As such, all analyses in this paper are limited to GPS referrals (i.e., of neglect), which make up the large majority of referrals

in our setting. Further details on the design and implementation of the algorithm are included in Subsection 2.3 below. After reviewing the allegations, as well as other historical information on the family from the Data Warehouse, the screener now runs the AFST, which shows a numerical score between 1 (lowest risk) and 20 (highest risk). Figure 2a shows an example of what the screener would see. The score is generated using only information that the screener has access to, but may not have time to review in detail and may not know how to incorporate into their assessment of safety and risk.

The use of the AFST is informed by two default decision protocols. For referrals with a score greater than 17 and at least one child aged 16 or under, the screener sees a "High Risk Protocol" notification, with no numeric score (see Figure 2b). These referrals are defaulted to be screened in for investigation, and require explicit supervisor approval to be screened out. For referrals with a score less than 11 and no children under the age of 12, the screener sees a "Low Risk Protocol" notification, again with no numeric score (see Figure 2c). These referrals are recommended to be screened out, but no supervisor approval is required to screen in these referrals.¹⁸ Low-risk protocols initially made up only 4% of referrals (Vaithianathan et al. 2019), and as such have a limited possible impact on overall screening rates. For referrals which do not meet the criteria for high- or low-risk protocols, the screener observes the numeric score and makes their own decision; that is, there is no default screening decision.

The score is not seen outside of the screening process. That is, investigators and caseworkers (who work with a family once a case is opened) do not have access to the results of the predictive risk model, and thus their downstream decisions should not be directly affected by the score. Screeners work in a centralized office, while investigators and caseworkers are based out of regional field offices, so the two groups have little chance to interact.

2.3 | Allegheny Family Screening Tool

For each child associated with a referral, the AFST predicts the risk that, if screened in for investigation, that child will experience

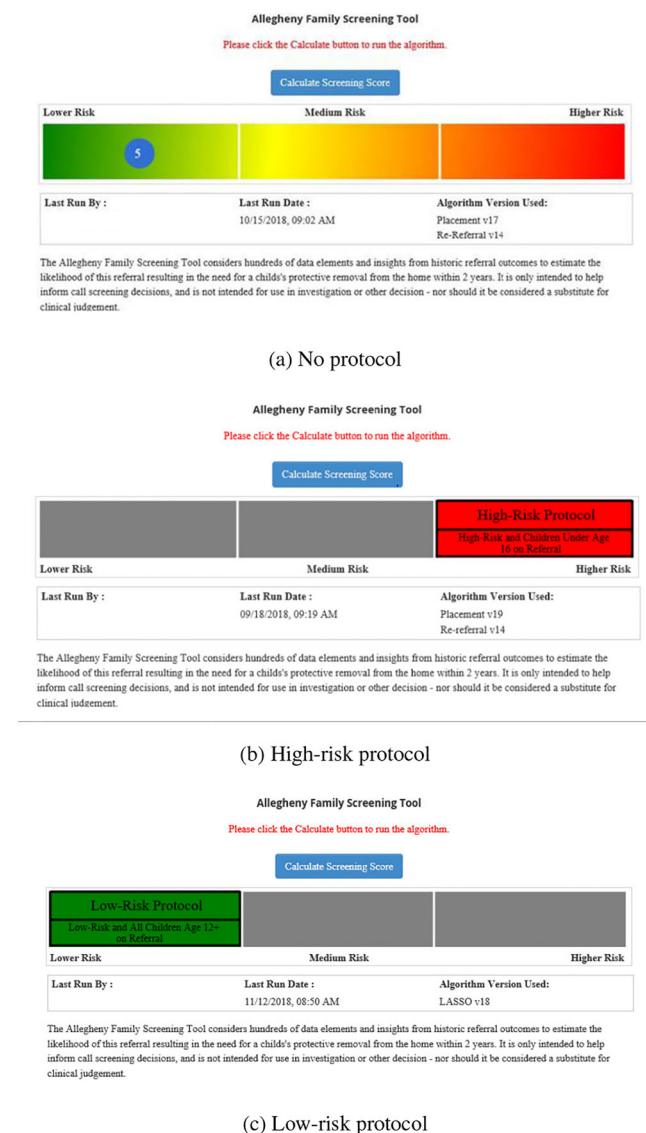


FIGURE 2 | Screener view of AFST output. *Note:* This figure presents the view that a screener has after running the algorithm, in three different scenarios. Panel (a) shows the output when neither the high-risk protocol nor the low-risk protocol is in place. Panel (b) shows the screener's view if the high-risk protocol is implemented (i.e., any child under age 16 and at least one score above 17). Panel (c) shows the screener's view if a low-risk protocol is implemented (i.e., all children are above a given age and all scores are below a given cutoff—the exact age and score cutoffs have changed over time).

a court-ordered removal from their home within 2 years. The model uses data associated with all individuals on the referral, including alleged victims and other children in the household, household members, parents, and alleged perpetrators. Data from past referrals and interactions with the child protection system, past and present involvement with the courts, jail and other county systems, as well as information from the child's birth record, are all used to generate a risk score. Note, race is not included as a predictor. A risk score is generated for each child living in the home of the alleged victim, but the screener only observes the maximum of these scores.¹⁹ Appendix B includes additional details on the construction and validation of the AFST.

Allegheny County Department of Human Services developed the AFST with the purpose of using existing data to improve the quality and consistency of screening decisions.²⁰ Importantly, the tool was never meant to replace human decision-making, but rather to inform and improve those decisions. That is, the tool was intended to be complementary to the call screener's professional judgment.

In August 2016, Allegheny County deployed the AFST. To the best of our knowledge, the timing of the implementation was not related to any other changes in the county's processes related to the screening or investigation of maltreatment allegations. Since then, they have updated the predictive risk model and the screening tool several times. In November 2018, the original model was updated with weights estimated using LASSO.²¹ In January 2019, the LASSO model was updated in response to a change in the data that were available to the model. Goldhaber-Fiebert and Prince (2019) were contracted to conduct an independent impact evaluation of the original AFST, and studied effects on accuracy, workload, disparities, and consistency.²²

3 | Data

We obtained de-identified administrative data from Allegheny County on the universe of child maltreatment referrals (both CPS and GPS) made between 2015 and 2020. For GPS referrals, we also obtain data on referrals made between 2010 and 2015.²³ We restrict our analysis to referrals made before October 2019, in order to observe at least a 6-month follow-up for each referral before the child protection system was severely impacted by COVID-19 in March 2020. Every referral links to unique IDs for each person living in the household, including the alleged victim, other children, parents and alleged perpetrators. For each alleged victim, the referral lists one or more allegations of abuse or neglect. For each individual, we observe demographic information including race, gender, and age. We also observe outcomes within the child protection system, including screening decision and foster care placements. We observe a unique ID for each call screener, but do not observe any demographic or other information about the workers.

Screening decisions occur at the referral, rather than individual, level. As such, we collapse data to the referral level in our main analyses. In order to study effects on removals (which occur at the individual level), we create a variable equal to one if any child on the referral is removed from their home within 3 months of the referral date, and zero otherwise. We choose 3 months since investigations are required to be completed within 60 days of referrals, and as such any removals associated with a given referral are likely to occur approximately within this time. We also define race at the referral level, classifying a referral as Black if at least one child on that referral is identified as Black or African American, and classifying a referral as White if there are no Black children and at least one White child on the referral. Referrals with no children identified as either Black or White make up approximately 6% of referrals from 2010 through 2020 and are excluded from our analysis sample.

Table 1 presents summary statistics separately for GPS (all and screened-in only) and CPS referrals.²⁴ We exclude referrals

TABLE 1 | Summary statistics.

	All GPS	Screened-in GPS	All CPS
Panel A: Demographics			
Black	0.498	0.552	0.510
Any infant	0.150	0.221	0.118
Any child age 1–5	0.454	0.476	0.443
Any child age 6–12	0.582	0.587	0.636
Any child age 13–17	0.388	0.391	0.438
Number of children	2.267	2.408	2.394
Panel B: Referral outcomes			
Screened in	0.444	1.000	0.997
Removed within 3 months	0.054	0.092	0.053
Panel C: Allegation categories			
Abandonment	0.009	0.011	0.003
Caregiver behavioral issues	0.035	0.047	0.012
Caregiver substance abuse	0.215	0.287	0.035
Causing death of child	0.004	0.006	0.007
Child behaviors	0.055	0.060	0.023
Domestic violence	0.068	0.081	0.026
Exposure to risk	0.096	0.105	0.022
Failure to protect	0.058	0.055	0.015
Imminent risk	0.035	0.040	0.015
Inadequate physical care	0.294	0.265	0.032
Medical neglect	0.036	0.046	0.017
Mental health	0.051	0.040	0.024
Mental injury	0.028	0.033	0.020
Neglect	0.108	0.108	0.014
No/inadequate home	0.108	0.139	0.012
Parent/child conflict	0.053	0.046	0.018
Physical altercation	0.015	0.015	0.023
Physical maltreatment	0.091	0.047	0.643
Sexual abuse or exploitation	0.020	0.010	0.177
Sexual contact between children	0.040	0.024	0.005
Truancy	0.042	0.068	0.006
Unwilling or unable to provide care	0.072	0.079	0.013
Youth substance abuse	0.011	0.009	0.003
Panel D: Reporter categories			
Agency	0.217	0.229	0.274
Anonymous	0.125	0.088	0.032
Community	0.050	0.042	0.022
Family	0.172	0.163	0.068
Law	0.120	0.187	0.086
Medical	0.093	0.106	0.158

(Continues)

TABLE 1 | (Continued)

	All GPS	Screened-in GPS	All CPS
School	0.134	0.126	0.159
Self	0.004	0.004	0.012
Therapist	0.078	0.047	0.186
<i>N</i>	79725	35416	13925

Note: This table reports means of referral characteristics separately for all GPS, screened-in GPS, and all CPS referrals. Note, since 100% of CPS referrals are screened in, averages for screened-in CPS referrals are identical to those in the full CPS sample. The sample is all referrals made between January 2010 and September 2019 to CYF, excluding active-family referrals, referrals involving neither Black children nor White children, and referrals stemming from truancy courts.

without a CPS or GPS designation (about 3% of referrals from 2010 to 2020). Our sample is comprised of 93,650 referrals, 85% of which are classified as GPS. We do not include referrals for families which, at the time of referral, have an active case with CYF (about 9% of referrals from 2010 to 2020). We also exclude referrals made by truancy courts (about 1% of referrals from 2010 to 2020). Note also that 100% of CPS referrals are investigated, reflecting the automatic screen in for this category.

