

NBER WORKING PAPER SERIES

REFORMING ROTATIONS

E. Jason Baron
Richard Lombardo
Joseph P. Ryan
Jeongsoo Suh
Quitze Valenzuela-Stookey

Working Paper 32369
<http://www.nber.org/papers/w32369>

NATIONAL BUREAU OF ECONOMIC RESEARCH

1050 Massachusetts Avenue
Cambridge, MA 02138
April 2024, Revised March 2026

We thank Peter Arcidiacono, Pat Bayer, Chris Campos, Sylvain Chassang, Allan Collard-Wexler, Michael Dinerstein, Laura Doval, Joseph Doyle, Federico Echenique, Natalia Emanuel, Francesco Fabbri, Maria Fitzpatrick, Ezra Goldstein, Aram Grigoryan, Peter Hull, Giacomo Lanzani, Jacob Leshno, Chris Mills, Matt Pecenco, Michael Pollmann, Canice Prendergast, Katherine Rittenhouse, Alex Teytelboym, and seminar participants at the ASSA Annual Meeting, BYU, CREST, Duke, ESSET 2024, MIT Sloan, NBER Market Design Meeting (Fall 2024), Opportunity Insights, Princeton, PSE, SAET 2024, SEA 2024, SITE 2024, UC Berkeley, UCL, UIUC, University of Chicago Booth, USC, Yale University, and the Workshop on Economics of Education (Valle Nevado, Chile) for helpful comments and suggestions. This paper subsumes previous drafts which circulated as “Mechanism Reform for Task Allocation” and “Mechanism Reform: An Application to Child Welfare.” The views expressed herein are those of the authors and do not necessarily reflect the views of the National Bureau of Economic Research.

NBER working papers are circulated for discussion and comment purposes. They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2024 by E. Jason Baron, Richard Lombardo, Joseph P. Ryan, Jeongsoo Suh, and Quitze Valenzuela-Stookey. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

Reforming Rotations

E. Jason Baron, Richard Lombardo, Joseph P. Ryan, Jeongsoo Suh, and Quitze Valenzuela-Stookey
NBER Working Paper No. 32369

April 2024, Revised March 2026

JEL No. D82, H75, J13, J45

ABSTRACT

In many settings, tasks are assigned to agents via a rotation. Such quasi-random allocation ignores heterogeneity in agents' performance and preferences. We study how a designer can exploit such heterogeneity to improve aggregate performance while simultaneously ensuring that no agent is worse off relative to the status-quo system. The key challenge is that, while the designer may be able to estimate agents' performance, their preferences are inherently unobservable. We develop a mechanism-design framework to study this problem and characterize optimal mechanisms in both static and dynamic settings. Optimal mechanisms can be interpreted as competitive equilibria in which agents trade tasks facing personalized, kinked budget sets. As an illustration, we apply our results to the assignment of Child Protective Services investigators to maltreatment cases. Simulations show that the mechanism reduces false positives (unnecessary foster care placements) by up to 14% while also lowering false negatives (missed maltreatment cases) and overall placements.

E. Jason Baron
Duke University
Economics Department
and NBER
jason.baron@duke.edu

Richard Lombardo
Harvard University
Department of Economics
richardlombardo@fas.harvard.edu

Joseph P. Ryan
University of Michigan
joryan@umich.edu

Jeongsoo Suh
Duke University
jeongsoo.suh@duke.edu

Quitze Valenzuela-Stookey
University of California, Berkeley
quitze@berkeley.edu

The need to divide heterogeneous tasks among agents arises frequently in organizations. In many settings where tasks arrive over time, the assignment is done via a rotation: when a task comes in, it is assigned to the agent at the top of a queue, and that agent then moves to the back of the queue. Rotations (or close variants thereof) are used to assign judges to criminal cases (Bhuller et al., 2020) and bankruptcy filings (Dobbie and Song, 2015); case managers to incarcerated individuals (Lee, 2023); and parole board members to parole hearings (Arbour and Marchand, 2024); among other examples.

The strength of rotations is that they smooth task-loads across agents: over any time horizon each agent receives the same number of tasks in expectation. Moreover, although some tasks may be more difficult to handle than others, in expectation (and over the long run) the composition of the assigned bundle of tasks is the same across agents. However, rotations have two important limitations. First, they ignore any heterogeneity that might exist in the performance of agents on different types of tasks. If the designer can estimate individual agents’ performance, they may wish to use this information to improve the overall efficiency of the system.¹ Second, they do not exploit potential preference heterogeneity among agents. The system equalizes workloads across agents only to the extent that all agents find any given task equally burdensome. If, instead, agents’ preferences differ, it may be possible to reassign tasks in a way that makes all agents better off.

In this paper, we study how a designer can exploit heterogeneity in performance and preferences across agents to improve upon the performance of a randomized task-allocation system while ensuring that no agents are made worse off.² For concreteness, we frame the discussion in terms of rotations—although, as clarified below, the analysis covers other instances of quasi-random task allocation. The key challenge is that agents’ preferences are generally not perfectly observable by the designer, and so assigning tasks based solely on the observable measures of their performance risks deteriorating agents’ welfare. We must therefore consider mechanisms that elicit and respond to agents’ preferences, while carefully structuring incentives to steer them towards the tasks that they are relatively good at.

We study this mechanism-design problem in a model with two types of tasks, referred to as

¹For concreteness, think of the performance of agent j on task i as the cost to the designer of having j complete this task.

²Or, more generally, while guaranteeing agents some fraction of their current welfare level. One interpretation of this restriction is that it allows us to assess the gains from exploiting performance information in a “revenue neutral” way, holding fixed the total use of inputs (agents’ time and energy). This constraint is also relevant in settings where the designer needs agents’ buy-in in order to implement a reform. Additionally, the designer may be concerned that reducing agents’ welfare could lead to a vicious cycle in which some agents quit, and the remaining agents are left to handle an increased workload, generating further turnover. See the empirical application in Section IV for further discussion.

“high”, h , and “low”, l .³ Each task must be assigned to one of J agents, each of whom has privately-observed preferences over tasks, represented by their relative cost of handling high-versus low-type tasks. The designer observes (or has an unbiased estimate of) each agent’s performance on each of the task types, and their objective is the expected sum of performance across task-agent pairs. We focus on settings in which transfers are not available as part of the design. While transfers play a role in many settings with rotations (e.g., employee salaries), they are often not directly tied to performance of specific tasks.⁴

In the language of mechanism-design, the general problem that we face is one of dynamic multi-item allocation with type-dependent welfare constraints. The problem is dynamic when tasks arrive over time and must be assigned as they come without knowledge of future arrivals.⁵ The welfare constraint is that no agent should be made worse off relative to the status-quo rotational system; it is type-dependent because an agent’s welfare under the status quo depends on their preferences. Each of these dimensions of the problem introduces its own set of technical challenges. Our approach is to study a sequence of relaxed problems, which allow us to isolate the key forces introduced by each dimension before arriving at a mechanism that can be implemented in the original problem of interest.

We first solve a version of the problem which is relaxed along two dimensions. The first relaxation is to a static setting, in which we have a fixed set of tasks to allocate among the agents (as is indeed the case in some applications). The second is that we require only that each task be assigned in expectation (taken over the profile of agents’ types) rather than ex-post (conditional on every realized type profile). Equivalently, this model represents a market with a continuum of agents, and we therefore refer to it as the Large-Market Static (LMS) problem. Conceptually, studying the large-market model allows us to focus clearly on the distortions introduced by asymmetric information and the status-quo welfare constraint.

Our main technical results characterize the optimal mechanism in the LMS problem, which we do via a two-step procedure. We first solve an “inner program” to characterize the support function of the set of expected assignments that can be induced for each agent by an incentive compatible mechanism that satisfies the status-quo welfare constraint (Theorem 1). We then reformulate the dual to the designer’s program in terms of these support functions (Theorem 2). This two-step procedure is more broadly applicable, and thus constitutes a

³In Section IV we discuss how the insights can be applied to settings with richer task heterogeneity.

⁴This is the case for all examples cited above. In restricting attention to counterfactual mechanisms which fix transfers, our approach can be viewed as an instance of minimalist market design (Sönmez, 2023).

⁵For example, in the application to Child Protective Services that we study below, an agent (investigator) must be assigned to a new task (case) within 24 hours. See Section IV for details.

potential methodological contribution of the paper.⁶

The optimal mechanism admits a simple indirect implementation: under a regularity condition on the distribution of preferences, it can be viewed as a generalization of competitive equilibrium in an exchange economy: each agent is endowed with their status-quo allocation and chooses from a personalized budget set that has a kink at the endowment (without regularity, at most three additional kinks are needed). The key distinction between this setting and typical applications of market-based mechanisms (e.g., [Hylland and Zeckhauser \(1979\)](#); [Budish \(2011\)](#)) is that the designer targets an objective based on agents’ performance, rather than efficiency with respect to agents’ preferences. This necessitates personalized (non-linear) pricing.

Additionally, we show that while our model treats agents’ performance as a policy-invariant parameter (e.g. ability) that can be estimated by the designer, the optimal mechanism for the LMS program does, in a local sense, reward agents with higher performance ([Theorem 3](#)). Moreover, the mechanism satisfies a notion of fairness, in that differences in agents’ assignments can be justified on the basis of their relative performance.

Having understood the distortions introduced by agents’ private information in a large-market, we next reimpose the constraint that every task be assigned ex-post. This defines a Small-Market Static (SMS) problem. We modify our mechanism from the LMS problem to approximately solve the SMS problem ([Proposition 3](#)). The basic idea is to use the indirect implementation of the mechanism that was optimal in the LMS problem (via non-linear pricing) but restrict the trades that can be executed to ensure that the market clears ex-post. This mechanism is asymptotically optimal in the SMS problem as the number of agents grows. It is also obviously strategy-proof ([Li, 2017](#)) and robust to both the form of agents’ preferences and to misspecification of the preference distribution on the part of the designer.

Finally, we convert the previous mechanism into one that works in the original setting, which we refer to as the Small-Market Dynamic (SMD) problem ([Proposition 5](#)). We do this by using the SMS mechanism to specify “target allocations” for each agent, and then assign cases as they arrive based on which agent is furthest from their target. This gives us a mechanism that can be applied in practice (indeed, a rotation is a special case of such a mechanism). It is asymptotically optimal and strategy-proof as the number of agents and the time horizon grows. Furthermore, this mechanism can be implemented without ex-ante knowledge of the distribution of agents’ preferences. Thus, this simple and easily-implementable mechanism provides a plausible lower bound on the potential gains from exploiting measures of agent

⁶The technique is subsequently applied in [Valenzuela-Stookey \(2025\)](#) to study automation, and in [Valenzuela-Stookey et al. \(2026\)](#) to study task-allocation in higher-dimensional settings.

performance, and can serve as a starting point for practical reforms.

To illustrate the potential for using performance estimates to improve upon rotational task-allocation systems in practice, we study an application to the assignment of Child Protective Services (CPS) workers (agents) to investigate reported cases of child maltreatment (tasks). This is an important high-stakes setting where preserving agent welfare is a primary concern for the designer (CPS) due to high rates of burnout, turnover, and chronic under-staffing.⁷ Moreover, this setting serves to motivate some of our main modeling assumptions, while also highlighting potentially interesting avenues for further study.

In order to apply our mechanism-design results in the CPS context, we first need to estimate agents' performance. The primary role of CPS investigators is to make a recommendation as to whether the child in question should remain in the home or be placed in foster care. The objective of CPS is to place a child in foster care if and only if they would experience subsequent maltreatment if left in the home. Performance is therefore measured by the propensity of an investigator to make correct placement decisions. Because counterfactual outcomes are inherently unobservable, this propensity cannot be non-parametrically identified. However, we show that the *difference* in performance between two investigators for a given category of cases is non-parametrically identified, exploiting the fact that investigators are quasi-randomly assigned to cases according to a rotation (Doyle, 2007).⁸ The identification results are more generally applicable to settings in which agents take a binary action and outcomes are selectively observed, and may therefore be of independent interest.

To apply the SMD mechanism in this setting, we categorize cases as high- or low-risk, based on the estimated probability of subsequent maltreatment. High-risk cases tend to be associated with greater workloads and psychological burdens. To inform the design of the new mechanism, we conduct a statewide survey of CPS investigators in Michigan, which reveals substantial heterogeneity in how the investigators perceive the burden of high- and low-risk cases. We combine this with an administrative dataset from Michigan covering all child maltreatment investigations between 2008 and 2016, which we use to estimate performance.

This analysis shows that there are potential gains to be had by exploiting performance data: in our baseline specification, assigning investigators to cases according to our mechanism lowers investigators' false positives (children placed in foster care who would have been safe in their homes) by 14%, while also reducing false negatives (children left at home who are

⁷Contact with CPS is surprisingly common in the U.S.: 37% of children are the subject of a maltreatment investigation by age 18 and 5% spend time in foster care (Doyle et al., 2025; Kim et al., 2017).

⁸In fact, the identification results continue to hold under the new mechanism that we introduce. Thus, performance estimates could be updated even after implementation of a reform.

subsequently maltreated) and overall placements. This result is robust to several alternative specification assumptions and measures of subsequent maltreatment, and the gains are even larger for the largest counties in our sample. Importantly, the mechanism involves reallocating only existing resources and ensures that no investigator is made worse off. In fact, it improves investigator welfare relative to the status quo by at least 10% for 26% of investigators in our sample. In contrast, the “naive” mechanism that simply fixes the total number of cases handled by each investigator generates better aggregate performance, but reduces welfare by at least 10% for 27% of investigators. Thus, failing to consider preferences leads to an incorrect estimate of the true potential gains, and implementing such a mechanism could significantly harm recruitment and increase turnover in an already strained system.

Related literature: This study’s primary contribution is a general mechanism-design framework that can be applied in a range of task-allocation settings in which (i) agents’ performance can be estimated, (ii) the designer seeks to reform a quasi-random status-quo allocation system to optimize aggregate performance, and (iii) the designer aims to ensure that no agent is made worse off relative to the status quo. This constraint is especially relevant in public sector settings like CPS, where high turnover and staff shortages make additional turnover infeasible. It may also bind in environments with hold harmless clauses ([Dinerstein and Smith, 2021](#)) or union contracts that explicitly prevent agents from being made worse off compared to the status quo.

In developing the framework, we contribute to several strands of literature. First and foremost, this paper belongs to the theoretical mechanism-design literature. Consider first the static versions of the problem (LMS/SMS). Broadly, these are problems of organizational economics in which tasks are allocated to agents whose cost for performing the task is unknown (e.g. [Grossman and Hart \(1983\)](#); [Holmstrom \(1989\)](#); [Baker et al. \(2001\)](#)). Unlike the bulk of this literature, agents’ private information matters not because we hope to influence their effort or minimize input cost, but because we need to ensure that they are not made worse off.

In the mechanism-design problem, the status-quo constraint is equivalent to a type-dependent participation constraint, which makes this a problem of countervailing incentives ([Lewis and Sappington, 1989](#); [Maggi and Rodriguez-Clare, 1995](#); [Jullien, 2000](#); [Dworczak and Muir, 2025](#)). Our solution uses an ironing approach similar to that of [Dworczak and Muir \(2025\)](#). However, the multi-item, multi-agent allocation problem that we study is new in this literature and requires the development of new techniques.⁹ Our work also relates to the growing literature on mechanism design with generalized social objectives (e.g., [Dworczak et al. \(2021\)](#)

⁹Less directly related are papers that consider Pareto-improving policies in other contexts such as contests ([Krishna et al., 2026](#)), or in an axiomatic matching framework (e.g., [Combe et al. \(2022\)](#)).

and Akbarpour et al. (2024)).

Within the literature on combinatorial allocation problems, in which indivisible objects must be assigned to agents with unknown preferences, the natural benchmark is the class of market-type mechanisms based on ideas from competitive equilibrium. Seminal contributions include Varian (1973) and Hylland and Zeckhauser (1979). This literature focuses on identifying mechanisms with various desirable properties, usually including a notion of Pareto efficiency or fairness with respect to agents’ preferences. In contrast, the current study is concerned with maximizing a social objective based on performance, and the preferences of the agents (investigators) enter only as a constraint. This is an important conceptual distinction, and leads to an optimal mechanism involving non-linear, personalized exchange rates for cases.

We focus here on a model with two types of tasks. However, as we discuss in the empirical application (Section IV), the mechanism can be applied in settings with richer task heterogeneity by partitioning the set of tasks into two categories. In such settings, the choice of the partition becomes an additional design parameter. Valenzuela-Stookey et al. (2026) extend this approach to allow for richer partitions (as well as non-partitional mechanisms). The focus there is on characterizing obviously strategy-proof mechanisms; in higher dimensions it is not in general possible to solve for the optimal mechanism.

The problem that we ultimately address (SMD) is inherently dynamic, in that cases must be assigned as they arrive. As with the static problem, we differ from the literature on dynamic combinatorial allocation (Combe et al., 2021; Nguyen et al., 2023) and matching (Karp et al., 1990; Mehta et al., 2007; Aggarwal et al., 2011; Baccara et al., 2020) in both the constraints that we face and the objective.

Finally, the empirical application relates to a structural and empirical literature on matching in foster care systems. Most of this work focuses on assigning children to foster or adoptive families (e.g., Baccara et al. (2014), MacDonald (2024), Robinson-Cortes (2019), and Slaugh et al. (2016)). Our work complements this literature by analyzing a different allocation margin—matching investigators to cases—and by taking a mechanism-design approach to propose a data-driven improvement to the existing assignment process.

I Model

We first introduce the static version of the model in which there is a fixed set of tasks to be allocated among a set $\mathcal{J} = \{1, \dots, J\}$ of agents. In Section III.B we build on insights from the static model to extend the analysis to the dynamic setting in which tasks arrive over time and must be assigned as they come. However, the static model may be of independent

interest for settings where tasks are assigned quasi-randomly in batches (e.g., [Kling 2006](#)).

Let n^h and n^l be the per-agent number of high-type (type- H) and low-type (type- L) tasks, respectively, so the total number of tasks of each type is Jn^h and Jn^l . Agent j has preferences over bundles of tasks $(\hat{n}^h, \hat{n}^l) \in \mathbb{R}_+^2$ represented by their (subjective) *workload* $p_j^h \hat{n}^h + p_j^l \hat{n}^l$, where $p_j^k \geq 0$ is the marginal cost to the agent of handling a type- k task (so agents prefer lower workloads). The rotational system amounts (in the long run and in expectation) to a random allocation of tasks. Thus, in the static model we refer to the equal split of tasks, in which each agent receives (n^h, n^l) , as the *status-quo allocation*. The constraint that agent j be made no worse off under the new assignment $(\hat{n}_j^h, \hat{n}_j^l)$ is thus given by $p_j^h \hat{n}_j^h + p_j^l \hat{n}_j^l \leq p_j^h n^h + p_j^l n^l$. Agents' preferences are their private information.

Remark 1. It is straightforward to extend the model to allow for an asymmetric status quo in which assignments differ across agents. In other words, the status-quo system need not be precisely a rotation. The important feature is that the status-quo allocation is independent of agents' preferences.¹⁰ We focus on the symmetric case because it is the relevant one for standard rotational systems (as in the CPS application) and to simplify notation.

Each task must be assigned to one agent, but agents can receive multiple tasks. The designer's objective is to minimize the social cost of completing the tasks, where the social cost of an allocation $(\hat{n}_j^h, \hat{n}_j^l)_{j \in \mathcal{J}} \in \mathbb{R}_+^{2 \times \mathcal{J}}$ is given by $\sum_{j \in \mathcal{J}} c^h(j) \hat{n}_j^h + c^l(j) \hat{n}_j^l$, where $c^h, c^l \geq 0$. We refer to $c^k(j)$ as j 's *performance* on type- k tasks, so that a lower $c^k(j)$ means better performance.¹¹ The performance functions c^h, c^l are known to the designer.¹²

Observe that no mechanism will be able to distinguish between types $(p_j^h, p_j^l), (\hat{p}_j^h, \hat{p}_j^l)$ such that $\frac{p_j^h}{p_j^l} = \frac{\hat{p}_j^h}{\hat{p}_j^l}$, as these agents have identical preferences over assignments. Thus, it is without loss of generality to consider mechanisms which elicit only the relative cost of high-type tasks, $p_j := \frac{p_j^h}{p_j^l}$.¹³ Henceforth, we refer to this ratio as the agent's type and write mechanisms simply as a function of the one-dimensional type. We maintain the following assumption:

Assumption. The type $p_j := \frac{p_j^h}{p_j^l}$ has a full-support distribution on a bounded interval $[\underline{p}_j, \bar{p}_j]$ with CDF F_j and continuous density f_j . Types are independent across agents.

¹⁰Additionally, it is useful, although not essential, that the status-quo not involve perfect specialization. Quantitatively, this helps ensure that there are "gains from trade" to be had by reallocating tasks. Depending on the setting, it may also be needed for estimating performance (see Section [IV.D](#)).

¹¹Modeling the designer as minimizing rather than maximizing their objective, and agents as minimizing workload, is merely a normalization.

¹²Because the designer's objective is linear, it is equivalent to say that c^h, c^l are the expected costs. In other words, we can allow for noise in the designer's cost estimates.

¹³This also means that the designer does not benefit from making different assignments to agents with the same preferences; for a formal proof, see [Dworczak et al. \(2021\)](#) Theorem 8, where the same observation on the reduction of a two-dimensional to a one-dimensional type appears, albeit in a different setting.

In the baseline model, we assume that the distributions $(F_j)_{j \in \mathcal{J}}$ are known to the designer. In Section III, we show how our solution can be used to construct a mechanism which does not depend on knowledge of these distributions.

I.A Discussion of the modeling assumptions

The model described here is a reasonable fit for a range of settings where there is heterogeneity in tasks and performance, and the designer is able to estimate performance. We do, however, abstract from some dimensions which may be important in certain applications. We therefore discuss the modeling assumptions in some detail before proceeding.

Linearity of preferences: We maintain throughout the linear specification of agents' preferences. This clearly rules out settings in which there are important spillovers between tasks. We find this assumption palatable in some settings, such as the CPS application, for a few reasons. First, the status-quo constraint ensures that aggregate workloads do not increase, which is conceptually similar to allowing the agents' preferences to be highly convex in workload. Second, the problem we ultimately want to study is dynamic: an agent's task-load consists of tasks assigned at different points in time. Indeed, if each task were resolved before the next one began, to assume linear costs would just be to assume that payoffs are separable across periods. While there are certainly valid critiques of time separability, it is a standard assumption on preferences in dynamic settings (of course, tasks for a given agent may overlap, in which case linear costs are not precisely equivalent to time separability). Finally, the mechanism that we ultimately propose is robust to the assumption of linear preferences. It can be implemented regardless of the form that agents' preferences take. It still has the potential to improve upon the status quo, and at worst achieves the same social welfare. See Section III for details.

Agents' performance: Our framework assumes that the designer has access to some (perhaps noisy) estimate of performance. For the purposes of designing the mechanism, we treat agents' performance, as captured by the function c , as fixed. A concern is that changes to the mechanism might affect performance. One channel for this effect is that, if agents' workload increases, they may perform worse on each individual task. However, our status-quo constraint ensures that this does not occur. Another channel would be through the agents' incentives for effort. This would be especially concerning if agents could reduce their workload by degrading their performance. Fortunately, we show that in our proposed mechanism the scope for such manipulation is limited (Theorem 3).¹⁴ Moreover, our mechanism can

¹⁴This property of the mechanism is also intimately related to the perceived fairness of the mechanism, in that agents who receive more favorable assignments are precisely those whose performance is higher. Formally, our mechanism possesses an *envy-freeness* property that can help justify to agents why their task-loads are

accommodate performance which evolves over time, as well as the arrival of new agents, because we can continue to update the performance estimates even after the mechanism has been implemented (see Appendix B.1).

Two types of tasks: While our analysis assumes that there are just two types of tasks, the mechanism can be applied in settings with richer task heterogeneity by grouping tasks into any two categories (see Section IV.C). Still, this approach is of course most relevant when there is in fact one such categorization that is particularly salient for both performance and preferences. An alternative approach is to allow for more than two categories. In that case the problem becomes one of multi-dimensional screening, and it is not in general possible to identify an optimal mechanism (e.g., Hart and Reny (2015)). Moreover, mechanisms are more complicated, potentially diminishing their practical relevance. Nonetheless, this approach may be valuable in some settings, and is explored in Valenzuela-Stookey et al. (2026).

II Large-Market Static problem

We first consider a relaxed problem that abstracts from the combinatorial dimension of the original SMD problem. We do this by allowing for fractional assignments of tasks and by requiring only that each task be assigned to some agent in expectation (taken over agents' types) rather than ex-post (for every realized type profile). Formally, in the LMS program the designer chooses a $(H^j, L^j) : [\underline{p}_j, \bar{p}_j] \rightarrow \mathbb{R}_+^2$ to minimize

$$\sum_{j \in \mathcal{J}} \mathbb{E}_{p_j \sim F_j} [c^h(j)H^j(p_j) + c^l(j)L^j(p_j)]$$

subject to

$$pH^j(p) + L^j(p) \leq pH^j(p') + L^j(p') \quad \forall j \in \mathcal{J} \quad p, p' \in [\underline{p}_j, \bar{p}_j] \quad (\text{IC})$$

$$pH^j(p) + L^j(p) \leq pn^h + n^l \quad \forall j \in \mathcal{J} \quad p \in [\underline{p}_j, \bar{p}_j] \quad (\text{SQ})$$

$$\sum_{j \in \mathcal{J}} \mathbb{E}_{p_j \sim F_j} [H^j(p_j)] = Jn^h \quad (H\text{-capacity})$$

$$\sum_{j \in \mathcal{J}} \mathbb{E}_{p_j \sim F_j} [L^j(p_j)] = Jn^l \quad (L\text{-capacity})$$

The fact that the capacity constraints are required to hold only in expectation is what makes this a “large-market” relaxation; this formulation is equivalent to assuming that each agent j is actually a unit-mass population of agents with identical performance and preference types distributed according to F_j .¹⁵ The IC constraint says that agents are better off (receive a

no longer identical (see Section II.C).

¹⁵Note that under the large-market assumption it is without loss of generality for the mechanism to

lower workload) if they report their type truthfully (the standard revelation principle of Myerson (1981) applies here). The SQ (status-quo) constraint ensures that every agent is weakly better off relative to the original allocation.¹⁶

II.A Solution preview

Before proceeding to the solution, it is useful to take a step back. The problem that we ultimately want to solve is one of combinatorial allocation. For such problems, the natural benchmark is the class of market-based mechanisms based on competitive equilibrium (CE) (Varian, 1973; Hylland and Zeckhauser, 1979; Budish, 2011; Nguyen and Vohra, 2021; Prendergast, 2022; Nguyen et al., 2023). To build intuition for our solution, we first evaluate a standard market-based solution applied to this context. This comparison is instructive for understanding the distinctive features of the current problem.

