

DEFERRED ACCEPTANCE WITH NEWS UTILITY

Bnaya Dreyfuss

Ofer Glicksohn

Ori Heffetz

Assaf Romm

WORKING PAPER 30635

NBER WORKING PAPER SERIES

DEFERRED ACCEPTANCE WITH NEWS UTILITY

Bnaya Dreyfuss
Ofer Glicksohn
Ori Heffetz
Assaf Romm

Working Paper 30635
<http://www.nber.org/papers/w30635>

NATIONAL BUREAU OF ECONOMIC RESEARCH
1050 Massachusetts Avenue
Cambridge, MA 02138
November 2022

We thank Ophir Averbuch, Yannai Gonczarowski, Shengwu Li, Vincent Meisner, Matthew Rabin, Alex Rees-Jones, Alvin Roth, Ran Shorrer, Davide Tagliatela, Jonas von Wangenheim, JESC-lab group and seminar participants at Cornell, Harvard, Hebrew University, LSE, Princeton, UPenn, INFORMS 2022, and MATCH-UP 2022 for invaluable comments; Tal Asif, Ronny Gelman, Ilana Lavie, Lev Maresca, Yonatan Rahimi, and Shira Zadik for excellent research assistance, and especially Guy Lichtinger for thoughtful and creative research assistance in the development of the project. This research was supported by the Israel Science Foundation (grant No. 2968/21), the Maurice Falk Institute, and the Johnson School at Cornell. This research was supported by the Israel Science Foundation (grant No. 2968/21), the Maurice Falk Institute, and the Johnson School at Cornell. 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 They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2022 by Bnaya Dreyfuss, Ofer Glicksohn, Ori Heffetz, and Assaf Romm. 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.

Deferred Acceptance with News Utility
Bnaya Dreyfuss, Ofer Glicksohn, Ori Heffetz, and Assaf Romm
NBER Working Paper No. 30635
November 2022
JEL No. D47,D82,D84,D91

ABSTRACT

Can incorporating expectations-based-reference-dependence (EBRD) considerations reduce seemingly dominated choices in the Deferred Acceptance (DA) mechanism? We run two experiments (total $N = 500$) where participants are randomly assigned into one of four DA variants—{static, dynamic} $\times$ {student proposing, student receiving}—and play ten simulated large-market school-assignment problems. While a standard, reference-independent model predicts the same straightforward behavior across all problems and variants, a news-utility EBRD model predicts stark differences across variants and problems. As the EBRD model predicts, we find that (i) across variants, dynamic student receiving leads to significantly fewer deviations from straightforward behavior, (ii) across problems, deviations increase with competitiveness, and (iii) within specific problems, the specific deviations predicted by the EBRD model are indeed those commonly observed in the data.

Bnaya Dreyfuss
Department of Economics, Littauer Center
Harvard University
Cambridge, Mass 02138
United States
bdreyfuss@g.harvard.edu

Ofer Glicksohn
Mt. Scopus
Department of Economics
Jerusalem
Israel
ofer.glicksohn@mail.huji.ac.il

Ori Heffetz
S.C. Johnson Graduate School of Management
Cornell University
324 Sage Hall
Ithaca, NY 14853
and The Hebrew University of Jerusalem
and also NBER
oh33@cornell.edu

Assaf Romm
Hebrew University of Jerusalem
assaf.romm@mail.huji.ac.il

A growing body of empirical evidence documents puzzling behavior in strategyproof matching mechanisms: many participants play seemingly dominated strategies, in both real, large-stakes applications and controlled lab experiments. A prominent example of this behavior is documented in centralized clearinghouses that use rank-order lists (ROLs) submitted by schools and students to match them together based on the deferred-acceptance algorithm (DA; [Gale and Shapley 1962](#)). Even though the mechanism is strategyproof for students, i.e., submitting a straightforward (“truthful”) ranking of schools is a weakly dominant strategy, a nontrivial share of supposedly informed students who participate in real-world matches appear to make dominated choices. Recent field evidence from various countries shows that students do not rank a study program with funding above the *same* program without the funding; see [Artemov et al. \(2017\)](#), [Hassidim et al. \(2021\)](#) and [Shorrer and Sóvágó \(2022\)](#). Similarly, in simple allocation experiments, a substantial fraction of participants rank smaller amounts of money above larger ones; see [Rees-Jones and Skowronek \(2018\)](#) for the first lab-in-the-field evidence from the US, and [Hakimov and Kübler \(2021\)](#) for a recent survey of lab evidence.

Recently, [Dreyfuss et al. \(2022\)](#) and [Meisner and von Wangenheim \(2021\)](#) showed that expectations-based reference dependence (EBRD; also referred to as EBLA, for expectations-based loss aversion) could potentially explain such non-straightforward behavior. Intuitively, participants with EBRD preferences may intentionally downrank (or, in some cases, completely omit) a high-value, low-probability school to avert the likely disappointment from rejection. Both papers show how an EBRD model, as formulated by [Kőszegi and Rabin \(2006, 2007, 2009\)](#), is consistent with various empirical patterns from the lab and the field. The first paper also structurally analyzes existing lab data from [Li \(2017\)](#) and shows that the EBRD model indeed fits the data significantly better than the classical, reference-independent-preferences benchmark.

In this paper we ask: is there a DA implementation that induces straightforward behavior even when students are loss averse? This question has both theoretical and practical, real-world implications.

To answer this question, we study, both theoretically and empirically, the potential role of EBRD in four different variants of DA. As summarized in Table 1, we analyze the predicted effect of varying two features of DA: (i) *static* (normal-form) versus *dynamic* (extensive-form) implementation, and (ii) *student-proposing* versus *student-receiving* (equivalently, school-proposing) role designation. In the industry-standard DA variant, Static student Proposing (SP), students propose to schools in an order determined by rank-order

Table 1: Four Deferred Acceptance Variants

	Static: Submit list in advance	Dynamic: Decide at each step
Proposing: Students apply to schools, schools retain the highest ranked applications.	SP	DP
Receiving: Schools send admission offers to students by ranking, students can retain at most one offer in each step.	SR	DR

lists (ROLs) they submit in advance. In the dynamic version of this variant, Dynamic student Proposing (DP), students actively propose to schools, in real-time, in any order they choose. In the Static student Receiving (SR) variant, students respond to school admission offers according to ROLs they submit in advance. Finally, in the Dynamic student Receiving (DR) variant, students actively respond to admission offers from schools, in real-time, by retaining at most one offer at any given moment.

We make four contributions: theoretical, empirical, methodological, and practical. First, we theoretically show that in large markets (with a continuum of students and a finite number of schools), the DR variant is “EBRD-strategyproof” in that it eliminates non-straightforward behavior for any level of loss aversion.¹ In contrast, EBRD-driven non-straightforward behavior is predicted to be quite prevalent under the other three variants, especially in highly competitive settings where the chances of admission into top schools are low.

Second, in what we view as our main contribution, we experimentally test the model’s predictions, both (a) across the four algorithm variants, by randomly assigning different participants to different variants, and (b) across matching settings, by exogenously varying the degree of competitiveness across ten matching problems that each participant faces.² In addition to providing strong empirical support for our DA-specific theoretical results above, our experiment is also unique in providing one of the sharpest tests to date of the EBRD model more generally. In particular, by holding everything fixed other than the

¹As discussed below, in concurrent work, [Meisner and von Wangenheim \(2021\)](#) prove a related result.

²In related work, [Klijn et al. \(2019\)](#) run a similar four-treatment experiment with the four DA variants. In contrast to our large-market framework, their experiment has four student subjects with full information about others’ preferences, and four (non-strategic) schools. This design allows for neither comparisons across role designation (proposing/receiving) nor testing our EBRD model predictions, for two main reasons. First, in small markets, switching roles drastically changes the (classical-preferences) incentive structure. Second, under full information, the uncertainty—the essential part of the EBRD theoretical framework—is about other subjects’ strategies, and is hence not easily estimated or modeled. Other related work with similar caveats includes [Echenique et al. \(2016\)](#), who run full-information, dynamic one-to-one DA. See [Hakimov and Kübler’s \(2021\)](#) review for further details.

timing of commitments and uncertainty resolution, (a) above provides a clean test of the model’s timing predictions; and by holding everything fixed other than the competitiveness of the setting, (b) above directly tests a central prediction of the model regarding how changes in probability beliefs affect behavior.

Third, our experiment uses a novel design that simulates DA in a continuum economy. Each subject draws a uniformly distributed *priority score* in every school, and can get matched with the school only if their priority score is greater than the school’s *threshold*. The subject learns each school’s threshold but not their own priority scores at the different schools.³ This design has two main advantages: first, our single-player framework allows us to directly (and exogenously) vary subjects’ beliefs about admission probabilities—a central component of the EBRD model—with minimal assumptions on probabilistic reasoning and no assumptions on beliefs about others’ play. Second, this design allows us to compare behavior across the four DA variants while keeping incentives constant. Methodologically, this design can be easily adapted to test other belief-based theories.

Finally, for practitioners, who may care less about theoretical results and their empirical tests and more about pragmatics, our finding that of the four mechanism variants we examine, DR in fact stands out in minimizing (though far from completely eliminating) non-straightforward behavior has practical implications. As discussed below, we view this result as a valuable, “bottom-line” policy-relevant empirical finding regardless of its theoretical drivers.