Screen-in rates and racial disparities in Allegheny are reflective of national data. Among our combined sample of CPS and GPS referrals, 53% are screened in for investigation. Nationwide in 2019, the average screen-in rate was 54.5% (Administration on Children and Families 2021).²⁵ While only 13% of the population in Allegheny County is Black, approximately 50% of referrals involve a Black child. While we do not have state or national data on referrals by race, the rate of maltreatment victimization is consistently higher for Black children than other race groups reported by NCANDS. In 2019, 13.8 per 1000 Black children nationwide were found to be victims of child maltreatment, relative to 7.8 per 1000 White children and 8.9 per 1000 children overall.

For each referral, we observe one or more algorithm-generated scores. First, we observe the score as calculated by the algorithm in use at the time of referral. In addition, we observe retroactively calculated scores generated by AFST V3, the most recent version of the model as of this writing, for all referrals made after January 2013. Going forward, we primarily focus on the V3 score in order to allow for comparability across years. For referrals made between January 2013 and July 2019, this comparable score was retroactively calculated, and can be thought of as the score that the screener *would have* seen, had this version of the algorithm been deployed.²⁶ Scores are strongly correlated across AFST versions. Figure A1 shows a heatmap of the correlation between retroactively calculated V3 scores and each of V1 and V2 scores.

We separately calculate summary statistics by child race and referral risk bin in Table A1. We classify scores as “low risk” (1–12), “medium risk” (13–17), and “high risk” (18–20). The sample is limited to GPS referrals. Note that across the risk bins, referrals involving Black children are more likely to be screened in, and more likely to be removed within 3 months. This may in part reflect differences in referral characteristics across race, within risk bin. For example, referrals involving Black children tend to involve more children, including more young children. To account for these differences across race, we control for referral-level characteristics in our regression analysis.

4 | Empirical Framework

4.1 | Screening

To test whether the algorithm changed disparities in screening rates, we exploit the timing of the AFST’s introduction, comparing GPS referrals involving Black children to those involving White children before and after the algorithm was implemented. While both groups were “treated” by the AFST, this approach tests the hypothesis that they were differentially affected.

We begin by estimating the following regression equation:

$$y_{it} = \beta_0 Black_{it} + \beta_1 Post_{it} + \beta_2 Black \times Post_{it} + \beta_3 X_{it} + \gamma_m + \gamma_y + \epsilon, \quad (1)$$

where y_{it} is an indicator equal to one if referral i , made in month t , is screened in for an investigation; $Black_{it}$ is an indicator equal to one if there are any Black children listed on referral i , and zero if there are no Black children, and at least one White child, listed on referral i ; and $Post_{it}$ is an indicator variable equal to one if referral i in month t was made after the implementation of the AFST, and zero otherwise. We include fixed effects for Year (γ_y) and Month-of-Year (γ_m), to control for variation across time and seasonality. Finally, X_{it} is a vector of referral-level controls, which includes allegation and reporter category indicators, indicators for child age groups, the number of children associated with the referral, and an indicator for any drug or alcohol concerns.²⁷ We include these variables to control for any differential trends across race and time. For example, substance exposure might differ across both race and time, as the opioid crisis has affected primarily White communities.

The identification assumption required for β_2 to be interpreted as the causal effect of the AFST on disparities is that there are no differential trends in screening rates for referrals involving Black versus White children. This assumption would be violated, for example, if there are other efforts to reduce disparities in screening decisions around the time that AFST was introduced, or if Black and White families are differently affected by shifting local conditions. One salient change around this time period was the 2014 amendment to Pennsylvania’s CPS Law. Among other changes, this amendment expanded both the definition of child abuse and the categories of mandated reporters. While we do not expect this change to differentially impact trends in screening decisions across race, we explore the possibility of differential pre-trends both in the raw data (Figure A2) and in a more formal event study (discussed below in Section 5). Moreover, we include in our main specification indicators for allegation and reporter

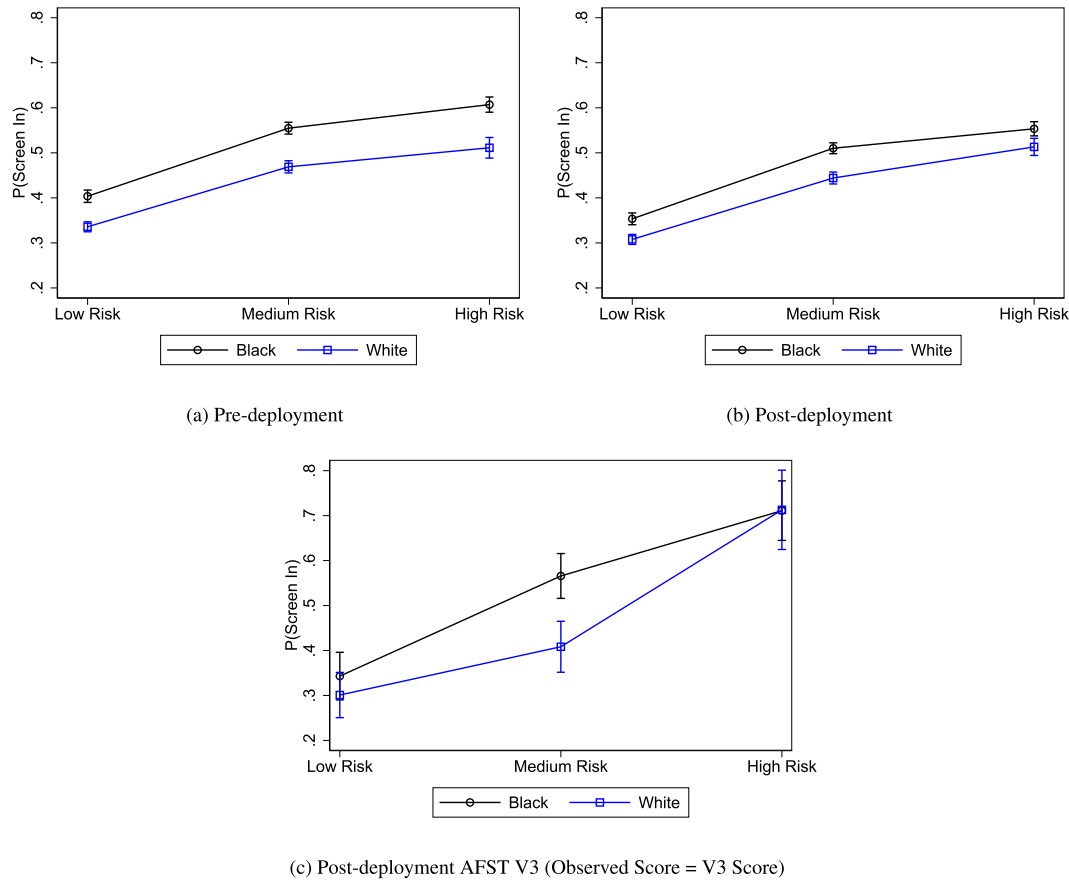


FIGURE 3 | Screen in disparities by score bin. *Note:* This figure reports average screen-in rates by algorithm score bin and race. Screen-in rate is defined as the share of referrals which are screened in for an investigation. See the notes to Table 1 for a description of the analysis sample. The sample is further restricted to referrals made in or after January 2013, due to limited availability of retroactively calculated scores. For each score bin and race, 95% confidence intervals are shown. Each sub-figure presents results from a different time period. Panel (a) presents screen-in rates from the period before the algorithm was implemented, from January 2013 through June 2016. The algorithm score bin in this panel comes from retroactively calculated AFST scores using AFST V3. Panel (b) presents screen-in rates from the period after the algorithm was implemented, from July 2016 through Sept. 2019. The algorithm score bin in this panel was also calculated using AFST V3. Panel (c) presents screen-in rates from the period after AFST V3 was implemented, that is, for July 2019 through September 2019. The algorithm score bin in this panel uses the AFST V3 score as seen by the call screeners.

categories, which control for changes in screening rates driven by changes in the composition both of mandated reporters and maltreatment types.

Finally, effects may differ across the range of algorithm scores. Given the low-risk and high-risk protocols, we might expect that effects on screening would be concentrated at these ends of the risk score distribution. To explore this heterogeneity, we classify scores as “low risk” (1–12), “medium risk” (13–17), and “high risk” (18–20).²⁸ This classification aligns with the risk score requirements of the most recent low- and high-risk protocols. We plot screening decisions for low-, middle-, and high-risk scores, before and after the AFST deployment and separately for referrals involving Black versus White children, in Figure 3. We show this comparison first using the entire post-period in Figure 3b, and then defining the post-period as only the time when AFST V3 was in place (i.e., after July 2019) in Figure 3c. Note that the change to the racial gap in screen-in rates seems particularly stark for high-risk referrals. Motivated by this observation, we look for heterogeneous effects according to algorithm risk scores, by interacting the indicator variables in Equation (1) with indicators

for each algorithm risk bin S^i , $i \in \{1, 3\}$, estimating the equation:

$$y_{it} = \sum_{j \neq 2}^3 \beta_j S_{it}^j + \sum_{j \neq 2}^3 \beta_{j+2} S_{it}^j \times Black_i + \sum_{j \neq 2}^3 \beta_{j+5} S_{it}^j \times Post_t + \sum_{j \neq 2}^3 \beta_{j+8} S_{it}^j \times Black_i \times Post_t + \beta_{11} X_{it} + \gamma_m + \gamma_y + \epsilon, \quad (2)$$

where y_{it} is an indicator equal to one if a referral is screened in for investigation, and zero otherwise. This approach allows us to test how the algorithm affected disparities in screen-in rates differently by comparable risk score.

4.2 | Removals

We next turn to the downstream outcome of home removal. Although home removals are not directly impacted by the AFST, they may be indirectly impacted through a change in the composition of screened-in referrals. For this analysis we conduct two empirical exercises, with different strengths and weaknesses.

TABLE 2 | Screen in.