To apply a CE mechanism to the current setting we would grant each agent an “endowment” equal to the expected status-quo assignment, denoted by (n^h, n^l) , and set a “price” p for high-type tasks in terms of low-type tasks. We then allow agent j to choose their favorite bundle from the budget set $\{(\hat{n}_j^h, \hat{n}_j^l) : p\hat{n}_j^h + \hat{n}_j^l \geq pn^h + n^l\}$. The price p should be set so that the market clears, i.e., all tasks are assigned. Given such a price, the allocation is efficient and fair; no agent can be made better off without making some other agent worse off, and no agent would prefer another’s assignment to their own (Varian, 1973). In finite markets, agents may be able to influence the price by distorting their demand, but this incentive disappears as the market grows (Roberts and Postlewaite, 1976).¹⁷ By construction, the status quo is always affordable, so no agent is worse off.

The downside of this market mechanism is that while it respects the preferences of agents, it ignores those of the designer: the efficiency of CE concerns the preferences of agents, while the designer’s objective concerns agents’ performance, which does not enter into the construction. This is the key distinction between our setting and those to which market-type mechanisms are typically applied (e.g., Budish (2011); Prendergast (2022)). The natural response is to introduce personalized prices. Suppose that agent j performs well on high-type tasks and poorly on low-type ones. Intuitively, we might try to steer j towards the former by increasing p^j , the number of additional low-type tasks j must take on in exchange for one fewer high-type

specify each agent’s assignment as a function of their own type, rather than the entire type profile. This is precisely the benefit of the large-market formulation, as explained in detail below.

¹⁶An alternative would be to maximize a weighted sum of social welfare and agent costs. This approach (or the dual of minimizing cost subject to a social welfare constraint) may be suitable in some settings. One drawback, however, is that the designer would need to take a stand on the relative weights of social cost and agent welfare; an unfamiliar, difficult, and politically fraught exercise in many settings (such as CPS).

¹⁷Existence of market clearing price is not guaranteed with indivisible goods. However, the CE mechanism can be well approximated in a way that preserves the efficiency and incentive properties (Budish, 2011).

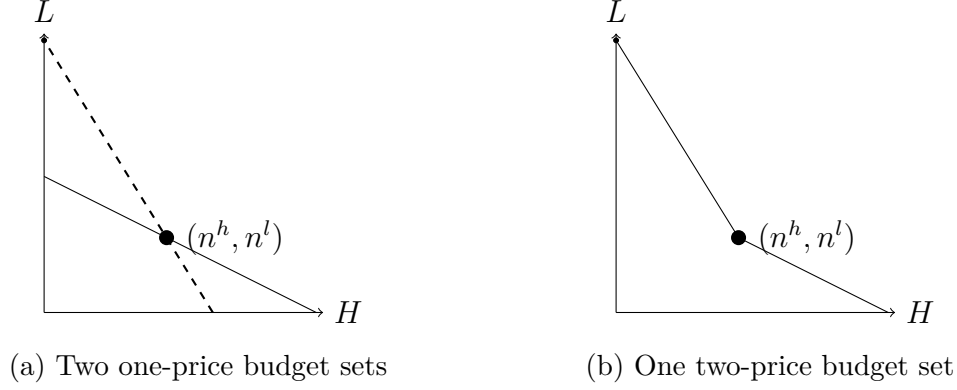


Figure 1: Budget sets

task, and allowing j to choose from the budget set $\{(\hat{n}_j^h, \hat{n}_j^l) : p^j \hat{n}_j^h + \hat{n}_j^l \geq p^j n^h + n^l\}$.

As j 's budget is determined by the value of the endowment, increasing p^j rotates the budget set around (n^h, n^l) . Figure 1a depicts the budget line. An increase in p^j corresponds to a rotation from the solid to the dotted budget set. The higher is p^j , the less attractive it is for j to take on additional low-type tasks. However, a larger p^j also means that j will perform fewer additional high-type tasks for each low-type task they give up. Thus there is a trade-off between increasing the probability that j specializes in high-type tasks and ensuring that j handles their fair share of the overall workload. This trade-off suggests that non-linear pricing could be useful. Suppose that in order to trade *away* a high-type task, j is forced to take on an additional p_2^j low-type tasks, while if j wants to trade *for* a high-type task they can give up p_1^j low-type tasks, where $p_1^j < p_2^j$. The induced budget set is depicted in Figure 1b. By increasing p_2^j above p_1^j we make it less attractive for j to give up high-type tasks, without affecting the rate at which they can give up low-type tasks. Instead, the kink in the budget set at (n^h, n^l) increases the likelihood that j opts to retain the status quo.

Given the potential value of non-linear pricing, we can consider even more flexible schemes. In the extreme we could allow the exchange rate between high- and low-type tasks to vary continuously in the space of task bundles. Surprisingly, additional flexibility is not needed. We say that F_j is *doubly regular* if $p \mapsto \phi_j(p) := p - \frac{1-F_j(p)}{f_j(p)}$ and $p \mapsto p + F_j(p)/f_j(p)$ are strictly increasing, and has *low asymmetry* if $1 - \frac{p_j}{p} \leq F_j(p) \leq \frac{p}{\bar{p}_j} \quad \forall p \in [\underline{p}_j, \bar{p}_j]$. We call F_j *strongly regular* if it is doubly regular and has low asymmetry.¹⁸

Proposition 1. If type distributions are strongly regular, then a personalized two-price

¹⁸Double regularity is just the usual buyer and seller Myerson regularity conditions. Intuitively, low asymmetry bounds the ratio of $p/\bar{p}_j$; if F is uniform, it is satisfied if and only if $\bar{p}_j/\underline{p}_j \leq 2$. See Valenzuela-Stookey (2025) for further discussion of this condition.

mechanism, i.e., one in which each agent receives a budget set as in Figure 1b, is optimal in the LMS problem, and the solution is unique if no two agents have the same performance for either task type. Absent any regularity conditions, at most four prices are needed for each agent (i.e., three kinks in the budget set).

This result is implied by Theorem 1 below. It only describes the qualitative features of the mechanism. The final crucial step (Theorem 2) is to derive the optimal prices for each agent.

II.B Solving the LMS program

To solve the LMS program, we make use of the fact that both the objective and the market-clearing conditions depend only on the expected task-loads for each agent. Thus, we can solve the problem in two parts. First, in an “inner problem” we characterize for each agent j the expected task-loads, $(\mathbb{E}_j[H^j(p)], \mathbb{E}_j[L^j(p)])$, that j can be assigned by some IC and SQ mechanism. We refer to this as the set of *incentive-feasible* pairs, denoted by $\mathcal{F}_j$. We then solve an “outer problem” in which for each j we choose an incentive-feasible expected task-load $(\hat{n}_j^h, \hat{n}_j^l) \in \mathcal{F}_j$ to minimize the designer’s objective, subject to the market-clearing conditions. The inner problem deals with IC and SQ, and the outer with market clearing.

It turns out that the dual to the outer problem then has a convenient formulation in terms of the support functions of the sets $(\mathcal{F}_j)_{j \in \mathcal{J}}$. This dual formulation is computationally useful and also facilitates analytical comparative statics. We view this solution technique (separating inner and outer programs and solving the outer using the support function of the incentive-feasible set) as one of the paper’s key methodological contributions.

Step 1: LMS inner problem: The inner problem concerns the design of a mechanism for a single agent, and so we drop the dependence on j in the notation here. It is easy to see that the set of incentive-feasible pairs, $\mathcal{F}$, is convex, since the IC and SQ constraints are linear in the mechanism (H, L) . This set can thus be described by its support function $S : \mathbb{R}^2 \rightarrow \mathbb{R}$ defined by $S(a, b) := \max\{a\hat{n}^h + b\hat{n}^l : (\hat{n}^h, \hat{n}^l) \in \mathcal{F}\}$, which yields the dual characterization $\mathcal{F} = \{(\hat{n}^h, \hat{n}^l) : a\hat{n}^h + b\hat{n}^l \leq S(a, b) \forall (a, b) \in \mathbb{R}^2\}$. In order to calculate the support function, we need to maximize over the set $\mathcal{F}$. The way to do this is to maximize directly over the set of IC and SQ mechanisms. That is, for arbitrary $(a, b) \in \mathbb{R}_+^2$, we solve

$$\begin{aligned}
S(a, b) = \max_{H, L \geq 0} \quad & a \int H(p) dF(p) + b \int L(p) dF(p) & (1) \\
s.t \quad & -pH(p) - L(p) \geq -pH(p') - L(p') \quad \forall p, p' \in [\underline{p}, \bar{p}] & (\text{IC}) \\
& -pH(p) - L(p) \geq -pn^h - n^l \quad \forall p \in [\underline{p}, \bar{p}] & (\text{SQ})
\end{aligned}$$

Note that while there are no transfers in our setting, we can think of H as playing the role of the physical allocation and L that of transfers. Thus, the program defining S shares many similarities with the classic monopoly pricing problem of Myerson (1981). The two essential distinctions between this program and the monopoly-pricing problem are (i) a non-negativity constraint on L , and (ii) a type-dependent participation constraint determined by the need to respect the status quo. Lower bounds on transfers, L in the current context, are studied by Loertscher and Muir (2021). However, their problem does not feature a type-dependent participation constraint. Technically, the closest existing paper is the contemporaneous work of Dworczak and Muir (2025); translated into the current context, their problem restricts H to lie in $[0, 1]$, and imposes no constraints on L .

As in most problems of countervailing incentives, the main challenge lies in identifying for which types the SQ constraint should bind. It turns out that SQ binds on an interval of intermediate types, and the solution takes the following form.¹⁹

Theorem 1. For any $(a, b) \in \mathbb{R}^2$ there is an optimal mechanism (H^*, L^*) defining the support function $S(a, b)$ which takes the following form: there exist three thresholds $\underline{p} \leq p_1 \leq p_2 \leq p_3 \leq \bar{p}$ and a level $n^h \leq H_2 \leq n^h + \frac{1}{p_2}n^l$ such that

$$(H^*(p), L^*(p)) = \begin{cases} \left(H_2 + \frac{n^l - p_2(H_2 - n^h)}{p_1}, 0 \right) & \text{if } p \in [\underline{p}, p_1) \\ (H_2, n^l - p_2(H_2 - n^h)) & \text{if } p \in [p_1, p_2] \\ (n^h, n^l) & \text{if } p \in (p_2, p_3] \\ (0, p_3 n^h + n^l) & \text{if } p \in (p_3, \bar{p}] \end{cases}$$

Proof. Proof in Appendix A.1. □

Theorem 1 says that the mechanism maximizing a weighted sum of expected H and L task-loads takes on no more than four distinct values; one intermediate set of types (between p_2 and p_3) who retain the status-quo assignment, one set above who get only low-type tasks, and two sets below. The reason there are two assignment levels below the status quo, as opposed to only one above, is that in this region the non-negativity constraint on L may bind, and ironing under this additional constraint can give rise to an additional assignment level. However, under strong regularity of F it is without loss to consider only two-part mechanisms when maximizing a non-decreasing objective.

¹⁹That SQ binds on an intermediate interval is perhaps not surprising to those familiar with Jullien (2000), given the linearity of the SQ constraint. However, the proof of Theorem 1 uncovers additional structure of the solution, which is useful elsewhere (e.g. Corollary 1 and Theorem 3).

Corollary 1. If F is doubly regular then for any $(a, b) > 0$ there is a unique (up to zero-measure perturbations) optimal mechanism defined by thresholds $p_1 \leq p_2$ such that

$$(H^*(p), L^*(p)) = \begin{cases} (n^h + \frac{1}{p_1}n^l, 0) & \text{if } p \leq p_1 \\ (n^h, n^l) & \text{if } p \in (p_1, p_2) \\ (0, p_2n^h + n^l) & \text{if } p \geq p_2. \end{cases}$$

Proof. Proof in Appendix A.1. □

As discussed above, such a mechanism has a simple indirect implementation, in which the agent is presented with a kinked budget set as in Figure 1b and allowed to exchange tasks. Without regularity, the budget set features two kinks: one at the endowment, and one at a point of more high-type and fewer low-type tasks.

Theorem 1 greatly simplifies the problem of solving for the value $S(a, b)$. Moreover, it tells us what a mechanism that achieves the value $S(a, b)$ will look like, which allows us to characterize $N^*(a, b) := \arg \max\{a\hat{n}^h + b\hat{n}^l : (\hat{n}^h, \hat{n}^l) \in \mathcal{F}\}$. Define the *frontier* to be the set $\{N^*(a, b) : \neg\{(a, b) \leq 0\}\}$, i.e., the set of solutions when at least one of a, b is strictly positive. The set $\mathcal{F}$ is the subset of the positive orthant that lies within the frontier. Abusing terminology, we say that $\mathcal{F}$ is *strictly convex* if the mixture of any two points on the frontier lies in the interior of $\mathcal{F}$. Say that $\mathcal{F}$ is *downward closed* if $(x, y) \in \mathcal{F}$ and $0 \leq (x', y') \leq (x, y)$ implies $(x', y') \in \mathcal{F}$. So $\mathcal{F}$ is downward closed if and only if the frontier is downward sloping.

Corollary 2. If F is strongly regular then $\mathcal{F}$ is strictly convex and downward closed.

Proof. Proof in Appendix C.1. □

Remark 2. Theorem 1 states that the value of $S(a, b)$ can be obtained with a mechanism involving at most 4 allocations. However, Proposition 1 allowed that the optimal mechanism could involve a mechanism involving 4 prices, and thus 5 allocations, for some agents. The discrepancy arises because while the value $S(a, b)$ can be obtained with a 3-price mechanism, 4 prices might be needed to implement some points in $\mathcal{F}$. This occurs precisely when the frontier of $\mathcal{F}$ has linear segments. However, our primary focus is on the strongly regular task, where $\mathcal{F}$ is strictly convex.

Step 2: LMS outer problem: Theorem 1 shows us how to characterize the incentive-feasible set for any type distribution. In particular, it allows us to easily compute the support function for this set. We now use this characterization to identify the optimal mechanism for the LMS problem. We begin with a convex set $\mathcal{F}_j \subset \mathbb{R}_+^2$ with a support function $S_j : \mathbb{R}^2 \rightarrow \mathbb{R}_+$ for each j . Let $N_j^*(a, b) := \arg \max\{a\hat{n}^h + b\hat{n}^l : (\hat{n}^h, \hat{n}^l) \in \mathcal{F}_j\}$. We study the optimal division of

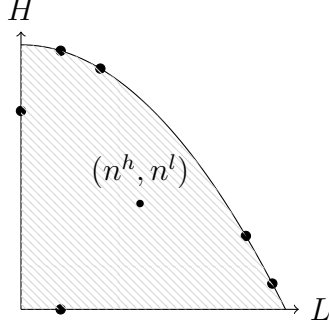


Figure 2: Outer problem with identical (strongly regular) type distributions

tasks among the agents, such that the task-load for each j is an element of $\mathcal{F}_j$. That is, we want to solve:

$$\begin{aligned} \min_{(\hat{n}_j^h, \hat{n}_j^l)_{j=1}^J} \sum_{j=1}^J c^h(j) \hat{n}_j^h + c^l(j) \hat{n}_j^l \quad s.t. \quad & (\hat{n}_j^h, \hat{n}_j^l) \in \mathcal{F}_j \quad \forall 1 \leq j \leq J \\ & \sum_{j=1}^J \hat{n}_j^h = Jn^h \quad , \quad \sum_{j=1}^J \hat{n}_j^l = Jn^l \end{aligned} \quad (2)$$

If type distributions are symmetric, so that $\mathcal{F}_j = \mathcal{F}$ for all j , then we can think of this as the problem of choosing a set of points $\mathcal{F}$ with barycenter equal to (n^h, n^l) , as illustrated in Figure 2, where $\mathcal{F}$ is depicted as the shaded region.

We solve the outer problem by studying its dual. Let λ_h, λ_l be the dual variables for the market-clearing constraints. Then, the previous program is equivalent to

$$\begin{aligned} \min_{(\hat{n}_j^h, \hat{n}_j^l)_{j=1}^J} \max_{\lambda_h, \lambda_l} \sum_{j=1}^J c^h(j) \hat{n}_j^h + c^l(j) \hat{n}_j^l + \lambda_h \left(Jn^h - \sum_{j=1}^J \hat{n}_j^h \right) + \lambda_l \left(Jn^l - \sum_{j=1}^J \hat{n}_j^l \right) \\ s.t. \quad (\hat{n}_j^h, \hat{n}_j^l) \in \mathcal{F}_j \quad \forall 1 \leq j \leq J \quad \text{(incentive-feasible)} \end{aligned}$$

Strong duality holds (see Appendix A.2), so this is equivalent to

$$\max_{\lambda_h, \lambda_l} \lambda_h Jn^h + \lambda_l Jn^l - \sum_{j=1}^J \max \{ (\lambda_h - c^h(j)) \hat{n}^h + (\lambda_l - c^l(j)) \hat{n}^l : (\hat{n}^h, \hat{n}^l) \in \mathcal{F}_j \} \quad (3)$$

In words, the dual variables λ_h, λ_l are the marginal social costs of adding high- and low-type tasks, respectively. Alternatively, we can view λ_h, λ_l as the “benefit” paid to the designer for completing each task. Then, the dual program can be interpreted as a competitive market in which firms seek to maximize cost minus benefit, and λ_h, λ_l adjust to clear the market.

Using the definition of the support function, we can rewrite the dual as

$$\max_{\lambda_h, \lambda_l} \lambda_h Jn^h + \lambda_l Jn^l - \sum_{j=1}^J S_j((\lambda_h - c^h(j)), (\lambda_l - c^l(j))) \quad (4)$$

Support functions are convex and continuous. Thus, the objective in (4) is concave in (λ_h, λ_l) . Using this formulation, we can simplify the outer problem of choosing incentive-feasible pairs $(\hat{n}_j^h, \hat{n}_j^l)$ for each agent, to the much simpler two-dimensional dual. Moreover, this formulation allows us to identify quantitative features of the solution and perform comparative statics. We say that the performance profile is *generic* if $c^k(j) \neq c^k(j')$ for all $j, j' \in \mathcal{J}$ and $k \in \{h, l\}$.

Theorem 2. Let (λ_h, λ_l) solve the dual program in (4). Then, there exist selections from $N_j^*(\lambda_h - c^h(j), \lambda_l - c^l(j))$ such that:

$$\sum_{j=1}^J \hat{n}_j^h = Jn^h \quad \text{and} \quad \sum_{j=1}^J \hat{n}_j^l = Jn^l,$$

and these constitute a solution to the social-cost minimization program. If the performance profile is generic then in any solution at most one agent is off the boundary of $\mathcal{F}_j$, and at most two agents have non-zero allocations that are off the frontier. If, moreover, every agent also has a strongly regular type distribution then there is a unique solution, $(\hat{n}_j^h, \hat{n}_j^l)_{j=1}^J$, to the social-cost minimization problem.

Proof. Proof in Appendix A.2. □

The optimal mechanism is determined by the performance parameters (through the objective function), and the type distributions (which determine the incentive-feasible sets $\mathcal{F}_j$). The allocation rule for agent j is the solution to the LMS inner problem for j for weights $(a, b) = (\lambda_h - c^h(j), \lambda_l - c^l(j))$. As discussed above, if each F_j is strongly regular then this allocation is characterized by a pair of prices (p_1^j, p_2^j) at which j is allowed to trade given their induced kinked budget set. We refer to the optimal mechanism under strong regularity as the *LMS two-price (LMS-TP) mechanism*, and focus primarily on this case.

We refer to agents whose assignment is off the frontier as *remedial*. Such agents belong to one of four categories. If $\lambda_h - c^h(j) = 0$ and $\lambda_l - c^l(j) < 0$ then j receives only type- h tasks, if $\lambda_l - c^l(j) = 0$ and $\lambda_h - c^h(j) < 0$ then they receive only type- l tasks, and if both $\lambda_l - c^l(j) < 0$ and $\lambda_h - c^h(j) < 0$ then j receives no tasks. The remaining remedial agents have $\lambda_l - c^l(j) = \lambda_h - c^h(j) = 0$, and only these can receive allocations in the interior of $\mathcal{F}_j$. Except for this last group, the assignment of remedial agents is independent of their type.

By the envelope theorem (Milgrom and Segal, 2002), S is differentiable almost everywhere, and if $N_j^*(a, b) = (x^*, y^*)$ then the right derivative of $S_j(a, b)$ with respect to the first argument is $\max\{x^*\}$, and with respect to the second argument is $\max\{y^*\}$. Along with

Theorem 2, having access to this derivative facilitates efficient computation of the optimal mechanism. Moreover, the dual formulation of the outer problem allows us to perform comparative statics and study agents’ incentives for effort.

Remark 3. The same approach of dividing the problem into inner and outer components can be applied regardless of the number of task types. Indeed, the outer problem is not fundamentally different in higher dimensions. The challenge lies in solving the inner problem when preferences are multi-dimensional. However, an approximate solution to the inner problem could be used to characterize a lower bound on the support function. This in turn could be used to identify an approximate solution to the outer problem, and thus to the full program. This is the approach employed by Valenzuela-Stookey et al. (2026).

II.C Fairness and incentives for effort

Our approach treats $c^h(j)$ and $c^l(j)$ as policy-invariant parameters. However, a natural concern in any performance-based assignment mechanism is whether it gives agents the right performance incentives. The concern is that agents might intentionally perform worse today if they expect their performance data to be used in the future to re-design the mechanism, particularly if low-performing workers are rewarded with lower task-loads. While not our primary focus, it turns out that in our mechanism the scope for such manipulation is limited.

To talk about the agents’ incentives to perform, we need to make explicit the mechanism’s dependence on the performance parameters. We thus think of the mapping from $(c^h(j), c^l(j))_{j \in J}$ to the LMS-TP mechanism as itself a meta-mechanism mapping performance parameters and type reports to allocations. We refer to this simply as the *optimal LMS-TP mechanism*.

The question of whether the optimal LMS-TP mechanism delivers the correct incentives for effort is fundamentally about its comparative statics in $(c^h(j), c^l(j))_{j \in J}$. Consider, for example, an agent who improves their performance on type- h tasks, holding that on type- l tasks fixed. Intuitively, the mechanism should try to assign this agent to more type- h tasks in expectation. From an ex-ante perspective, i.e., without knowing the agent’s type, there are two ways to do this: (i) ensure that the agent receives a large number of high-type tasks in the event that they report a low type, or (ii) increase the size of this event, i.e., the probability that the agent specializes in high-type tasks. These two objectives are at odds: we must lower p_1^j to achieve (i), and raise it to achieve (ii).²⁰ Notice, however, that if the second objective dominates, the agent receives better terms for exchanging for high-type tasks the better their performance on these tasks. Indeed, this is the case under the optimal

²⁰The intuition here is incomplete, since we also need to consider changes in p_2^j . The proof in Appendix A.3 deals with this additional complication.

LMS-TP mechanism under a slight strengthening of strong regularity.

Proposition 2. Assume F_j is strongly regular and $F(p)/p$ is increasing. Then p_1^j is decreasing in $c^h(j)$ and p_2^j is increasing in $c^l(j)$ for any j on their frontier.

In other words, conditional on specializing in type- k tasks, j benefits from reducing $c^k(j)$. Proposition 2 is proven along with Theorem 3 below using the dual in eq. (4). The intuitive connection between Proposition 2 and effort incentives can be formalized as follows.

Definition. A mechanism is *locally effort-inducing* for j if the following holds: if j reports their type truthfully and receives a type- k case, then j 's payoff must be locally increasing in their performance on type- k cases (i.e., decreasing in $c^k(j)$).

To understand this definition, consider a two-period model of mechanism design. In the first period, a mechanism is designed and implemented for the LMS problem, based on some initial estimates of $(c^h(j), c^l(j))_{j \in \mathcal{J}}$. In the second period, the outcomes from the first period are used to update the performance estimates, and a new mechanism is designed and implemented. Suppose that agent j expects the performance of other agents to remain unchanged from the first to the second period, but in the first period can choose to degrade their own performance when assigned a type- k case, so as to increase $c^k(j)$. That the mechanism implemented in both periods is locally effort-inducing means precisely that j cannot gain by such degradation (for small increases in $c^k(j)$), conditional on having truthfully reported their type in the first period. To be clear, this does not rule out the possibility that j could profit from the double deviation of misreporting their type in period 1 and degrading their performance on the cases they are assigned. However, such deviations are costly, since j must take on a less-preferred caseload in period 1 in order to potentially improve their assignment in period 2. We therefore view local effort-inducing as a real, albeit qualified, restriction on the potential scope for manipulation by degrading performance.

A closely related condition concerns the fairness of a mechanism. Say that an agent j with type p_j *envies* agent j' if j' is not excluded (i.e., j' receives some cases), and j would prefer to be offered the mechanism $(H^{j'}, L^{j'})$ rather than (H^j, L^j) . We say that agent j 's envy is *justified* if, moreover, they have the same type distribution and either $H^j(p_j) > 0, L^j(p_j) = 0, c^h(j) < c^h(j')$, and $c^l(j) = c^l(j')$; or $H^j(p_j) = 0, L^j(p_j) > 0, c^l(j) < c^l(j')$, and $c^h(j) = c^h(j')$. To understand this definition, suppose j and j' are both asked to specialize in type- h cases, but j has justified envy for j' . This means that j' handles fewer type- h cases than j , despite the fact that j performs better on these cases, and both perform the same on type- l

cases.²¹ Such an outcome is arguably unfair to agent j .²²

Definition. A mechanism is *locally fair* for agent j with type p_j if there exists $\epsilon > 0$ such that there is no j' with $|c^h(j) - c^h(j')| + |c^l(j) - c^l(j')| < \epsilon$ for which j has justified envy.

While this is a limited notion of fairness, it implies that disparities in task-loads can at least be well-justified between similar agents on the basis of performance.

Theorem 3. Assume F_j is strongly regular and $F(p)/p$ is increasing. Then, for each agent j , the optimal LMS-TP mechanism is locally fair for agent j , regardless of their type. Moreover, the optimal LMS-TP mechanism is locally effort-inducing for all agents except (potentially) those that receive non-zero allocations off the frontier (of which there are at most two for generic performance profiles).