We report our theoretical analysis in Section 1. We focus on large-market economies, with a continuum of students and a finite set of capacity-constrained schools, using the characterization of [Azevedo and Leshno \(2016\)](#) and [Abdulkadiroğlu et al. \(2015\)](#). We show that in such economies, DR is EBRD-strategyproof: for any degree of loss aversion and any belief distribution, the EBRD model predicts that participants will always play what we call *straightforward strategies* (“truthful strategies”). Straightforward strategies are strategies consistent with a ranking of alternatives by consumption value. In DR, straightforwardness implies that the student always picks the highest-consumption-value offer from the set of available offers in each step. In contrast, the other three variants are not EBRD-strategyproof: in highly competitive settings, the model predicts loss-averse individuals under SP, SR, and DP to behave non-straightforwardly.⁴

³[Rees-Jones et al. \(2020\)](#) use a similar design in a static, non-strategyproof setting.

⁴We prefer the more normatively neutral terms “straightforward” ([Roth and Sotomayor, 1990](#)) and “non-straightforward” rather than “truthful” and “non-truthful” because it is unclear how to define “truthful reporting” for an agent whose preferences over schools depend on an endogenously determined reference point.

The theory in Section 1 thus leads to three sets of testable predictions under the EBRD model. First, *within* each variant, non-straightforward behavior should increase with the competitiveness of the matching setting in the three non-EBRD-strategyproof variants, but not in DR. Second, *between* the four variants, DR should stand out in having the lowest share of non-straightforward behavior in competitive environments, but not in noncompetitive ones (where shares should be similarly low in all variants). Third, non-straightforward behavior should consist of *specific strategies in specific environments* (precisely identified by the EBRD model).

These three sets of qualitative predictions, together with a given distribution of the degree of loss aversion in a population—which we plug in from previous work—yield specific quantitative predictions both across matching settings and across DA variants. Our experiment is designed to evaluate them.

We outline our experimental design in Section 2. Subjects are randomly assigned to one of the four variants, and play the same ten matching problems (but in different random orders). Each problem simulates a different large-market matching economy, with five schools whose values to participants are given in dollar amounts of takeaway money. Using the population distribution of loss aversion from a previous study, the ten matching problems are designed to create variation in the predicted prevalence, according to the EBRD model, of straightforward behavior in the three non-DR treatments. Three of the ten problems (“weak student” problems) are designed to simulate highly competitive settings, and predict a moderately high fraction of non-straightforward behavior (24–31 percent); another four (“medium student”) problems predict a lower fraction of such behavior (12–24 percent); and the remaining three (“strong student”) problems predict no non-straightforward behavior (0 percent). (Recall that in comparison, the standard, no-EBRD model predicts 0 percent of non-straightforward behavior under all four treatments and all ten matching problems.)

In Section 3 we report experimental results from two independent samples, pooled together ($N = 500$): Cornell students ($N = 196$) and Israeli psychology graduate-school candidates ($N = 304$) who participated in the (DA-based) Israeli Psychology Master’s Match (IPMM) and are therefore a highly relevant population with strong incentives to understand the mechanism. Our main predictions and results are summarized in Figure 3 (page 26).

Our findings support all three sets of theoretical predictions. First, within the three non-DR treatments, the observed shares of non-straightforward behavior monotonically

decrease from weak- (33–48 percent) to medium- (30–37 percent) to strong-student problems (19–22 percent). In contrast, within the (EBRD-strategyproof) DR treatment, the shares are essentially flat (19 to 16 to 16 percent). Second, across treatments, these shares also imply that we find substantially lower levels of observed non-straightforward behavior in DR compared to the other three variants in competitive environments, and essentially no difference in noncompetitive environments. Third, moving beyond mere shares of non-straightforward behavior, we further show that (i) the specific non-straightforward strategies that the model predicts and does not predict, are indeed those we commonly observe and do not observe in the data, respectively; and (ii) in specific *problems* in which the model predicts specific non-straightforward strategies (e.g., ranking the third-highest-value school on the top of the list), these strategies are indeed the most common.

A potential worry when comparing behavior across static and dynamic variants (in our case, the static variants vs. DR) is that while in static variants we observe submitted rank-order lists—i.e., complete strategies—in dynamic variants we only observe *actions*. In theory, this could mechanically downward-bias observed non-straightforward behavior in DR. In practice, however, our evidence suggests that such potential bias is unlikely. Indeed, we find that in DR, (a) observed non-straightforward behavior primarily consists of rejecting offers from low-value schools—offers that we typically observe; and (b) observed non-straightforward behavior is no more prevalent in rounds where subjects receive more offers.

We close Section 3 by discussing other theories that might explain our data. We make two observations. First, we find across all problems and variants a baseline share of 16–22 percent observed non-straightforward behavior that the model does not predict and cannot explain, and that is likely explained by strategic confusion, misunderstanding, noisy decision-making, or other theories. Second, no model of misunderstanding or noise that we are aware of predicts or explains the three sets of qualitative predictions that our model predicts and that are borne out in the data: *variation* within and across variants, and prevalence of *specific* types of non-straightforward behavior.

We conclude in Section 4, where we discuss the potential implications, as well as limitations of our findings. From a theoretical point of view, our findings are unambiguously more consistent with the EBRD than with a no-EBRD model—importantly, without adding degrees of freedom in the EBRD model (as we fixed parameter values in advance at a previously estimated distribution). At the same time, we always find 10–20 percent more non-straightforward behavior than the (parameter-constrained) EBRD model predicts—

strongly suggesting that loss aversion cannot account for all such observed behavior. From a policy-maker’s point of view, our findings suggest that one of the four implementations of DA that we examine—dynamic student receiving (DR)—minimizes non-straightforward behavior. Of course, our study leaves open the pragmatic question of whether DR is feasible to implement in practice (however, see [Grenet et al. 2022](#) for a recent real-world example of quasi-dynamic implementation of student-receiving DA).

1 Theoretical Analysis

This section is notation-heavy and technical. Readers with powerful intuition may wish to skip it, moving directly to the mostly self-contained Section 2 (p. 16).

1.1 General Framework

Consider a two-sided market with n capacity-constrained schools and a unit mass continuum of atomistic students, governed by a DA mechanism that matches (masses of) students to schools. Throughout this section, we rely on the large-market characterization of [Abdulkadiroğlu et al. \(2015\)](#) and [Azevedo and Leshno \(2016\)](#) and their adaptation of DA to a continuum economy.⁵

In what follows we focus on the decision from student i ’s perspective, fixing others’ (i.e., schools’ and other students’) behavior. We start by describing the fundamentals of the economy underlying that perspective. The set of schools is denoted by $S = \{s_1, \dots, s_{n+1}\}$, where s_{n+1} represents the outside option. Each school j has a capacity $c_j \in (0, 1]$ that represents the share of students it can admit and assigns each student i with a relative rank $\rho_{ij} \in [0, 1]$, which we call a *priority score*.

Each student i has a strict (reference-independent) preference over schools, represented by the vector of utilities $\mathbf{m}_i = (m_{i,1}, \dots, m_{i,n+1})$. We index WLOG schools by student i ’s preferences so that $m_{i,j} > m_{i,k}$ for all $j < k \leq n$. We call this utility component *consumption utility*. Later in this section we add a *news-utility* ([Kőszegi and Rabin, 2009](#)) component.

As shown in [Azevedo and Leshno \(2016\)](#), the match resulting from running DA in this economy (given some students’ and schools’ preferences) is characterized by a vector of thresholds $\mathbf{T} = (T_1, \dots, T_{n+1})$, one for each school, where each student i is assigned to

⁵While in the continuum economy DA converges to a well-defined allocation in the limit, the process might not complete in finite time. We ignore this issue and assume that a matching is always reached in finite time.

her highest ranked school such that $\rho_{ij} \geq T_j$. Moreover, student-proposing and student-receiving DA result in the same match (and associated thresholds vector $\mathbf{T}$).⁶ We assume that the outside option is always available with certainty, i.e., $c_{n+1} = 1$, which then implies $T_{n+1} = 0$.

Students know the thresholds $\mathbf{T}$ and—since students are measure zero, correctly—treat them as exogenous. However, we assume that students do not know their own priority score at each school, ρ_{ij} , but rather perceive a probability distribution over their score, which we denote by the CDF G_{ij} .⁷ This assumption captures the idea that while students know how selective a school is, they do not necessarily know their exact ranking relative to other students.⁸ We denote student i 's joint distribution of priority scores in all schools by G_i . Since we focus on the single-agent problem, we often suppress the index i . To simplify the presentation, we restrict our attention to schools that are acceptable to i , and normalize $m_{i,n+1} = 0$, but none of our results meaningfully depend on this restriction.

As explained in the introduction, we focus on four different *variants* of the DA algorithm, resulting from the product $\{\text{proposing, receiving}\} \times \{\text{static, dynamic}\}$. We denote the DA variant governing the market by M . A *matching problem* from student i 's perspective can therefore be summarized by $\langle M, G_i, \mathbf{T}, \mathbf{m}_i \rangle$, where M defines the matching process, G_i is the joint distribution over priority scores, $\mathbf{T}$ is a vector of thresholds, and $\mathbf{m}_i$ is a vector of school-consumption utilities.

1.2 Timing, Beliefs, and Strategies

Fixing all other players' actions, we can treat each DA variant M as a game student i plays against Nature, where G governs Nature's moves. We define periods for each variant below, but generally, we think of periods as decision nodes where either the student or Nature makes a move.

A history $h \in \mathcal{H}$ is a sequence of actions by the student and Nature prescribed to every period t it includes. We define $G(\cdot | h)$ as the student's joint belief over priority scores given history h . A terminal history (in the set of terminal histories) $z \in \mathcal{Z}$ is a history

⁶For the stable match to be unique, the joint distribution of preferences and priority scores needs to satisfy some regularity conditions (see [Azevedo and Leshno 2016](#), Theorem 1). For simplicity we restrict student i 's beliefs to satisfy those regularity conditions throughout. Notice, however, that we do not assume that others' behavior is straightforward.