	(1) DD	(2) +FE	(3) +Controls	(4) Restricted sample	(5) +Risk score
Post	−0.00652 (0.00506)	0.0252** (0.0120)	0.00988 (0.0109)	0.0137 (0.0112)	0.0139 (0.0112)
Black	0.109*** (0.00446)	0.108*** (0.00446)	0.0852*** (0.00419)	0.0873*** (0.00558)	0.0641*** (0.00563)
Black × Post	−0.0334*** (0.00720)	−0.0331*** (0.00720)	−0.0280*** (0.00662)	−0.0317*** (0.00753)	−0.0310*** (0.00751)
Year FE	No	Yes	Yes	Yes	Yes
Month-of-year FE	No	Yes	Yes	Yes	Yes
Score FE	No	No	No	No	Yes
Controls	No	No	Yes	Yes	Yes
Data span	2010–2019	2010–2019	2010–2019	2013–2019	2013–2019
Mean	0.444	0.444	0.444	0.444	0.446
Obs.	79725	79725	79725	58273	57825

Note: Standard errors in parentheses. This table reports coefficients and standard errors from estimating four specifications of Equation (2). The sample is described in the notes to Table 1. In Columns 4 and 5, the sample is further restricted to referrals made in or after January 2013, when retrospective AFST V3 scores are available. The outcome variable in each regression is an indicator equal to one for referrals which are screened in, and zero otherwise.

* $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

We start by estimating Equation (1), where y_{it} is an indicator equal to one if any child associated with referral i , made in month t , is removed from their home within 3 months. The identification strategy requires the parallel trends assumption that removals evolve over time similarly for referrals involving Black versus White children. We test this assumption directly using an event-study specification of Equation (1).

We also estimate a difference-in-differences model, using CPS referrals as the control group. For this analysis, we use a subsample of referrals made between 2015 and 2020, as CPS referrals made before 2015 are expunged in our data. This model is estimated with the following equation:

$$\begin{aligned}
 y_{it} = & \beta_0 Black_{it} + \beta_1 GPS_{it} + \beta_2 Post_{it} + \beta_3 Black \times GPS_{it} + \beta_4 Black \\
 & \times Post_{it} + \beta_5 GPS \times Post_{it} + \beta_6 Black \times Post \times GPS_{it} + \beta_7 X_{it} \\
 & + \gamma_m + \gamma_y + \epsilon,
 \end{aligned}
 \tag{3}$$

where everything is defined as in Equation (1), and GPS_{it} is an indicator equal to one if referral i falls under GPS, and zero if referral i falls under CPS. The coefficient of interest, β_6 , tells us how removal rates change for Black children, relative to White children, in the treated group relative to the control group of referrals. This specification relaxes the required assumption from the previous specification, by allowing us to control for changes in racial disparities in removals which affect both CPS and GPS referrals. For example, if CYF implemented other policies to reduce racial disparities around the same time as the algorithm was introduced, Equation (1) would not be able to distinguish the effect of the algorithm from that of other efforts. By using CPS referrals as a control group unaffected by the algorithm,

Equation (2) allows us to distinguish those effects. Note, this specification is not possible for estimating effects on screening decisions, as all CPS referrals are screened in both before and after algorithm implementation.

The identification assumption for this specification requires common trends in racial disparities for CPS and GPS referrals.²⁹

5 | Results and Discussion

5.1 | Screening

In Table 2, we report results from estimating Equation (1), or the effects of the algorithm on disparities in screen-in rates. Column 1 reports results from a specification excluding fixed effects and controls, Column 2 adds year and month-of-year fixed effects, Column 3 adds referral-level controls, Column 4 replicates Column 3 for the time period in which we have AFST V3 scores, and Column 5 adds an additional fixed effect for the underlying AFST V3 score. First, note that referrals involving Black children are more likely to be screened in than referrals involving White children. There is a 10.9 percentage point gap in screen-in rates before controlling for referral-level characteristics, which is reduced to 6.4 percentage points after controlling for referral-level characteristics and underlying risk scores. The algorithm reduces the unconditional gap by 3.3 percentage points (Column 1), or 31%. It reduces the conditional gap by 3.1 percentage points (Column 5), or 48%.

We next assess the robustness of our results to two additional sample restrictions. First, we limit our sample to referrals in which no child has been referred in the previous 12 months.

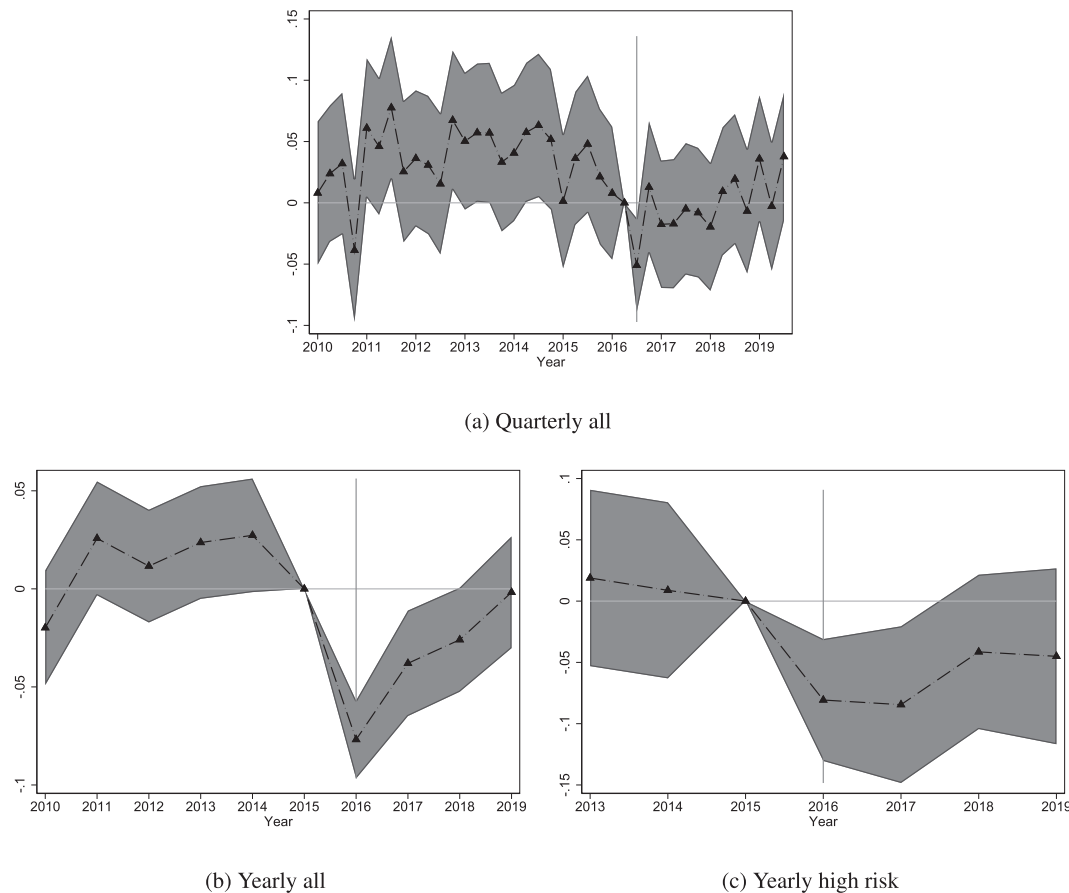


FIGURE 4 | Event study: Screening. *Note:* This figure presents coefficients and 95% confidence intervals from estimating an event study version of Equation (1), where the outcome variable is an indicator equal to one if the referral was screened in and zero otherwise. Panel (a) presents a quarterly event study, where 2016Q2 is the excluded quarter. Panels (b) and (c) present annual event studies, where the omitted year is 2015. Panel (c) restricts the sample to high-risk referrals (i.e., with AFST scores above 17). Retroactively scores are only available beginning in 2013.

Second, we limit the sample to referrals *not* made by CYF workers. Results are shown in Panel A of Table A1. Column 1 presents our baseline results, and Columns 2 and 3 present results for each sample restriction. Effects are similar and remain statistically significant across samples.

Figure 4a presents coefficients from a quarterly event study version of Equation (1), where $Black_i$ is interacted with quarterly indicator variables.³⁰ While the estimated coefficients are noisy, there appears to be a shift in average screen-in disparities around the time of the AFST deployment.³¹ A version of the event study aggregated to the annual level presents a less-noisy story, suggesting screening disparities dropped sharply and significantly the year the AFST was introduced, and rebounded over time (see Figure 4b).³²

Motivated by the patterns in screening disparities observed in Figure 3, we next study how effects on screening decisions vary across algorithm scores. Figure 5 shows coefficients and 95% confidence intervals from estimating Equation (2). Figure 5a shows a positive relationship between algorithm score bin and screen-in rate. Figure 5b shows that referrals involving Black children are more likely to be screened in at every risk score. Figure 5c suggests

that the algorithm primarily increased screen-in rates for referrals with the highest risk scores, consistent with the intended effect of the high-risk protocol. Finally, Figure 5d graphically presents the coefficients on $Black \times Post \times Score^i$. The effect on screening disparities is negative for every score bin, with varying magnitudes. For referrals in the highest risk bin, the algorithm reduced the gap in screening rates across race by 6.0 percentage points, or 65% of the pre-existing disparity in this score bin.³³ An annual event study restricted to the high-risk bin is presented in Figure 4c suggests that these reductions in disparities were sustained over time.³⁴

Referrals within this high-risk bin are most often defaulted to be screened in through the high-risk protocol under AFST (see Section 2.2).³⁵ Recall, although referrals in this category may still be screened out, this decision requires a supervisor's override. This result suggests that the protocol plays an important role in changing screening outcomes.

The differential effects observed across risk bins provide additional evidence for the required parallel trends assumption. That is, any differential trends across race would have to differ across risk bins in order to explain our pattern of findings.