Proof. Proof in Appendix A.3. □

This result depends on restrictions on the type distributions. If fairness and effort concerns are important in practice, the designer can impose these conditions on the distribution. If the distributions are misspecified, the solutions to the SMS and SMD problems derived from the LMS-TP mechanism will be sub-optimal, but they continue to satisfy IC and SQ.

Remark 4. Theorem 3 is stated for the LMS-TP mechanism, but these properties translate approximately to the SMS and SMD mechanisms described below. In fact, we can say a bit more: both the SMS and SMD mechanisms are constructed using the prices $(p_1^j, p_2^j)_{j=1}^J$ defined in the LMS-TP mechanism, and so Proposition 2 remains directly relevant.

III Small-Market Static and Dynamic problems

III.A Small market static

In the previous section, we solved a relaxed problem in which we only required that all tasks be assigned in expectation. The problem is more difficult if we require that all tasks be assigned ex-post, i.e., conditional on each realized type profile—rather than just in expectation. While it may be possible to solve for the optimal Bayesian incentive compatible (BIC) mechanism in the SMS problem directly, doing so would sacrifice the simplicity of

²¹We require equal performance for cases that j is not assigned to guarantee that j and j' are roughly comparable agents. Strict equality is not important, it would suffice for their performance on these cases to be similar. The reason we need the agents to be similar has to do with comparative advantage. Suppose j is assigned only type- h cases. If j' performs worse than j for both case types, but is significantly worse for the type- l cases, then j' may still have a comparative advantage for type- h cases. Thus, it would be justifiable for the mechanism to match j' with no type- l cases and still give j' relatively few type- h cases.

²²The definition of justified envy does not cover the case where j retains the status quo. This is because the status quo is the same for all investigators, so j cannot have justified envy for another investigator who also retains the status quo. This is the only relevant comparison for local fairness, which is our focus here.

the LMS-TP mechanism.²³ We opt instead to modify the LMS solution to accommodate the ex-post market-clearing condition. This approach has the additional benefit that we obtain a mechanism which is strategy proof (in fact Obviously Strategy-Proof as in Li (2017)) and robust to misspecification of the type distribution. Moreover, the status-quo constraint will be satisfied ex-post (for any realized type profile) rather than just ex-interim. The trade-off is that our solution is only approximately optimal.

Assume that all agents' type distributions are strongly regular (we can apply a similar approach without strong regularity.) Let $\mathcal{E}$ be the set of non-remedial agents in the optimal LMS-TP mechanism, i.e., those agents whose expected allocation is on the frontier of their incentive-feasible set. Let $(p_1^j, p_2^j)_{j \in \mathcal{E}}$ be the prices defining the optimal LMS-TP mechanism for the non-remedial agents. Let $P = (p_j)_{j=1}^J$ be a type profile. Fixing the mechanism, $j \in \mathcal{E}$ is a *buyer* (of type- h tasks) if $p_j \leq p_1^j$, a *seller* if $p_j > p_2^j$, and retains the status quo otherwise. Let $\mathcal{B}$ be the set of buyers and $\mathcal{S}$ the set of sellers.

We wish to approximate the LMS-TP mechanism in the small market, while ensuring ex-post market clearing. The basic idea is to offer each non-remedial agent the same prices (p_1^j, p_2^j) that they would be given in LMS-TP, and ask them whether they would like to trade for high-type tasks, trade away high-type tasks, or retain the status quo. They are then guaranteed an allocation that lies on their budget set, in the direction that they choose to move. However, whereas in the LMS-TP mechanism all agents either retain the status quo or go to a corner, here we may restrict the volume of trade. For remedial agents the assignment is independent of their reported type. Subject to these constraints, and market clearing, the allocation is chosen to minimize social cost. Concretely, the mechanism is the following.

Small-market static two-price (SMS-TP) mechanism. Assign (n^h, n^l) to all $j \in \mathcal{E} \setminus (\mathcal{B} \cup \mathcal{S})$, i.e., the non-remedial agents who are neither buyers nor sellers. To allocate the remaining tasks, solve the linear program:

$$\begin{aligned} \min_{\{b_j\}_{j \in \mathcal{B}}, \{s_j\}_{j \in \mathcal{S}}, \{\dot{n}_j^h, \dot{n}_j^l\}_{j \notin \mathcal{E}}} & \sum_{j \in \mathcal{B}} b_j (c^h(j) - p_1^j c^l(j)) - \sum_{j \in \mathcal{S}} s_j (c^h(j) - p_2^j c^l(j)) + \sum_{j \notin \mathcal{E}} (\dot{n}_j^h c^h(j) + \dot{n}_j^l c^l(j)) \\ \text{s.t.} & \quad \bar{p}_j \dot{n}_j^h + \dot{n}_j^l \leq \bar{p}_j n^h + n^l \quad \forall j \notin \mathcal{E} \\ & \quad \underline{p}_j \dot{n}_j^h + \dot{n}_j^l \leq \underline{p}_j n^h + n^l \quad \forall j \notin \mathcal{E} \\ & \quad 0 \leq \dot{n}_j^h, 0 \leq \dot{n}_j^l \quad \forall j \notin \mathcal{E} \quad ; \quad 0 \leq b_j \leq \frac{n^l}{p_1^j} \quad \forall j \in \mathcal{B} \quad ; \quad 0 \leq s_j \leq n^h \quad \forall j \in \mathcal{S} \end{aligned}$$

²³The main challenge is that in the small-market problem there are additional constraints on the interim allocations and standard characterizations (e.g., Border (1991)) do not apply to this multi-item setting. More recent developments in Valenzuela-Stookey (2023) can be used to characterize interim allocations for this setting, but how to design the optimal mechanism remains an open question.

$$\begin{aligned} \sum_{j \in \mathcal{B}} (n^h + b_j) + \sum_{j \in \mathcal{S}} (n^h - s_j) + \sum_{j \notin \mathcal{E}} \dot{n}_j^h + |\mathcal{E} \setminus (\mathcal{B} \cup \mathcal{S})| n^h &= J n^h & (h\text{-capacity}) \\ \sum_{j \in \mathcal{B}} (n^l - p_1^j b_j) + \sum_{j \in \mathcal{S}} (n^l + p_2^j s_j) + \sum_{j \notin \mathcal{E}} \dot{n}_j^l + |\mathcal{E} \setminus (\mathcal{B} \cup \mathcal{S})| n^l &= J n^l & (l\text{-capacity}) \end{aligned}$$

Let $(b_j^*)_{j \in \mathcal{B}}, (s_j^*)_{j \in \mathcal{S}}, (\dot{n}_j^h, \dot{n}_j^l)_{j \notin \mathcal{E}}$ be the solution to this program. The assignment for $j \notin \mathcal{E}$ is $(\dot{n}_j^h, \dot{n}_j^l)$. For $j \in \mathcal{B}$ the assignment is $(n^h + b_j^*, n^l - p_1^j b_j^*)$, and for $j \in \mathcal{S}$ it is $(n^h - s_j^*, n^l + p_2^j s_j^*)$.

To assess the asymptotic performance of the SMS-TP mechanism in finite markets, consider a sequence of “replica economies” in which there are y copies of each agent. Let $V_{SMS}(\{F\}_{j=1}^J | y)$ be the expected social cost achieved by SMS-TP in the y -replica economy, given the profile of type distributions $\{F\}_{j=1}^J$. Let $V_{OPT}(\{F\}_{j=1}^J | y)$ be the cost achieved by the (unknown) optimal SMS mechanism. The source of the divergence between the small market and large market is that in the former we do not know ex-ante the mass of agents who will be buyers and sellers of high-type tasks. Unsurprisingly, the SMS-TP mechanism is a better approximation to the optimal mechanism as this aggregate uncertainty about the agents’ types decreases, so that the small market approaches the large-market idealization. A mechanism is *strategy-proof* if truthful reporting is optimal for each agent, regardless of the type reports made by others.

For nearly all agents, the allocation that they receive under LMS-TP will be (approximately) feasible in SMS-TP as the market grows. The only exception is if there is a j such that 1) $\lambda^h - c^h(j) = \lambda^l - c^l(j) = 0$, and 2) $p \hat{n}_j^h + \hat{n}_j^l > p n^h + n^l$ for some $p \in [\underline{p}_j, \bar{p}_j]$, where $\hat{n}_j^h, \hat{n}_j^l, \lambda^h, \lambda^l$ are solutions to the primal and dual LMS outer programs. If all agents have distinct performance (which holds generically in applications) then there is at most one such agent. If there is no such agent, say that the LMS solution is *standard*.

Proposition 3. In the small-market static setting, SMS-TP is strategy-proof and respects the status quo ex-post. Moreover, if F_j satisfies strong regularity for all $j \in \mathcal{J}$ and the LMS solution is standard, then the value $V_{SMS}(\{F\}_{j=1}^J | y)$ converges to $V_{OPT}(\{F\}_{j=1}^J | y)$ as either

- i. $y \rightarrow \infty$, and/or
- ii. F_j converges in distribution to a constant for all j .

Proof. Proof in Appendix C.2. □

It is possible to modify the SMS-TP solution to extend the asymptotic result to the case where the LMS solution is non-standard. However, the difference between the standard and non-standard cases concerns (generically) at most one agent, and in finite markets the modified mechanism is nearly identical to SMS-TP. We therefore focus on the standard case to avoid introducing additional notation.

Strategy-proofness follows from linearity of the agents' preferences and the fact that prices are fixed ex-ante. In fact, it is easy to see that SMS-TP is Obviously Strategy-Proof (Li, 2017). However, as an indirect mechanism, SMS-TP can be implemented even if agents' preferences are not linear. To do this we can ask each agent to choose between being a buyer, a seller, or retaining the status quo at the stated prices, rather than eliciting type reports directly. From the designer's perspective, the worst possible outcome is that all agents report that they would like to retain the status quo. Thus, the designer can never do worse than the status quo, and does strictly better as long as some agents are willing to buy and sell.

Additionally, the mechanism is robust to misspecification of the type distribution.

Proposition 4. Regardless of the true type distributions, the SMS-TP mechanism based on distributions $(F_j)_{j \in \mathcal{J}}$ is (obviously) strategy-proof and respects the status-quo constraint. Moreover, the expected social cost of the mechanism is no worse than that of the status quo.

Proposition 4 is an immediate implication of the fact that prices are fixed and the status quo is feasible in the linear program defining SMS-TP.

III.B Small market dynamic

In dynamic settings, tasks arrive over time and must be assigned “online” without knowledge of future arrivals. To go from the static to the dynamic setting, we develop a mechanism to approximately implement the SMS-TP mechanism, where the approximation in this case gets better the longer the time horizon.

The dynamic model is as follows. Time is continuous and runs from 0 to T .²⁴ In each period one task may arrive. Let $\tau_t \in \{h, l, 0\}$ be the type of the task in period t , where $\tau_t = 0$ if no task arrives in period t . Denote by $N^k(t)$ the number of type- k tasks which have arrived up to and including time t , and let $\bar{n}^k(t) = \frac{1}{T}N^k(t)$.

Agents report their type only once, at time zero. The payoff of agent j who receives a cumulative task-load of $(\hat{n}_j^h, \hat{n}_j^l)$ by time T is $p_j \hat{n}_j^h + \hat{n}_j^l$. That is, agents care about their total undiscounted workload.²⁵ We start by estimating the total number of type k tasks per agent that will arrive by time T ; denote the estimate by n^k for $k \in \{h, l\}$. Given (n^h, n^l) , we solve for the SMS-TP assignment, which we denote by $(\hat{n}_j^h, \hat{n}_j^l)_{j=1}^J$. For concreteness, assume that high- and low-type tasks arrive at Poisson rates ρ^h and ρ^l respectively, so $n^k = \frac{T}{J}\rho^k$. This

²⁴The assumption of continuous time simplifies the discussion but is not essential; in the empirical application we modify the mechanism to run in discrete time.

²⁵Ultimately, every task needs to be assigned as it comes, so there would be no scope for the designer to take advantage of agents' discounting of future payoffs by back-loading tasks. Moreover, the mechanism we propose here smooths each agent's workload evenly over time regardless of their type report, so discounting should not significantly affect incentives to report truthfully.

assumption is not necessary, but simplifies the exposition.²⁶

Index each task by the time at which it arrives. Let z_t be the agent to which task t is assigned. For each j we keep track of their running task count $\hat{n}_j^k(t)$, i.e. the number of type- k tasks that they have been assigned thus far. Define the *score* $r_j(t, k) = \hat{n}_j^k(t)/\dot{n}_j^k$, where $r_j(t, k) = \infty$ if $\dot{n}_j^k = 0$.

SMD-TP mechanism. For each time t at which a task arrives, assign it to the agent with the lowest value of $r_j(t, \tau_t)$ (using any tie-breaking rule).

We can view this mechanism as a generalization of a standard rotation: in the standard rotation agent j 's score is $r_j(t, h) = r_j(t, l) := \hat{n}_j^h(t) + \hat{n}_j^l(t)$. The SMD-TP mechanism modifies this rule to distinguish between the two task types and move agents towards their target allocation.

Let $V_{SMD}((F_j)_{j=1}^J, A, T)$ be the value of SMD-TP mechanism given a sequence of task arrivals A . Abusing notation, let $V_{SMS}((F_j)_{j=1}^J, A, T)$ be the value of the SMS-TP mechanism given the aggregate task counts from sequence A over time horizon T . Say that a mechanism is ε -IC (ε -SQ) if for any agent the ratio of the expected workload under truthful reporting to that of any deviation (to the status quo) is no more than $1 + \varepsilon$.

Intuitively, the SMD-TP algorithm tries to allocate a task of type- k so as to move each agent towards their target task-load $\dot{n}_j^k$ in a way that smooths assignments over time. How well the algorithm can do this depends on how far $N^h(T)$ and $N^l(T)$ are from their expected values $\rho^h T$ and $\rho^l T$. The algorithm improves as T increases, since by the strong law of large numbers, $\frac{1}{T} (N^k(T) - T\rho^k) \xrightarrow{a.s.} 0$ as $T \rightarrow \infty$.

Proposition 5. $\frac{V_{SMD}((F_j)_{j=1}^J, A, T)}{V_{SMS}((F_j)_{j=1}^J, A, T)} \xrightarrow{a.s.} 1$ as $T \rightarrow \infty$. Moreover, for any $\varepsilon > 0$ there exists $\bar{T}$ such that the SMD-TP mechanism is ε -IC and ε -SQ for any $T \geq \bar{T}$.

Proof. Proof in Appendix C.3. □

In addition to targeting the aggregate task-loads $(\dot{n}_j^h, \dot{n}_j^l)$, the SMD-TP mechanism also attempts to smooth the arrivals over time. This is the benefit of using the ratio $r_j(t, k) = \frac{\hat{n}_j^k(t)}{\dot{n}_j^k}$ to assign tasks, as opposed to the difference $\hat{n}_j^k(t) - \dot{n}_j^k$; the latter would front-load tasks to agents with high targets. On the other hand, this method is somewhat extreme in that it never assigns type- k tasks to an agent j with $\dot{n}_j^k = 0$. Assigning based on the difference between target and realized task-loads would ensure that this difference is small, even if

²⁶What we need for the results is that $\frac{1}{T} (N^k(T) - T\rho^k) \xrightarrow{a.s.} 0$ as $T \rightarrow \infty$.

it means giving a few type- k tasks to agents with $\dot{n}_j^k = 0$. In Appendix B.2, we discuss finite-sample adjustments to the assignment rule that move between these extremes.

The dynamic mechanism inherits, up to the approximation error associated with predicting aggregate task arrivals, the robustness properties of the SMS-TP mechanism (Proposition 4). In the dynamic setting, this allows us to implement our mechanism without relying on distributional assumptions. We can do this by implementing the mechanism first under some initial specification of the type distribution for some period, say a year. The types reported in the first year can then be used to update the estimate of the type distribution used in the second year, and so on. This updating need not interfere with incentives, as we can base the estimate of F_j only on the reports of other agents. Under mild assumptions, the estimated distributions converge to their true values.

IV Empirical Application

In this section, we estimate the magnitude of potential gains from reforming a rotational system in an important high-stakes setting: the assignment of CPS investigators to reported cases of child maltreatment.

IV.A Overview of CPS context

We begin with a brief introduction to the CPS context. Further details can be found in Appendix F. Contact with CPS is surprisingly common in the U.S.: 37% of children are the subject of a maltreatment investigation by age 18 and 5% spend time in foster care (Doyle et al., 2025; Kim et al., 2017). Moreover, foster care placement is one of the most far-reaching government interventions and previous work shows that it has large effects on children’s later-in-life outcomes (Bald et al., 2022). Thus, understanding whether there are simple mechanisms to improve the efficacy of CPS investigations is of policy importance.

At a high level, the system operates as follows: Cases are initiated through calls to a state-level hotline. After initial screening, cases that require further investigation are allocated to a regional office based on the child’s location. The case is then promptly (within 24 hours) assigned to one of several investigators through a rotational system: a new case gets assigned to the investigator at the top of the queue, and that investigator moves to the end of the queue. The investigator probes the allegations and determines whether the child should be placed in foster care. Under CPS guidelines, this decision should be based on the assessed probability that the child will experience subsequent maltreatment if left in the home (Baron et al., 2024).

Handling cases is costly for investigators—requiring time, energy, and imposing emotional

and psychological burdens—and we provide empirical evidence below that these burdens vary across case types and investigators and translate into significant turnover. In CPS agencies, where turnover and shortages are already chronic challenges (Casey Family Programs, 2023), these dynamics are especially salient: high turnover not only imposes heavy recruitment and training costs, but also risks a vicious cycle in which departing staff increase the caseloads of those who remain, further fueling exits. Absent a detailed understanding of this unraveling dynamic, it is prudent to approach investigator welfare with caution. This motivates the focus on mechanisms which generate a Pareto improvement for investigators.

IV.B Modeling the CPS objective

Let $\mathcal{I} = \{1, \dots, I\}$ be a set of cases (tasks), with typical element i .²⁷ Cases are differentiated by characteristics observable at the time of investigator assignment (due to initial screening). We allow both investigator preferences and performance to vary across cases. For each investigator j , preferences over cases are represented by a function $p_j : \mathcal{I} \rightarrow \mathbb{R}$, where $p_j(i)$ is the cost to j of handling case i . For any collection of cases $X \subset \mathcal{I}$ assigned to j , total workload is $\sum_{i \in X} p_j(i)$, which we refer to as j 's *workload*.

While CPS agencies may place value on investigator welfare, their primary directive is to make accurate foster care placement decisions, removing children from the home only if they would otherwise experience subsequent maltreatment. Let $D_{ij} \in \{0, 1\}$ denote investigator j 's potential placement decision for case i , where $D_{ij} = 1$ indicates foster care placement, and let $Y_i^* \in \{0, 1\}$ denote the child's maltreatment potential, where $Y_i^* = 1$ indicates maltreatment would occur absent removal.

The joint realization of (D_{ij}, Y_i^*) yields four outcomes: a false negative ($D_{ij} = 0, Y_i^* = 1$), a false positive ($D_{ij} = 1, Y_i^* = 0$), a true negative ($D_{ij} = 0, Y_i^* = 0$), and a true positive ($D_{ij} = 1, Y_i^* = 1$). Define $\text{FN}_{ij} = \Pr(D_{ij} = 0, Y_i^* = 1)$, $\text{FP}_{ij} = \Pr(D_{ij} = 1, Y_i^* = 0)$, $\text{TN}_{ij} = \Pr(D_{ij} = 0, Y_i^* = 0)$, and $\text{TP}_{ij} = \Pr(D_{ij} = 1, Y_i^* = 1)$, so that these probabilities sum to one.

The expected social cost of assigning case i to investigator j is

$$c(i, j) = \text{FN}_{ij}c_{\text{FN}} + \text{FP}_{ij}c_{\text{FP}} + \text{TN}_{ij}c_{\text{TN}} + \text{TP}_{ij}c_{\text{TP}},$$

where $(c_{\text{FN}}, c_{\text{FP}}, c_{\text{TN}}, c_{\text{TP}})$ represent the designer's valuation of the four possible outcomes. Thus, $(\text{FN}_{ij}, \text{FP}_{ij}, \text{TN}_{ij}, \text{TP}_{ij})$ characterizes the joint distribution of investigator j 's potential

²⁷As cases are assigned within CPS offices, the problem is separable across offices. The theoretical analysis describes the assignment mechanism for a given office.

decision and the maltreatment potential of case i , while $(c_{\text{FN}}, c_{\text{FP}}, c_{\text{TN}}, c_{\text{TP}})$ constitutes the designer’s Bernoulli utility over these outcomes.

An assignment is represented by a matrix $Z \in \{0, 1\}^{I \times J}$, where $Z_{ij} = 1$ if case i is assigned to investigator j . The total social cost induced by assignment Z is

$$C(Z) = \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} Z_{ij} c(i, j). \quad (5)$$

The parameters $(c_{\text{FN}}, c_{\text{FP}}, c_{\text{TN}}, c_{\text{TP}})$ are taken as given primitives of the model. In practice, they must be chosen by the designer (CPS); in the empirical application, we calibrate them. The central empirical challenge, addressed below, is identification of the investigator-specific performance measures $(\text{FN}_{ij}, \text{FP}_{ij}, \text{TN}_{ij}, \text{TP}_{ij})_{i \in \mathcal{I}, j \in \mathcal{J}}$.

IV.C Binary classification

To apply our results from the model with two task types to a setting with richer heterogeneity, we partition the set of cases into two groups, *high* (H) and *low* (L). We then consider mechanisms which are measurable with respect to this partition; that is, for which the probability of assigning cases i and i' to investigator j is the same if i and i' are of the same type. We refer to this as a *binary-classification mechanism*.

Fixing an arbitrary partition, define $c^k(j) := \hat{E}[c(i, j) | i \text{ is type } k]$ and $p_j^k := \hat{E}[p_j(i) | i \text{ is type } k]$ for $k \in \{H, L\}$, where $\hat{E}$ denotes the empirical expectation. Then, in any binary-classification mechanism in which investigator j is assigned H^j high-type cases and L^j low-type cases, the expected social cost across the population of cases is $\sum_{j \in \mathcal{J}} c^h(j) H^j + c^l(j) L^j$ and the expected workload of each j from this assignment is $p_j^h H^j + p_j^l L^j$. Within this class of mechanisms, the analysis is the same as in the two-type case.

A key question for the empirical simulation is how to partition cases. The theory allows for any binary partitioning of cases. In the context of CPS, this could mean distinguishing between abuse and neglect, or separating cases based on the gender or age of the children involved. Given CPS’s primary focus on preventing further child maltreatment, and guided by discussions with local agencies in Michigan, we partition cases based on the predicted risk of subsequent maltreatment. We construct this measure by training a machine learning algorithm to predict the risk of subsequent maltreatment in the home, which we discuss further in Appendix E.3. We define high-risk cases as those in the top quartile of predicted algorithmic risk, and low-risk cases as all other cases.²⁸

²⁸Predictive risk modeling and binary risk partitions are also used in recent CPS policy efforts. For example, the Los Angeles County Risk Stratified Supervision Model uses ML techniques to notify supervisors

Remark 5. Note that in a binary-classification mechanism agents face uncertainty about their workload (stemming from the inherent randomness in the mechanism) that is separate from their uncertainty about others’ preferences. This is also the case in the original rotational system. However, if the number of tasks assigned to each agent is large, as it is in the CPS context, then the difference between expected and realized workloads will be small.

IV.D Identification

Thus far we have treated investigator performance, $c^k(j)$, as observed. In practice, however, these parameters are not point-identified, even under random assignment of investigators to cases. This reflects the standard “selective labels” problem: for children placed in foster care, whether maltreatment would have occurred had they remained at home is unobserved, preventing separate identification of true and false positives.

Our strategy is therefore to identify differences in costs across investigators rather than their levels. These differences are sufficient to construct social preferences over mechanisms and hence to evaluate welfare comparisons. Although $c^k(j)$ is not nonparametrically identified, the difference $c^k(j) - c^k(j')$ is identified provided investigators j and j' are quasi-randomly assigned to cases and an exclusion restriction holds.

The selective labels problem arises because Y_i^* is observed only when the assigned investigator does not remove the child. Formally, we observe Y_i^* if and only if case i is assigned to investigator j with $D_{ij} = 0$. The realized outcome for observed subsequent maltreatment under assignment to j can therefore be written as $Y_{ij} := (1 - D_{ij})Y_i^*$.

Normalizing $c_{\text{TN}} = 0$ without loss of generality, identification of $c(i, j)$ would require recovering $(\text{FN}_{ij}, \text{FP}_{ij}, \text{TP}_{ij})$. Under random assignment, FN_{ij} and TN_{ij} are identified because Y_i^* is observed when $D_{ij} = 0$. When $D_{ij} = 1$, however, Y_i^* is unobserved, so FP_{ij} and TP_{ij} are not nonparametrically identified. While this precludes point identification of the social cost function $c(i, j)$ without further structure, we show next that social preferences—i.e., the ranking over the set $\mathcal{Z}$ of feasible assignments—are nevertheless identified under random assignment and an exclusion restriction.

Let $I \subset \mathcal{I}$ be a subset of cases. We say that assignment Z is *random conditional on I* if $(D_{ij}, Y_i^*) \perp\!\!\!\perp Z_{ij}$ conditional on $i \in I$, for all $j \in J$. Investigator j ’s assignment is *supported on I* if $\Pr(\{i \in I, Z_{ij} = 1\}) \neq 0$. Let $\text{FN}_j^I := \mathbb{E}[\text{FN}_{ij} | i \in I]$ and similarly for TN_j^I , TP_j^I , and FP_j^I .

when a new case, classified by the model as “complex-risk,” has been assigned to their office, so that they can devote additional time and attention to it. There, the binary partition has been justified by the need for simplicity when explaining the practical implications of high and low risk to supervisors and investigators.

Lemma 1. Assume that the observed assignment is random conditional on I . Then, for any $j, j' \in \mathcal{J}$ whose assignments are supported on I , the following three differences are identified: (i) $\text{FP}_j^I - \text{FP}_{j'}^I$, (ii) $\text{TP}_j^I - \text{TP}_{j'}^I$, and (iii) $\mathbb{E}[c(i, j) - c(i, j') | i \in I]$.