⁷We stress that in our model the thresholds can, but do not necessarily have to, arise from explicit beliefs on the joint distribution of other applicants' play and priority scores, and on school capacities.

⁸Notice that this framework can easily capture uncertainty about the thresholds by simply incorporating it into G_{ij} .

that includes a *terminal period*, i.e., the period in which no more actions are taken, and the student is informed about her final match from S . The index of that final match is represented by the outcome function $O(z) : \mathcal{Z} \rightarrow \{1, \dots, n+1\}$, and the terminal period is denoted by $\bar{t}_z$. Subhistories are denoted using subscripts: z_t (with $t \leq \bar{t}_z$) is the realization of the terminal history z up to period t .

A strategy $l \in \mathcal{L}$ for the student assigns an action to every period in which the student is called to act, for every possible history.⁹ We denote the set of continuation strategies consistent with (i.e., do not preclude reaching) history h by $\mathcal{L}_h$. Given a variant M and thresholds $\mathbf{T}$, the joint distribution of priority scores G combined with a strategy l jointly induce a distribution of terminal histories, which we denote by $\tilde{F}_l$, and also a distribution of final payoffs (i.e., over the support $\mathbf{m}_i$) which we denote by F_l . Similarly, given a history h , the conditional distribution $G(\cdot | h)$ and the strategy l jointly induce conditional distributions of terminal histories and final payoffs, denoted by $\tilde{F}_{l|h}$, and $F_{l|h}$, respectively.

In Appendix A we formally present the game our large-market framework induces under each of the four variants. In the static variants (SP and SR), students submit a rank-order list (ROL) in the first period and learn about the result in the next one; in DP, they apply to a school in one period and learn the result in the following period, until the first acceptance; and in DR they receive a set of offers in one period, and can keep one of them in the following period, until no more offers are received.

We now introduce our definition of non-straightforward behavior (“misrepresentation” or “non-truthfulness”), which applies to all four variants.

Definition 1. *In a DA variant M , a strategy l that can be represented by a ROL ordering schools by their consumption-utility value is a **straightforward** (SF) strategy. A **non-straightforward** (NSF) strategy is any strategy inconsistent with this ROL.*

In the static variants SP and SR, the SF strategy is simply the ROL that ranks schools by their values. In DP, the SF strategy is sequentially applying to schools by order of their values. In DR, it is retaining the highest value offer at any given decision node for any given history.

⁹Throughout this section, we limit our attention to pure strategies. While we do not formally prove this, we strongly suspect that none of our results would be affected if we allowed for mixed strategies. For a related discussion on mixed strategies under EBRD see [Dato et al. \(2017\)](#).

1.3 Expectations-Based-Reference-Dependent (EBRD) Preferences

As mentioned earlier, in addition to classical consumption utility, students' utility functions have a news-utility component. Absent this additional news component, given any history h , the student's objective function is simply $\mathbb{E}_{F_{l|h}}[m]$, i.e., the expected consumption utility under strategy l given history h . The additional news-utility component is aimed to capture a basic feature in people's preferences: the utility and disutility from belief updating. In particular, learning about a higher likelihood of a good outcome—a pleasant surprise relative to previously held beliefs—in itself entails positive utility, while learning about a lower likelihood—a disappointment—in itself entails disutility. *Loss aversion* in this model means that the disutility from a downward update of one's beliefs is larger than the utility from an equally sized upward update.

News utility is therefore defined over belief updates. Formally, during the matching process, when moving from a history h_t to h_{t+1} , the student rationally updates her beliefs over final outcomes from $F_{l|h_t}$ to $F_{l|h_{t+1}}$. Let F and $\hat{F}$ be, respectively, previously held and updated belief distributions over outcomes. The news utility function is given by:

$$N(\hat{F} | F) = \int_0^1 \mu(\hat{F}^p - F^p) dp, \quad (1)$$

where F^p denotes the consumption level at percentile p of F , and $\mu(\cdot)$ is defined as:

$$\mu(x) = \begin{cases} x & \text{if } x \geq 0 \\ \lambda x & \text{if } x < 0, \end{cases} \quad (2)$$

with $\lambda \geq 1$ representing an individual's loss-aversion parameter. The function N describes how updated beliefs are compared with previously held beliefs, percentile by percentile: in each period t , the student compares, for each percentile, the outcome under $\hat{F}$ to the outcome under F , with a higher weight λ on negative surprises.

The total utility from a strategy l is therefore given by:

$$U(l) = \mathbb{E}_{F_l}[m] + \mathbb{E}_{\tilde{F}_l} \left[N(F_{l|z_1} | F_{l|\emptyset}) \right] + \mathbb{E}_{\tilde{F}_l} \left[\sum_{t=2}^{\bar{t}_z} N(F_{l|z_t} | F_{l|z_{t-1}}) \right]. \quad (3)$$

The first term is expected consumption utility. The last two terms are the expected sum of news utility streams from a terminal history, given a probability distribution over possible

terminal histories $\tilde{F}_l$.¹⁰ Note that beliefs over payoff prior to period 1, (i.e., *before choosing a strategy*), $F_{l\emptyset}$, are not yet defined. To close the model, we assume that prior to entering the mechanism, the student believes that she will consume the outside option with certainty, i.e., $F_{l\emptyset} = F_0$, where $F_0(x) = 1$ for all $x \geq m_{n+1} = 0$ and 0 otherwise. This assumption is calibrationally (i.e., quantitatively), but not qualitatively, substantive. In particular, it does not affect our between-variant results, nor does it change our qualitative predictions within each mechanism. However, it is calibrationally substantive in that assuming more optimistic initial beliefs will require a higher λ for predicting NSF.

Notice that this formulation assumes unidimensional utility—all schools belong to the same consumption dimension, capturing situations where schools are vertically differentiated (rather than horizontally). Our between-variant result (Proposition 1) does not hold if we relax it. Dreyfuss et al. (2022) show a within-variant result (similar to Proposition 2) under complete horizontal differentiation.

To summarize, given a realization of a terminal history z and a chosen strategy l , timing in our model is defined as follows:

- **Period 0:** The student believes that she will consume the outside option with certainty. No action is taken.
- **Period 1:** The student learns about the mechanism and chooses a strategy l . If she is called to act, she takes the action prescribed to z_1 by l .¹¹ She updates her beliefs from F_0 to $F_{l|z_1}$ and receives $N(F_{l|z_1} | F_0)$.
- **Period $1 < t < \bar{t}_z$:** If called to act, the student takes the action prescribed to z_t by l . She updates from $F_{l|z_{t-1}}$ to $F_{l|z_t}$ and receives $N(F_{l|z_t} | F_{l|z_{t-1}})$.¹²
- **Period $\bar{t}_z$:** The student learns about her final match and updates to the degenerate CDF $F_{l|z}$. She receives $m_{O(z)} + N(F_{l|z} | F_{l|z_{\bar{t}_z-1}})$.¹³

¹⁰ Kőszegi and Rabin (2009) has two additional parameters: η , the weight on news relative to consumption utility, and γ , the discount on news about future consumption. We assume that $\eta = 1$ as well as $\gamma = 1$. The first assumption is essentially a normalization. The second assumption is more substantive and implies that the weight on news utility does not depend on when in the future said consumption occurs. See a related discussion in Dreyfuss et al. (2022).

¹¹In DR, Nature is the first to take an action, whereas in the other three variants, the student takes the first action.

¹²However, note that in equilibrium, new information is only learned after Nature's actions, in all variants. Therefore, action taking and non-degenerate belief updating do not occur in the same period.

¹³In all four variants, Nature is always the last to take an action: in the static variants, it responds to the submitted ROL with a final match; in DP, it accepts an application; and in DR it stops sending offers.

In both static variants, news utility reduces to $N(F_l | F_0) + \mathbb{E}_{\tilde{F}_l} [N(F_{l|z} | F_l)]$. The first term compares the belief over final matches induced by the submission of the ROL l in period 1 to an initial belief F_0 (consuming the outside option with certainty). The second term compares the result of the matching process with the belief over outcomes induced by the submitted ROL l (and recall that $F_{l|z}$ is degenerate). In the two dynamic variants, after the initial period-1 update, beliefs potentially update after each time Nature takes an action (or when the student changes her strategy, which does not happen in equilibrium). For example, in DP, the student updates after each rejection or acceptance, and in DR, she potentially updates after receiving each new set of offers.

1.4 Optimal Strategies

To derive predictions, we use Preferred Personal Equilibrium (PPE; [Kőszegi and Rabin, 2009](#)) as our solution concept for all variants.¹⁴ First, given a strategy l and a non-terminal, t -periods-long history h , utility from deviating to some strategy $l' \in \mathcal{L}_h$ is given by

$$U_h(l' | l) = \mathbb{E}_{F_{l'|h}}[m] + N(F_{l'|h} | F_{l|h}) + \mathbb{E}_{\tilde{F}_{l'|h}} \left[\sum_{s=t+1}^{\bar{t}_z} N(F_{l'|z_s} | F_{l'|z_{s-1}}) \right], \quad (4)$$

i.e., expected consumption utility under the deviating strategy plus news utility from deviating from l to l' (at time period t) plus the expected sum of streams of news utility given the distribution of terminal histories under l' .