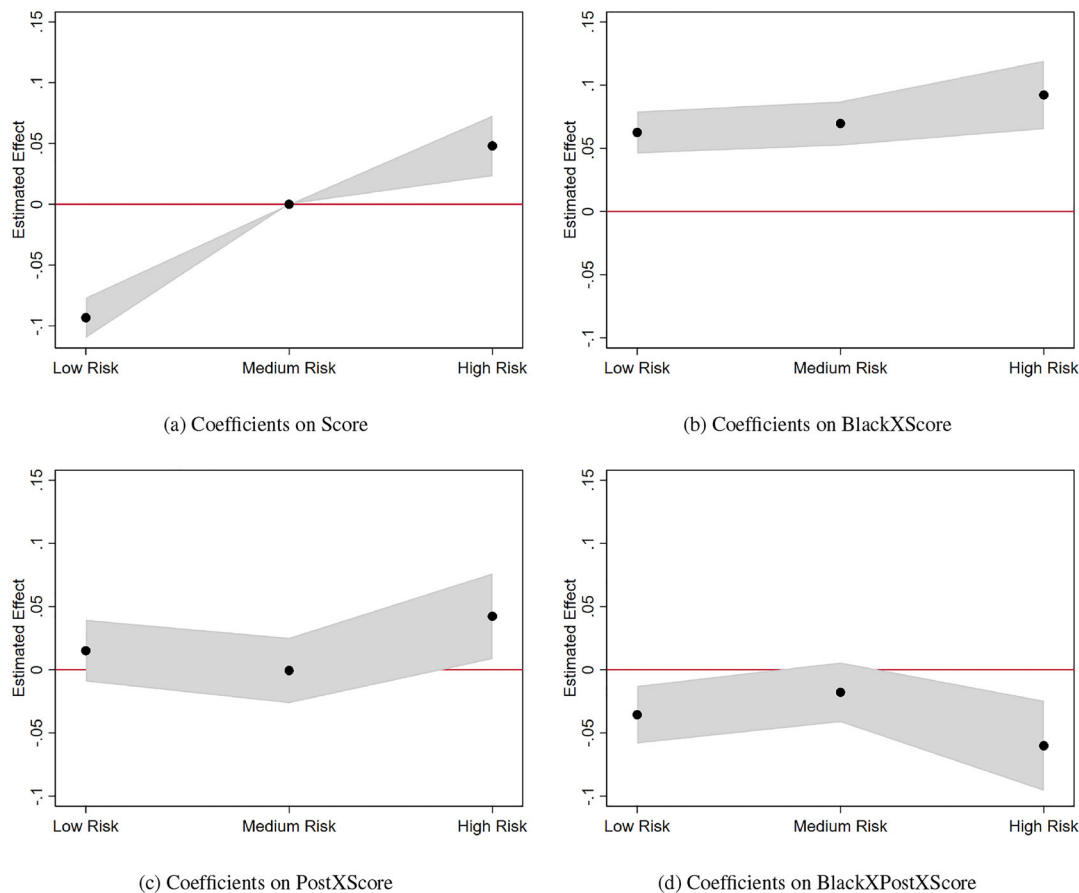


FIGURE 5 | Screening by score bin. *Note:* This figure graphically presents coefficients and 95% confidence intervals from estimating Equation (2), where the outcome variable is equal to one if the referral is screened in, and zero otherwise. Panel (a) plots coefficients β_1 and β_2 , representing the relationship between algorithm score and probability of screen in. Panel (b) plots coefficients β_3 through β_5 , on the interaction of algorithm score bin and $Black_i$. Panel (c) plots coefficients β_6 through β_8 , on the interaction of algorithm score bin and $Post_t$. Finally, Panel (d) plots coefficients β_9 through β_{11} , on the interaction of algorithm score bin, $Black_i$ and $Post_t$. See notes to Table 1 for a description of the analysis sample. The sample is further restricted to referrals made in or after January 2013, when retrospective AFST V3 scores are available.

5.2 | Removals

We next turn to the effects of the algorithm on home removals. We start by comparing referrals across race, before versus after the implementation of the AFST, using only the GPS referrals. Results from estimating Equation (1), where y_{it} is an indicator for home removal within three months of referral are reported in Column 1 of Table 3. Referrals involving Black children are on average 3.7 percentage points more likely to involve a child who is removed from their home within 3 months (72% of the sample mean). The coefficient on $Black \times Post$ suggests that the implementation of the algorithm reduced disparities in removal rates by 1.9 percentage points, or 51% of the pre-existing Black–White gap.

Figure 6a presents coefficients from a quarterly event study versions of Equation (1).³⁶ Again, there appears to be a shift in removal disparities around the time of AFST implementation. An annual version of the event study presents a less-noisy picture (Figure 6b), showing a sharp drop in disparities beginning in 2016.

Next, we incorporate CPS referrals as a control group in a difference-in-differences specification. For this analysis, we must

restrict our sample to referrals made in or after 2015. For easier comparison across specifications, we also estimate Equation (1) using this shorter time span. Results from estimating Equation (1) on this restricted sample are reported in Column 2 of Table 3, and are statistically indistinguishable from the baseline results in Column 1.

Results from estimating Equation (3), where y_{it} is an indicator for removal within 3 months of referral, are reported in Columns 3 through 5 of Table 3. Column 3 reports results without fixed effects or controls, Column 4 reports results from a regression which adds year and month-of-year fixed effects, and Column 5 reports results from a regression which adds referral-level controls (our preferred specification). Our coefficient of interest, on the triple interaction $Black \times Post \times GPS$, is significant at the 10% level, and stable across specifications. In our preferred specification (reported in Column 5), the coefficient implies that the introduction of the algorithm reduced the racial disparity in 3-month removal rates by 1.7 percentage points, or 48% of the pre-existing difference.³⁷

We again limit the sample to referrals with no CYF calls in the previous 12 months, and those *not* made by CYF workers. Results

TABLE 3 | Results: Placed within 3 months.

	(1) DD	(2) Restricted sample	(3) DDD base	(4) +FE	(5) +Controls
Post	0.0149*** (0.00527)	0.0182*** (0.00562)	0.00494 (0.00463)	0.0224*** (0.00620)	0.0194*** (0.00610)
Black	0.0371*** (0.00210)	0.0357*** (0.00383)	0.0294*** (0.00654)	0.0297*** (0.00653)	0.0264*** (0.00634)
Black × Post	−0.0190*** (0.00317)	−0.0179*** (0.00448)	−0.00217 (0.00796)	−0.00238 (0.00796)	−0.000880 (0.00777)
GPS			0.00434 (0.00427)	0.00444 (0.00426)	−0.00942** (0.00451)
Black × GPS			0.00961 (0.00758)	0.00896 (0.00758)	0.00893 (0.00739)
GPS × Post			−0.000242 (0.00531)	−0.000276 (0.00531)	0.000626 (0.00525)
Black × Post × GPS			−0.0164* (0.00916)	−0.0158* (0.00915)	−0.0170* (0.00896)
Year FE	Yes	Yes	No	Yes	Yes
Month-of-year FE	Yes	Yes	No	Yes	Yes
Controls	Yes	Yes	No	No	Yes
Data span	2010–2019	2015–2019	2015–2019	2015–2019	2015–2019
Mean	0.0543	0.0469	0.0463	0.0463	0.0459
Obs.	79725	43399	55262	55262	55134

Note: Standard errors in parentheses. This table reports coefficients and standard errors from five separate regressions estimating Equation (1) (Columns 1 and 2) and different specifications of Equation (3) (Columns 3–5). The sample is described in the notes to Table 1. The sample is further restricted to referrals made before October 2020 to allow for a 3-month follow-up for each referral. In Columns 2–5, the sample is further restricted to referrals made in or after January 2015. The outcome variable in each regression is an indicator equal to one for referrals which are associated with any child who is removed within 3 months of the referral date (regardless of whether the referral resulted in an investigation), and zero otherwise. Controls include allegation and reporter category indicators, indicators for child age groups, the number of children associated with the referral, and an indicator for any drug or alcohol exposure.

* $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

are presented in Panel B of Table A2. Effects are similar and statistically significant in each sample.

To study where in the risk distribution effects are concentrated we estimate Equation (2), setting the outcome variable equal to an indicator for removal within 3 months of referral. For this analysis, we restrict our sample to referrals made in or after 2013, as we have comparable scores for these dates. The coefficients on $Score_i$, $Black \times Score_i$, $Post \times Score_i$, and $Black \times Post \times Score_i$ are shown in Figure 7. In Figure 7b, the race differential in likelihood of removal is most stark for the highest risk referrals. The effect on disparities is also concentrated in the highest risk bin, as shown in Figure 7d. An annual event study specification restricted to high-risk referrals is presented in Figure 6c, and shows a sharp and sustained drop in disparities beginning in 2016.

Recall that the AFST cannot directly impact removal decisions. As such, these results imply that the introduction of the algorithm changes the composition of investigations, in such a way as to reduce the inequality in removal likelihood across race. This may occur either through a reduction in the investigation rate for Black children at high risk of removal, an increase in the

investigation rate for White children at high risk of removal, or some combination of these two channels. In Section 6, we explore the welfare implications of the AFST's impacts on screening decisions, and provide suggestive evidence on the nature of these compositional changes.

6 | Mechanisms and Welfare

One goal of using predictive risk models in the child welfare setting is to improve consistency across decision-makers. We explore how the AFST affected individual call screeners, studying racial disparities by worker before and after the AFST is implemented. We first limit our sample to call screeners who handle at least 20 calls in each of the pre- and the post-period. For each of these 29 screeners, we calculate the raw gap in screening rate by race.

Figure 9 plots the distribution of disparities across screeners, before and after the AFST is implemented. In the pre-period, there are several workers with large gaps in screening rates by race. In the post-period, there is less variation in disparities across