Proof. First, recall that Y_{ij} and D_{ij} are observed for the set of cases when $Z_{ij} = 1$. Then, under random assignment conditional on I , we have

$$\text{FN}_j^I := \mathbb{E}[\text{FN}_{ij} | i \in I] = \mathbb{E}[Y_i^*(1 - D_{ij}) | i \in I] = \mathbb{E}[Y_{ij} | i \in I] = \mathbb{E}[Y_i | i \in I, Z_{ij} = 1]$$

and $P_j^I := \mathbb{E}[D_{ij} | i \in I] = \mathbb{E}[D_i | i \in I, Z_{ij} = 1]$. Moreover, we can express TN_j^I as $\text{TN}_j^I = 1 - (\text{TP}_j^I + \text{FP}_j^I) - \text{FN}_j^I = 1 - P_j^I - \text{FN}_j^I$. Thus, FN_j^I , TN_j^I , and P_j^I are identified if j 's assignment is supported on I . Let $S_j^I = \text{TP}_j^I + \text{FN}_j^I$. Note that, under random assignment conditional on I , $S_j^I = S_{j'}^I = \mathbb{E}[Y_i^* | i \in I]$ for all $j, j' \in \mathcal{J}$. Then,

$$\begin{aligned} \text{FP}_j^I - \text{FP}_{j'}^I &= (1 - \text{TP}_j^I - \text{FN}_j^I - \text{TN}_j^I) - (1 - \text{TP}_{j'}^I - \text{FN}_{j'}^I - \text{TN}_{j'}^I) \\ &= (1 - S_j^I - \text{TN}_j^I) - (1 - S_{j'}^I - \text{TN}_{j'}^I) = -(\text{TN}_j^I - \text{TN}_{j'}^I). \end{aligned}$$

Similarly, $\text{TP}_j^I - \text{TP}_{j'}^I = -(\text{FN}_j^I - \text{FN}_{j'}^I)$. This is sufficient to identify cost differences. $\square$

Lemma 6, presented in Appendix E.5, generalizes Lemma 1 beyond binary Y_i^* .

Example. To build intuition for Lemma 1, consider two investigators j and j' in the same office. Assume that cases are assigned as if at random and that investigators cannot affect Y_i^* . This implies that both investigators face the same underlying rate of true maltreatment: $\mathbb{E}[Y_i^* | Z_{ij} = 1] = \mathbb{E}[Y_i^* | Z_{ij'} = 1] = S$. Suppose investigator j has a false negative rate $\text{FN}_j = 0.09$ and placement rate $P_j = 0.02$, while investigator j' has $\text{FN}_{j'} = 0.08$ and $P_{j'} = 0.03$. What can we infer about their false positive rates? Although S itself is unobserved, differences in false positive rates can still be recovered using the result above that $\text{FP}_j - \text{FP}_{j'} = (P_j + \text{FN}_j) - (P_{j'} + \text{FN}_{j'})$. Plugging in the numbers, $\text{FP}_j - \text{FP}_{j'} = 0.11 - 0.11 = 0$. The interpretation is that the additional placements by j' exactly offset the reduction in false negatives. Because both investigators see cases with the same underlying risk S , the two must have the same false positive rate. This illustrates the point of the lemma: while the absolute level of S is not observed, *differences* in error rates across investigators are identified under random assignment.

The identification results rest on two key assumptions. First, investigators must be quasi-randomly assigned to cases within a CPS office. This assumption has been carefully examined in the Michigan CPS context (e.g., see Baron et al. (2024)). As discussed in Appendix F,

investigators are assigned cases rotationally, following a queue of who is “next up,” rather than through matching of specific investigators to specific families. We discuss how we extend the results from Lemma 1, which requires random assignment, to a setting with quasi-random assignment within offices in Section E.2. The second key assumption is an exclusion restriction: investigators can only “reveal” Y_i^* through their placement decisions, and cannot otherwise directly affect it. This restriction has also been probed in this prior work in our context. The intuition is that investigators’ primary influence on children’s outcomes is the foster care decision. Other actions, such as referrals to preventive services, have been shown to have little impact on child maltreatment risk. Thus, conditional on placement, investigators are unlikely to directly shift Y_i^* .

Computing Welfare Differences Across Mechanisms. We next show how Lemma 1 can be used to recover social preferences over assignment mechanisms.²⁹ First, note that the parameter $c^k(j)$ is a natural measure of the performance of investigator j on cases of type k . While $c^k(j)$ is not non-parametrically identified, the difference $c^k(j) - c^k(j')$ is identified for any j, j' and is given by $c^k(j) - c^k(j') = c_{FP}(P_j^{I_k} - P_{j'}^{I_k}) + (c_{FP} + c_{FN} - c_{TP})(FN_j^{I_k} - FN_{j'}^{I_k})$. Defining $\gamma_j^k := c_{FP}P_j^{I_k} + (c_{FP} + c_{FN} - c_{TP})FN_j^{I_k}$, we have that $c^k(j) \leq c^k(j')$ if and only if $\gamma_j^k \leq \gamma_{j'}^k$. Intuitively, γ_j^k tells us the position of investigator j in the distribution of investigator performance among cases of type k . We therefore refer to γ_j^k as j ’s *performance score* on type- k cases, where a lower score corresponds to greater performance. As we show next, γ_j^k is a useful parameter because it can be used to compute social preferences.

Let $\{I_k\}_{k=1}^K$ be a partition of $\mathcal{I}$ into disjoint sets. Say that Z is *conditionally random* given partition $\{I_k\}_{k=1}^K$ if it is random conditional on I_k for all K . Say that it *has full support* if every agent’s assignment is supported on I_k , for all k .

Corollary 3. Let Z be an observed assignment that is conditionally random and has full support given a partition $\{I_k\}_{k=1}^K$ of $\mathcal{I}$. Then $\mathbb{E}[C(Z)] - \mathbb{E}[C(Z')]$ is identified for any other assignment Z' that is conditionally random given the same partition, and equal to:

$$\sum_{k=1}^K \sum_{j \in \mathcal{J}} \gamma_j^k \sum_{i \in I_k} Z_{ij} - \sum_{k=1}^K \sum_{j \in \mathcal{J}} \gamma_j^k \sum_{i \in I_k} Z'_{ij}.$$

That is, under the conditions of Corollary 3, the cardinal ranking over $\mathcal{Z}$ is non-parametrically identified. Note that from Lemma 1 we can also identify the expected difference in false

²⁹We are certainly not the first to note that differences in false positive rates can be expressed in terms of placement rates, baseline maltreatment risk S , and false negative rates. Our contribution in this section is to show that, for the purpose of identifying social preferences over assignment mechanisms, it is sufficient to focus on *differences* in false positives and true positives across investigators. This is in contrast to prior work, which has aimed to directly identify the *levels* of false and true positives—typically by structurally estimating S (Chan et al., 2022) or by relying on alternative strategies such as contraction (Kleinberg et al., 2018) or identification at infinity (Arnold et al., 2022; Angelova et al., 2023).

negatives, false positives, and the placement rate across the two assignments.

Remark 6. One useful application of Lemma 1 is to pick an arbitrary investigator, j' , and define $\tilde{c}(i, j) = c(i, j) - c(i, j')$. Then, we can replace the objective of the designer, $C(Z') := \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} Z'_{ij} c(i, j)$, with $\tilde{C}(Z') := \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} Z'_{ij} \tilde{c}(i, j) = \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} Z'_{ij} c(i, j) - \sum_{i \in \mathcal{I}} c(i, j')$, where the last equality follows from the fact that under each $Z' \in \mathcal{Z}$, each case is assigned to exactly one investigator. Since $\sum_{i \in \mathcal{I}} c(i, j')$ does not depend on Z' , $\tilde{C}$ and C represent the same preferences over $\Delta(\mathcal{Z})$. Moreover, if the observed assignment Z is conditionally random and has full support given a partition $\{I_k\}_{k=1}^K$ then $\mathbb{E}[\tilde{c}(i, j)|i \in I_k]$ is identified for all k , and $\mathbb{E}[\tilde{C}(Z')]$ is identified if Z' is conditionally random given the same partition.

Finally, we consider an investigator's *relative advantage* across case types, $c^l(j) - c^h(j) = \delta_j$. While relative advantage is not identified, differences in relative advantage are: define the *comparative advantage* of j relative to j' by $D(j, j') = \delta_j - \delta_{j'} = c^l(j) - c^l(j') - (c^h(j) - c^h(j'))$. When high- and low-type cases are equally costly for all investigators, comparative advantage is the sufficient statistic for the optimal assignment. Let $d_j := \gamma_j^l - \gamma_j^h$ be investigator j 's *comparative advantage score*. Below, we use this identification result to document evidence of comparative advantage across investigators within offices in our data.

IV.E Data and estimation

We use administrative data from the Michigan Department of Health and Human Services covering the universe of child maltreatment investigations between 2008 and 2016. The data include investigator identifiers, foster care placement decisions, child and allegation characteristics, and subsequent child welfare outcomes. Consistent with the institutional setting described above, cases are assigned to investigators via rotation within local offices, generating quasi-random assignment conditional on office-by-year.

We construct a sample of investigations that satisfy the quasi-random assignment assumption and for which subsequent outcomes can be measured. The final sample consists of 322,758 investigations involving 908 investigators. Following the literature (e.g., [Baron et al., 2024](#); [Putnam-Hornstein et al., 2021](#)), we measure subsequent maltreatment as a new CPS investigation within six months among children left at home following the focal investigation. Summary statistics, sample construction details, and robustness to alternative maltreatment measures are provided in Appendix [E.1](#).

Our identification results imply that while performance levels are not identified, differences across investigators are identified under quasi-random assignment. We therefore estimate investigator-specific performance scores by case type, γ_j^h and γ_j^l , as well as relative cost measures $\tilde{c}^k(j)$ and comparative advantage scores d_j . We also construct overall performance

measures that aggregate across case types. These objects are sufficient to rank investigators and evaluate welfare comparisons across assignment mechanisms. To mitigate overfitting, we implement a split-sample procedure: investigator performance measures are estimated in a training sample and welfare effects are evaluated in a hold-out sample. We apply empirical Bayes shrinkage to reduce noise in investigator-level estimates. Full estimation details and calibration of the social cost parameters are provided in Appendix E.2.

IV.F Main empirical results

IV.F.1 Motivating empirical facts

We begin by presenting empirical facts that motivate the use of our mechanism in this setting, including new evidence from a statewide survey of CPS investigators in Michigan.

Considerable variation in performance and comparative advantage: Intuitively, gains from investigator reassignment in our proposed mechanism can only occur if there is heterogeneity in investigators’ relative performance across cases and within offices. That is, it is not enough for investigators to differ in the level of their performance, they must also differ in their comparative advantage scores.

We find considerable variation in performance and comparative advantage. Table A2 estimates the relationship between performance and comparative advantage metrics (γ_j and d_j) on the training dataset and prediction error rates in the evaluation dataset. Panel A shows that investigators with a one standard deviation γ_j below the office mean achieve a 1.1pp [7.1%] reduction in false negatives and 1.6pp [70.4%] reduction in false positives. Panels B and C further regress prediction error rates on comparative advantage scores, d_j . Investigators with a one standard deviation greater comparative advantage score achieve 2.0pp [8.5%] lower false negative rates and 2.3pp [57.3%] lower false positive rates in high-risk cases, but 0.04pp [0.3%] *higher* false negative rates and 0.26pp [13.6%] *higher* false positive rates in low-risk cases. That is, investigators with greater comparative advantage in high-risk cases achieve lower prediction error rates in high-risk cases but higher error rates in low-risk cases—providing evidence of investigator specialization across case types within offices.

High-risk cases are costly to investigators: Under the status-quo rotational system, the composition of caseloads in expectation is equal across investigators within an office. In practice, however, there may be time periods in which some investigators receive larger numbers of high- or low-risk cases by random chance. In Table A3, we leverage this variation to examine whether greater exposure to high-risk cases leads to increased investigator turnover. To do so, we use survival analysis techniques to measure the effect of caseload risk composition on investigator career length. Column 1 shows that a one standard deviation increase in the

mean predicted risk of an investigator’s caseload increases turnover risk by 149%. Column 2 reports that being assigned to an above-median share of high-risk cases increases turnover risk by 54%. Thus, exposure to a greater share of high-risk cases leads to large increases in investigator turnover, suggesting that these cases are costlier to investigators.

Heterogeneity in perceived costs of high-risk cases: To complement the administrative data, we conducted a survey of Michigan CPS investigators. The survey probed how investigators perceive high- versus low-risk cases and elicited their marginal rates of substitution (MRS) between case types. For a detailed report of the responses, see Appendix E.4.

The survey provides direct evidence that investigators themselves perceive high-risk cases as disproportionately costly. Respondents rated high-risk cases as much more time-consuming (mean = 8.8 on a 0–10 scale), more emotionally taxing (8.6), and only modestly more satisfying (5.3) than low-risk cases. Open-ended responses reinforced these findings, with investigators frequently linking complex cases to burnout and turnover.

The survey also revealed significant heterogeneity in preferences. Investigators reported widely varying MRSs between high- and low-risk cases. On average, investigators were willing to trade one high-risk case for about 3.6 low-risk cases, with substantial heterogeneity across individuals (Figure A4). Importantly, investigators reported similar MRSs when asked about their coworkers’ preferences—an average MRS of 3.7.

Altogether, these patterns motivate the potential for the SMD-TP mechanism to achieve welfare gains in this context. There is evidence in the data of significant comparative advantage across high- and low- risk cases within offices. However, investigators themselves recognize sharp trade-offs between case types. Additional high-complexity cases are experienced as disproportionately costly, both in terms of job satisfaction and turnover, and these preferences are heterogeneous. This creates scope for welfare gains through reallocation that better matches investigators to cases, while accounting for their preferences.

IV.F.2 The role of investigator type distributions

As is standard in the Bayesian mechanism-design literature, our theoretical analysis takes the preference distributions, F_j , as an input. However, our mechanism can in fact be implemented without prior knowledge of these distributions. As shown above, if F_j is misspecified, the SMD-TP mechanism retains its incentive properties and continues to respect the status quo. The downside of using incorrect type distributions is that the mechanism may not converge to the optimal outcome in the large market. Fortunately, information about type distributions can be realistically obtained in practice, unlike knowledge of each investigator’s realized type. For example, the MRS distribution from our survey can be used to inform an initial

specification for F_j . We can then run the SMD-TP mechanism for an initial trial period, using this estimated distribution. The trial run can then generate further data on individual investigators’ preferences, allowing refinement of estimates of their preference distributions.³⁰ Therefore, the SMD-TP mechanism can be implemented without prior knowledge of the investigators’ preference distribution.

Simulating welfare gains: While it is feasible to eventually gather information about F_j , we would like to understand now whether the mechanism could generate welfare gains across a range of distributions. Our survey provides a natural baseline for specifying investigator type distributions. We use the distribution of MRSs reported by investigators as the baseline specification of F_j . We estimate hazards from a kernel-smoothed empirical distribution from the survey and impose monotone hazards by choosing the nondecreasing hazard sequence that minimizes the discrete sum of squared deviations in hazard rate. We refer to this as the “regularized own-type empirical distribution” and apply the same procedure to the coworker-reported responses (see Appendix E.4 for details).³¹ In our baseline analysis, we assume a common type distribution across investigators; in Appendix E.7 we also consider an extension that allows for correlation between types and performance, in which we continue to see significant welfare gains.

While we do not claim that the survey responses reflect the “true” distribution of types, they provide a useful benchmark to illustrate the potential welfare gains achievable under the SMD-TP mechanism in a data-driven way. To assess robustness, however, we also simulate outcomes under several alternative distributions that seem like natural candidates, such as uniform and truncated normal specifications. Across all of these cases, we consistently find welfare improvements relative to the status quo, demonstrating that the mechanism’s potential gains do not depend on any single distributional assumption.

³⁰To avoid introducing additional agency problems by making the mechanism in future periods dependent on type reports in the trial period, we cannot use the type report of agent j to learn about types for agents in the same office. However, we can only use information about j ’s type to learn about the distribution for agents in other offices. Given the large number of investigators involved, this should be sufficient to generate significant learning. This also raises the question of how the trial period should be chosen, trading off the benefits of experimentation versus exploitation. Resolving this trade-off is beyond the scope of the current paper, but recent work in [Nguyen et al. \(2023\)](#) provides a potential path forward.

³¹The regularized empirical distribution may not satisfy strong regularity. As a result, the optimal LMS mechanism could involve a four-prices for some investigators (Proposition 1). Nonetheless, for implementation we restrict attention to two-price mechanisms. As discussed above, in the small market these mechanisms have desirable incentive properties, and their simplicity means that they could plausibly be implemented in practice. Moreover, as we show below, two-price mechanisms generate sizable welfare gains.

IV.F.3 Social welfare gains

Corollary 3 shows that the difference in social cost between assignments is identified using investigator performance scores and their caseload composition in the two mechanisms. Using this result, we report differences between the SMD-TP mechanism and a “status-quo counterfactual” which splits high- and low-risk cases equally across investigators within counties.³² Our strategy to estimate welfare gains accounts for (i) uncertainty in investigator type distributions and (ii) over-fitting concerns. Under an initial distributional assumption for F_j , we use a split-sample strategy that combines investigator performance measures with their p_j draws to compute the assignment generated by the SMD-TP mechanism for that draw in the training set. We then calculate the realized welfare gains for the given type profile in the evaluation set. We summarize welfare gains for a given specification of the type distribution as the average welfare gain across 100 draws of types. Finally, we estimate standard errors, clustered by investigators, using a bootstrapping procedure that accounts for uncertainty in estimates of both investigator performance and their type draw.

Table 1 presents estimated welfare gains across different distributional assumptions. Each column compares outcomes when investigators are assigned to cases using the SMD-TP mechanism versus a counterfactual in which high- and low-risk cases are evenly split. For example, when F_j is specified using the regularized own-type empirical distribution, we estimate a decline in social costs of 1,189 [5.9%] driven by a reduction in both 755 false negatives [1.5%] and 1,017 false positives [13.9%] (Column 1).³³ The mechanism also reduces the number of total placements by 262 [2.6%]. Results are nearly identical when F_j is instead based on investigators’ reports of their coworkers’ preferences (Column 2).

To estimate the importance of investigator private information for welfare gains, in Column 3 we again assume that F_j is distributed as in Column 1, but that the designer observes each investigator’s type directly and implements the first-best assignment for each realized type profile. In this simulation the welfare gains increase dramatically—social costs decline by 16.7%, driven by a false negative decline of 4.3% and false positive decline of 39%. The comparison with Column 1 shows that information rents are significant when types are unobserved. As discussed in Appendix B.1, over time the designer can use the data generated by the mechanism to learn about investigators’ preferences. Column 3 represents an upper bound on welfare gains as the designer can better predict preferences.

³²The equal split of cases is the expected assignment under the rotational system. Table A6 shows that we obtain similar welfare gains when we benchmark the SMD-TP mechanism to the observed status quo.

³³Baseline false positive counts are unobserved. To express welfare changes in percent terms, we use extrapolation-based estimates of the false positive rate, which we describe in Appendix E.6.

Table 1: Gains from Investigator Reallocation

	(1)	(2)	(3)
	Own Type	Coworker Type	Known Type, Own Type
Social Costs	-1,189.2*** (300.7) [-5.9%]	-1,111.2*** (274.3) [-5.5%]	-3,351.7*** (454.9) [-16.7%]
False Negatives	-754.9*** (219.5) [-1.5%]	-713.3*** (213.4) [-1.4%]	-2,209.0*** (372.9) [-4.3%]
False Positives	-1,017.0*** (239.7) [-13.9%]	-950.1*** (226.8) [-13.0%]	-2,865.3*** (385.6) [-39.2%]
Placements	-262.1** (106.4) [-2.6%]	-236.7** (97.8) [-2.4%]	-656.3*** (190.2) [-6.6%]

Notes. This table reports the welfare gains derived from the SMD-TP mechanism. Each column corresponds to a different distributional assumption for p_j . Column 1 presents gains from the regularized own-type empirical distribution: we estimate hazards from a kernel-smoothed empirical distribution from the survey and impose monotone hazards (Myerson regularity) by choosing the nondecreasing hazard sequence that minimizes the discrete sum of squared deviations in hazard rate (see Appendix E.4 for details). Column 2 performs the same procedure but using investigator reports of their coworkers' preferences. Column 3 uses the same type distribution as in Column 1, but assumes that p_j is known to the designer. We estimate welfare changes and robust standard errors (in parenthesis), clustered by investigator, using the approach described in the text. Percent effects relative to the baselines are presented in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$

Robustness checks: We conduct two exercises to assess the robustness of our main findings. First, Table A4 demonstrates that the welfare gains persist across a wide range of distributional assumptions.³⁴ The largest gains arise when p_j follows a degenerate distribution, in which case the investigator type is known and invariant across investigators, and thus further shows that a naive analysis which ignores investigators' private information would significantly overstate welfare gains. Second, Table A8 shows that our results are robust to alternative proxies for maltreatment. While we use a subsequent investigation within six months as our preferred measure, since it is external to the local CPS office, we obtain similar results when contracting the time horizon or when restricting to subsequent *substantiated* investigations.³⁵

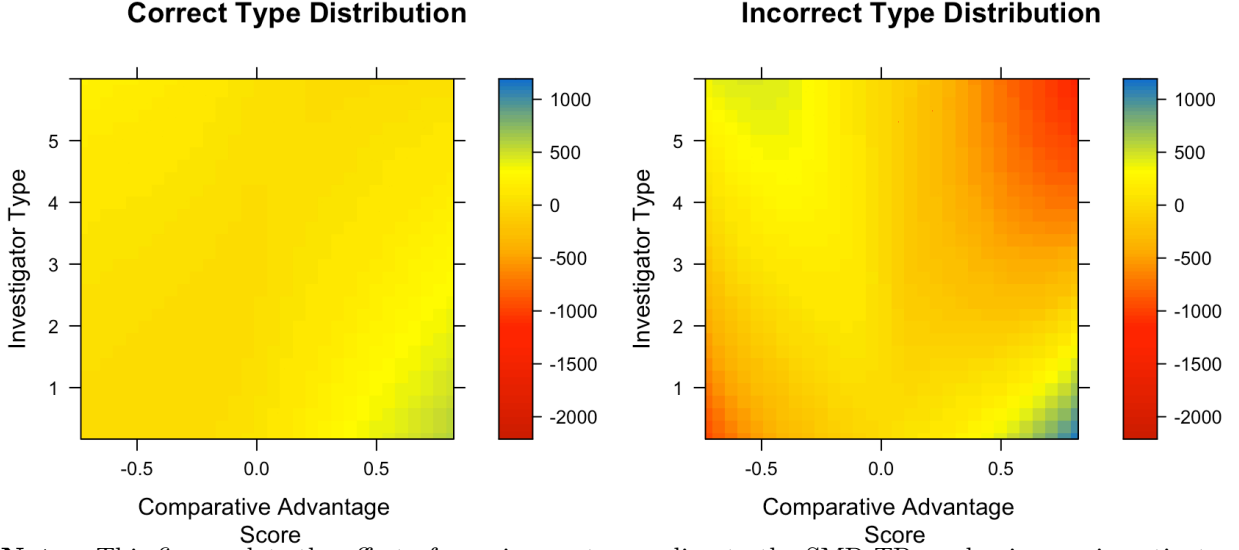
IV.F.4 Investigator preferences

Figure 3 demonstrates the importance of considering investigators' heterogeneous preferences in the SMD-TP mechanism. We define investigator welfare for an investigator with type

³⁴Figure A2 further illustrates this point by showing reductions in social costs under truncated normal distributions with varying means μ and standard deviations σ . The mechanism consistently achieves welfare improvements relative to the status quo.

³⁵For computational purposes, we compare the welfare gains across different proxies for subsequent maltreatment in the LMS-TP mechanism, for which the SMD-TP is an approximation.

Figure 3: The Importance of Accounting for Investigator Preferences



Notes. This figure plots the effect of reassignment according to the SMD-TP mechanism on investigators' welfare by their comparative advantage score, d_j and their type, p_j . Investigator welfare for an investigator type p_j and assigned to caseload (n^h, n^l) is $-(p_j n^h + n^l)$. We report the difference between investigator welfare under the SMD-TP mechanism and a counterfactual in which cases are equally split within counties. The left panel assumes that the true distribution of investigator types is the regularized own-type empirical distribution. The right panel calculates changes in investigator welfare under an assignment that assumes $p_j = 1 \forall j \in \mathcal{J}$, but where the true p_j is distributed according to the regularized own-type empirical distribution. We present results averaged across the 100 investigator-type draws.

p_j and caseload (n^h, n^l) as $-(p_j n^h + n^l)$, the negative of their price-weighted caseload. In the left panel of Figure 3, we derive the optimal allocation of cases assuming that the true distribution of types is the regularized own-type empirical distribution. We then compute the difference between investigator welfare under the SMD-TP mechanism relative to the status quo. Under the SMD-TP mechanism, which accounts for investigator type heterogeneity, investigator welfare is improved by approximately 51 price-weighted cases, on average. Average investigator welfare is -562 in the equal-split counterfactual, so this represents a modest welfare improvement. Importantly, the reassignment makes no investigator substantively worse off: no investigator experiences a welfare loss greater than 10 price-weighted cases, and the 1st percentile of investigator welfare change is a loss of 3.4 cases.³⁶ In fact, 36% of investigators experience welfare gains under the correct SMD-TP mechanism of greater than 10 price-weighted cases and 26% of investigators experience welfare gains of at least 10%, which could in turn improve recruitment and retention.

The right panel of Figure 3 instead assigns cases without considering heterogeneity in

³⁶The constraint that no investigator is made worse off by the mechanism is imposed exactly in the static model. The SMD-TP mechanism approximates the SMS-TP mechanism, and this approximation improves as the time horizon grows. Thus, in the dynamic version, some investigators can be made slightly worse off than the equal-split counterfactual.

investigator preferences. Formally, the mechanism assigns cases assuming that $p_j = 1$ for all investigators. But, when computing investigator welfare, we assume that their types truly follow the regularized own-type empirical distribution. Under this scenario, 27% of investigators experience welfare losses of at least 10%. Moreover, Figure 3 shows that there is significant heterogeneity in investigator welfare loss by comparative advantage and type. The investigators experiencing the largest losses are those with a large comparative advantage on high-risk cases as well as high p_j —investigators above the median in both their comparative advantage score, d_j , and p_j experience an average welfare loss of 214 price-weighted cases (38%), and those in the top quartile of both experience an average welfare loss of 340 (60%). On the other hand, investigators with low comparative advantage on high-risk cases and high p_j are made better off under the mechanism that ignores preference heterogeneity.

Figure 3 highlights why considering investigator preferences in the assignment problem is paramount. If the mechanism ignores types, investigators with large comparative advantage in high-risk cases receive more of these cases, but are made substantially worse off if assignment to high-risk cases is costly relative to low-risk cases. This would likely create greater turnover or worsened performance among such investigators, a particularly negative outcome in a system that already suffers from staff shortages.