We adapt the backward-recursive definition of [Kőszegi and Rabin \(2009\)](#) of the PPE consumption plan to our setup:

Definition 2. Let z be a terminal history. We define the set of z_t -credible strategies in the following backward-recursive way: the credible set $\mathcal{L}_{z_{t-1}}^*$ includes all strategies that maximize utility at the last action period given the expectation they induce, i.e., all strategies $l \in \mathcal{L}$ that satisfy $U_{z_{t-1}}(l | l) \geq U_{z_{t-1}}(l' | l)$ for all $l' \in \mathcal{L}$. Then the set $\mathcal{L}_{z_t}^*$ contains all strategies $l \in \cap_{h \in \mathcal{H}_{z_t}} \mathcal{L}_h^*$ satisfying $U_{z_t}(l | l) \geq U_{z_t}(l' | l)$ for all $l' \in \cap_{h \in \mathcal{H}_{z_t}} \mathcal{L}_h^*$, where $\mathcal{H}_{z_t}$ is the set of histories that strictly contain the subhistory z_t . I.e., all strategies that maximize utility in the current period given the expectation they induce out of the strategies that prescribe a credible

¹⁴More precisely, we use a slightly modified solution concept, from [Kőszegi and Rabin's \(2009\)](#) Online Appendix, called *optimal consistent plan* (OCP). The main difference between the two solution concepts is that while in PPE the DM holds correct beliefs about all future contingencies from the moment of birth (i.e., before the first period), in OCP the DM has some initial prior beliefs when she forms her plan (which we fixed to F_0 above).

continuation plan.

A z_t -credible strategy maximizes utility given the expectation induced by the strategy: a student that expects to play l must find it optimal to follow through with l at every decision node that originates in z_t , assuming any future self also plays optimally.¹⁵ The definition of a z_t -credible strategy allows us to define our solution concept:

Definition 3. Let $\mathcal{L}^* \equiv \mathcal{L}_0^*$, i.e., the set of strategies that are credible in the empty history. A PPE strategy l^* is a strategy that satisfies

$$l^* \in \arg \max_{l \in \mathcal{L}^*} U(l).$$

In PPE, the core difference between static and dynamic variants is commitment: the submission of a ROL in period 1 in a static variant can be seen as committing in advance to a strategy in its dynamic counterpart.¹⁶ In the dynamic variants, deviations are possible in each period, and therefore a strategy has to be optimal in all decision nodes (not just the first one). For example, in DP, on-path equivalent strategies can be described as a ROL. However, for a ROL to be a PPE, there can be no profitable deviations at any possible point where a decision can be made: conditional on the first application prescribed by l (and given the belief induced by the continuation of l), the student must find it optimal to follow through and send the second application prescribed by l , and so on.

1.5 Strategyproofness

In our large-market framework, all four variants are trivially strategyproof for classical-preferences students. The following definition extends strategyproofness to markets with agents with EBRD preferences.

Definition 4. A DA variant M is **EBRD-strategyproof** if for any degree of loss aversion λ , thresholds vector $\mathbf{T}$ and belief G , an optimal (PPE) strategy is SF.

It is worth emphasizing that EBRD-strategyproofness implies standard strategyproofness. Our analysis therefore imposes a large-market structure that completely shuts down the (standard) strategic channel that could potentially drive non-straightforward

¹⁵The forward-looking definition implies that when evaluating a strategy's credibility, only credible deviations are considered, i.e., potential deviations that prescribe non-credible continuation strategies are not considered.

¹⁶In these variants, PPE strategy coincides with the definition of *Choice-Acclimating Personal Equilibrium* (Kőszegi and Rabin, 2007).

behavior in student-receiving DA. EBRD-strategyproofness requires that, in addition to classical-preferences considerations, the mechanism must be such that under *any* belief about others’ behavior (which in our framework is summarized by T and G), SF behavior is optimal for applicant i for any λ .

Equipped with our new notion of strategyproofness, the following proposition differentiates between the model’s prediction in DR and the other three variants.

Proposition 1. *The Dynamic student-Receiving (DR) variant is EBRD-strategyproof, while the other three variants (SP, SR, DP) are not EBRD-strategyproof.*

The proof is relegated to Appendix B. The second part of the proposition has been discussed in Dreyfuss et al. (2022) and Meisner and von Wangenheim (2021) and is easily proved by a counterexample (see below). The first part of the proposition is proved by backward induction and has a simple intuition. The model predicts that, no matter what she had planned to do, when a student is presented with a choice set of alternatives that she can get with certainty, she will always choose the one with the highest consumption utility. In DR, the student chooses only between schools that have in fact sent her an offer—only between “sure things”—and hence cannot do better than keeping the highest value one. Applying this argument iteratively shows that in DR, the straightforward strategy is indeed the only optimal one in any possible history. This highlights the importance of dynamic implementation: in (the static) SR, the student can lower her expectations and reduce potential disappointment by committing (via ROL submission) to reject some offers. In contrast, in DR, the student anticipates that she will keep desirable offers and is therefore forced to choose an SF strategy.¹⁷

We note that in concurrent work, Meisner and von Wangenheim (2021) (Proposition 4) show a similar, albeit different result about DR, using Rosato’s (2014) model of dynamic EBRD preferences. Assuming a unique stable matching for any realization of preferences, they show the existence of a Bayesian Nash equilibrium in which all players play SF. The stable matching is always unique in large markets; see Karpov (2019) for necessary and sufficient conditions in non-large markets. We assume large markets and show EBRD-strategyproofness, i.e., for any profile of other players’ strategies, it is optimal for an arbitrarily loss averse student to play SF. We view the two results as complementary:

¹⁷This also illustrates a subtle point about welfare: while DR maximizes the students’ ex-ante *consumption* utility, it actually *decreases* their overall (consumption and news) utility. We still find DR appealing since we suspect that EBRD-driven behavior in these settings may be a mistake. For further discussion, see Dreyfuss et al. (2022).

assuming large markets implies (EBRD) strategyproofness, generalizing to all markets in which a unique stable matching is guaranteed implies existence of an SF equilibrium.

1.6 Varying Competitiveness

Proposition 1 makes a stark prediction about observed behavior in DR vs. the other variants. We now formalize the model’s second prediction, which relates observed behavior in the static variants to competitiveness.

In each non-DR variant, the threshold λ above which NSF behavior is optimal depends on the matching environment. In particular, since NSF behavior reflects an attempt to avoid disappointment from future rejections, this threshold λ increases as disappointment becomes less likely, all else equal. This makes an additional testable comparative static: NSF behavior increases with competitiveness.

Formally, denote the admission probability to school j by $q_j \equiv \Pr(\rho_j > T_j)$. Competitiveness is defined by the probability of acceptance at the highest value school, q_1 . The next proposition is an extension of Meisner and von Wangenheim (2021)’s Proposition 2. Before stating it, we impose the following assumption:

Assumption 1.

1. $\rho_j \perp \rho_k$ for all j, k .
2. $q_j < q_k$ for all $j < k$.

In words, we assume that (1) priority score is independent across schools, and (2) admission probability is decreasing with school value. While its part (1) in particular may not hold in important real-life deployments of DA, Assumption 1 yields sharper theoretical predictions that our experiment—which satisfied it by design—can cleanly test. In particular, it generates a tighter upper bound on the threshold λ , and it generates an upper bound in the case of $q_1 \geq 0.5$.¹⁸

Proposition 2 (Based on Meisner and von Wangenheim 2021).

1. If $\lambda < 1 + \frac{2}{1-q_1} \equiv \underline{\lambda}$, the SF ROL is strictly optimal in SP and SR.
2. Suppose that assumption 1 holds. Then:
If $\lambda > 1 + \frac{2}{(1-q_1)^2} \equiv \bar{\lambda}$, the SF ROL is strictly suboptimal in SP and SR.

¹⁸In the parameters of our model, Meisner and von Wangenheim’s (2021) upper bound $\bar{\lambda}$ translates to $1 + \frac{1}{1-2q_1}$, and does not exist for $q_1 \geq 0.5$.

In the following sections, we test specific predictions implied by this proposition (see, e.g., Figure 2 and the related explanation in the next section). We do so by experimentally varying q_1 . Specifically, subjects in our experiment face matching environments (i.e., rounds) with three levels of competitiveness: high, medium, and low, with $q_1 = 0.05$, $q_1 = 0.2$ – 0.3 , and $q_1 = 0.6$ – 0.65 , respectively. In Proposition 2 these q_1 values translate to the following non-overlapping bounds around λ : $\underline{\lambda} = 3.11$, $\bar{\lambda} = 3.22$ (high competitiveness); $\underline{\lambda} = 3.50$ – 3.85 , $\bar{\lambda} = 4.13$ – 5.08 (medium); and $\underline{\lambda} = 6.00$ – 6.71 , $\bar{\lambda} = 13.50$ – 17.32 (low competitiveness). As long as the population distribution of λ has a sufficiently large mass outside the bounded intervals (i.e., inside the ranges 3.22 – 3.5 and 5.08 – 6.00), the proposition implies that the prevalence of NSF behavior should increase with round competitiveness.¹⁹

2 Experimental Design and Predictions

2.1 The Experiment

The experiment consists of simulating a large matching market.²⁰ Each subject is randomly assigned to one of 2×2 ($\{\text{Static vs. Dynamic}\} \times \{\text{Proposing vs. Receiving}\}$) treatments denoted SP, SR, DP, and DR. This 2×2 design allows us to independently explore the difference both between proposing and receiving mechanisms and between static and dynamic implementations, by varying each of the two features while holding the other one fixed at both states. However, since the theory sets the EBRD-proof DR variant apart from the other three, it is assigned 40% of the subjects, while the other three are each assigned 20% of subjects.