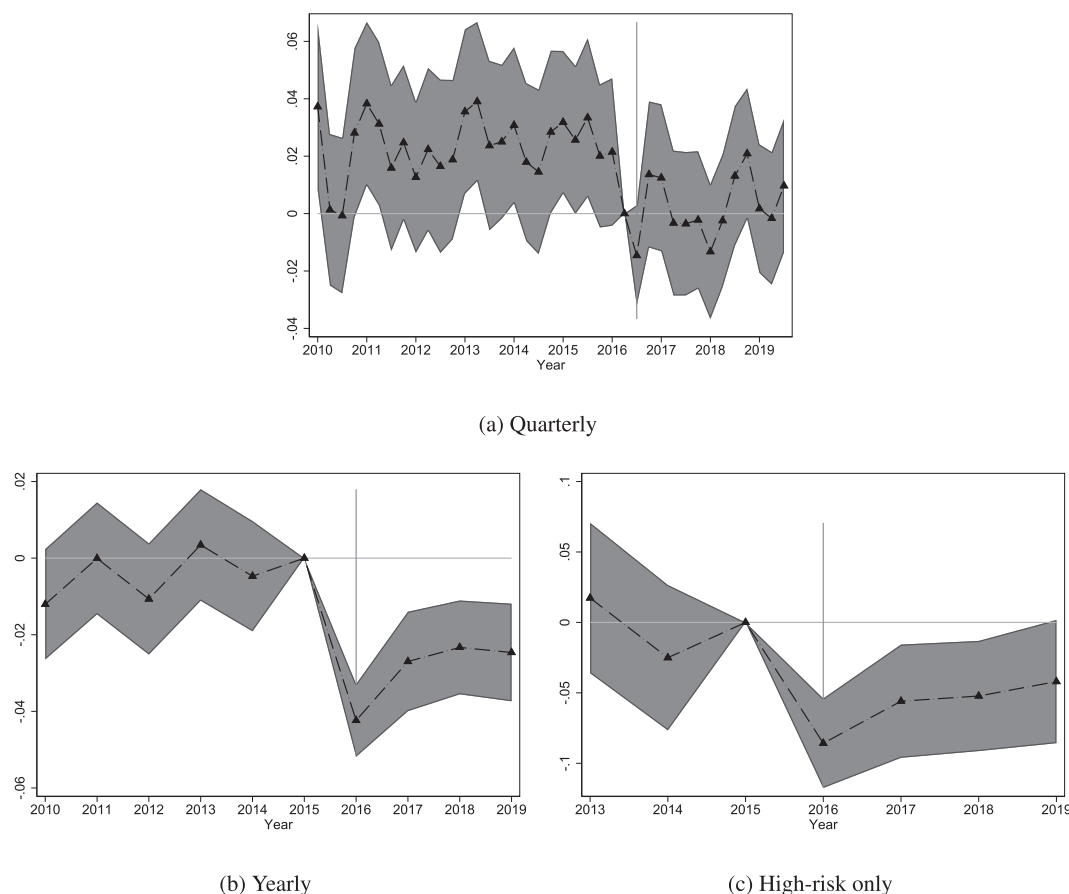


FIGURE 6 | Event study: Removals. *Note:* This figure presents coefficients and 95% confidence intervals from estimating an event study version of Equation (1), where the outcome variable is an indicator equal to one if any child on the referral was removed within three months and zero otherwise. Panel (a) presents a quarterly event study, where 2016Q2 is the excluded quarter. Panels (b) and (c) present annual event studies, where the omitted year is 2015. Panel (c) restricts the sample to high-risk referrals (i.e., with AFST scores above 17). Retroactively scores are only available beginning in 2013.

screeners. This pattern is consistent with the AFST increasing consistency across workers, and reining in screeners whose disparities are furthest from the mean.³⁸

One may be concerned that a reduction in observed disparities is harmful if the disparities were in fact warranted. To address this concern, we adopt a measure of unwarranted disparities proposed in Baron et al. (2026). We estimate unwarranted disparities in screening rates before and after AFST implementation.³⁹ We define unwarranted disparity as differences in screening rates across race conditional on potential for future maltreatment, where future maltreatment is proxied for by foster care placement within 6 months if left at home.⁴⁰ Given that screeners do not appear to be randomly assigned in our setting, we estimate bounds for unwarranted disparities, again following the approach in Baron et al. (2026).⁴¹

Prior to the implementation of the AFST, referrals involving Black children were 8.8%–10.4% more likely to be screened in than referrals involving White children with the same potential for future maltreatment. After the implementation of the AFST this rate dropped to 6.2%–7.2%. Figure A4 shows the ranges of unwarranted disparities in the pre- and post-periods, for each

risk bin. Note that pre-AFST unwarranted disparities are largest in the high-risk bin, and reductions in unwarranted disparities after the implementation of the AFST are largest in this risk bin as well, consistent with the effects on raw screening disparities. This analysis again suggests that the high-risk protocol, which defaults referrals above a certain risk score to an investigation, plays an important part in minimizing the role of human bias in decision-making.

Still, the welfare effects of reductions in disparities are a priori ambiguous. A reduction in unwarranted disparities implies more consistent, but not necessarily more accurate decision-making. If this result is achieved through a reduction in screen-ins for children who would benefit from an investigation, or an increase in screen-ins for children who do not benefit from an investigation, children may be harmed even by a reduction in unwarranted disparities. We do not know the costs and benefits of investigation, or how those values differ across race and other observable characteristics. Instead, we turn to proxy measures in order to classify screening decisions as “false negatives” (referrals which were screened out, but would have led to benefits for the children if screened in) and “false positives” (referrals which were screened in unnecessarily).

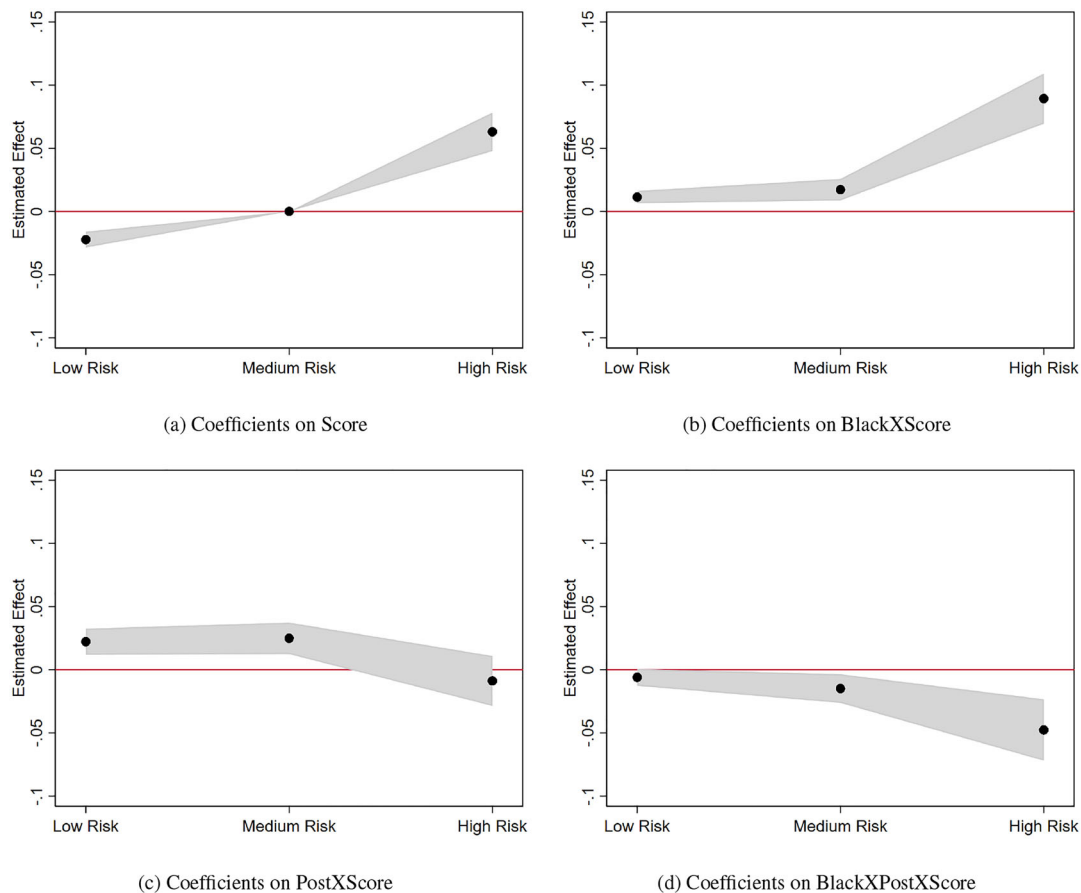


FIGURE 7 | Results: Removal (3m) by score bin. *Note:* This figure graphically presents coefficients and 95% confidence intervals from estimating Equation (2), where the outcome variable is equal to one if any child associated with the referral is placed within 3 months, and zero otherwise. Panel (a) plots coefficients β_1 and β_2 , representing the relationship between algorithm score and probability of removal. Panel (b) plots coefficients β_3 through β_5 , on the interaction of algorithm score bin and $Black_i$. Panel (c) plots coefficients β_6 through β_8 , on the interaction of algorithm score bin and $Post_i$. Finally, Panel (d) plots coefficients β_9 through β_{11} , on the interaction of algorithm score bin, $Black_i$ and $Post_i$. See notes to Table 1 for a description of the analysis sample. The sample is further restricted to referrals made before October 2020, to allow for a 3-month follow-up, and to referrals made in or after January 2013, when retrospective AFST V3 scores are available.

We define a referral as a false negative if it is screened out and any child from the referral is placed in foster care within 6 months. This implies that the child was left at home in circumstances that led to severe maltreatment.⁴² We use two definitions of false positives. For consistency with our definition of false negatives, we first include as false positives referrals which are screened in with no child from the referral placed in foster care within 6 months. This measure is somewhat difficult to interpret, as it may be the case that children benefit from an investigation even if they are not placed in foster care. Earlier intervention may even help to reduce the need for home removal. A reduction in false positives as defined here would suggest that fewer non-severe cases of maltreatment are investigated, while an increase in false positives would indicate that a wider net is being cast. To help address these limitations, we also analyze an alternative definition of false positives: referrals which are screened in but have no substantiated allegations. Table 4 presents false negative and positive rates, before and after the implementation of the AFST and separately for referrals involving Black versus White children. The table also reports p -values from a one-sided t -test of the null hypothesis that false positive or negative rates are

higher in the period after AFST implementation. False negative rates fall by 22% for referrals involving White children, and by 11% for referrals involving Black children. False positive rates (defined using foster care placements) fall by 4% for both referrals involving White and Black children. False positive rates (defined using substantiations) fall by 18% for both referrals involving White and Black children. Each of these differences is statistically significant. We are also able to study false positive rates for our control group of CPS referrals. Among this group, false positive rates (under both definitions) actually increased for both referrals involving Black and White children. This is reassuring in that the patterns observed among the treated GPS referrals cannot be explained by overall trends in reporting or removals that would affect both GPS and CPS referrals. Unfortunately, we are not able to conduct a similar placebo analysis for false negative rates, as no CPS referrals are screened out.

These patterns provide suggestive evidence on the nature of the compositional changes brought about by the AFST. Although all children see a reduction in false negative rates, the change is larger for White children (22%) than for Black children (11%).

TABLE 4 | False negatives and positives.