IV.F.5 Dynamic nature of the mechanism

Figure 3 considers investigators’ welfare over their cumulative caseloads, but does not consider how their caseloads are spread over time. While smoothing caseloads over time does not directly enter the mechanism-design problem as a constraint, the SMD-TP solution attempts to do so by allocating cases based on the percent of the target level for each case type that each investigator has completed thus far. Thus at any point in time, each investigator should have completed approximately the same percentage of their high- or low-risk cumulative caseload. Figure A3 describes how cumulative price-weighted workloads vary over time. This figure shows that the SMD-TP mechanism is successful in spreading caseloads.

IV.F.6 Welfare comparisons across mechanisms

For reference, we also estimate welfare gains under the SMS-TP and LMS-TP mechanisms and compare these to outcomes under the SMD-TP mechanism in Table A5. Columns 1–3 report results when p_j is drawn from the regularized own-type empirical distribution, while Columns 4–6 repeat the exercise assuming $p_j \sim \text{Unif}[1, 2]$. Comparing Columns 1 and 2, we find that moving from LMS-TP to SMS-TP significantly reduces welfare gains. The fact that LMS yields significantly larger gains than SMS in the baseline specification suggests that the benefits of reassignment may be greater in larger offices. Consistent with this,

Table A7, shows that focusing on the five largest counties in Michigan produces considerably higher welfare gains relative to our baseline estimates in Table 1. Finally, differences between SMS-TP and SMD-TP are small, indicating that the “online” nature of case assignment—the need to allocate cases without knowing which cases will arrive in the future—is not a first-order concern.

Finally, we investigate the potential gains from going beyond the binary-classification framework. Recall that our main binary split considered high-risk cases to be those in the top quartile of the predicted risk distribution, and low-risk as all other cases. Suppose instead that we allow the mechanism to condition on each of the four quartiles of this distribution. In general, solving for the optimal mechanism with four categories is beyond the scope of the current study. However, we can get a sense of the potential gains by considering the special case in which the designer observes investigators’ preferences, in which case the SMS allocation problem reduces to a simple linear program. Table A9 reports this comparison. We find that the welfare gains from considering quartile mechanisms only increase social welfare gains from 1,506 to 1,711. This provides some evidence that the gains from considering more complicated mechanisms may be limited in this setting.

V Conclusion

The ultimate objective of this work is a tractable framework for reforming institutions which allocate tasks among agents. Our main contribution is a mechanism-design analysis of this problem. Moreover, we demonstrate empirically that the simple mechanism we identify has the potential to improve outcomes in the CPS context. The primary takeaway of this empirical analysis is that there are gains to be had, and further study is warranted on the use of performance data in assigning CPS investigators to cases. With the goal of providing a starting point for reform (in this and other contexts) the theoretical analysis sought to address what we view as the primary challenges facing such reforms:

1. *Unobservable agent preferences and status-quo constraints.* The designer may wish to preserve agents’ welfare (perhaps motivated by recruitment and turnover considerations or internal politics), yet agents’ preferences are unobservable.
2. *Effort incentives.* While we do not explicitly model the decision to exert effort, under our mechanism agents’ payoffs improve with their performance, at least locally (Theorem 3).
3. *Perceived fairness of the mechanism.* Even if every agent is better off relative to the current system, they may resent the new disparity in task-loads. We show that such disparity can at least be justified on the basis of performance: agents who receive fewer type- k tasks for

the same type report are those who perform better on type- k tasks (Theorem 3).

Several additional considerations might need to be addressed for such a mechanism to be implemented in practice. For example, a key challenge is effectively communicating the mechanism to agents and establishing a clear protocol for reporting their preferences. In the direct implementation of the mechanism, agents only need to report a single number—their MRS between high- and low-type cases. However, they would likely require guidance on how to understand this parameter. Once types are elicited, the SMD-TP mechanism can operate with no further input from agents and requires minimal changes to current office procedures.

A Proofs of main results

A.1 Proof of Theorem 1 and Corollary 1

Proof. We begin, as in Myerson (1981), by using the envelope condition to simplify the IC constraints. First, note that in any IC mechanism H must be non-increasing. Also, by the envelope theorem (Milgrom and Segal, 2002)

$$-pH(p) - L(p) = -\underline{p}H(\underline{p}) - L(\underline{p}) - \int_{\underline{p}}^p H(z)dz$$

in any IC mechanism. Moreover, if H is non-increasing and H, L satisfy the envelope condition, then the mechanism is IC. From the envelope condition and monotonicity of H , we then have that L is non-decreasing. Thus non-negativity of $L(\underline{p})$ is sufficient for non-negativity of L . Note also that

$$\begin{aligned} \int L(p)dF(p) &= \underline{p}H(\underline{p}) + L(\underline{p}) - \int_{\underline{p}}^{\bar{p}} \left(pH(p) - \int_{\underline{p}}^p H(z)dz \right) dF(p) \\ &= \underline{p}H(\underline{p}) + L(\underline{p}) - \int_{\underline{p}}^{\bar{p}} H(p) \left(p - \frac{1 - F(p)}{f(p)} \right) dF(p) \end{aligned}$$

We can use the above IC characterization to simplify the SQ constraint to

$$\underline{p}n^h + n^l - \underline{p}H(\underline{p}) - L(\underline{p}) - \sup_p \left\{ \int_{\underline{p}}^p (H(z) - n^h)dz \right\} \geq 0.$$

Let $\phi(p) := p - \frac{1-F(p)}{f(p)}$ be the *virtual type* of p . Putting together our previous observations, the program defining the support function $S(a, b)$ becomes

$$S(a, b) = \max_{H, L} b(\underline{p}H(\underline{p}) + L(\underline{p})) + \int_{\underline{p}}^{\bar{p}} H(p)(a - b\phi(p))dF(p) \quad (6)$$

$$s.t \quad H \text{ is non-increasing} \quad (\text{IC}')$$

$$\underline{p}n^h + n^l - \underline{p}H(\underline{p}) - L(\underline{p}) - \sup_p \left\{ \int_{\underline{p}}^p (H(z) - n^h)dz \right\} \geq 0 \quad (\text{SQ}')$$

$$H(p) \geq 0, \quad L(p) \geq 0 \quad \forall p \in [\underline{p}, \bar{p}]$$

By inspection of the program in eq. (6), it is optimal to choose $L(\underline{p})$ so that the (SQ') constraint binds. Then the program becomes

$$\begin{aligned} \max_{H \geq 0} \quad & b \left(\underline{p} n^h + n^l - \sup_p \left\{ \int_{\underline{p}}^p (H(z) - n^h) dz \right\} \right) + \int_{\underline{p}}^{\bar{p}} H(p) (a - b\phi(p)) dF(p) \\ \text{s.t.} \quad & H \text{ is non-increasing} \end{aligned} \quad (\text{IC}')$$

$$\underline{p} n^h + n^l - \underline{p} H(\underline{p}) - \sup_p \left\{ \int_{\underline{p}}^p (H(z) - n^h) dz \right\} \geq 0 \quad (\text{non-negative } L)$$

Now notice that because H is non-increasing, $\sup_p \left\{ \int_{\underline{p}}^p (H(z) - n^h) dz \right\} = \int_{\underline{p}}^{p^*} (H(z) - n^h) dz$, for any $\sup\{p : H(z) > n^h\} \leq p^* \leq \inf\{p : H(z) < n^h\}$. So we can solve the above program in two steps. First, for any fixed $\underline{p} \leq p^* \leq \bar{p}$ we solve

$$\begin{aligned} \max_{H \geq 0} \quad & b \left(\underline{p} n^h + n^l - \int_{\underline{p}}^{p^*} (H(z) - n^h) dz \right) + \int_{\underline{p}}^{\bar{p}} H(p) (a - b\phi(p)) dF(p) \\ \text{s.t.} \quad & H \text{ is non-increasing} \end{aligned} \quad (\text{IC}')$$

$$\underline{p} n^h + n^l - \underline{p} H(\underline{p}) - \int_{\underline{p}}^{p^*} (H(z) - n^h) dz \geq 0 \quad (\text{non-neg. } L)$$

$$H(p) \geq n^h \quad \forall p \in [\underline{p}, p^*] \quad , \quad H(p) \leq n^h \quad \forall p \in [p^*, \bar{p}]$$

then we can optimize over p^* . We can solve this program separately for H on $[\underline{p}, p^*]$ and H on $[p^*, \bar{p}]$. First, fix H on $[\underline{p}, p^*]$. Then we choose H on $[p^*, \bar{p}]$ to solve

$$\begin{aligned} \max_{H \geq 0} \quad & \int_{p^*}^{\bar{p}} H(p) (a - b\phi(p)) dF(p) \\ \text{s.t.} \quad & H \text{ is non-increasing} \\ & H(p) \leq n^h \quad \forall p \in [p^*, \bar{p}] \end{aligned} \quad (\text{IC}')$$

This looks exactly like a standard monopoly pricing problem. The extreme points of the set of feasible functions are step functions taking values in $\{0, n^h\}$. As the objective is linear, is a solution in this set. If $(a, b) > 0$ and ϕ is strictly increasing, the solution is unique.

Now consider the other half of the problem, choosing H on $[\underline{p}, p^*]$. Rearranging the non-negative L constraint, we have

$$\begin{aligned} \max_H \quad & b(p^* n^h + n^l) + \int_{\underline{p}}^{p^*} H(p) \left(a - b \left(p + \frac{F(p)}{f(p)} \right) \right) dF(p) \\ \text{s.t.} \quad & H \text{ is non-increasing} \end{aligned} \quad (\text{IC}')$$

$$p^*n^h + n^l - \underline{p}H(\underline{p}) \geq \int_{\underline{p}}^{p^*} H(z)dz \quad (\text{non-neg. } L)$$

$$H(p) \geq n^h \quad \forall p \in [\underline{p}, p^*]$$

Fix $H(\underline{p}) > n^h$. The standard ironing argument implies that there is an optimal mechanism taking at most three values in $\{n^h, x, H(\underline{p})\}$ for some $x \in (n^h, H(\underline{p}))$. To be precise, if we ignore the non-negative L constraint then the extreme points of the feasible set are step functions taking values in $\{n^h, H(\underline{p})\}$, and to satisfy the non-negative L constraint we need to take a mixture between at most two such functions. The solution is unique and takes on only values in $\{n^h, H(\underline{p})\}$ if $p + F(p)/f(p)$ is strictly increasing and $(a, b) > 0$.

Consider now the choice of $H(\underline{p})$. If the non-negative L constraint is slack, it is optimal to increase the value of $H(\underline{p})$ since doing so relaxes the monotonicity constraint (IC'). More explicitly, by the ironing argument we know that whenever the optimal mechanism given a fixed $H(\underline{p})$ takes three values, it must be that the non-negative L constraint binds. Thus (given the fixed $H(\underline{p})$) the non-negative L constraint is slack if and only if it is satisfied when we maximize over simple step functions, which means

$$(H(\underline{p}) - n^h) \min \left\{ \arg \max_{z \in [\underline{p}, p^*]} \left\{ \int_{\underline{p}}^z \left(a - b \left(p + \frac{F(p)}{f(p)} \right) \right) dF(p) \right\} \right\} < n^l$$

However if this holds then it would be optimal to increase $H(\underline{p})$. Thus the non-negative L constraint always binds (meaning $L(\underline{p}) = 0$) under the optimal mechanism.

Combining the solutions above and below p^* yields the general solution described in Theorem 1. When ϕ and $p + F(b)/f(b)$ are strictly increasing there is an optimal mechanism must use only two prices, which is unique if $(a, b) > 0$. $\square$

A.2 Proof of Theorem 2

We first verify that strong duality holds. We know $(n^h, n^l) \in \mathcal{F}_j$ for all j . If (n^h, n^l) is, moreover, in the interior of $\mathcal{F}_j$ for all j then Slater's condition is satisfied, so we are done. Otherwise, define a perturbed problem by letting $S_j^\varepsilon(a, b) = S_j(a, b)(1 + \varepsilon)$ for all j , and $\mathcal{F}_j^\varepsilon = \{(\hat{n}^h, \hat{n}^l) \in \mathbb{R}^2 : a\hat{n}^h + b\hat{n}^l \leq S_j^\varepsilon(a, b) \forall (a, b) \in \mathbb{R}^2\}$. Then $(\hat{n}^h, \hat{n}^l)$ is in the interior of $\mathcal{F}_j^\varepsilon$ for all j , so Slater's condition is satisfied for any $\varepsilon > 0$. The value of the dual in the perturbed program is

$$\max_{\lambda_h, \lambda_l} \lambda_h J n^h + \lambda_l J n^l - (1 + \varepsilon) \sum_{j=1}^J S_j((\lambda_h - c^h(j)), (\lambda_l - c^l(j)))$$

which converges to the value of the original dual in eq. (4) as $\varepsilon \rightarrow 0$. Moreover $\varepsilon \mapsto \mathcal{F}_j^\varepsilon$ is upper- and lower-hemicontinuous; and compact and convex valued, so the value of the primal program is continuous in ε by Berge's maximum theorem. Because strong duality holds for all $\varepsilon > 0$, we conclude that it holds in the limit.

The first part of the theorem, up to and including the claim that each agent receives a taskload on the boundary of $\mathcal{F}_j$, is immediate from the dual formulation. Suppose now that no two agents are identical, in the sense stated in the result. We first show that at most two agents have non-zero allocations that are off of the frontier. Note that if $a, b < 0$ then $N_j^*(a, b) = \{(0, 0)\}$, and $N_j^*(a, b)$ contains non-zero points that are off of the frontier if and only if either $a \leq b = 0$ or $b \leq a = 0$. If agents are not identical, for any λ_h, λ_l there is at most one j such that $\lambda_h - c^h(j) = 0$, and one j' such that $\lambda_l - c^l(j') = 0$.

It remains to prove the stated implications of strong regularity. There are two cases to consider. First, suppose there exists an optimal mechanism such that some agent, j , receives a strictly positive quantity of both tasks. Let (λ_h, λ_l) be the corresponding dual solution. Suppose there exists another dual solution (λ'_h, λ'_l) . Then there exists α_j such that $(\lambda_h - c^h(j), \lambda_l - c^l(j)) = \alpha_j(\lambda'_h - c^h(j), \lambda'_l - c^l(j))$; otherwise S_k is strictly convex over some portion of the segment between (λ'_h, λ'_l) and (λ_h, λ_l) (given that F_k is doubly regular), so a mixture obtains a larger value in the dual program. Suppose $\alpha_j > 1$ ($\alpha_j < 1$ is symmetric). Then $\lambda_h/\lambda'_h < \alpha_j$ and $\lambda_l/\lambda'_l < \alpha_j$, so in fact (λ'_h, λ'_l) yields a strictly higher value in the dual program. Thus the dual solution is unique. Moreover, under strong regularity $N_j^*(\lambda_h - c^h(j), \lambda_l - c^l(j))$ is single-valued for all j such that either $\lambda_h - c^h(j) \neq 0$ and/or $\lambda_l - c^l(j) \neq 0$. The solution for the remaining two agents is uniquely pinned down by market clearing.

Alternatively, suppose that in all solutions every agent receives one type of task. Then in any solution there is a set $A \subset \mathcal{J}$ of agents who receive no l tasks, and a set $B \subset \mathcal{J}$ of agents who receive no h tasks. For each pair of sets (A, B) there is a unique allocation of the tasks (under the non-identical $c^k(j)$ assumption): among A give as many tasks as possible to the agents with lower $c^h(j)$, and similarly for B . Suppose that there two solutions, say (A, B) and (A', B') , such that $j \in A \cap B'$ and j receives a non-zero allocation in both. As the objective is linear, the half-half mixture of these two assignments is also a solution. However in that case j gets some of both types of tasks, contradicting our initial assumption.

A.3 Proof of Theorem 3

Proof. We begin with some preliminary comparative statics observations.

Lemma 2. If $c^k(j)$ increases (fixing $c^{-k}(j)$) then in expectation j receives fewer type- k tasks in the optimal LMS-TP mechanism (where the expectation is taken over p_j). Similarly, if $c^k(j) > c^k(j')$, $c^{-k}(j) = c^{-k}(j')$, and $F_j = F_{j'}$ then j receives fewer type- k tasks than j' in expectation.

Proof. The first case is easily seen by observing that the objective function in the program in eq. (2) has increasing differences in $c^k(j)$ and $\hat{n}_j^k$. The second case is immediate from eq. (4) and the definition of S^j . $\square$

Given Lemma 2, the remaining question is how changes in the optimal expected taskloads translate into changes in the prices offered to each agent.

Consider now the claim about local fairness. If j is remedial then any agent with worse performance is excluded, so j cannot have justified envy. Assume therefore that j is on their frontier. The prices p_1^j, p_2^j defining the optimal LMS-TP mechanism solve

$$\max_{p_j \leq p_1 \leq p_2 \leq \bar{p}^j} (\lambda_h - c^h(j)) \left(F_j(p_2) n^h + \frac{F_j(p_1)}{p_1} n^l \right) + (\lambda_l - c^l(j)) ((1 - F_j(p_1)) n^l + (1 - F_j(p_2)) p_2 n^h).$$

Fixing λ_h, λ_l , the solution p_1^j is decreasing in $c^h(j)$ if $p \mapsto F_j(p)/p$ is increasing. Similarly, p_2^j is increasing in $c^l(j)$ if $p \mapsto (1 - F_j(p))p$ is decreasing, which is equivalent to $f(p)p \geq (1 - F(p))$. This holds under strong regularity.

Suppose $p_j < p_1^j$, so $H^j(p_j) = n^h + \frac{1}{p_1^j} n^l$ and $L^j(p_j) = 0$. If $c^h(j') > c^h(j)$ and $c^l(j') = c^l(j)$ then by the previous claim $p_1^j \geq p_1^{j'}$. Similarly, if $p_j > p_2^j$, $c^h(j') = c^h(j)$ and $c^l(j') > c^l(j)$ then $p_2^j < p_2^{j'}$. The workload of agent j is

$$\min \left\{ p_j(n^h + \frac{1}{p_1^j} n^l), (p_j n^h + n^l), (p_2^j n^h + n^l) \right\}.$$

If $p_j \leq p_1^j$ then the marginal impact of increasing p_1^j is $-(p_1^j)^{-2} p_j n^l < 0$. Similarly if $p_j \geq p_2^j$ then the marginal impact of reducing p_2^j is $-n^h < 0$. Finally, if $p_j \in (p_1^j, p_2^j)$ then the agent's welfare is invariant to local perturbations of $c^h(j), c^l(j)$. This proves local fairness for j .

We now prove that the mechanism is locally effort inducing for non-remedial agents. Following the same argument used to prove local fairness above, it suffices to show that p_j^1 is decreasing in $c^h(j)$ and p_j^2 is increasing in $c^l(j)$. Unlike comparisons across agents, however, to evaluate changes in an individual agent's performance we need to know how λ^h, λ^l respond to changes in $c^h(j), c^l(j)$. To this end, recall the conclusions of Lemma 2. From eq. (3) we can see that in order for the expected allocation of type- k tasks for j to strictly increase, it must be that

$(\lambda_h - c^h(j)/(\lambda_l - c^l(j)))$ increases. This follows from convexity of $\mathcal{F}_j$. Given this observation, the comparative statics above (for fixed λ_h, λ_l) go through unchanged. $\square$

References

- Aggarwal, Gagan, Gagan Goel, Chinmay Karande, and Aranyak Mehta (2011) “Online vertex-weighted bipartite matching and single-bid budgeted allocations,” in *Proceedings of the twenty-second annual ACM-SIAM symposium on Discrete Algorithms*, 1253–1264, SIAM.
- Akbarpour, Mohammad, Eric Budish, Piotr Dworczak, and Scott Duke Kominers (2024) “An economic framework for vaccine prioritization,” *The Quarterly Journal of Economics*, 139 (1), 359–417.
- Angelova, Victoria, Will Dobbie, and Crystal Yang (2023) “Algorithmic Recommendations and Human Discretion,” Working paper.
- Arbour, William and Steeve Marchand (2024) “Can parole reduce both time served and crime,” *Working paper*.
- Arnold, David, Will S Dobbie, and Peter Hull (2022) “Measuring Racial Discrimination in Bail Decisions,” *American Economic Review*, 112 (9), 2992–3038.
- Baccara, Mariagiovanna, Allan Collard-Wexler, Leonardo Felli, and Leeat Yariv (2014) “Child-Adoption Matching: Preferences for Gender and Race,” *American Economic Journal: Applied Economics*, 6 (3), 133–58.
- Baccara, Mariagiovanna, SangMok Lee, and Leeat Yariv (2020) “Optimal dynamic matching,” *Theoretical Economics*, 15 (3), 1221–1278.
- Baker, George, Robert Gibbons, and Kevin J Murphy (2001) “Bringing the market inside the firm?” *American Economic Review*, 91 (2), 212–218.
- Bald, Anthony, Joseph Doyle Jr., Max Gross, and Brian Jacob (2022) “Economics of Foster Care,” *Journal of Economic Perspectives*, 36 (2), 223–246.
- Baron, E. Jason, Joseph J. Doyle, Natalia Emanuel, Peter Hull, and Joseph Ryan (2024) “Discrimination in Multi-Phase Systems: Evidence from Child Protection,” *The Quarterly Journal of Economics*, Forthcoming.
- Bhuller, Manudeep, Gordon B Dahl, Katrine V Løken, and Magne Mogstad (2020) “Incarceration, recidivism, and employment,” *Journal of Political Economy*, 128 (4), 1269–1324.
- Border, Kim C (1991) “Implementation of reduced form auctions: A geometric approach,” *Econometrica*, 1175–1187.
- Budish, Eric (2011) “The combinatorial assignment problem: Approximate competitive equilibrium from equal incomes,” *Journal of Political Economy*, 119 (6), 1061–1103.
- Casey Family Programs (2023) “Turnover Costs and Retention Strategies,” <https://www.casey.org/turnover-costs-and-retention-strategies/>, Accessed: 2024-08-27.
- Chan, David C, Matthew Gentzkow, and Chuan Yu (2022) “Selection with Variation in Diagnostic Skill: Evidence from Radiologists,” *The Quarterly Journal of Economics*, 137 (2), 729–783.
- Combe, Julien, Vladyslav Nora, and Olivier Tercieux (2021) “Dynamic assignment without money: Optimality of spot mechanisms,” *Working paper*.

- Combe, Julien, Olivier Tercieux, and Camille Terrier (2022) “The design of teacher assignment: Theory and evidence,” *The Review of Economic Studies*, 89 (6), 3154–3222.
- Dinerstein, Michael and Troy D Smith (2021) “Quantifying the supply response of private schools to public policies,” *American Economic Review*, 111 (10), 3376–3417.
- Dobbie, Will and Jae Song (2015) “Debt relief and debtor outcomes: Measuring the effects of consumer bankruptcy protection,” *American Economic Review*, 105 (3), 1272–1311.
- Doyle, Joseph J (2007) “Child Protection and Child Outcomes: Measuring the Effects of Foster Care,” *American Economic Review*, 97 (5), 1583–1610.
- Doyle, Joseph J, Natalia Emanuel, and Raimundo Eyzaguirre (2025) “Cumulative Risks of Foster Care Placement for US Children: Comparing Birth Cohort and Synthetic Cohort Approaches,” *NBER Working Paper #w34069*.
- Dworczak, Piotr, Scott Duke Kominers, and Mohammad Akbarpour (2021) “Redistribution through markets,” *Econometrica*, 89 (4), 1665–1698.
- Dworczak, Piotr and Ellen Muir (2025) “A mechanism-design approach to property rights,” Working paper.
- Grossman, Sanford J and Oliver D Hart (1983) “An analysis of the principal-agent problem,” *Econometrica*.
- Hart, Sergiu and Philip J Reny (2015) “Maximal revenue with multiple goods: Nonmonotonicity and other observations,” *Theoretical Economics*, 10 (3), 893–922.
- Holmstrom, Bengt (1989) “Agency costs and innovation,” *Journal of Economic Behavior & Organization*, 12 (3), 305–327.
- Hylland, Aanund and Richard Zeckhauser (1979) “The efficient allocation of individuals to positions,” *Journal of Political Economy*, 87 (2), 293–314.
- Jullien, Bruno (2000) “Participation constraints in adverse selection models,” *Journal of Economic Theory*, 93 (1), 1–47.
- Karp, Richard M, Umesh V Vazirani, and Vijay V Vazirani (1990) “An optimal algorithm for on-line bipartite matching,” in *Proceedings of the twenty-second annual ACM symposium on Theory of computing*, 352–358.
- Kim, Hyunil, Christopher Wildeman, Melissa Jonson-Reid, and Brett Drake (2017) “Lifetime Prevalence of Investigating Child Maltreatment Among US Children,” *American Journal of Public Health*, 107 (2), 274–280.
- Kleinberg, Jon, Himabindu Lakkaraju, Jure Leskovec, Jens Ludwig, and Sendhil Mullainathan (2018) “Human Decisions and Machine Predictions,” *Quarterly Journal of Economics*, 133 (1), 237–293.
- Kling, Jeffrey R (2006) “Incarceration length, employment, and earnings,” *American Economic Review*, 96 (3), 863–876.
- Krishna, Kalyan, Sergey Lychagin, Wojciech Olszewski, Ron Siegel, and Chloe Tergiman (2026) “Pareto Improvements in the Contest for College Admissions,” *Review of Economic Studies*, 93 (1), 629–663.
- Lee, Logan M (2023) “Halfway home? residential housing and reincarceration,” *American Economic Journal: Applied Economics*, 15 (3), 117–149.
- Lewis, Tracy R and David EM Sappington (1989) “Countervailing incentives in agency problems,” *Journal of Economic Theory*, 49 (2), 294–313.