Our goal was to design four treatments as similar to each other as reasonably possible in terms of instructions structure, length and language, and, more generally, all aspects of the user-interface look and feel. (Online Appendix A provides screenshots for all treatments, and uses a four-color coding scheme to indicate all cross-treatment differences.) In all treatments, following a tutorial and an attention check, each subject participates in ten order-randomized incentivized *matching problems*, each simulating a different large

¹⁹Specifically, these analytical bounds imply that the increase in NSF share between high vs. medium round competitiveness is determined by the mass inside the range 3.22 – 3.5 and the increase in NSF share between medium vs. low competitiveness is determined by the mass inside the range 5.08 – 6.00 . Notice that in the next section, we use specific matching problems and therefore get specific thresholds (rather than bounds).

²⁰The experiment was programmed using the oTree platform (Chen et al., 2016).

matching market. We then collect demographic variables and feedback. Each subject who successfully finishes the experiment receives a participation fee and the sum of matched-school values from every matching problem.

2.1.1 A Matching Problem

A matching problem consists of five schools; the subject can be matched with at most one at the end of the matching process. Each school is given a dollar value—the payment the subject will gain if matched with said school.

The experimental design closely follows the theory section: The subject is assigned a *priority score* at every school, which is a random, uniformly and independently distributed integer between 0 and 99. Each school has a *threshold*, the minimal accepted priority score. That is, only candidates whose priority score is above a school’s threshold can be accepted to that school.

At the beginning of each problem, the subject learns each school’s threshold but not her own priority scores at the different schools. This creates an *unconditional* probability of acceptance at each school ($= 1 - \frac{\text{threshold}}{100}$). Figure 1 reproduces the example matching problem given in the tutorial, the way it appears on subjects’ screen.²¹

Figure 1: Screenshot of an example matching problem

School	Threshold	Value	Your Priority Score	Chance that Your Priority Score $\geq$ Threshold
Pine Peak	60	\$0.75	?	40%
Birch Hill	0	\$0.25	?	100%
Hickory Bridge	50	\$0.50	?	50%
Maplecrest	80	\$1.25	?	20%
Elm South	70	\$1.00	?	30%

Notes: Example screenshot of the table describing a matching problem. The table’s structure is identical in all treatments and rounds. School names and parameters differ by round. See Table 2 for the parameters used in the ten paying rounds (the shown screenshot is taken from the tutorial round).

To make the process transparent and easier to monitor, at the end of every matching process subjects receive full information: they learn their match and priority score at each

²¹We randomized the order of columns in the interface. Half the subjects had, for the ten matching problems, “Threshold” to the left of “Value” (as in Figure 1). The other half had “Value” to the left of “Threshold.” As shown in Online Appendix B, the order of columns had no effect on behavior under SP, SR and DR. We do not have data on column order for DP due to a data-logging error.

school (the “?”s in Figure 1 change to the realized priority scores).

2.1.2 The Matching Process

The matching process varies by treatment, closely resembling the structure outlined in the theory section. In both *static* treatments (SP and SR), subjects submit a rank-order list (ROL) of schools in advance and are matched with the highest ranked school in which their priority score exceeds its threshold. In DP, subjects sequentially apply to schools, and get accepted to the first school in which their priority score exceeds its threshold. In DR, subjects sequentially receive offers from schools in which they exceed the threshold, with scores higher above the threshold resulting in earlier-arriving offers, and can keep at most one offer at any point.²²

2.2 Theoretical Predictions

2.2.1 Matching Problems

As explained in the theory section, a subject with loss aversion λ is faced with a matching problem $\langle M, G, \mathbf{T}, \mathbf{m} \rangle$, where M is the matching process, G is the joint distribution over priority scores (recall that priority scores in the experiment are uniform and independent), $\mathbf{T}$ is a vector of thresholds and $\mathbf{m}$ is a vector of school consumption utilities. We assume a linear consumption utility, i.e., if s_j is worth v dollars, then $m_j = v$. Last, schools are denoted by their indices (i.e., s_1 is denoted by 1), and ROLs are represented by sequences of numbers denoting the order in the list (i.e., the ROL ranking five schools in descending order of value is 12345).

When designing our experiment, our goal was to create variation in NSF predictions not only across treatments—DR vs. the other three variants, testing Proposition 1—but also across problems—more vs. less competitive, testing Proposition 2. To create variation in competitiveness, we searched for settings we could classify into three levels of predicted NSF: high, medium, and low. We now briefly describe our search process; for further details, see Online Appendix C.

We calculated optimal-strategy predictions for a range of λ s, for each of the non-DR

²²The stream of offers a subject i receives in DR contains at most five periods and is determined as follows. For each school s_j in which the subject exceeds the threshold, the segment $[\text{threshold}, 99]$ is divided into five quintiles, $Q_1, \dots, Q_5$ (where Q_1 is the bottom quintile). The subject then receives an offer from s_j in period $6 - Q_{ij}$, where Q_{ij} is the quintile in which her priority score lies. This induces a stream of offers from different sets of schools at different periods, where periods with no offers are eliminated.

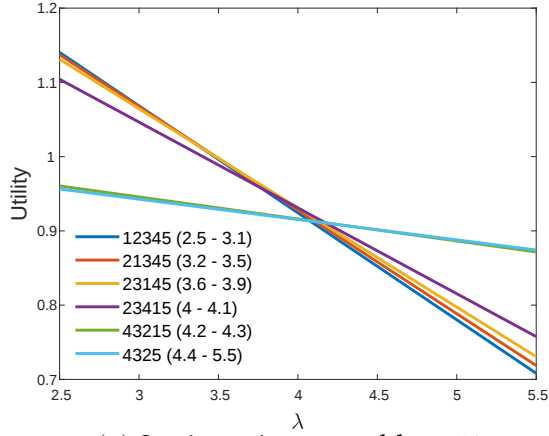
variants, in each matching problem from a large pool of candidate problems. Figure 2 illustrates, for two example matching problems (#1 and #7), the prediction over the range $\lambda \in [2.5, 5.5]$. For each possible strategy—which, for these three variants, can be represented as a ROL—we calculated the resulting expected overall (consumption + news) utility at a range of λ s. Each subfigure shows all the ROLs that maximize utility for some λ in the range for a given problem and variant. For the dynamic variant DP, we also verified that any point on the envelope represents a consistent strategy, i.e., it is immune to profitable surprise deviations, and is thus a PPE. Therefore, the envelope in each subfigure represents the model’s optimal-strategy predictions by λ .

As the figure demonstrates, an optimal strategy in the non-DR variants indeed depends on λ . For low levels of loss aversion, the SF strategy (i.e., the ROL 12345) is optimal. However, there exists a threshold λ above which NSF strategies are optimal. Because the gradual resolution of uncertainty in DP yields a different expected news utility, said threshold differs across the variants—compare subfigures (a) and (c) with subfigures (b) and (d). Moreover, because the environment in problem #7 is less competitive than in problem #1, the threshold is higher in problem #7, implying less prevalent NSF behavior—compare subfigures (a) and (b) with subfigures (c) and (d). Importantly, the envelopes consist of *specific* NSF strategies that the model predicts as optimal under the different problems and variants, yielding an additional prediction that we assess empirically below.

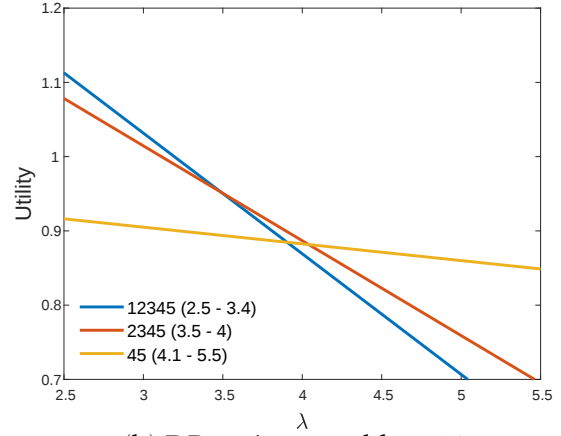
Having created, for each candidate matching problem, optimal-strategy predictions as a function of an individual’s loss aversion λ , we generated population predictions. We wanted to create ex-ante, no-degrees-of-freedom predictions but, naturally, did not have a previously estimated distribution of λ in our subject populations. We opted for basing our population predictions on past estimates from a somewhat similar population: estimates from (Dreyfuss et al., 2022) among lab-experiment participants in a related context (Li, 2017). In those estimates, 67 percent have $1 \leq \lambda \leq 3$, and are therefore never predicted to deviate from SF behavior; 24 percent have $3 < \lambda \leq 5$; 7 percent have $5 < \lambda \leq 7$; and 2 percent have $7 < \lambda \leq 10$; for details, see Online Appendix C. We note that any distribution with a tail of $\lambda > 3$ will result in qualitatively similar directional predictions.