	Pre	Post	p-Value
Panel A: GPS referrals			
False negative (White)	0.0150	0.0117	0.0068
False negative (Black)	0.0302	0.0270	0.047
False positive FC (White)	0.3703	0.3558	0.0053
False positive FC (Black)	0.4282	0.4120	0.0026
False positive Subst. (White)	0.2169	0.1782	0.0000
False positive Subst. (Black)	0.2479	0.2022	0.0000
Panel B: CPS referrals			
False negative (White)	N/A	N/A	N/A
False negative (Black)	N/A	N/A	N/A
False positive FC (White)	0.9423	0.9570	0.995
False positive FC (Black)	0.8927	0.9143	0.998
False positive Subst. (White)	0.7926	0.8823	1.000
False positive Subst. (Black)	0.7619	0.8718	1.000

Note: This table reports false positive and false negative rates. As described in Section 6, false negatives are defined as referrals which are screened out and involve a child who is removed from their home in the following 6 months. We show two definitions of false positives. False positive (FC) describes referrals which are screened in and for which no child is removed from their home within the following 6 months. False positive (Subst.) describes referrals which are screened in and which do not result in any substantiated allegations. Panel A shows rates of false negatives and positives among GPS referrals, separately for referrals involving Black versus White children, before and after AFST implementation. Panel B reports false positive rates for CPS referrals, separately for referrals involving Black versus White children, before and after AFST implementation. The third column presents *p*-values from a one-sided *t*-test of the null hypothesis that false positive (or negative) rates are higher in the post-period.

This is consistent with the AFST disproportionately adding White referrals at high risk of removal to the investigation pool.

We further investigate these patterns by risk bin in Figure 8. Panels (a) and (b) show false negative rates for referrals involving Black and White children, respectively. Among low- and medium-risk referrals, false negative rates are similar before and after the AFST for both referrals involving Black and White children. However, false negative rates are reduced among high-risk referrals for both race groups. This aligns with the purpose of the high-risk protocol, which defaults referrals above a certain score to be screened in for investigation. Indeed, it seems that the algorithm is catching cases of severe maltreatment that were previously missed. Panels (c) and (d) show false positive (foster care definition) rates for referrals involving Black and White children, respectively. The patterns are similar across race. For low- and medium-risk referrals, false positives decline in the post-AFST period. This suggests that within these risk bins, the AFST is improving targeting of investigations toward those who are at risk of foster care. However, false positive rates increase for those in the highest risk bin. Again, this is consistent with the high-risk protocol increasing screen-ins above a given score. The observed increase may represent a decline in welfare among this group if investigations with no subsequent foster care are costly

to children and households. However, it may be the case that investigations among this high-risk group are providing other benefits or interventions that reduce the need for foster care. When false positives are instead defined using substantiations, the rate of false positives declines for each risk group in the post-period, as depicted in panels (e) and (f).

Welfare implications depend on the relative costs of false negatives and false positives among low-, medium-, and high-risk referrals. In order for the observed patterns to be consistent with a reduction in welfare, the cost of the increase in false positives among high-risk referrals would need to outweigh the benefit of reductions in false positives among low- and medium-risk referrals, combined with the benefit of reductions in false negatives among high-risk referrals. This is unlikely to be the case. First, false positives as defined may be less costly for higher risk referrals. Children on high-risk referrals may see benefits even from investigations that do not result in home removal, if they are provided with in-home services.⁴³ Even if these children do not benefit from an investigation which does not result in foster care, it is not clear why they would be more harmed by an investigation than lower risk children. Second, false negatives as defined are likely to be highly costly, as they require that a child be screened out and left in a home where they experience subsequent maltreatment severe enough to warrant foster care placement. Given that observed patterns are similar across race, this analysis suggests that the AFST reduced disparities while improving welfare in each group.

Together, these patterns suggest that the AFST reduced the role of human bias in decision-making, particularly among referrals subject to the high-risk protocol. Crucially for understanding welfare effects, children on these high-risk referrals are exactly those who are most likely to benefit from intervention.

7 | Conclusion

Predictive risk models are increasingly used to assist decision-makers in a wide variety of settings. Academics, activists, and policymakers have rightfully raised concerns that such algorithms may exacerbate existing biases, or even create new ones. We show that predictive risk models can also serve to reduce racial disparities, relative to human decision makers.

We study how the introduction of the AFST affected racial disparities in child protection system involvement. Comparing outcomes for Black versus White children before versus after the implementation of the AFST, we find that the AFST reduced disparities in both screening decisions and home removals, with effects concentrated among the highest risk referrals subject to the high-risk protocol. While we do not speak to welfare-relevant outcomes beyond the child protection system, an analysis of false positive and false negative rates suggests that these reductions in disparities were achieved while improving decision quality for children of both races. In particular, the high-risk protocol appears to have caught cases of severe maltreatment that were previously missed, reducing false negatives among high-risk referrals for both Black and White children.

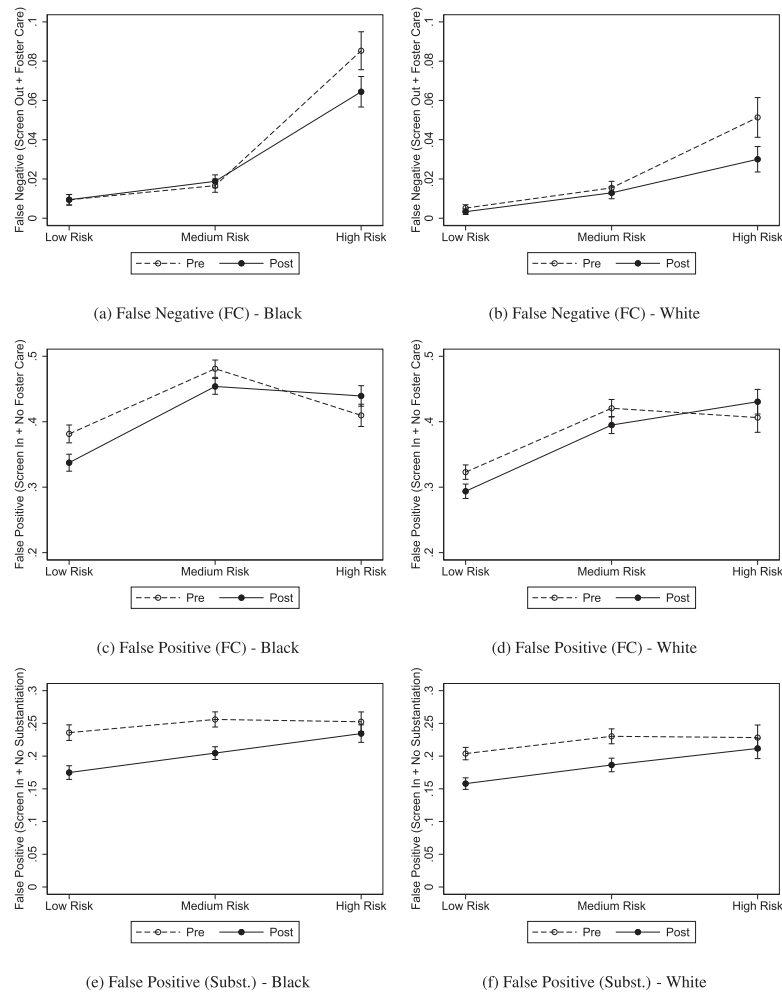


FIGURE 8 | False positives and false negatives. *Note:* This figure presents rates of false positives and false negatives. As described in Section 6, false negatives are defined as referrals which are screened out and involve a child who is removed from their home in the following 6 months. We show two definitions of false positives. False positive (FC) describes referrals which are screened in and for which no child is removed from their home within the following 6 months. False positive (Subst.) describes referrals which are screened in and which do not result in any substantiated allegations. In each subfigure, AFST risk bin is shown on the x-axis and false positive/negative rate is shown on the y-axis. Hollow circles and dashed lines represent the period prior to AFST implementation, while solid circles and solid lines represent the period after AFST implementation. Panels (a), (c), and (e) present patterns for referrals involving Black children, while panels (b), (d), and (f) present patterns for referrals involving White children. The time period is limited to referrals made between January 2013 and September 2019.

Our findings should be interpreted with important caveats. First, while the AFST reduced disparities, unwarranted disparities in screening rates remained after implementation. Moreover, the algorithm operates only at one of many decision points in the child welfare system. Disparities in how cases come to the attention of the child protection system in the first place—through differential rates of reporting by race—are outside the reach of a screening tool. Similarly, disparities may persist in subsequent decisions such as substantiation, case opening, and home removal.

Second, the effects of any given algorithm are likely to be context dependent. The AFST operates in part through a high-risk protocol that defaults certain referrals to investigation, a design feature that appears to drive much of the equity gain we document. Algorithms that provide information without a binding default recommendation may have different effects on disparities and

decision quality. Similarly, the magnitude and direction of effects may differ in settings with different baseline rates of disparity, different compositions of referred families, or different institutional constraints on screener discretion. Policymakers considering the adoption of predictive risk models in other jurisdictions should not assume that the results documented here will generalize without further evaluation.

With these caveats in mind, our findings offer meaningful guidance for policy. Racial disparities in child protection system involvement are a longstanding and high-priority concern for agencies, advocates, and policymakers. Reducing disparities while maintaining standards of child safety is a key challenge for policymakers. Our work shows that carefully designed predictive risk models can make progress toward this goal. Our findings suggest that equity concerns should not be an automatic disqualifier for algorithm adoption. The evidence from Allegheny County

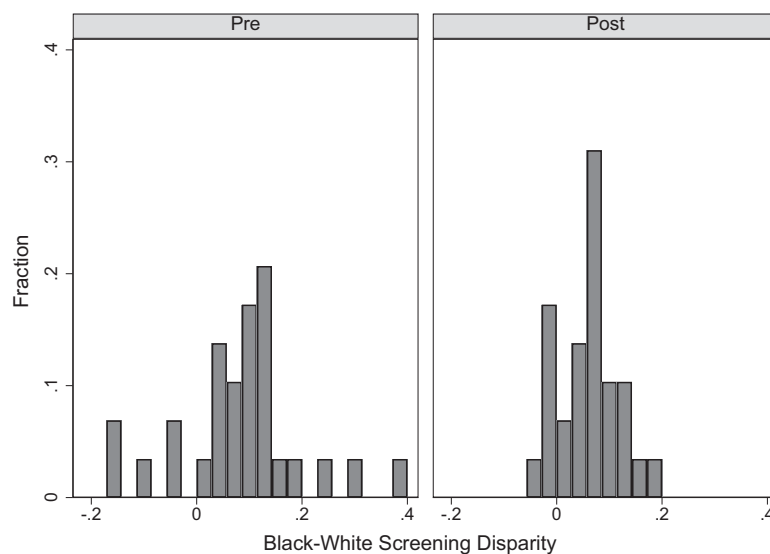


FIGURE 9 | Distribution of screening gaps across workers. *Note:* This figure presents the distribution of worker-specific racial disparities in screening rates, before and after AFST implementation. Worker-specific disparities are measured as the raw gap in screen-in rates, that is, the average screen-in rate for referrals involving Black children minus the average screen-in rate for referrals involving White children. The sample is limited to 29 screeners who handle at least 20 referrals in each of the pre- and post-period.

suggests that the right algorithm, in the right context, can reduce disparities while improving decision quality.