- Li, Shengwu (2017) “Obviously strategy-proof mechanisms,” *American Economic Review*, 107 (11), 3257–3287.
- Loertscher, Simon and Ellen V Muir (2021) “Wage dispersion, minimum wages and involuntary unemployment: A mechanism design perspective,” *Working paper*.
- MacDonald, Diana E. (2024) “Foster Care: A Dynamic Matching Approach,” *Working paper*.
- Maggi, Giovanni and Andres Rodriguez-Clare (1995) “On countervailing incentives,” *Journal of Economic Theory*, 66 (1), 238–263.
- Mehta, Aranyak, Amin Saberi, Umesh Vazirani, and Vijay Vazirani (2007) “Adwords and generalized online matching,” *Journal of the ACM (JACM)*, 54 (5), 22–es.
- Milgrom, Paul and Ilya Segal (2002) “Envelope theorems for arbitrary choice sets,” *Econometrica*, 70 (2), 583–601.
- Myerson, Roger B (1981) “Optimal auction design,” *Mathematics of operations research*, 6 (1), 58–73.
- Nguyen, Thành, Alexander Teytelboym, and Shai Vardi (2023) “Dynamic Combinatorial Assignment,” *Working paper*.
- Nguyen, Thành and Rakesh Vohra (2021) “ Δ -Substitute Preferences and Equilibria with Indivisibilities,” in *Proceedings of the 22nd ACM Conference on Economics and Computation*, 735–736.
- Prendergast, Canice (2022) “The allocation of food to food banks,” *Journal of Political Economy*, 130 (8), 1993–2017.
- Putnam-Hornstein, Emily, John Prindle, and Ivy Hammond (2021) “Engaging Families in Voluntary Prevention Services to Reduce Future Child Abuse and Neglect: A Randomized Controlled Trial,” *Prevention Science*, 22 (7), 856–865.
- Roberts, Donald John and Andrew Postlewaite (1976) “The incentives for price-taking behavior in large exchange economies,” *Econometrica*, 115–127.
- Robinson-Cortes, Alejandro (2019) “Who gets placed where and why? An empirical framework for foster care placement,” *Working paper*.
- Slaugh, Vincent W, Mustafa Akan, Onur Kesten, and M Utku Ünver (2016) “The Pennsylvania Adoption Exchange improves its matching process,” *Interfaces*, 46 (2), 133–153.
- Sönmez, Tayfun (2023) “Minimalist market design: A framework for economists with policy aspirations,” *arXiv preprint arXiv:2401.00307*.
- Valenzuela-Stookey, Quitzé (2023) “Interim Allocations and Greedy Algorithms,” *Working paper*.
- Valenzuela-Stookey, Quitzé (2025) “Automation and Task Allocation Under Asymmetric Information,” *Working paper*.
- Valenzuela-Stookey, Quitzé, Jason Baron, and Richard Lombardo (2026) “Obviously Strategy-Proof Multi-Dimensional Allocation and the Value of Choice,” *Working paper*.
- Varian, Hal R (1973) “Equity, envy, and efficiency,” *Journal of Economic Theory*.

Online Appendix

B Extensions and additional properties

B.1 Learning about $c^h(j), c^l(j)$

In Section IV.D, we leverage the quasi-random nature of the observed assignment to identify the cost parameters $c^k(j)$. A natural concern is that if one were to implement the SMD-TP mechanism, we would lose the ability to continue to learn about the performance, $c^h(j)$ and $c^l(j)$, of agents. Fortunately, what matters for identification is that the assignment be quasi-random *conditional on task type*, which the SMD-TP mechanism is. The only remaining challenge to continued learning about agent performance is that the SMD-TP assignment may violate the full-support condition of Lemma 1. In other words, if agent j never receives any type- k tasks, then we cannot hope to learn about $c^k(j)$.

A simple way to solve this problem is to introduce some additional randomness into the mechanism, so that every agent receives at least some of each task type. This is also how the proposed mechanism can accommodate the arrival of new agents: by keeping them on the status-quo track until we have enough information to estimate their performance. In essence, we face the familiar experimentation-exploitation trade-off (Weitzman, 1978; Bolton and Harris, 1999). A more sophisticated solution would involve explicitly modeling this trade-off as part of the mechanism design problem, as in Kasy and Teytelboym (2023).

B.2 Finite-sample adjustments to SMD-TP mechanism

To move between the two extremes of assigning based on the difference between realized and target taskloads, versus assigning based on the ratio, we can modify the algorithm by adjusting the score as follows for some $\varepsilon > 0$

$$\tilde{r}_j(t, k) = \frac{\hat{n}_j^k(t) + \varepsilon}{\dot{n}_j^k + \varepsilon}.$$

For large ε the assignments generated by using the ratio $\tilde{r}$ converge to those generated by using the difference $\hat{n}_j^k(t) - \dot{n}_j^k$.

Lemma 3. For any $\hat{n}_j^k, \hat{n}_m^k$ and $\dot{n}_j^k, \dot{n}_m^k$, there exists x large enough such that

$$\frac{\hat{n}_j^k + \varepsilon}{\dot{n}_j^k + \varepsilon} < \frac{\hat{n}_m^k + \varepsilon}{\dot{n}_m^k + \varepsilon} \Leftrightarrow \dot{n}_j^k - \hat{n}_j^k > \dot{n}_m^k - \hat{n}_m^k$$

for all $\varepsilon > x$.

Proof.

$$\begin{aligned}
\frac{\hat{n}_j^k + \varepsilon}{\dot{n}_j^k + \varepsilon} < \frac{\hat{n}_m^k + \varepsilon}{\dot{n}_m^k + \varepsilon} &\Leftrightarrow (\hat{n}_j^k + \varepsilon)(\dot{n}_m^k + \varepsilon) < (\hat{n}_m^k + \varepsilon)(\dot{n}_j^k + \varepsilon) \\
&\Leftrightarrow \varepsilon(\hat{n}_j^k + \dot{n}_m^k) + \hat{n}_j^k \dot{n}_m^k < \varepsilon(\hat{n}_m^k + \dot{n}_j^k) + \hat{n}_m^k \dot{n}_j^k \\
&\Leftrightarrow \hat{n}_m^k \dot{n}_j^k + \varepsilon(\dot{n}_j^k - \hat{n}_j^k) > \hat{n}_j^k \dot{n}_m^k + \varepsilon(\dot{n}_m^k - \hat{n}_m^k).
\end{aligned}$$

Taking ε large yields the result. $\square$

Thus, by adjusting ε we can smoothly move between the two extremes of assigning based on ratios and assigning based on differences. More generally, in finite samples we can balance the desire to smooth agents' taskloads over time on the one hand, accomplished by assigning based on the ratio, versus ensuring that the difference between target and realized taskloads is small, by using a generalized scoring rule of the form

$$\tilde{r}_j(t, k) = \frac{\hat{n}_j^k(t) + x(t)}{\dot{n}_j^k + x(t)}.$$

for some increasing function $f > 0$. The asymptotic properties of the SMD-TP mechanism are preserved, but it may be possible to adjust f to improve finite sample performance. We leave this as a topic for future work.

C Further omitted proofs

C.1 Proof of Corollary 2

Proof. We want to show that under strong regularity, the frontier of $\mathcal{F}$ is equal to $\{N^*(a, b) : (a, b) \geq 0\}$, i.e. it suffices to consider only non-negative weights. To prove this we show that if $a = 0, b > 0$ then $N^*(a, b) = (0, x)$ for $x > 0$, and if $a > 0, b = 0$ then $N^*(a, b) = (y, 0)$ for $y > 0$. This implies that the entire frontier is downward sloping, and is traced out by considering $(a, b) \geq 0$.

Consider first $a = 0, b > 0$. Then under strong regularity, $(a - b\phi(p))$ and $(a - b(p + F(p))/f(p))$ are negative for all p . Thus from the proof of Theorem 1, the optimal mechanism sets

$$H(p) = \begin{cases} n^h & \text{on } [\underline{p}, p^*] \\ 0 & \text{on } (p^*, \bar{p}] \end{cases}$$

for some $p^* \in [\underline{p}, \bar{p}]$. The value obtained is $bn^l F(p^*) + b(p^* n^h + n^l)(1 - F(p^*))$. Under low asymmetry this is maximized at $p^* = \underline{p}$. Suppose instead that $a > 0, b = 0$. Then from the

proof of Theorem 1 it is optimal to set

$$H(p) = \begin{cases} n^h + \frac{1}{p^*}n^l & \text{on } [\underline{p}, p^*] \\ n^h & \text{on } (p^*, \bar{p}] \end{cases}$$

for some $p^* \in [\underline{p}, \bar{p}]$. The value obtained is $a \left(n^h + \frac{1}{p^*}n^l \right) F(p^*) + an^h(1 - F(p^*)) = an^h + an^l \frac{F(p^*)}{p^*}$ which is maximized at $p^* = \bar{p}$ under low asymmetry.

We have shown that it suffices to consider $(a, b) \geq 0$. Then if F is strongly regular, Theorem 1 tells us that for any point $(\hat{n}^h, \hat{n}^l)$ on the frontier, the only way to implement $(\hat{n}^h, \hat{n}^l)$ is with a two-price mechanism. Suppose there is a linear segment of the frontier which contains distinct points $(\hat{n}_1^h, \hat{n}_1^l)$ and $(\hat{n}_2^h, \hat{n}_2^l)$. Then the mixture $\alpha(\hat{n}_1^h, \hat{n}_1^l) + (1 - \alpha)(\hat{n}_2^h, \hat{n}_2^l)$ can be induced by the α mixture of the two-price mechanisms that induce $(\hat{n}_1^h, \hat{n}_1^l)$ and $(\hat{n}_2^h, \hat{n}_2^l)$. However since such a mixture is not itself a two-price mechanism, $\alpha(\hat{n}_1^h, \hat{n}_1^l) + (1 - \alpha)(\hat{n}_2^h, \hat{n}_2^l)$ cannot be on the frontier. $\square$

C.2 Proof of Proposition 3

Consider first the case of $y \rightarrow \infty$. First, notice that in the large-market problem, there is an optimal mechanism which gives all identical agents the same allocation. This follows from eq. (4). We focus on this mechanism, and show that SMS-TP approximates it as $y \rightarrow \infty$.

In the replica economy, we index the k^{th} copy of agent j as (j, k) , so for example $p_{j,k}$ is the type of this agent. In theory, we could treat each (j, k) as an separate agent. However in order to obtain a lower bound for V_{SMS} , we assume that if (j, k) and (j, k') are both buyers (or both sellers) then they receive the same allocation. With a slight abuse of notation, denote the allocation for j 's copies who are buyers as b_j , and for those who are sellers as s_j .

Given a realized type profile P , let $\hat{F}_j(\cdot|P, y)$ be the empirical CDF of types among the y copies of agent j . So $y \cdot \hat{F}_j(p_1^j|P, y)$ is the number of the j -replica agents who are buyers, and $y \left(1 - \hat{F}_j(p_2^j|P, y) \right)$ is the number of these agents who are sellers. Then for a given type profile P we can write the program defining SMS-TP in the replica economy as

$$\begin{aligned} \min_{\{b_j\}_{j \in \mathcal{B}}, \{s_j\}_{j \in \mathcal{S}}, \{\dot{n}_j^h, \dot{n}_j^l\}_{j \notin \mathcal{E}}} \quad & \sum_{j \in \mathcal{B}} \hat{F}_j(p_1^j|P, y) b_j (c^h(j) - p_1^j c^l(j)) \\ & - \sum_{j \in \mathcal{S}} \left(1 - \hat{F}_j(p_2^j|P, y) \right) s_j (c^h(j) - p_2^j c^l(j)) + \sum_{j \notin \mathcal{E}} \dot{n}_j^h c^h(j) + \dot{n}_j^l c^l(j) \\ \text{s.t.} \quad & \bar{p}_j \dot{n}_j^h + \dot{n}_j^l \leq \bar{p}_j n_j^h + n_j^l \quad \forall j \notin \mathcal{E} \\ & \underline{p}_j \dot{n}_j^h + \dot{n}_j^l \leq \underline{p}_j n_j^h + n_j^l \quad \forall j \notin \mathcal{E} \end{aligned} \tag{7}$$

$$\begin{aligned}
& 0 \leq \dot{n}_j^h, \quad 0 \leq \dot{n}_j^l \quad \forall j \notin \mathcal{E} \quad ; \quad 0 \leq b_j \leq \frac{n^l}{p_1^j} \quad \forall j \in \mathcal{B} \quad ; \quad 0 \leq s_j \leq n^h \quad \forall j \in \mathcal{S} \\
& \sum_{j \in \mathcal{B}} \hat{F}_j(p_1^j | P, y) b_j - \sum_{j \in \mathcal{S}} \left(1 - \hat{F}_j(p_2^j | P, y)\right) s_j + \sum_{j \notin \mathcal{E}} \dot{n}_j^h + |\mathcal{E}| n^h = J n^h \quad (h\text{-capacity}) \\
& - \sum_{j \in \mathcal{B}} \hat{F}_j(p_1^j | P, y) p_1^j b_j + \sum_{j \in \mathcal{S}} \left(1 - \hat{F}_j(p_2^j | P, y)\right) p_2^j s_j + \sum_{j \notin \mathcal{E}} \dot{n}_j^l + |\mathcal{E}| n^l = J n^l \quad (l\text{-capacity})
\end{aligned}$$

Now notice if we replace $\hat{F}_j$ with F_j for all j in this program, the solution is precisely the LMS-TP allocation. This holds because the fact that the LMS solution is standard implies that the LMS allocation is feasible in the SMS-TP program for all $j \notin \mathcal{E}$ (even if the solution is non-standard, this is true for all j such that either $\lambda^h \neq c^h(j)$ and/or $\lambda^l \neq c^l(j)$). Moreover, by the strong law of large numbers $\hat{F}_j(p_1^j | P, y) \xrightarrow{a.s.} F_j(p_1^j)$ and $\hat{F}_j(p_2^j | P, y) \xrightarrow{a.s.} F_j(p_2^j)$ as $y \rightarrow \infty$. Thus to show that V_{SMS} converges to V_{OPT} , we just need to show that the value of the program in eq. (7) continuous in $\hat{F}_j$, which is what we proceed to demonstrate.

Let $\hat{F}_j^1 = \hat{F}_j(p_1^j | P, y)$ and $\hat{F}_j^2 = \hat{F}_j(p_2^j | P, y)$. Let $R \left((\hat{F}_j^1, \hat{F}_j^2)_{j=1}^J \right)$ be the set of feasible choice variables $((b_j)_{j \in \mathcal{B}}, (s_j)_{j \in \mathcal{S}}, (\dot{n}_j^h, \dot{n}_j^l)_{j \notin \mathcal{E}})$ in the above program given parameters $(\hat{F}_j^1, \hat{F}_j^2)_{j=1}^J$.

Lemma 4. R is upper and lower hemicontinuous.

Proof. Define $\varphi \left((\hat{F}_j^1, \hat{F}_j^2)_{j=1}^J \right)$ as the set of $((b_j)_{j \in \mathcal{B}}, (s_j)_{j \in \mathcal{S}}, (\dot{n}_j^h, \dot{n}_j^l)_{j \notin \mathcal{E}})$ satisfying

$$\begin{aligned}
& \bar{p}_j \dot{n}_j^h + \dot{n}_j^l \leq \bar{p}_j n_j^h + n_j^l \quad , \quad \underline{p}_j \dot{n}_j^h + \dot{n}_j^l \leq \underline{p}_j n_j^h + n_j^l \quad , \quad 0 \leq \dot{n}_j^h, \quad 0 \leq \dot{n}_j^l \quad \forall j \notin \mathcal{E} \\
& 0 \leq b_j \leq \frac{n^l}{p_1^j} \quad \forall j \in \mathcal{B} \quad , \quad 0 \leq s_j \leq n^h \quad \forall j \in \mathcal{S} \\
& \sum_{j \in \mathcal{B}} \hat{F}_j(p_1^j | P, y) b_j - \sum_{j \in \mathcal{S}} \left(1 - \hat{F}_j(p_2^j | P, y)\right) s_j + \sum_{j \notin \mathcal{E}} \dot{n}_j^h = J n^h - |\mathcal{E}| n^h \quad (h\text{-capacity})
\end{aligned}$$

and $\eta \left((\hat{F}_j^1, \hat{F}_j^2)_{j=1}^J \right)$ as the set satisfying

$$\begin{aligned}
& \bar{p}_j \dot{n}_j^h + \dot{n}_j^l \leq \bar{p}_j n_j^h + n_j^l \quad , \quad \underline{p}_j \dot{n}_j^h + \dot{n}_j^l \leq \underline{p}_j n_j^h + n_j^l \quad , \quad 0 \leq \dot{n}_j^h, \quad 0 \leq \dot{n}_j^l \quad \forall j \notin \mathcal{E} \\
& 0 \leq b_j \leq \frac{n^l}{p_1^j} \quad \forall j \in \mathcal{B} \quad , \quad 0 \leq s_j \leq n^h \quad \forall j \in \mathcal{S} \\
& - \sum_{j \in \mathcal{B}} \hat{F}_j(p_1^j | P, y) p_1^j b_j + \sum_{j \in \mathcal{S}} \left(1 - \hat{F}_j(p_2^j | P, y)\right) p_2^j s_j + \sum_{j \notin \mathcal{E}} \dot{n}_j^l = J n^l - |\mathcal{E}| n^l \quad (l\text{-capacity})
\end{aligned}$$

so that $R = \varphi \cap \eta$. Both ϕ and η are given by the intersection of the polytope defined by

$$\begin{aligned} \bar{p}_j \dot{n}_j^h + \dot{n}_j^l &\leq \bar{p}_j n_j^h + n_j^l \quad , \quad \underline{p}_j \dot{n}_j^h + \dot{n}_j^l \leq \underline{p}_j n_j^h + n_j^l \quad , \quad 0 \leq \dot{n}_j^h, 0 \leq \dot{n}_j^l \quad \forall j \notin \mathcal{E} \\ 0 \leq b_j &\leq \frac{n^l}{p_1^j} \quad \forall j \in \mathcal{B} \quad , \quad 0 \leq s_j \leq n^h \quad \forall j \in \mathcal{S} \end{aligned}$$

with a hyperplane the normal vector of which is a continuous function of $(\hat{F}_j^1, \hat{F}_j^2)_{j=1}^J$. Thus both φ and η are upper and lower hemicontinuous. Since both are also convex and compact valued, upper and lower hemicontinuity of $R = \varphi \cap \eta$ follows.³⁷ $\square$

Given Lemma 4, Berge's Maximum Theorem implies that the value of the program defining SMS-TP for the replica economy is continuous in $(\hat{F}_j^1, \hat{F}_j^2)_{j=1}^J$. Finally, by the strong law of large numbers $\hat{F}_j(p_1^j|P, y) \xrightarrow{a.s.} F_j(p_1^j)$ and $\hat{F}_j(p_2^j|P, y) \xrightarrow{a.s.} F_j(p_2^j)$ as $y \rightarrow \infty$. Combined with continuity of the program defining SMS-TP, this implies convergence of the expected cost to V_{SMS} . The case of F_j converging in distribution to a constant for all j is similar. Let $n \mapsto (F_j^n)_{j=1}^J$ be a sequence of distributions which converge in distribution to a vector of constants $(x_j)_{j=1}^J \in [\underline{p}, \bar{p}]^J$. (Note that we maintain the assumption that each F_j^n is regular.) In the limit, i.e. when each agent's type is known, V_{SMS} and V_{OPT} coincide. We now make use of the following intermediate result.

Lemma 5. $(F_j)_{j=1}^J \mapsto V_{OPT}((F_j)_{j=1}^J|y)$ and $(F_j)_{j=1}^J \mapsto (p_1^j, p_2^j)_{j=1}^J$ are continuous.

Proof. Recall that $(p_1^j, p_2^j)_{j=1}^J$ are defined from the solutions to eq. (4). S^j is continuous in F_j . Moreover, if F^j satisfies strict regularity for all j then the objective in eq. (4) is unique. The lemma follows from Berge's maximum theorem. $\square$

Moreover, following the same argument as that of Lemma 4, we can show that V_{SMS} is continuous in $(p_1^j, p_2^j)_{j=1}^J$. Combined with Lemma 5, this implies that V_{SMS} converges to V_{OPT} along any sequence of strictly regular $(F_j^n)_{j=1}^J$, as desired.

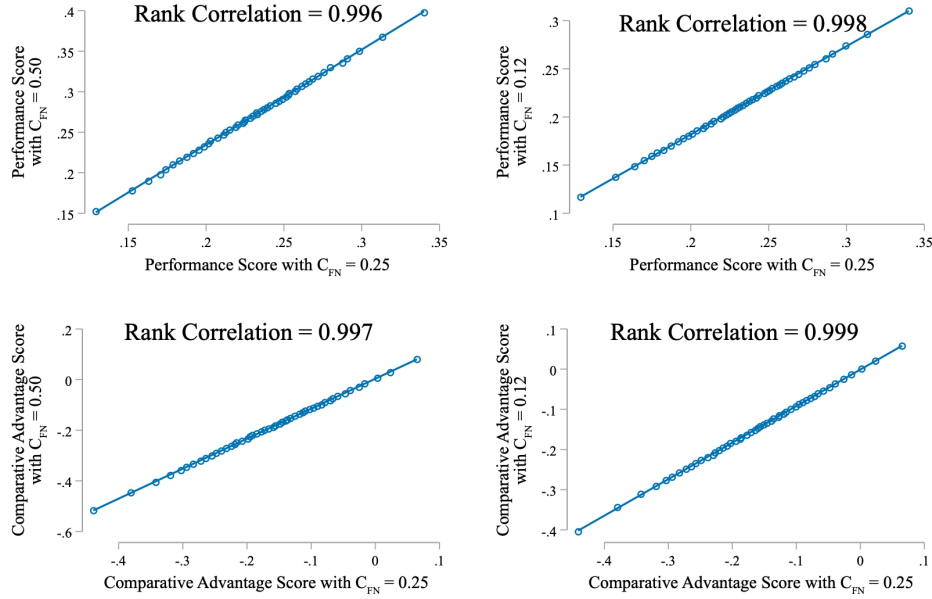
C.3 Proof of Proposition 5

Proof. By the strong law of large numbers, $\frac{1}{T} (N^k(T) - T\rho^k) \xrightarrow{a.s.} 0$ as $T \rightarrow \infty$. Then, by construction, for each $j \in J$ and $k \in \{h, l\}$, we have $r_j(T, k) \xrightarrow{a.s.} 1$ as $T \rightarrow \infty$, and so the mechanism is ε -IC for large enough T . Note also that $\frac{1}{T} V_{SMD}((F_j)_{j=1}^J, A, T)$ is just a weighted sum of $(\hat{n}_j^h, \hat{n}_j^l)_{j=1}^J$, and $\frac{1}{T} V_{SMS}((F_j)_{j=1}^J, A, T)$ is a weighted sum of $(\dot{n}_j^h, \dot{n}_j^l)_{j=1}^J$. Convergence of the ratio of values follows from convergence of $r_j(T, k)$ for all $j \in \mathcal{J}$ and $k \in \{h, l\}$. $\square$

³⁷See for example [Border \(2013\)](#) Proposition 24 (for upper hemicontinuity) and [Lechicki and Spakowski \(1985\)](#) (for lower hemicontinuity).

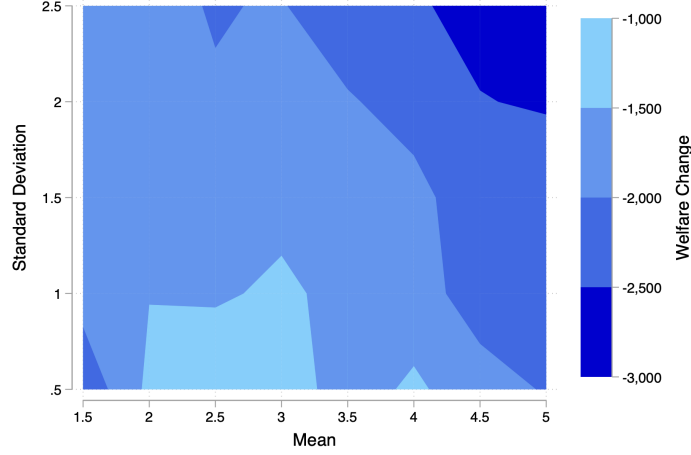
D Supplemental Figures and Tables

Figure A1: Robustness to Different Choices of Social Costs



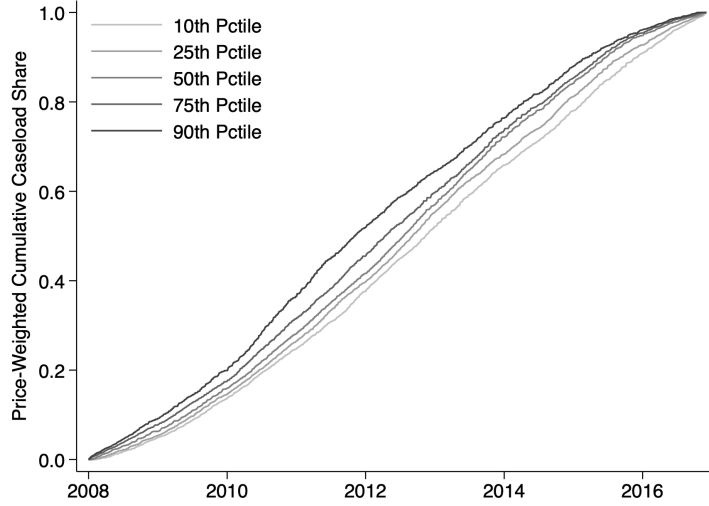
Notes. This figure plots the relationship between performance scores, γ_j , and comparative advantage scores, d_j , as we vary the choice of social costs. Benchmark estimates of γ_j and d_j are reported in the x-axis of each subfigure, with $(c_{TP}, c_{FN}, c_{FP}) = (0, 0.25, 1)$. In the left subfigures, we re-estimate γ_j and d_j with $c_{FN} = 0.50$. In the right subfigures, we re-estimate γ_j and d_j with $c_{FN} = 0.12$. Binned scatter plot estimates of the new performance score versus the benchmark performance score are displayed with 50 bins in each figure. We also report the Spearman's rank correlation coefficient between the new performance score measure and the benchmark measure. To minimize noise, for the comparative advantage estimates, the sample is limited to investigators that were assigned to at least 50 high-risk and low-risk cases across the sample. Investigator-specific and case type-specific estimates of subsequent maltreatment and placement rates are estimated via a regression adjustment for zipcode-by-year fixed effects.

Figure A2: Social Welfare Gains Across Distributional Assumptions



Notes. This figure presents the reduction in social costs (relative to the status quo) from implementing the LMS-TP mechanism under the assumption that F_j is truncated normal with mean, μ , and standard deviation, σ (shown in the horizontal and vertical axis, respectively). The distribution is truncated to $[1, \mu + 2\sigma]$. For computational purposes, we implement this exercise in the LMS-TP mechanism, for which the SMD-TP is an approximation.