Table 2 presents the ten matching problems we selected for our experiment. Each problem is defined as five school values and (unconditional) probabilities of acceptance. The three columns under “NSF (%)” report the population prediction for the prevalence of NSF strategies under each of the four treatments (recall that our large market property implies that the two static treatments are equivalent, so their predictions coincide). For

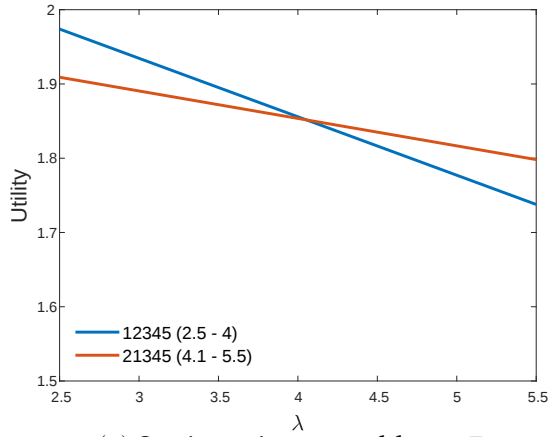
Figure 2: Candidate optimal strategies in example matching problems



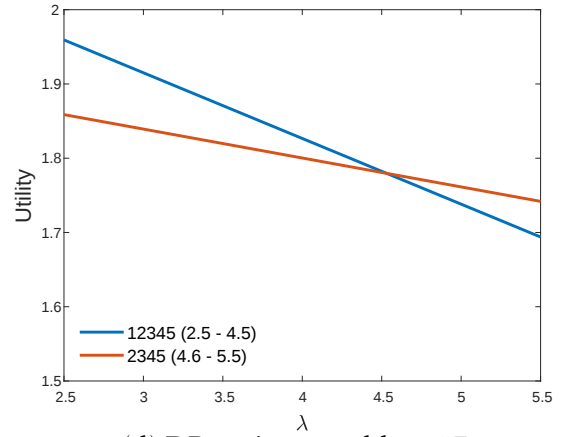
(a) Static variants, problem #1



(b) DP variant, problem #1



(c) Static variants, problem #7



(d) DP variant, problem #7

Notes: Candidate optimal strategies for matching problems #1 and #7. Schools are denoted by their indices (i.e., s_1 is denoted by 1), and ROLs are represented by a sequences of numbers denoting the order in the list. Matching problem #1 is defined by the vectors of dollar school values (1.5, 1, 0.75, 0.5 0.25) and thresholds (0.95, 0.8, 0.75, 0.1, 0), and problem #7 is defined by school values (1.25, 1, 0.75, 0.5 0.25) and thresholds (0.7, 0.15, 0.1, 0.05, 0). (See Table 2 for school values and thresholds for all problems #1–#10.) Subfigures (a) and (c) plot, for the static variants, utilities from all strategies that are optimal for some loss-aversion parameter in the range $2.5 \leq \lambda \leq 5.5$ (grid step = 0.1). Panels (b) and (d) do the same for DP.

example, the 31% prediction in problem #1 in the SP/SR column reflects the estimated population share with $\lambda > 3.1$, which in that problem implies that the ROL 12345 is no longer an optimal strategy (see Figure 2(a)). We chose these ten matching problems, which each subject in our experiment faces, based on the population predictions in these “NSF (%)” columns. (The four rightmost columns report, for the relevant treatments, population predictions using alternative measures; we discuss them below.)

2.2.2 Within-treatment Predictions

The most important feature for predicting NSF shares in a matching problem is the probability of acceptance at the highest value school (see Proposition 2). Specifically, this probability determines the predicted behavior of a subject with a moderately high loss-aversion parameter (say, $3 < \lambda < 6$).

We use the label “weak student” problems for matching problems 1–3, where the probability of acceptance at the highest value school is small (5%), which implies a large predicted share of NSF strategies (24%–31% in the non-DR “NSF (%)” columns). We use the label “medium student” problems for problems 4–7, where the probability of acceptance at the highest value school is larger (20%–30%), implying a lower predicted share of NSF strategies (12%–24%). Finally, we use the label “strong student” problems for problems 8–10, where the probability of acceptance at the highest value school is high (60%–65%), which implies no predicted NSF strategies (under the past empirically estimated distribution of λ we use). Online Appendix C further describes the process and the criteria we used for choosing the sets of values and probabilities in the ten problems.

In summary, we predict that under the three non-DR treatments, the share of NSF strategies will monotonically decrease from weak- to medium- to strong-student problems. We also predict *how* agents will depart from SF strategies. For example, as Figure 2 implies, the second and third most prevalent NSF ROLs in the static treatments are predicted to be 21345 and 23145, respectively.

2.2.3 Between-treatments Predictions

As reported in the three columns under “NSF (%)” in Table 2, we predict varying shares of NSF strategies under the three non-DR treatments, and none under DR. It is important to note that these are predictions on chosen *strategies*; however, in our experiment we only collect data on observed *behavior*. Under the static treatments (SP, SR) there is no difference between strategies and behavior, as we fully observe subjects’ submitted ROLs.

Table 2: Matching problems and EBRD predictions

Matching problem		Predictions						
#	\$ Values (s_1, s_2, s_3, s_4, s_5)	NSF (%)			Costly NSF (%)		DO NSF (%)	
	Probs. (q_1, q_2, q_3, q_4, q_5)	DR	SP/SR	DP	SP/SR	DP	SP	SR
1	(1.50, 1.00, 0.75, 0.50, 0.25) (0.05, 0.20, 0.25, 0.90, 1.00)	0%	31%	26%	7%	7%	31%	7%
2	(1.50, 1.00, 0.75, 0.50, 0.25) (0.05, 0.20, 0.30, 0.40, 1.00)	0%	31%	24%	7%	9%	31%	7%
3	(1.50, 1.00, 0.75, 0.50, 0.25) (0.05, 0.10, 0.85, 0.9, 1.00)	0%	29%	29%	4%	3%	29%	4%
4	(1.25, 1.00, 0.75, 0.50, 0.25) (0.20, 0.80, 0.90, 0.95, 1.00)	0%	24%	18%	4%	3%	24%	4%
5	(1.50, 1.00, 0.75, 0.50, 0.25) (0.25, 0.80, 0.85, 0.95, 1.00)	0%	21%	18%	4%	4%	21%	4%
6	(1.25, 1.00, 0.75, 0.50, 0.25) (0.25, 0.75, 0.80, 0.85, 1.00)	0%	19%	12%	4%	2%	19%	4%
7	(1.25, 1.00, 0.75, 0.50, 0.25) (0.30, 0.85, 0.90, 0.95, 1.00)	0%	18%	14%	5%	3%	18%	5%
8	(1.25, 1.00, 0.75, 0.50, 0.25) (0.65, 0.70, 0.85, 0.95, 1.00)	0%	0%	0%	0%	0%	0%	0%
9	(1.50, 1.25, 1.00, 0.50, 0.25) (0.65, 0.75, 0.80, 0.85, 1.00)	0%	0%	0%	0%	0%	0%	0%
10	(1.25, 1.00, 0.75, 0.50, 0.25) (0.60, 0.65, 0.80, 0.90, 1.00)	0%	0%	0%	0%	0%	0%	0%

Notes: Parameters and population predictions for each of the ten matching problems. Matching problem columns describe schools' money values (top row) and unconditional probabilities of acceptance (bottom row) in each problem. Prediction columns are based on an empirically estimated distribution of λ from [Dreyfuss et al. \(2022\)](#).

However, under the dynamic treatments (DP, DR), we only observe actions dynamically taken by subjects, which only reveal partial information on full strategies.

Specifically, under DP, we only observe the set and order of schools a subject applied to by the time the process ends (when either some school accepts, the subject decides to stop the process, or every school rejected the subject). Similarly, under DR we only observe the sequence of offer sets a student received and the schools the student decided to keep from those sets. Such observed decisions can be either *SF-consistent* or *SF-inconsistent*.

For example, suppose a subject's priority score at the highest value school is higher than the threshold (so once the subject applies to that school, the process terminates). In that case, a SF-consistent behavior under DP consists of the subject only applying to that school. However, such observed behavior could also result from a NSF strategy such as 14235. On the other hand, suppose a subject's priority score at the highest value school is *lower* than the threshold. Under DR, the school never sends an offer to that subject. In that case—by far the most common case in weak-student matching problems (see Table 2)—we will never observe the subject's choice regarding that highest value school in DR. More generally, a subject's behavior may be SF-consistent merely because we had limited opportunity to observe deviations. As a result, SF-inconsistent behavior in dynamic treatments is only a lower bound on NSF strategies.

Table 2's four rightmost columns report theoretical predictions using two alternatives to our primary ("NSF (%)") measure. These alternative measures allow for a more direct comparison across treatments. However, as discussed below, the theory predicts only relatively small cross-treatment differences in these measures—differences that our experiment is not optimized to detect.

The first measure is *costly* NSF behavior, which measures NSF behavior that is also payoff relevant (i.e., it affected the subject's final match). This measure allows for direct comparisons between all four treatments. Clearly, under DR, we predict no such behavior; however, as the two columns under "Costly NSF (%)" in Table 2 show, the predicted share for this measure varies little across both matching problems and treatments (and always remains below 10%).²³

The second measure is *dynamically-observable* (DO) NSF behavior, which counts NSF ROLs in the static treatments that would have been observed as SF-inconsistent behavior if implemented dynamically (i.e., as a sequence of applications to schools in SP, and as keep/reject responses to sequentially arriving school offers in SR).

²³A potential alternative to costly NSF is actual earnings; see Online Appendix B. For further details on generating predictions for costly and dynamically observable NSF, see Supplementary Material Appendix A.

As suggested by the two columns under “DO NSF (%)” in Table 2, the theory predicts that under SP, every NSF strategy is also dynamically observable. In contrast, the theory predicts that under SR, all dynamically observable NSF strategies are costly. These two predicted identities limit the usability of this measure as an alternative for cross-treatment comparisons; however, they provide additional within-treatment predictions that we investigate in the next section.²⁴

3 Results

3.1 Sample

We ran the experiment on two different samples:²⁵ participants in Cornell’s BSL (Business Simulation Lab) SONA-system, recruited from June 30 to July 9, 2021, and participants in the 2020 and 2021 Israeli Psychology Matching Mechanism (IPMM)—a DA-based clearinghouse that matches students and graduate programs in psychology—recruited from August 26 to September 5, 2021. We used identical experimental interfaces, except for language (English vs. Hebrew) and currency (1 USD = 4 NIS in the experiment and \$4 vs. 15 NIS as a show-up fee).