Acknowledgments

We are grateful for helpful comments from David Arnold, Jason Baron, Julie Cullen, Natalia Emanuel, Mary Evans, Sarah Font, Marie-Pascale Grimon, Max Gross, Lindsey Lacey, Katherine Meckel, Chris Mills, and David Simon. We also thank seminar participants at the 2024 ASSA conference and the 2024 Atlanta Workshop on Public Policy and Child Well-Being.

Funding

Putnam-Hornstein acknowledges support from philanthropic grant funding for research projects (dating back to 2014) related to the use of algorithms in child protection in California, Colorado, and Pennsylvania. Funding has been received from the following foundations and agencies: The Ballmer Group, the Conrad N. Hilton Foundation, the California Office of Child Abuse Prevention, the Reissa Foundation, First 5 Los Angeles, and the Ralph M. Parsons Foundation. Vaithianathan acknowledges support from philanthropic grant funding for research projects (dating back to 2014) related to the use of algorithms in child protection in California, Colorado, and Pennsylvania. Funding has been received from the following foundations and agencies: the Conrad N. Hilton Foundation, the Reissa Foundation, and the Ralph M. Parsons Foundation. Rittenhouse acknowledges support from the Institute for Practical Ethics at UC San Diego.

Disclosure

Vaithianathan and Putnam-Hornstein wish to disclose that they were contracted by Allegheny County to build the AFST and continue to work with the county on other projects. Additionally, Vaithianathan and Putnam-Hornstein disclose interests in Social Data Analytics Ltd, which inter alia contracts for data analytics consulting. The opinions, findings, and conclusions or recommendations expressed in the paper are those

of the authors and do not necessarily reflect the view of any agency or funding partner. All errors are our own.

Data Availability Statement

The data used in this study were received through a data use agreement with Allegheny County Department of Human Services. While they cannot be shared publicly, researchers may request access by following instructions available here: <https://www.alleghenycountyanalytics.us/requesting-data/>. The authors are willing and able to provide additional guidance as necessary.

Endnotes

¹Discussing lessons learned from the movement toward algorithms in the justice system, Ludwig and Mullainathan (2021) note that: “the optimism for machine learning in criminal justice did not last long... Reports emerged of algorithms that were themselves discriminatory, producing racially disparate outcomes at a high enough rate that the phrase ‘algorithmic bias’ has entered the lexicon.”

²See, for example, Goncalves and Mello (2021), Antonovics and Knight (2009), Rehavi and Starr (2014), Arnold et al. (2018), Abrams et al. (2012), and Baron et al. (2026).

³Oregon’s Department of Human Services halted use of an algorithm after concerns were raised about its potential to disproportionately affect Black families (Ho and Burke 2022b). California’s Department of Social Services abandoned plans to implement a predictive risk tool in part due to worries about the effects on racial equity (Ho and Burke 2022a). Allegheny County, PA (the context for this study), has faced widespread criticism for its use of a predictive risk model and the potential disparate impacts by race.

⁴See, for example, Gateway (2021) and Thomas and Halbert (2021).

⁵As discussed further below, call screen workers only have discretion in the decision to investigate for a subset of referrals. Allegations of abuse or more severe forms of maltreatment are mandated by state law to be investigated. As such, our analysis is limited to the subset of allegations which are investigated at the call screener’s discretion (85% of referrals in our sample).

- ⁶We use as a proxy for future maltreatment foster care placement within 6 months if left at home. This definition is discussed further in Section 6 below.
- ⁷See Barry Jester et al. (2015) and Angwin et al. (2016).
- ⁸Albright (2026) shows that algorithmic recommendations may have importantly different effects from algorithmic predictions.
- ⁹Our main analysis is not affected by this glitch, as we use comparable retroactively calculated scores from a newer version of the AFST, rather than the observed scores for each given referral. Moreover, the glitch only affected a small fraction of referrals, and in general the shown score was close to the true score. See De-Arteaga et al. (2020) for more details on the technical glitch, as well information on the correlation between observed and true scores.
- ¹⁰See, for example, Berger et al. (2017); Raissian and Bullinger (2017); Cancian et al. (2013); Kovski et al. (2022); Rittenhouse (2025); Bullinger et al. (2023).
- ¹¹See Putnam-Hornstein et al. (2013) and Maloney et al. (2017).
- ¹²The CPS Law is the relevant Pennsylvania statute which defines child abuse and prescribes the counties' responsibility. Allegheny County provides a brief overview of the two types of referrals here: <https://www.alleghenycounty.us/Human-Services/Programs-Services/Children-Families/Protective-Services.aspx>.
- ¹³In some cases, there might be multiple referrals made by different people about the same allegation or incident. The county may in these instances, combine all of these referrals into one referral, that is, requiring just one investigation.
- ¹⁴Importantly for the context of this paper, screeners can often observe child race either through a reporter's physical description of the child, or through the linked data system.
- ¹⁵Allegheny County Data Warehouse (2024) provides a detailed description on the linked data systems, and how they feed into the AFST.
- ¹⁶Based on conversations with call screeners and supervisors, they primarily focus on age and allegation in order to determine the likely safety of children on referrals.
- ¹⁷In some very rare cases where a fatality or near-fatality occurred, screening staff might learn about any mistakes that had been made, for example, in screening out an at-risk child.
- ¹⁸The low risk protocol has changed over time. Prior to 2018 there was no low risk protocol. From November 2018 through October 2019, referrals fell under the low risk protocol if the maximum score was less than 10 and all children were over age 11. In October 2019, the protocol criteria was expanded to include referrals with a maximum score less than or equal to 12 and no children aged 6 or younger.
- ¹⁹For example, two children in the same household may have different histories with child protective services, which leads to different risk scores. However, the screener will only see one score per referral.
- ²⁰See Vaithianathan et al. (2017) for an overview of the development of the original AFST. Additional background and documents related to the AFST are available at: <https://www.alleghenycounty.us/Human-Services/News-Events/Accomplishments/Allegheny-Family-Screening-Tool.aspx>.
- ²¹See Vaithianathan et al. (2019).
- ²²This evaluation used interrupted time-series methods, and a follow-up period of 17 months. Although limited in their power to detect statistically significant effects, their analysis suggests that the AFST increased the accuracy of screen-in decisions, in particular for referrals involving White children.
- ²³CPS referrals made prior to 2015 were expunged under former practices and are not available for analysis. In 2018, Pennsylvania law changed, allowing counties to maintain records in their own databases which were expunged from the statewide database (23 PA.C.S.; P.L. 375, No. 541).
- ²⁴Note that a small share of GPS referrals include allegations which should qualify for CPS status. To the best of our understanding, this may occur if allegations are added to a referral after the initial screening decision.
- ²⁵Note, national data are for the 45 states which report both screened-in and screened-out referrals to the National Child Abuse and Neglect Data System (NCANDS).
- ²⁶The way in which predictive features are retrospectively coded, it is as close to what those features would have been at the time that the call came in. It is possible that some features may have changed due to subsequent data entering the data-warehouse—for example, demographic data might be updated, but these are in the minority.
- ²⁷Allegation categories are listed in Panel C of Table 1. Reporter categories are listed in Panel D of Table 1. The indicator for exposure to drugs or alcohol is equal to one if any allegation on the referral mentions drugs or alcohol, and zero otherwise.
- ²⁸We use the retroactively generated V3 score for this classification, which allows us to compare like with like.
- ²⁹Removal trends are plotted for each of GPS and CPS referrals in Figure A3.
- ³⁰Raw screening trends are shown in Figure A2.
- ³¹The omitted period (2016q2) has disparities that are lower than average for the pre-period, giving the appearance of significant effects in the pre-period, and null effects in the post-period. However, we do not observe a consistent trend in the pre-period.
- ³²Note, in the annual event study, the excluded year is 2015 and the first treated year (2016) includes several untreated months (January to July).
- ³³In unreported results from this regression, the coefficient on $Black \times Post \times Score^3$ is -0.0601 and the coefficient on $Black \times Score^3$ is 0.0919 (each significant at the 1% level).
- ³⁴Note that this event study is restricted to years in which we have retroactively calculated AFST scores (2013 on).
- ³⁵Since we use the comparable AFST V3 score, the high-risk bin does not correspond exactly with the high-risk protocol. That is, there may be referrals with an AFST V3 score of 18–20, but a screener-observed score (from an earlier algorithm version) below 17.
- ³⁶Raw removal trends are reported in Figure A3.
- ³⁷To calculate the pre-existing difference in 3-month removal rates for GPS referrals involving Black versus White children, we sum the coefficients on $Black$ and $Black \times GPS$. $0.0264 + 0.00893 = 0.03533$.
- ³⁸An F -test comparing variance in disparities across the two distributions results in an f -statistic of 4.7, strongly rejecting the null hypothesis that variance is equal across the two distributions.
- ³⁹We calculate unwarranted disparities following the methodology of Baron et al. (2026).
- ⁴⁰This is a different definition than Baron et al. (2026), who use investigation within 6 months as a proxy for subsequent maltreatment. We chose this outcome as it is not directly influenced by the AFST. Results are similar when using future investigations.
- ⁴¹Details of this calculation are included in the notes to Figure A4.
- ⁴²Implicit in this definition is the assumption that children are only removed from their home when they are truly victims of maltreatment. In practice, this may not always be the case, as home removal decisions may be subject to error and biases. However, we argue that the introduction of the AFST is unlikely to impact this error rate and as

such we can still learn about the welfare implications of the AFST by studying the changes in false negative and false positive rates across time.

⁴³In ongoing work, Rittenhouse et al. (2025) study the impacts of investigations for children at the high-risk protocol threshold.