Figure A3: Smoothing Investigator Caseloads Over Time



Notes. This figure reports the distribution of cumulative price-weighted caseloads assigned by the SMD-TP mechanism over time. The sample is limited to the 34 counties that appear in every sample year. We re-estimate welfare gains under SMD-TP for this sample and find very similar results relative to those in Table 1. Cases are assigned according to the SMD-TP mechanism for one investigator type draw where types are drawn from the regularized own-type empirical distribution. We compute cumulative price-weighted caseloads in day t for investigator j as $\frac{\hat{n}_j^l(t) + p_j \hat{n}_j^h(t)}{\hat{n}_j^l(T) + p_j \hat{n}_j^h(T)}$, where T is the last day of the sample period. We then report percentiles of this statistic for each day of our sample.

Table A1: Summary Statistics

<i>Panel A: Child Socio-Demographics</i>	
White	0.597
Black	0.266
Female	0.482
Child had a previous investigation	0.445
Number of previous investigations	1.024
Age at investigation	6.791
<i>Panel B: Investigation Traits</i>	
Alleged perpetrator is the mother/stepmother	0.772
Alleged perpetrator is the father/stepfather	0.328
Alleged perpetrator is a non-parent relative	0.053
Investigation included a domestic violence allegation	0.103
Investigation included an improper supervision allegation	0.530
Investigation included a medical neglect allegation	0.046
Investigation included a physical abuse allegation	0.290
Investigation included a physical neglect allegation	0.435
Investigation included a substance abuse allegation	0.170
<i>Panel C: Outcome, if left at home</i>	
Re-investigated for child maltreatment within 6 months	0.164
Foster care rate	0.032
Number of investigations	322,758
Number of children	261,021
Number of investigators	908

Notes. This table summarizes the analysis sample. The sample consists of maltreatment investigations of children in MI between 2008 and 2016, assigned to investigators who handled at least 200 cases during this period. The sample excludes repeat investigations and investigations of sexual abuse, as discussed in the main text. The final sample consists of 322,758 unique investigations of 261,021 children assigned to 908 investigators. Investigations can include multiple allegations and perpetrators, so these categories are not mutually exclusive.

Table A2: Estimates of Measures of Performance on Investigator Prediction Errors

	(1) False Negative	(2) Foster Care Placement	(3) False Positive
<i>Panel A: Across all Cases</i>			
Standardized Performance Score	1.13*** (0.13)	0.49*** (0.09)	1.62*** (0.14)
<i>Panel B: Across High-Risk Cases</i>			
Standardized Comparative Advantage Score	-2.04*** (0.32)	-0.25 (0.37)	-2.29*** (0.56)
<i>Panel C: Across Low-Risk Cases</i>			
Standardized Comparative Advantage Score	0.04 (0.18)	0.23 (0.24)	0.26 (0.22)

Notes. This table reports the results of OLS regressions of the investigator's false negative, foster care, and false positive rates on measures of their performance, γ_j and comparative advantage on high-risk cases, d_j . The independent variables are standardized to mean 0 and variance 1, and are estimated only in the randomized 50% training set. False negative rates and placement rates are estimated on the evaluation set, and are computed using a standard empirical Bayes shrinkage procedure. Implied false positive changes in Column 3 are estimated as the sum of coefficient estimates from Column 1 and Column 2, as $FP_j - FP_{j'} = (FN_j - FN_{j'}) + (P_j - P_{j'})$ by Lemma 1. In Panel A, we estimate this specification across all cases, in Panel B only among high-risk cases, and in Panel C among low-risk cases. All regressions are weighted by estimates of the inverse variance (clustered by investigator) of the investigator's performance or comparative advantage score. Baseline false negative and false positive rates are 15.9% and 2.3% over all cases, 24.0% and 4.0% over high-risk cases, and 13.1% and 1.9% over low-risk cases. False positive rates are identified via an identification-at-infinity strategy, described in Appendix E.6. Robust standard errors are reported in parentheses.

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

Table A3: Hazard Ratio Estimates of Risky Caseload Effect on Investigator Turnover

	(1)	(2)
	Career Length	Career Length
Mean Risk Level (Normalized)	2.488*** (0.380)	
Above Median High-Risk Share		1.538*** (0.158)
Investigator Count	908	908

Notes. This table reports the results of estimating a Cox proportional hazards model of investigator career length on caseload risk measures. We record an investigator’s career length as the distance (in days) between their first and last observed CPS case assignment, and denote that this length is censored if the investigator is working in 2016 (the final year of the sample). Column 1 uses mean risk level—the average algorithmic predicted risk score across all of this investigator’s cases, normalized to mean 0 variance 1 within each sample. Column 2 uses an indicator recording whether the share of an investigator’s cases that are high-risk is above the median for this sample. All estimates include a modal county fixed effect. We report the point estimates in terms of hazard ratios, with robust standard errors in parentheses.

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

Table A4: Gains from Investigator Reallocation, Alternative Type Distributions

	(1)	(2)	(3)	(4)	(5)	(6)
	Unif[1,2]	Unif[1,3]	$\mathcal{N}(2, 0.5^2)$	$\mathcal{N}(2, 1^2)$	$p_j = 2$	Known type, Unif[1,2]
Social Costs	-947.1*** (259.7) [-4.7%]	-977.6*** (232.0) [-4.9%]	-925.5*** (229.2) [-4.6%]	-940.6*** (246.0) [-4.7%]	-1,716.5*** (163.3) [-8.5%]	-1,860.1*** (285.5) [-9.2%]
False Negatives	-624.3*** (211.2) [-1.2%]	-641.6*** (180.1) [-1.3%]	-618.0*** (170.5) [-1.2%]	-618.1*** (191.7) [-1.2%]	-1,140.2*** (103.7) [-2.2%]	-1,282.3*** (234.0) [-2.5%]
False Positives	-805.2*** (224.4) [-11.0%]	-828.5*** (197.1) [-11.3%]	-785.6*** (179.3) [-10.7%]	-798.5*** (202.8) [-10.9%]	-1,460.7*** (124.7) [-20.0%]	-1,571.4*** (248.7) [-21.5%]
Placements	-180.9* (95.7) [-1.8%]	-186.9** (81.3) [-1.9%]	-167.6** (78.5) [-1.7%]	-180.4** (83.3) [-1.8%]	-320.5*** (62.2) [-3.2%]	-289.1*** (100.1) [-2.9%]

Notes. This table reports the welfare gains derived from the SMD-TP mechanism under different distributional assumptions. We estimate welfare changes and robust standard errors (in parenthesis), clustered by investigator, using the approach described in the text. Percent effects relative to the baselines are presented in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$

Table A5: Importance of Small Market and Dynamic Considerations

	(1)	(2)	(3)	(4)	(5)	(6)
	Regularized MRS			Unif[1,2]		
	LMS-TP	SMS-TP	SMD-TP	LMS-TP	SMS-TP	SMD-TP
Social Costs	-2,363.7*** (114.9) [-11.8%]	-1,189.0*** (311.4) [-5.9%]	-1,189.2*** (324.2) [-5.9%]	-1,850.8*** (93.4) [-9.2%]	-948.4*** (273.0) [-4.7%]	-947.1*** (266.3) [-4.7%]
False Negatives	-1,516.1*** (77.8) [-3.0%]	-755.1*** (236.1) [-1.5%]	-754.9*** (252.9) [-1.5%]	-1,228.4*** (70.4) [-2.4%]	-625.1*** (204.5) [-1.2%]	-624.3*** (195.8) [-1.2%]
False Positives	-2,028.8*** (93.4) [-27.7%]	-1,016.8*** (245.2) [-13.9%]	-1,017.0*** (262.9) [-13.9%]	-1,563.2*** (77.4) [-21.4%]	-806.2*** (229.2) [-11.0%]	-805.2*** (220.6) [-11.0%]
Placements	-512.6*** (49.3) [-5.2%]	-261.7*** (99.1) [-2.6%]	-262.1** (104.7) [-2.6%]	-334.8*** (38.4) [-3.4%]	-181.1* (94.5) [-1.8%]	-180.9* (96.0) [-1.8%]

Notes. This table reports the welfare gains derived from the LMS-TP, SMS-TP, and SMD-TP mechanisms. Columns 1-3 estimate welfare gains under the distribution assumption that investigator type is drawn from the regularized own-type empirical distribution: we estimate hazards from a kernel-smoothed empirical distribution from the survey and impose monotone hazards (Myerson regularity) by choosing the nondecreasing hazard sequence that minimizes the discrete sum of squared deviations in hazard rate. Columns 4 to 6 estimate welfare gains from uniform distributions with support [1, 2]. We estimate welfare changes and robust standard errors (in parenthesis), clustered by investigator, using the approach described in the text. Percent effects relative to the baselines are presented in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$

Table A6: Welfare Gains from Observed Counterfactual

	(1)	(2)	(3)
	Own Type	Coworker Type	Known Type, Own Type
Social Costs	-1,277.0*** (298.5) [-6.3%]	-1,198.9*** (271.1) [-5.9%]	-3,439.5*** (441.7) [-16.9%]
False Negatives	-815.4*** (242.7) [-1.6%]	-773.8*** (215.1) [-1.5%]	-2,269.5*** (383.9) [-4.4%]
False Positives	-1,091.7*** (260.3) [-14.5%]	-1,024.7*** (221.2) [-13.6%]	-2,939.9*** (418.3) [-39.1%]
Placements	-276.3** (108.9) [-2.8%]	-250.8** (105.9) [-2.5%]	-670.4*** (203.5) [-6.7%]

Notes. This table reports the welfare gains derived from the SMD-TP mechanism compared to a counterfactual approximating the observed assignment matrix that strictly restricts investigators to one county. Some cases are handled by investigators outside their modal counties—for such cases, we randomly reassign these cases to an investigator working in the focal county. Each column corresponds to a different distributional assumption for p_j . Column 1 presents gains from the regularized own-type empirical distribution: we estimate hazards from a kernel-smoothed empirical distribution from the survey and impose monotone hazards by choosing the nondecreasing hazard sequence that minimizes the discrete sum of squared deviations in hazard rate (see Appendix E.4 for details). Column 2 performs the same procedure but using investigator reports of their coworkers’ preferences. Column 3 uses the same type distribution as in Column 1, but assumes that p_j is known to the designer. We estimate welfare changes and robust standard errors (in parenthesis), clustered by investigator, using the approach described in the text. Percent effects relative to the baselines are presented in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$

Table A7: Welfare Gains for Five Largest Counties

	(1)	(2)	(3)
	Own Type	Coworker Type	Known Type, Own Type
Social Costs	-846.5*** (211.7) [-9.7%]	-790.2*** (172.4) [-9.1%]	-2,005.7*** (340.3) [-23.0%]
False Negatives	-509.0*** (171.0) [-2.3%]	-476.5*** (143.9) [-2.1%]	-1,222.2*** (255.6) [-5.5%]
False Positives	-728.0*** (182.6) [-23.1%]	-681.0*** (151.1) [-21.6%]	-1,730.7*** (261.4) [-54.9%]
Placements	-219.0*** (84.8) [-5.4%]	-204.5*** (76.3) [-5.0%]	-508.5*** (143.9) [-12.5%]

Notes. This table reports the welfare gains derived from the SMD-TP mechanism, using a sample of the 5 largest counties in the data. Each column corresponds to a different distributional assumption for p_j . Column 1 presents gains from the regularized own-type empirical distribution. Column 2 performs the same procedure but using investigator reports of their coworkers’ preferences. Column 3 uses the same type distribution as in Column 1, but assumes that p_j is known to the designer. We estimate welfare changes and robust standard errors (in parenthesis), clustered by investigator, using the approach described in the text. Percent effects relative to the baselines are presented in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$

Table A8: Welfare Gains Under Alternative Proxies for Subsequent Maltreatment

	(1) Baseline	(2) Inv. within 5 months	(3) Inv. within 4 months	(4) Inv. within 3 months	(5) Inv. within 2 months	(6) Inv. within 1 months	(7) Subst. Inv. within 6 months
Social Costs	-2,363.7*** (114.1) [-11.8%]	-2,014.7*** (103.0) [-10.8%]	-1,851.1*** (105.9) [-10.9%]	-1,659.7*** (97.1) [-10.6%]	-1,556.1*** (86.3) [-10.9%]	-1,477.8*** (77.1) [-11.8%]	-2,048.8*** (74.7) [-16.1%]
False Negatives	-1,516.1*** (83.5) [-3.0%]	-1,275.6*** (69.7) [-2.9%]	-1,170.4*** (80.2) [-3.2%]	-869.0*** (77.0) [-3.0%]	-788.4*** (54.6) [-3.9%]	-502.0*** (43.6) [-4.6%]	-924.2*** (54.5) [-6.1%]
False Positives	-2,028.8*** (96.9) [-27.7%]	-1,728.5*** (88.3) [-22.8%]	-1,591.0*** (96.6) [-20.5%]	-1,473.8*** (89.4) [-17.6%]	-1,388.6*** (70.3) [-15.2%]	-1,355.5*** (68.3) [-13.7%]	-1,814.3*** (71.7) [-20.3%]
Placements	-512.6*** (46.2) [-5.2%]	-452.9*** (49.0) [-4.6%]	-420.6*** (45.6) [-4.2%]	-604.8*** (40.7) [-6.1%]	-600.2*** (44.0) [-6.0%]	-853.5*** (51.3) [-8.6%]	-890.1*** (49.4) [-9.0%]

Notes. This table reports the welfare gains derived from the LMS-TP mechanism for different definitions of maltreatment risk under the distributional assumption that investigator types are drawn from the regularized own-type empirical distribution. We estimate welfare changes and robust standard errors (in parenthesis), clustered by investigator, using the approach described in the text. Percent effects relative to the baselines are presented in brackets. For computational purposes, we compare the welfare gains across different proxies for subsequent maltreatment in the LMS-TP mechanism, for which the SMD-TP is an approximation. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$

Table A9: Welfare Gains from Binary versus Quartile Case Splits

	(1) Binary Split	(2) Quartile Split
Social Costs	-1,506.1*** (350.9) [-14.0%]	-1,711.5*** (421.8) [-16.3%]
False Negatives	-922.2*** (289.8) [-3.5%]	-1,004.7*** (289.1) [-3.8%]
False Positives	-1,328.5*** (326.2) [-32.0%]	-1,477.0*** (303.8) [-37.7%]
Placements	-406.3*** (152.5) [-7.6%]	-472.3*** (160.8) [-9.1%]

Notes. This table reports the welfare gains derived from the SMD-TP mechanism using a quartile split versus a binary split. To define performance for these new categories we re-estimate Equations 8 and 9, but allowing for interactions for 2nd risk quartile by worker, 3rd risk quartile by worker, and 4th risk quartile by worker. We estimate this separately in the training and test samples. This table is limited to workers assigned to 50+ cases of each risk quartile group for all analysis. In both columns, it is assumed that the designer knows worker preferences. In Column 1, preferences on high-risk (top quartile risk) cases are drawn according to regularized own-type empirical distribution: we estimate hazards from a kernel-smoothed empirical distribution from the survey and impose monotone hazards (Myerson regularity) by choosing the nondecreasing hazard sequence that minimizes the discrete sum of squared deviations in hazard rate (see Appendix E.4 for details). In Column 2, worker performance is estimated across risk quartiles. Worker preferences for the riskiest quartile, p_{j4} are drawn from the regularized own-type empirical distribution as above. Worker preferences on the 3rd quartile, p_{j3} are distributed uniformly between $[1.2, 1.5]$. Worker preferences on the 2nd quartile, p_{j2} are distributed uniformly between $[0.9, 1.2]$. We then set $p_{j1} = 3 - p_{j2} - p_{j3}$ so that the mean of p_{j1}, p_{j2}, p_{j3} is equal to 1 for comparability with binary results. We estimate welfare changes and robust standard errors (in parenthesis), clustered by investigator, using the approach described in the text. Percent effects relative to the baselines are presented in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$

E Empirical Appendix

E.1 Data sources and analysis sample

We obtained the universe of child maltreatment investigations in Michigan between January 2008 and November 2016 from the Michigan Department of Health and Human Services. The dataset includes the allegation report date as well as child and investigation traits including the child’s zip code, age, gender, race, relationship to the alleged perpetrator, and maltreatment type (e.g., physical abuse versus neglect). It also includes indicators for whether the child was placed in foster care following the investigation and investigator numeric identifiers.

To construct the analysis sample, we begin with the set of child maltreatment investigations that did not involve either sexual abuse or repeat reports (involving children who had been

the subject of an investigation in the past year) as these cases are not quasi-randomly assigned to investigators. Given that foster care placement rates are low, we drop cases assigned to investigators who handled fewer than 200 investigations to minimize noise in our estimates of investigator placement rates ($N = 152,686$). We then drop observations in rotations (zip code by year pairs) with fewer than four investigators to compare investigators in a given “office” by year ($N = 22,201$). Furthermore, we drop cases for which we cannot observe subsequent child welfare outcomes for at least six months after the focal investigation ($N = 20,462$), as this will be the primary outcome of interest. We next drop a relatively small number of cases with missing child zip code information ($N = 4,856$), since quasi-random assignment of investigators is conditional on a zip code by year fixed effect. Finally, we limit to investigators assigned to at least 50 high- and low-risk cases, defined below, to limit noise in estimates of investigator comparative advantage ($N = 50,386$).

The final sample consists of 322,758 investigations involving 261,021 children assigned to 908 investigators. 3.2% of investigations result in foster care placement. Table A1 presents summary statistics: 60% of children in our sample are white, 48% are female, 45% have had a CPS investigation prior to their focal one, and the average child is nearly seven years old (Panel A). Investigations in our sample tend to include at least one allegation of improper supervision (53%), physical neglect (44%), and physical abuse (29%). In 77% of investigations, at least one of the alleged perpetrators of maltreatment is the child’s mother or step-mother (Panel B).

Panel C summarizes rates of “subsequent maltreatment” for children left at home following the focal investigation. Our primary maltreatment measure considers whether a child was re-investigated within six months of the focal investigation. This is a common proxy for subsequent maltreatment in the child welfare literature (e.g., [Baron et al. \(2024\)](#); [Putnam-Hornstein et al. \(2021\)](#)). Nevertheless, re-investigations are imperfect proxies for actual child maltreatment, as they only account for cases that are re-reported to CPS. While there are other potential proxies, such as a subsequent *substantiated* investigation, we prefer re-investigation because re-investigations within a few months may be assigned to the initial investigator who will again make substantiation decisions. In contrast, both the decision to re-report and to screen-in a case, the two steps required for a re-investigation, are outside of the initial investigator’s control. Still, we show below that our findings are robust when using these alternative proxies for subsequent maltreatment. With these caveats in mind, we refer to a re-investigation within six months as “subsequent maltreatment” throughout the manuscript for ease of exposition. Note that this maltreatment outcome is mechanically missing for children placed in foster care, which is the primary identification challenge in this study.

16.4% of children experience subsequent maltreatment in the home within six months of the focal investigation.³⁸

E.2 Estimation strategy

The results of Lemma 1 allow us to identify differences in social cost in low- and high-risk cases between investigators. Our estimation approach accounts for (i) over-fitting concerns and measurement error in investigator moment estimates, and (ii) the fact that investigators are quasi-randomly assigned only within offices. Define $\tilde{c}^k(j) := c^k(j) - c^k(j^0)$, where j^0 is the benchmark investigator used for all social cost comparisons.³⁹ Then $\tilde{c}^k(j)$ is identified and equal to $\tilde{c}^k(j) = \gamma_j^k - \gamma_{j^0}^k$.

To avoid concerns that our estimates of the benefits of reassignment are overstated due to over-fitting, we follow a split-sample strategy. Specifically, we randomize within the set of cases that each investigator was assigned into a “training” set (50%) and an “evaluation” set (50%). We use the training set to derive the optimal investigator assignment mechanism, and then test its effectiveness on the evaluation set.

We first estimate j ’s performance scores across case types: γ_j^l for low-risk cases and γ_j^h for high-risk cases. This requires investigator-specific estimates of placement and false negative rates. Let $D_i = \sum_j D_{ij}Z_{ij}$ and $\text{FN}_i = \sum_j \text{FN}_{ij}Z_{ij}$, so that D_i is an indicator for whether case i resulted in placement, and FN_i an indicator for whether the case is a false negative.

Investigators in Michigan are rotationally assigned to cases within CPS offices. Typically, each county in the state has its own office, but some large counties have multiple offices, and many offices split investigators into geographic-based teams (Baron and Gross, 2022). As such, we define an “office” throughout based on the child’s zip code. Our identification results apply separately to each office. To compare investigators across offices, we use a linear adjustment to estimate investigator placement and false negative rates.⁴⁰ That is, we estimate regressions of the form:

$$D_i = \sum_j (\beta_{j1}^D Z_{ij} + \beta_{j2}^D \text{High-Risk}_i Z_{ij}) + \mathbf{X}_i' \alpha^D + u_i \quad (8)$$

³⁸Moreover, while we maintain the assumption that Y_i^* is binary throughout, our identification and theoretical results can readily accommodate richer partitions of Y_i^* (see Appendix E.5).

³⁹The results of Section IV.D apply if j and j^0 are in the same office-by-year. However, under an additional linearity assumption introduced below, we can use a single reference investigator across all offices. The choice of j^0 has no impact on the mechanism results. We choose j^0 as the investigator with greatest caseload over our sample.

⁴⁰As discussed in Arnold et al. (2022), this approach tractably incorporates the large number of zip code-by-year fixed effects, under an additional assumption that placement and false negative rates are linear in the zip code-by-year effects for each investigator and case type.

$$\text{FN}_i = \sum_j (\beta_{j1}^{\text{FN}} Z_{ij} + \beta_{j2}^{\text{FN}} \text{High-Risk}_i Z_{ij}) + \mathbf{X}_i' \boldsymbol{\alpha}^{\text{FN}} + \epsilon_i \quad (9)$$

where High-Risk_i is an indicator equal to one if case i is high-risk. We estimate Equations 8 and 9 separately in the training and evaluation samples. $\mathbf{X}_i$ is a vector of office-by-year fixed effects to account for the level of randomization. We demean $\mathbf{X}_i$ so that β_{j1} represents strata-adjusted investigator-specific estimates of each outcome for low-risk cases and $\beta_{j1} + \beta_{j2}$ represents strata-adjusted investigator estimates of each outcome for high-risk cases. We use our estimates of these parameters to estimate performance scores among low- and high-risk cases, separately in the training and evaluation dataset, as:

$$\begin{aligned} \widehat{\gamma}_j^l &= c_{FP} \widehat{\beta}_{j1}^D + (c_{FN} + c_{FP} - c_{TP}) \widehat{\beta}_{j1}^{\text{FN}} \\ \widehat{\gamma}_j^h &= c_{FP} (\widehat{\beta}_{j1}^D + \widehat{\beta}_{j2}^D) + (c_{FN} + c_{FP} - c_{TP}) (\widehat{\beta}_{j1}^{\text{FN}} + \widehat{\beta}_{j2}^{\text{FN}}) \end{aligned}$$

Following Chan et al. (2022), we assume $c_{TP} = c_{TN} = 0$, so that the welfare measure is focused only on prediction mistakes. As mentioned above, the value of c_{FN}, c_{FP} must ultimately be chosen by the agency. To bring our mechanism to data, we assume that $c_{FP} = 1$ and $c_{FN} = 0.25$, though we show below that our results are robust to this choice of values. To motivate this choice, note that CPS investigators in our context place 3.2% of children but 16.4% of children face subsequent maltreatment when left at home. This mismatch may imply that CPS views $c_{FN} < c_{FP}$. Normalizing $c_{FP} = 1$ suggests that $c_{FN} \in (0, 1)$. For our benchmark estimates, the ratio between placement rates and subsequent maltreatment rates suggests that c_{FN} is roughly 0.25. We explore robustness to this assumption in Figure A1, where we show that the ranking of investigators is well-preserved if we instead assign, for example, $c_{FN} = 0.12$ or $c_{FN} = 0.5$. To reduce noise in the estimates of the investigator moments, we follow Arnold et al. (2022) and use empirical Bayes shrinkage estimates of $\widehat{\gamma}_j^l, \widehat{\gamma}_j^h$ to adjust for finite sample error. We then estimate $\tilde{c}^k(j)$ as $\widehat{\gamma}_j^k - \widehat{\gamma}_{j0}^k$ and d_j as $\widehat{\gamma}_j^l - \widehat{\gamma}_j^h$.⁴¹

E.3 Maltreatment risk prediction

We estimate an algorithmic risk prediction that child i will face subsequent maltreatment if left at home, $Pr(Y_i^* = 1|X_i)$, where X_i includes case and child attributes of case i available to the investigator at the time of the placement decision. Following Kleinberg et al. (2018), we use a gradient boosted decision tree to predict $Pr(Y_i^* = 1|X_i)$. We hypertune the algorithm to select for optimal parameters using a 5-fold cross-validation technique. Only children left at home are used to train the model since Y_i^* is unobserved for children placed in foster care. The features used to train the algorithm, X_i , are coded by the initial screener and include: the type of allegations in the investigation (physical abuse, medical neglect, physical neglect,

⁴¹We also estimate $\widehat{\gamma}_j$, the average performance of investigator j across all cases, by following this same methodology but omitting the interaction terms in Equations 8 and 9.

domestic violence, substance abuse, improper supervision), the relationship of the alleged perpetrator to the child, prior child welfare investigation history, the gender and age of the child, and their residing county. The out-of-sample AUC of the algorithm on left-at-home cases is 0.634.⁴²

E.4 Investigator Survey

We partnered with the Michigan Department of Health and Human Services (MDHHS) to design and administer a statewide survey of CPS investigators. Our goal was to better understand how investigators perceive their work—how they evaluate different types of investigations, how they manage caseload tradeoffs, and whether our modeling assumptions reflect how they themselves think about assignment and workload. The survey served four main purposes and yields several insights:

1. **Testing core modeling assumptions.** A central goal of the survey was to assess whether our model reflects how investigators perceive their work. Specifically, we sought to understand whether the nature of the cases—rather than just the number—shapes preferences over caseloads. To test this assumption, we asked respondents to compare high- and low-risk investigations along three dimensions—time demand, emotional toll, and satisfaction—using 0–10 scales. The results suggest that investigators perceive high-risk cases as substantially more burdensome, and only modestly more satisfying.
2. **Estimating heterogeneity in preferences.** A core motivation for our mechanism design is the hypothesis that investigators may have idiosyncratic preferences over high- and low-type cases. The survey allows us to test this. By estimating each respondent’s MRS between high- and low-risk cases, we recover a distribution of p across the workforce. The results show substantial variation—some investigators are highly averse to high-risk cases, while others are more willing to take them on. As we discuss in the main text, estimating the distribution of investigators’ MRS values also inform our empirical simulations of the potential welfare gains from implementing the mechanism.
3. **Connecting case complexity to outcomes.** In our administrative data, a one standard deviation increase in caseload risk is associated with a large increase in turnover risk. The survey provides complementary, qualitative evidence for this finding. Investigators consistently report that high-risk cases are more stressful and emotionally draining. Many explicitly link complexity to burnout and turnover.
4. **Probing the focus on binary-classification mechanisms.** The survey allows us to

⁴²The model is trained on an 80% random sample and evaluated on the remaining 20%. Because Y_i^* is observed only for children left at home, evaluation is restricted to that sample.

assess whether our focus on binary classification mechanisms can broadly capture how investigators perceive variation in their work in the real world. Across both structured and open-ended responses, we observe consistent separation between the two case types as defined in the paper (high- and low-risk). Investigators broadly prefer low-risk cases and describe high-risk cases as more demanding. These patterns suggest that our partition captures meaningful differences in how cases are experienced and valued. Moreover, the relatively high MRS estimates—along with open-text responses that we summarize below—indicate that a binary classification can capture significant heterogeneity in preferences.