At Cornell, 223 subjects clicked on the experiment’s link, and we closed the experiment after 206 subjects completed it. They earned on average \$13.23 (including the show-up fee); the median completion time was 16.8 minutes. As preregistered, we dropped eight subjects who failed the attention check and another two who completed the experiment itself (after the tutorial) in more than an hour. The remaining 196 subjects’ assignment is: 77 (39.3 percent) DR, 39 (19.9) SR, 40 (20.4) SP and 40 (20.4) DP.

In Israel, we first emailed 1,095 invites to the 2021 IPMM participant pool, of whom 225 clicked on the experiment and 171 completed it. To meet our preregistered 200-subjects target, we emailed 1,206 additional invites to the 2020 pool, of whom 215 clicked on the experiment and 138 completed it. Together, they earned an average of \$16.43 (including the show-up fee) in a median completion time of 22.2 minutes. After dropping five subjects who failed the attention check, the remaining 304 subjects’ assignment is: 126 (41.4 percent) DR, 60 (19.7) SR, 59 (19.4) SP, and 59 (19.4) DP.

²⁴The driver behind these predicted identities is the fact that the NSF strategies predicted by the model in the non-DR treatments always involve the highest value school. See [Meisner and von Wangenheim \(2021\)](#)’s characterization.

²⁵Preregistered separately at <https://aspredicted.org/rd2y5.pdf> and <https://aspredicted.org/4qc8q.pdf>. Online Appendix D discusses our main-text analysis in light of the prespecified analysis plan.

In the rest of this section we pool the samples ($N = 500$; 203 DR and 99 each SP, SR, and DP). Online Appendix B replicates the main analysis by sample. While the general level of observed NSF behavior is markedly lower in the IPMM sample, our main findings remain similar across the samples.

3.2 NSF Shares

Panel (a) of Figure 3 presents the EBRD model’s predictions regarding NSF shares (based on said past estimated population distribution of loss aversion; $\blacktriangle$) compared to classical-preferences predictions ($\blacktriangledown$), by treatment and problem type. Within treatments, in the three non-DR treatments, the EBRD-predicted NSF share declines as competitiveness decreases. Across treatments, the EBRD-predicted NSF share is lower (0 percent) in DR compared to the other three treatments, in all but strong-student problem types. Trivially, classical preferences predict 0 percent NSF in all treatments and problems.

Panel (b) presents empirical shares of NSF behavior ($\blacksquare$), as well as p -values from equality-of-coefficients tests.

We make three observations. First, looking at *levels*, there appears to be an across-the-board minimum “baseline” of 16–22 percent NSF behavior, in all treatments, that neither classical nor EBRD preferences can explain.

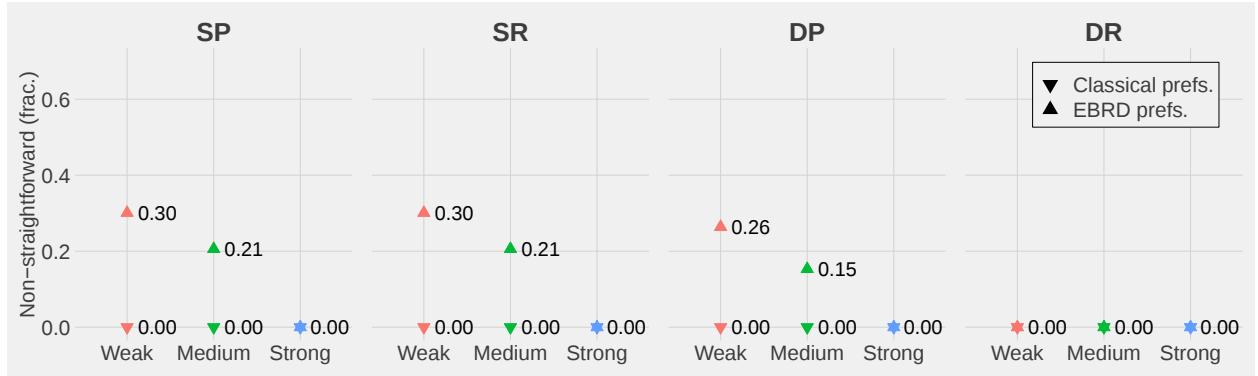
Second, looking at within-treatment *trends*, the data strongly support the EBRD prediction of a notable decrease in the share of observed NSF behavior in all non-DR treatments when moving from weak- to medium- to strong-student problem types: from 33–48 to 30–37 to 19–22 percent. (The flat, zero-NSF predictions of classical preferences are easily rejected.²⁶) In contrast, the trend in DR is much flatter (from 19 to 16 to 16 percent NSF), consistent with our EBRD-model predictions (which, for DR, coincide with those of classical preferences).²⁷

Third, comparing across treatments, the EBRD predictions are also supported by the data. In contrast with the constant no-behavior-difference prediction of classical preferences, NSF shares in DR are substantially lower than in other treatments in weak-

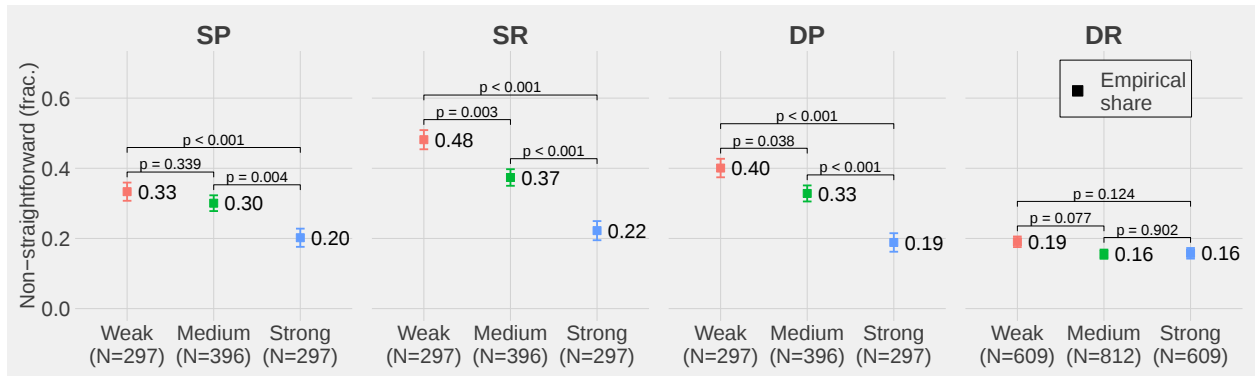
²⁶The classical-vs.-EBRD comparison is “fair” in that both models have zero degrees of freedom. (While the EBRD model has a free parameter (λ), our predictions have no free parameters, as they were generated, prior to data collection, based on a previously estimated population distribution of λ .)

²⁷In addition to these aggregate shares, in Online Appendix B we look at pairwise correlations in NSF across the ten different problems. We find that individual behavior appears fairly consistent across the ten problems and, in line with the model’s predictions, behavior is substantially more consistent within weak- or medium-student problems (correlations = 0.55–0.70 for pairs within problems #1–7) than across strong-student problems (corr. = 0.35–0.51 for pairs that include at least one problem from #8–10) in non-DR treatments, but not in DR (corr. = 0.19–0.57 and 0.15–0.45, respectively).

Figure 3: Non-straightforward (NSF) behavior
(a) Theoretical Predictions: Classical & EBRD Preferences



(b) Results



Notes: Panel (a): NSF-share predictions by treatment and problem type under classical preferences (downwards triangle) and EBRD preferences (upwards triangle). Panel (b): empirical shares. Error bars: standard errors from a regression of NSF behavior on problem type, clustering at the individual level. p -values: Wald tests.

and medium-student problems, where EBRD predicts a difference; but are essentially indistinguishable in strong-student problems, where the model predicts no difference. (In Online Appendix B we test this formally by regressing NSF on treatment, by problem type; we strongly reject the null $SP = SR = DP = DR$ in weak- and medium- ($p = 0.000$), but not in strong-student problem types ($p = 0.32$).)

Online Appendix B reproduces panel (b) of Figure 3 twice: for subjects' first vs. last five rounds. While the above three observations generally hold in each of the two subsamples, NSF shares noticeably drop from the first to last five rounds—suggesting that experience-based learning during the experiment may make our results still more consistent with the EBRD model's predictions. In particular, regarding *levels*, the minimum baseline that neither model can explain decreases in the last five rounds to 10–17 percent; and, regarding *trends*, the declines in the non-DR treatments from weak- to medium- to strong-student problems become 30–37 to 23–36 to 12–16.