References

- Abrams, D. S., M. Bertrand, and S. Mullainathan. 2012. "Do Judges Vary in Their Treatment of Race?" *Journal of Legal Studies* 41, no. 2: 347–383.
- Administration on Children, Y., and C. B. Families. 2021. *Child Maltreatment 2019*. Tech. rep., U.S. Department of Health & Human Services, Administration for Children and Families.
- Albright, A. 2019. "If You Give a Judge a Risk Score: Evidence From Kentucky Bail Decisions." Working Paper.
- Albright, A. 2026. "The Hidden Effects of Algorithmic Recommendations." OIGI Working Paper #104. <https://www.minneapolisfed.org/research/institute-working-papers/the-hidden-effects-of-algorithmic-recommendations>.
- Allegheny County Data Warehouse. 2024. Tech. rep., Allegheny County Department of Human Services. <https://analytics.alleghenycounty.us/wp-content/uploads/2024/02/24-ACDHS-03-Datawarehouse.pdf>.
- Angwin, J., J. Larson, S. Mattu, and L. Kirchner. 2016. "Machine Bias." *ProPublica*. <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing>.
- Antonovics, K., and B. G. Knight. 2009. "A New Look at Racial Profiling: Evidence From the Boston Police Department." *Review of Economics and Statistics* 91, no. 1: 163–177.
- Arnold, D., W. Dobbie, and P. Hull. 2021. "Measuring Racial Discrimination in Algorithms." *AEA Papers and Proceedings* 111: 49–54.
- Arnold, D., W. Dobbie, and C. S. Yang. 2018. "Racial Bias in Bail Decisions." *Quarterly Journal of Economics* 133, no. 4: 1885–1932.
- Baron, E. J., J. J. Doyle Jr, N. Emanuel, P. Hull, and J. Ryan. 2024. "Discrimination in multiphase systems: Evidence from child protection." *The Quarterly Journal of Economics* 139, no. 3: 1611–1664.
- Baron, E. J., J. J. Doyle Jr, N. Emanuel, and P. Hull. 2026. "Unwarranted racial disparity in US foster care placement." *Journal of Policy Analysis and Management* 45, no. 2: e70030.
- Barry Jester, A. M., B. Casselman, and D. Goldstein. 2015. "The New Science of Sentencing." The Marshall Project.
- Berger, L. M., S. A. Font, K. S. Slack, and J. Waldfogel. 2017. "Income and Child Maltreatment in Unmarried Families: Evidence From the Earned Income Tax Credit." *Review of Economics of the Household* 15, no. 4: 1345–1372.
- Biden, J. R. 2021. A Proclamation on National Foster Care Month, 2021. <https://www.whitehouse.gov/briefing-room/presidential-actions/2021/04/30/a-proclamation-on-national-foster-care-month-2021/>.
- Bullinger, L., A. Packham, and K. Raissian. 2023. "Effects of Universal and Unconditional Cash Transfers on Child Maltreatment." National Bureau of Economic Research Working Paper #31733.
- Cancian, M., M.-Y. Yang, and K. S. Slack. 2013. "The Effect of Additional Child Support Income on the Risk of Child Maltreatment." *Social Service Review* 87, no. 3: 417–437.
- Cheng, H.-F., L. Stapleton, A. Kawakami, et al. 2022. "How Child Welfare Workers Reduce Racial Disparities in Algorithmic Decisions." In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems*, 1–22.
- Chohlas-Wood, A. 2020. *Understanding Risk Assessment Instruments in Criminal Justice*. Tech. rep., Brookings Institution.
- Chouldechova, A., D. Benavides-Prado, O. Fialko, and R. Vaithianathan. 2018. "A Case Study of Algorithm-Assisted Decision Making in Child Maltreatment Hotline Screening Decisions." In *Conference on Fairness, Accountability and Transparency* PMLR, 134–148.
- Concluding Observations on the Combined Tenth To Twelfth Reports of the United States of America. 2022. Tech. rep., United Nations Committee on the Elimination of Racial Discrimination.
- De-Arteaga, M., R. Fogliato, and A. Chouldechova. 2020. "A Case for Humans-in-the-Loop: Decisions in the Presence of Erroneous Algorithmic Scores." In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, 1–12.
- Drake, B., J. M. Jolley, P. Lanier, J. Fluke, R. P. Barth, and M. Jonson-Reid. 2011. "Racial Bias in Child Protection? A Comparison of Competing Explanations Using National Data." *Pediatrics* 127, no. 3: 471–478.
- Drake, B., M. Jonson-Reid, H. Kim, C.-J. Chiang, and D. Davalishvili. 2021. "Disproportionate Need as a Factor Explaining Racial Disproportionality in the CW System." In *Racial Disproportionality and Disparities in the Child Welfare System*, 159–176. Springer.
- Gateway, C. W. I. 2021. *Child Welfare Practice to Address Racial Disproportionality and Disparity*. Tech. rep., Children's Bureau.
- Goldhaber-Fiebert, J. D., and L. Prince. 2019. *Impact Evaluation of a Predictive Risk Modeling Tool for Allegheny County's Child Welfare Office*. Tech. rep., Allegheny County Analytics.
- Goncalves, F., and S. Mello. 2021. "A Few Bad Apples? Racial Bias in Policing." *American Economic Review* 111, no. 5: 1406–1441.
- Grimon, M. P., and C. Mills. 2022. "The Impact of Algorithmic Tools on Child Protection: Evidence From a Randomized Controlled Trial." Working Paper.
- Hampton, R. L., and E. H. Newberger. 1985. "Child Abuse Incidence and Reporting by Hospitals: Significance of Severity, Class, and Race." *American Journal of Public Health* 75, no. 1: 56–60.
- Ho, S., and G. Burke. 2022a. *An Algorithm That Screens for Child Neglect Raises Concerns*. Associated Press.
- Ho, S., and G. Burke. 2022b. *Oregon Dropping AI Tool Used in Child Abuse Cases*. Associated Press.
- Howell, S. T., T. Kuchler, D. Snitkof, J. Stroebe, and J. Wong. 2021. *Racial Disparities in Access to Small Business Credit: Evidence from the Paycheck Protection Program*. Tech. rep., National Bureau of Economic Research.
- Jenny, C., K. P. Hymel, A. Ritzen, S. E. Reinert, and T. C. Hay. 1999. "Analysis of Missed Cases of Abusive Head Trauma." *Jama* 281, no. 7: 621–626.
- Kim, H., C. Wildeman, M. Jonson-Reid, and B. Drake. 2017. "Lifetime Prevalence of Investigating Child Maltreatment Among US Children." *American Journal of Public Health* 107, no. 2: 274–280.
- Kleinberg, J., H. Lakkaraju, J. Leskovec, J. Ludwig, and S. Mullainathan. 2018. "Human Decisions and Machine Predictions." *Quarterly Journal of Economics* 133, no. 1: 237–293.
- Kovski, N. L., H. D. Hill, S. J. Mooney, F. P. Rivara, and A. Rowhani-Rahbar. 2022. "Short-Term Effects of Tax Credits on Rates of Child Maltreatment Reports in the United States." *Pediatrics* 150, no. 1.
- Lane, W. G., D. M. Rubin, R. Monteith, and C. W. Christian. 2002. "Racial Differences in the Evaluation of Pediatric Fractures for Physical Abuse." *Jama* 288, no. 13: 1603–1609.
- Ludwig, J., and S. Mullainathan. 2021. "Fragile Algorithms and Fallible Decision-Makers: Lessons From the Justice System." *Journal of Economic Perspectives* 35, no. 4: 71–96.
- Maloney, T., N. Jiang, E. Putnam-Hornstein, E. Dalton, and R. Vaithianathan. 2017. "Black-White Differences in Child Maltreatment Reports and Foster Care Placements: A Statistical Decomposition Using Linked Administrative Data." *Maternal and Child Health Journal* 21, no. 3: 414–420.
- Obermeyer, Z., B. Powers, C. Vogeli, and S. Mullainathan. 2019. "Dissecting Racial Bias in an Algorithm Used to Manage the Health of Populations." *Science* 366, no. 6464: 447–453.

- Price, W. N. 2019. *Risks and Remedies for Artificial Intelligence in Health Care*. Tech. rep., Brookings Institution.
- Putnam-Hornstein, E., B. Needell, B. King, and M. Johnson-Motoyama. 2013. "Racial and Ethnic Disparities: A Population-Based Examination of Risk Factors for Involvement With Child Protective Services." *Child Abuse & Neglect* 37, no. 1: 33–46.
- Raghavan, M., S. Barocas, J. Kleinberg, and K. Levy. 2020. "Mitigating Bias in Algorithmic Hiring: Evaluating Claims and Practices." In *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency*, 469–481.
- Raissian, K. M., and L. R. Bullinger. 2017. "Money Matters: Does the Minimum Wage Affect Child Maltreatment Rates?" *Children and Youth Services Review* 72: 60–70.
- Rehavi, M. M., and S. B. Starr. 2014. "Racial Disparity in Federal Criminal Sentences." *Journal of Political Economy* 122, no. 6: 1320–1354.
- Rittenhouse, K. 2025. "Income and Child Maltreatment: Evidence From a Discontinuity in Tax Benefits." Review of Economics and Statistics.
- Rittenhouse, K., S. David, and L. Lindsey. 2025. "Child Maltreatment Investigations and Family Well-being." National Bureau of Economic Research Working paper 34334.
- Stevenson, M. T., and J. L. Doleac. 2021. "Algorithmic Risk Assessment in the Hands of Humans." *American Economic Journal: Economic Policy* 16, no. 4: 382–414.
- Thomas, K., and C. Halbert. 2021. *Transforming Child Welfare: Prioritizing Prevention, Racial Equity, and Advancing Child and Family Well-Being*. Tech. rep., National Council on Family Relations.
- Vaithianathan, R., E. Kulick, E. Putnam-Hornstein, and D. B. Prado. 2019. *Allegheny Family Screening Tool: Methodology, Version 2*. Tech. rep., Centre for Social Data Analytics.
- Vaithianathan, R., E. Putnam-Hornstein, N. Jiang, P. Nand, and T. Maloney. 2017. *Developing Predictive Models to Support Child Maltreatment Hotline Screening Decisions: Allegheny County Methodology and Implementation*. Tech. rep., Centre for Social Data Analytics.
- Wildeman, C., and N. Emanuel. 2014. "Cumulative Risks of Foster Care Placement by Age 18 for US Children, 2000–2011." *PloS One* 9, no. 3: e92785.

Supporting Information

Additional supporting information can be found online in the Supporting Information section.

Data S1