The survey was distributed to CPS investigators via an MDHHS-managed email. While the listserv included approximately 1,200 to 1,400 recipients, the exact number of investigators who received the email is uncertain due to turnover and staff on leave. We received 459 responses, and, after filtering out incomplete entries, retained approximately 380 responses for descriptive sections and 322 for the MRS-elicitation questions. This yields a response rate between 23–27%. The survey link is available [here](#).

E.4.1 Survey Design and Key Measures

The survey consisted of four sections: open-ended reasoning, preferences over current workloads, perceptions of case types, and MRS elicitation.

In order to derive a practical policy, the paper focuses on binary-classification mechanisms. From a theoretical perspective, this partition can be arbitrary; in the empirical application, we define it using the predicted risk of subsequent maltreatment and split cases into high- and low-risk. In this section we sometimes refer to them as “high-” and “low-complexity” cases, as this terminology is more commonly used within Michigan CPS.

High-risk investigations are those in the top quartile of predicted risk. These cases typically involve younger children, families with prior CPS involvement, and multiple or severe allegations (e.g., substance abuse, domestic violence, physical neglect).

Low-risk investigations fall into the bottom three quartiles and typically involve: older children, families with little or no CPS history, and less severe or isolated allegations (e.g., minor neglect or improper supervision).

These definitions, along with a summary table of case characteristics, were presented to respondents at the beginning of the survey.

E.4.2 Qualitative Explanations for Tradeoff Choices

Before delving into responses to structured questions, we begin by presenting responses to open-ended questions about why investigators prefer certain bundles of cases over others.

We analyzed these open-text responses to identify recurring themes:

1. High-Complexity Cases Are More Demanding. Respondents consistently described high-complexity cases as more time-consuming, emotionally taxing, and harder to manage. One wrote: “High complexity investigations are more time consuming, are far more taxing to manage, [and] involve multiple households.” Another said: “Just one more high-complexity case can be the time equivalent of three or four low-complexity cases.”

2. Emotional Toll and Burnout Are Salient. Burnout and emotional exhaustion were common themes. “Secondary trauma is exhausting,” one respondent wrote. Another added: “CPS’s turnover rate is not because it’s ‘not for everyone’... it’s because once people go into full rotation, the workload is too high.”

3. Low-Complexity Cases Help Balance the Load. Respondents emphasized that low-complexity cases allow time to complete other work and decompress. One wrote, “Low-complexity cases are a bit of a mental break. You need them to stay afloat.”

4. Tradeoff Logic Mirrors Our Model. Many respondents articulated reasoning directly in line with the structure of our model. One wrote: “Two low complexity investigations normally equal one high complexity,” while another noted: “I’d trade three low complexity cases for one high one, maybe.”

E.4.3 Testing Core Modeling Assumptions

To test whether high-risk cases are perceived as more burdensome, we asked respondents to evaluate them relative to low-risk cases across three dimensions: time demand, emotional toll, and satisfaction. Each was rated on a 0–10 scale, with higher values indicating more time, more stress, or more satisfaction associated with high-risk cases.

The average scores were 8.8 for time, 8.6 for emotional toll, and 5.3 for satisfaction—suggesting that investigators experience high-risk cases as significantly more demanding, but only modestly more rewarding. These responses offer support for our empirical finding that a one standard deviation increase in caseload risk is associated with a significant increase in turnover risk. Investigators consistently linked complexity to emotional burden, and, as we discussed above, several linked high-complexity cases to secondary trauma and turnover in open-ended responses.

E.4.4 Eliciting the Marginal Rate of Substitution

To quantify tradeoffs between case types, we asked investigators to choose between hypothetical caseloads composed of different mixes of high- and low-complexity cases. Each respondent began by choosing between 12 high-complexity and 12 low-complexity cases. Based on their

answer, we offered a sequence of comparisons in which one type was incrementally added and the other removed, until a preference switch was observed. These responses allowed us to estimate each investigator’s MRS. Approximately 87% of respondents preferred 12 low-complexity cases to 12 high-complexity cases, implying an MRS of at least 1.

Figure A4 reports the CDF and density of the survey responses, along with the kernel density smoothed estimate. The average MRS in the sample was approximately 3.6—suggesting that, on average, investigators would be willing to trade one high-complexity case for 3.6 low-complexity ones. We also asked respondents to complete the same exercise from the perspective of a typical coworker. Results were similar, with a slightly higher average MRS of approximately 3.7. The figure also shows the regularized own-type empirical distribution. As discussed in the main text, these survey responses provide a natural estimate of F_j to use in our policy simulations. We regularize the own-type distribution by estimating hazards from the kernel-smoothed distribution and imposing monotone hazards. Specifically, we select the nondecreasing hazard sequence that minimizes the discrete sum of squared deviations in hazard rates.

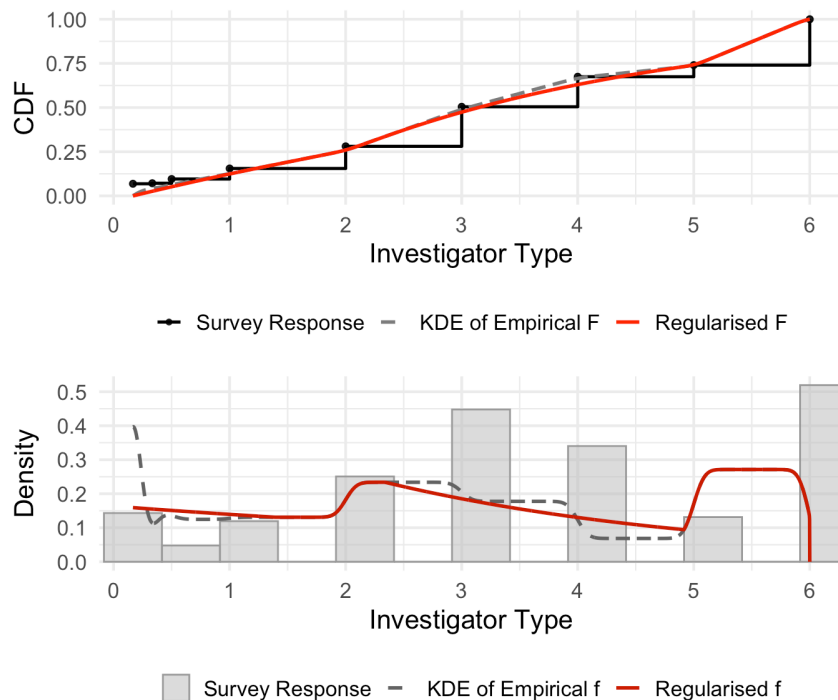


Figure A4: Investigator Own-Type Distribution

E.4.5 Field Validation of Complexity Definitions

The survey also allows us to assess whether there is variation in the way that investigators experience and value case difficulty across the binary classification that we use. Across both structured and open-ended responses, we observe clear separation in preferences between case types. Most respondents preferred low-complexity bundles when asked to choose between hypothetical caseloads. When choosing bundles that reveal an MRS of greater than one, investigators tend to cite emotional toll and resource demands—validating that our partition can broadly capture how complexity is perceived. Taken together with the observed heterogeneity in MRS estimates—many of which are relatively high—the results suggest that our binary-classification captures meaningful variation in preferences.

E.4.6 Summary

We believe that the survey broadly provides support for the assumptions and simplifications used in our model. The results also reinforce our empirical finding that complex caseloads are associated with increased turnover risk. Investigators clearly distinguish between high- and low-complexity cases, consistently describe high-complexity cases as more demanding, and exhibit substantial heterogeneity in preferences. Estimated MRS values vary meaningfully across the sample, and qualitative explanations confirm that investigators think about case assignments in terms of tradeoffs.

E.5 Generalizing beyond binary outcomes

As discussed in the main text, the assumption that Y_i^* is binary valued is innocuous. Suppose Y_i^* takes values in a finite set $\mathcal{X}$. Maintain the assumption that when case i is assigned to investigator j , Y_i^* is observed if and only if $D_{ij} = 0$. For $X \in \mathcal{X}$ define $PX_{ij} = Pr(\{Y_i^* = x, D_{ij} = 1\})$ and $NX_{ij} = Pr(\{Y_i^* = x, D_{ij} = 0\})$. The joint distribution of Y_i^* and D_{ij} is described by the vector $(PX_{ij}, NX_{ij})_{X \in \mathcal{X}}$. The cost of assigning i to j is, as in the binary case, a linear function of the joint distribution, denoted by $c(i, j)$. Lemma 1 generalizes immediately to this setting.

Lemma 6. Assume that the observed assignment is random conditional on I . Then for any $j, j' \in \mathcal{J}$ whose assignments are supported on I , the following are identified:

- the difference $PX_j^I - PX_{j'}^I$,
- the cost difference $\mathbb{E}[c(i, j) - c(i, j') | i \in I]$.

Proof.

$$PX_j^I - PX_{j'}^I = S^I(X) - NX_j^I - (S^I(X) - NX_{j'}^I)$$

$$= - (NX_j^I - NX_{j'}^I).$$

Given that we can identify the cost differences $\mathbb{E}[c(i, j) - c(i, j') | i \in I]$, the remainder of the mechanism-design analysis is unchanged. $\square$

E.6 Identification of false positive rates

Suppose we wish to identify $\mathbb{E}[FP_{ij}] = \mathbb{E}[D_{ij}] - \mathbb{E}[Y_i^*] + \mathbb{E}[FN_{ij}]$. In that expression, $\mathbb{E}[FN_{ij}] = \mathbb{E}[FN_i | Z_{ij} = 1]$ and $\mathbb{E}[D_{ij}] = \mathbb{E}[D_i | Z_{ij} = 1]$ are identified under random assignment by the observed false negative rate and placement rate of each investigator (where $Z_{ij} = 1$ if investigator j were assigned to case i). However, $\mathbb{E}[Y_i^*]$ is not identified as Y_i^* is not measured when $D_{ij} = 1$, or when the investigator places the child in foster care. Therefore, the identification challenge reduces to the challenge of identifying $\mathbb{E}[Y_i^*]$.

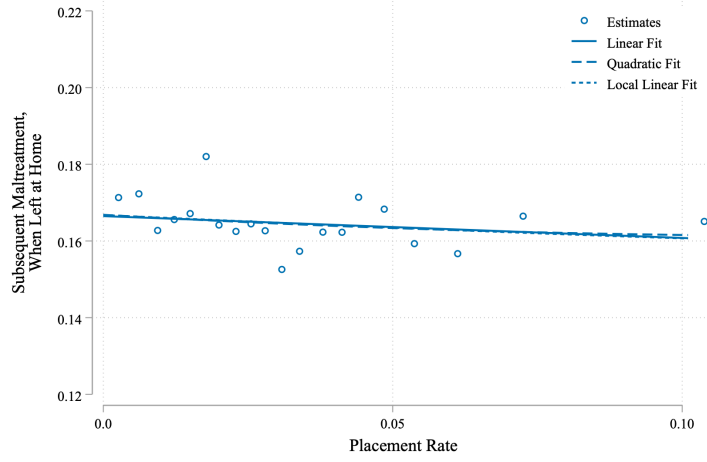
To identify this parameter, we follow [Arnold et al. \(2022\)](#) and use an extrapolation-based identification strategy. To build intuition, suppose there exists an “infinitely lenient” investigator j^* with a placement rate of zero and that cases are randomly assigned to investigators. Then, $\mathbb{E}[Y_i^*]$ of such an investigator would not suffer from selective observability concerns, since D_{ij^*} would equal zero for all i . Because cases are randomly assigned to investigators, the average subsequent maltreatment rate of cases assigned to this supremely lenient investigator would be close to the overall average: $\mathbb{E}[Y_i^* | D_{ij^*} = 0] \approx \mathbb{E}[Y_i^*]$.

Without a supremely lenient investigator, this parameter can be estimated via extrapolation. Estimates of $\mathbb{E}[Y_i^*]$ may come, for example, from the vertical intercept at zero of a linear, quadratic, or local linear regression of investigators’ subsequent maltreatment rates (among children left at home) on their placement rates. As [Arnold et al. \(2022\)](#) discuss, this approach is similar to extrapolations of average potential outcomes near a treatment cutoff in a regression discontinuity design. Here, we extrapolate across randomly assigned investigators with very low placement rates. This method is related to “identification at infinity” approaches in sample selection models ([Andrews and Schafgans, 1998](#); [Chamberlain, 1986](#); [Heckman, 1990](#)) and has been used to identify selectively observed parameters in several recent studies ([Arnold et al., 2021, 2022](#); [Angelova et al., 2023](#); [Baron et al., 2024](#)). In practice, this approach works well whenever there are many decision-makers with low treatment rates. Because foster care placement rates are low (3% in our sample), the CPS setting is particularly well-suited to this approach, yielding limited extrapolation and precise estimates.

We use the strata-adjusted investigator-specific placement and subsequent maltreatment rates from Section IV to extrapolate toward the unselected first moment, $\mathbb{E}[Y_i^*]$. Figure A5 reports the investigator-specific estimates that are used for the extrapolation, with a binned scatter plot of estimates of each investigator’s placement and subsequent maltreatment rate

(net of zip code by year fixed effects). The large mass of investigators with placement rates near zero suggests the extrapolation may be reliable in this context. We show extrapolations from linear, quadratic, and local linear regressions of each investigator’s subsequent maltreatment rate among children left at home on their placement rate. The vertical intercept at zero is the estimate of the unselected first moment of subsequent maltreatment. The most flexible local linear extrapolation yields an estimate of 0.167 (SE=0.001). Figure A5 shows that alternative extrapolation specifications yield nearly identical point estimates.

Figure A5: Extrapolation Estimates of Average Subsequent Maltreatment Potential



Notes. This figure presents the results of the extrapolation strategy used to estimate $\mathbb{E}[Y_i^*]$. Binned scatter plot estimates of investigator-specific placement rates versus conditional subsequent maltreatment rates are displayed, with 20 bins. All estimates adjust for zipcode-by-year fixed effects, and are obtained from investigator-level regressions that inversely weight observations by variance of estimated subsequent maltreatment rate among children not placed in foster care. The local linear regression uses a Gaussian kernel with a rule-of-thumb bandwidth.

E.7 Correlation between Preferences and Performance

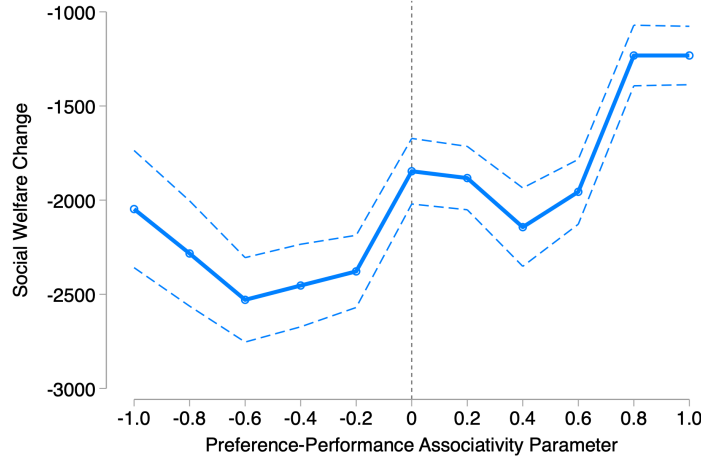
Finally, we consider how social welfare gains change if investigator preferences are correlated with their comparative advantage score, d_j . Let $\underline{d}_j, \bar{d}_j$ be the minimum and maximum d_j within counties, respectively. Assume that investigator types are drawn from a uniform distribution with full support on $[g(d_j) + 1, g(d_j) + 2]$, where $g(d_j) = b \frac{d_j - \underline{d}_j}{\bar{d}_j - \underline{d}_j}$ for $b \geq 0$ and $g(d_j) = -b(1 - \frac{d_j - \underline{d}_j}{\bar{d}_j - \underline{d}_j})$ for $b < 0$.⁴³ Then the associativity parameter b captures the strength and direction of the correlation between comparative advantage and preferences, where $b > 0$ indicates that investigators who are relatively good at high-risk cases tend to

⁴³Note that if $b = 0$, this reduces to the Unif[1, 2] setting. This construction of $g(\cdot)$ is also symmetric: for any $b \geq 0$, the investigator with maximal d_j in the county has the same distribution under associativity parameter b as the investigator with minimal d_j under parameter $-b$.

find such cases more costly.

For computational purposes, we compare the welfare gains for different values of b in the LMS-TP mechanism, of which SMD-TP is an approximation. Figure A6 reports the results of this exercise. When investigators with high d_j tend to have lower type draws, we find that welfare gains are significantly larger: for $b = -1$, the welfare gains are 2,047 relative to the expected social cost of the status quo. Compared to the $b = 0$ case where there is no association between preferences and performance, this is a 11% increase in the welfare gains. This confirms the intuition that when investigators with a comparative advantage in high-risk cases also relatively prefer these cases, the mechanism can achieve larger welfare gains. On the other hand, if investigators with a comparative advantage in high-risk cases tend to dislike such cases, the welfare gains are attenuated. The largest reduction in Figure A6 occurs when $b = 1$, in which case welfare gains are 1,232, a 33% decline compared to the $b = 0$ case. Thus, while a strong positive correlation between d_j and p_j may reduce the potential welfare gains, there still exists a significant potential for welfare improvement even under this scenario.⁴⁴

Figure A6: LMS Welfare Changes, Correlation Between Preference and Performance



Notes. This figure presents the welfare gains from the LMS-TP mechanism under distributional assumptions that allow for correlation between investigator type distributions, $F(p)$, and their comparative advantage score, d_j . Investigator types are drawn from a uniform distribution $[g(d_j) + 1, g(d_j) + 2]$, where $g(d_j) = b \frac{d_j - \underline{d_j}}{d_j - \underline{d_j}}$ for $b \geq 0$ and $g(d_j) = -b(1 - \frac{d_j - \underline{d_j}}{d_j - \underline{d_j}})$ for $b < 0$. 95% confidence intervals are reported.

⁴⁴Strong positive correlation appears unlikely in the current context: We find no evidence that investigators with an above-median comparative advantage in high-risk cases are differentially likely to quit when their caseload includes an above-median share of high-risk cases.

F Description of the CPS and Foster Care Systems

This section describes the CPS and foster care systems in Michigan, which work similarly to other states. The process begins when someone calls the state’s child abuse hotline to report an allegation of child abuse (e.g., bruises or burns) or neglect (e.g., inadequate supervision due to substance abuse). While anyone can call the hotline, the most frequent reporters are those legally required to do so, such as educational personnel ([Benson et al., 2022](#)). There are two central hotline call centers in Michigan, one in Detroit and one in Grand Rapids, but they share the same hotline number. When a new call comes in, it is quasi-randomly routed to the screener who has been on queue the longest, with no exceptions. Screeners have discretion on whether to “screen-in” the call: about 60% of all initial calls are screened-in, which launches a formal CPS investigation. A screened-out call concludes CPS involvement.

Once a call is screened-in, the screener transfers all relevant paperwork to the alleged victim’s local child welfare office, including the alleged maltreatment type (e.g., physical abuse versus physical neglect), and basic demographics of the child such as age, gender, and race. Each county in Michigan has its own office and some larger counties can have multiple offices. When the local office receives the report, it assigns the case to a CPS investigator based on a rotational system rather than their particular skill set or characteristics. There are two exceptions to rotational assignment, both of which we exclude from the analysis. First, given their sensitivity, reports of sexual abuse tend to be assigned to more experienced investigators. Second, new reports involving a child for whom there was a very recent prior investigation are usually assigned to the original investigator given the investigator’s familiarity with the case. Accordingly, we exclude cases involving sexual abuse and those involving children who had been the subject of an investigation in the year before the report.

The investigator has 24 hours to begin an investigation, 72 hours to establish face-to-face contact with the alleged child victim, and 30 days to complete the investigation. The investigator makes two sequential decisions that determine the outcome of the investigation. First, the investigator interviews the people involved, reviews any relevant police or medical reports, and decides whether there is enough evidence to “substantiate” the allegation. In Michigan, 74 percent of investigations were unsubstantiated during our sample period. In these cases, CPS concludes the investigation and there is no further contact with the family.

Conditional on a substantiated investigation, the investigator must also decide whether to temporarily place the child in foster care. Under CPS investigator guidelines in Michigan, the primary justification for foster care placement is a potential for subsequent maltreatment in the home: Investigators are instructed to recommend placement if the child is in imminent

danger of maltreatment in the home, but to keep the child with their family otherwise.⁴⁵ While there is technically a standardized 22-question risk assessment that helps the investigator determine whether foster care placement is appropriate, in practice investigators have immense discretion over placement. Many of the questions in the assessment are inherently subjective and previous research suggests that investigators tend to manipulate responses in order to match their priors (Gillingham and Humphreys, 2010; Bosk, 2015).

If the investigator believes there is a potential for subsequent maltreatment in the home, they request to the office’s supervisor to file a petition with the local court to temporarily place the child in foster care. In practice, it is rare for either the supervisor or the judge asked to sign the petition to disagree with investigators’ recommendations. Regardless of the placement recommendation, investigators can also recommend prevention-focused services. These services range from referrals to food pantries or support groups to substance abuse or parenting classes. Nevertheless, families are usually not mandated by the courts to engage in these services. Previous research conducted in our setting has indicated that the preventive services’ impact on subsequent maltreatment within the home and other outcomes is generally small (Baron et al., 2024; Gross and Baron, 2022; Baron and Gross, 2022).

The foster care system in Michigan is similar to the rest of the country. Children are temporarily placed with either an unrelated foster family, relatives, or (in about 10% of cases) in a group home or residential setting. During our sample period, children spend roughly 17 months in foster care on average; most children are reunified with their birth parents once the court decides that the parents have made the necessary changes in their lives to get their children back.

References

- Andrews, Donald and Marcia M. A. Schafgans (1998) “Semiparametric Estimation of the Intercept of a Sample Selection Model,” *Review of Economic Studies*, 65 (3), 497–517.
- Angelova, Victoria, Will Dobbie, and Crystal Yang (2023) “Algorithmic Recommendations and Human Discretion,” Working paper.
- Arnold, David, Will Dobbie, and Peter Hull (2021) “Measuring Racial Discrimination in Algorithms,” *AEA Papers and Proceedings*, 111, 49–54.
- Arnold, David, Will S Dobbie, and Peter Hull (2022) “Measuring Racial Discrimination in Bail Decisions,” *American Economic Review*, 112 (9), 2992–3038.

⁴⁵As an example, Michigan’s Department of Health and Human Services’ *Children’s Protective Services Policy Manuals* reads: “placement of children out of their homes should occur only if their well-being cannot be safeguarded with their families” (p.3). It also directs investigators to recommend placement “in situations where the child is unsafe, or when there is resistance to, or failure to benefit from, CPS intervention and that resistance/failure is causing an imminent risk of harm to the child” (p.5).

- Baron, E. Jason, Joseph J. Doyle, Natalia Emanuel, Peter Hull, and Joseph Ryan (2024) “Discrimination in Multi-Phase Systems: Evidence from Child Protection,” *The Quarterly Journal of Economics*, Forthcoming.
- Baron, E Jason and Max Gross (2022) “Is There a Foster Care-To-Prison Pipeline? Evidence from Quasi-Randomly Assigned Investigators,” NBER Working Paper 29922.
- Benson, Cassandra, Maria D Fitzpatrick, and Samuel Bondurant (2022) “Beyond Reading, Writing, and Arithmetic: The Role of Teachers and Schools in Reporting Child Maltreatment,” *Journal of Human Resources*, forthcoming.
- Bolton, Patrick and Christopher Harris (1999) “Strategic experimentation,” *Econometrica*, 67 (2), 349–374.
- Border, Kim C (2013) “Introduction to correspondences.”
- Bosk, Emily Adlin (2015) *All Unhappy Families: Standardization and Child Welfare Decision-Making* Ph.D. dissertation, University of Michigan.
- Chamberlain, Gary (1986) “Asymptotic efficiency in semi-parametric models with censoring,” *Journal of Econometrics*, 32 (2), 189–218.
- Chan, David C, Matthew Gentzkow, and Chuan Yu (2022) “Selection with Variation in Diagnostic Skill: Evidence from Radiologists,” *The Quarterly Journal of Economics*, 137 (2), 729–783.
- Gillingham, Philip and Cathy Humphreys (2010) “Child Protection Practitioners and Decision-Making Tools: Observations and Reflections from the Front Line,” *British Journal of Social Work*, 40 (8), 2598–2616.
- Gross, Max and E. Jason Baron (2022) “Temporary Stays and Persistent Gains: The Causal Effects of Foster Care,” *American Economic Journal: Applied Economics*, 14 (2), 170–99.
- Heckman, James (1990) “Varieties of Selection Bias,” *American Economic Review*, 80 (2), 313–18.
- Kasy, Maximilian and Alexander Teytelboym (2023) “Matching with semi-bandits,” *The Econometrics Journal*, 26 (1), 45–66.
- Kleinberg, Jon, Himabindu Lakkaraju, Jure Leskovec, Jens Ludwig, and Sendhil Mullainathan (2018) “Human Decisions and Machine Predictions,” *Quarterly Journal of Economics*, 133 (1), 237–293.
- Lechicki, Alojzy and Andrzej Spakowski (1985) “A note on intersection of lower semicontinuous multifunctions,” *Proceedings of the American Mathematical Society*, 95 (1), 119–122.
- Putnam-Hornstein, Emily, John Prindle, and Ivy Hammond (2021) “Engaging Families in Voluntary Prevention Services to Reduce Future Child Abuse and Neglect: A Randomized Controlled Trial,” *Prevention Science*, 22 (7), 856–865.
- Weitzman, Martin (1978) *Optimal search for the best alternative*, 78: Department of Energy.