3.3 ROL Types

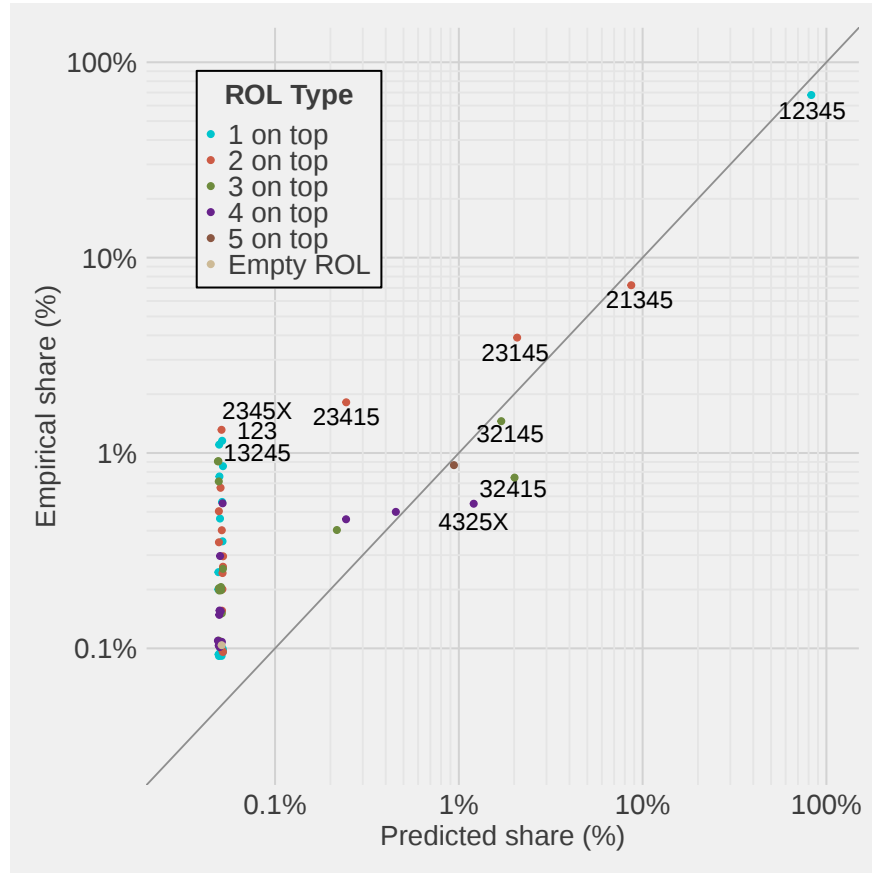
Moving beyond NSF shares, the model (with our distribution of λ) predicts a specific distribution of NSF ROLs in both static treatments.²⁸ Figure 4 compares the predicted distribution (horizontal axis) versus the observed distribution (vertical axis), pooling together all ten matching problems in the two static treatments.

The figure shows that NSF ROLs that are predicted to be prevalent—dots closer to the right side of the figure—are indeed roughly as empirically common—close to the 45° line. In particular, the two NSF ROLs with the first and second highest predicted shares—21345 and 23145, respectively—are also the first and second most empirically prevalent NSF ROLs. More generally, all predicted NSF ROLs are observed in the data, and the differences between prediction and empirical prevalence are never much higher than one percentage point.

Finally, zooming further in, we examine specific ROL-type predictions in specific *problems*. We focus on weak problems: problems #1–3 in Table 2. In addition to having a higher predicted NSF share, we intentionally designed these three problems to show variation in the top-ranked school among loss-averse subjects. Specifically, in problem #3, $q_2 = 0.1$ is relatively low, while $q_3 = 0.85$ is relatively high, yielding a prediction that the most prevalent NSF ROL will rank the third-highest-value school on top. In contrast,

²⁸The same is true for DP; however, in DP we do not observe complete strategies but rather incomplete sequences of applications.

Figure 4: Predicted vs. observed frequency of ROLs



Notes: Theoretical EBRD predictions and empirical shares of ROLs in the static treatments. $N = 1,980$. Points are jittered for readability. Logarithmic scale; $1/1980$ is added to each observation's predicted and observed shares to allow for the inclusion of observations with zero predicted/observed shares. All ROLs are color coded by the first-ranked school. ROLs with at least 1% predicted or empirical share are labeled. Since the probability of admission to school 5 is one, ROL sets that are identical up to school 5 are grouped together, with "X" denoting the payoff-irrelevant part of the list.

in problems #1 and #2, $q_2 = 0.2$ is higher, and $q_3 = 0.25\text{--}0.3$ is only slightly above q_2 , yielding a prediction that the most prevalent NSF ROL will rank the second-highest-value school on top. Since problems #1–3 are otherwise similar to each other, we view the test of these predictions as a particularly sharp test of the theory.

Figure 5 has similar structure to Figure 3, but it shows predictions and results by *problem*, only for problems #1–3 and only for NSF behavior with s_2 on top in panels (a) and (b), and with s_3 on top in panels (c) and (d).

Comparing panel (a) vs. (b) and panel (c) vs. (d) suggests that the data qualitatively track most patterns predicted by the EBRD model. In contrast, the non-trend predicted by classical preferences is mostly rejected by the data.

3.4 NSF Behavior in DR

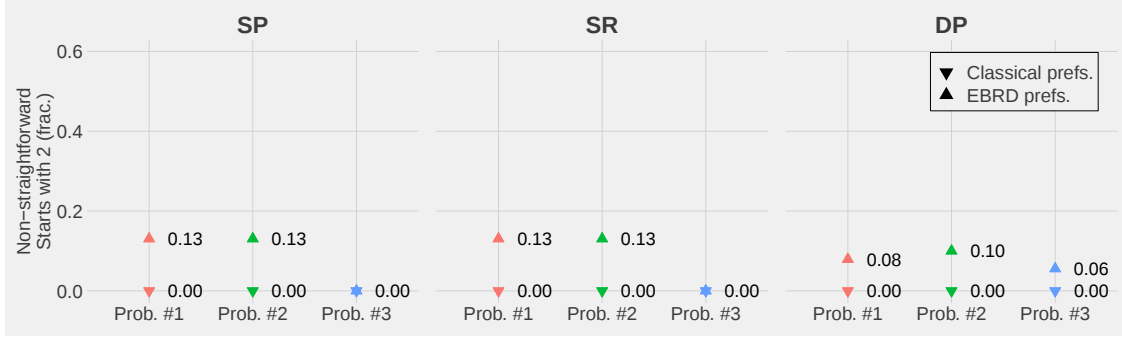
Our main between-treatments prediction compares DR to the other three treatments. However, as discussed in Section 2, such comparisons may favor DR because our main outcome, NSF shares, only counts *actions* taken by subjects, not full strategies. The general worry is that unlike in the static treatments, in DR, some NSF strategies would not lead to observed NSF behavior unless the subject received sufficiently many offers. In particular, if NSF behavior in DR is driven by strategies consistent with ROLs prevalent in the static treatments—e.g., ROLs like 21345—then observations in which the subject does not receive an offer from the highest value school, which are quite common given our design, would not be classified as SF-inconsistent even if the subject played a non-straightforward strategy. In short, subjects may have fewer opportunities in DR to display SF-inconsistent behavior—especially if they play NSF strategies consistent with common NSF ROLs observed in the static treatments.

We now present two types of evidence suggesting that this is not likely the case. First, we show that the most prevalent NSF behavior in DR does not reflect NSF strategies that would often fail to appear as SF-inconsistent. Second, we show that observed NSF does not increase with the number of offers, or with the presence of high-value offers.

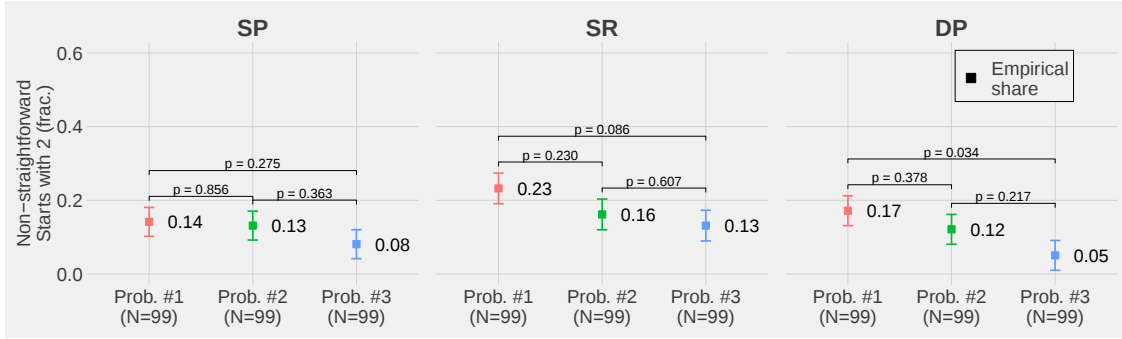
First, across all ten problems, 338 out of 2,030 subject-round observations (17 percent) in DR are classified as NSF. Of these 338 NSF observations, 249 (74 percent) involve rejecting offers received in the round’s first period. Of these first-period offer sets, 215 (86 percent) contain only the lowest-, or second-to-lowest-value school, or both. In other words, most observed NSF behavior in DR involves getting offers only from low-value schools in the round’s first period, and rejecting them. This behavior appears unique to

Figure 5: NSF by ROL Type

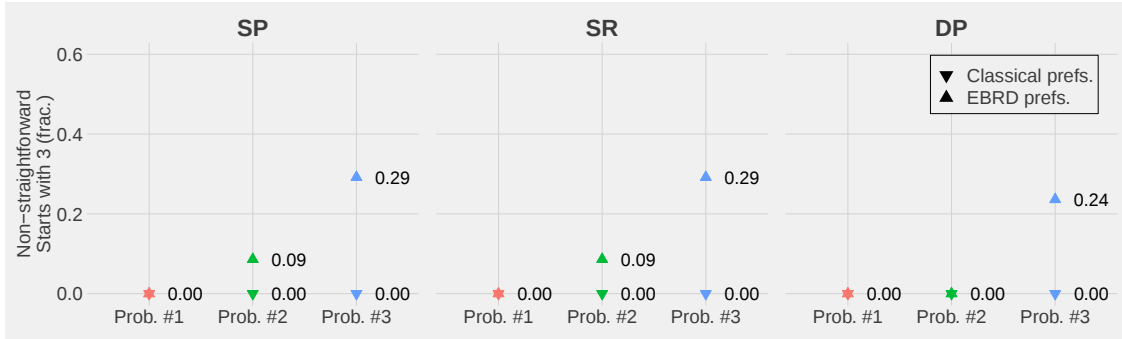
(a) Theoretical Predictions (2 on Top): Classical & EBRD Preferences



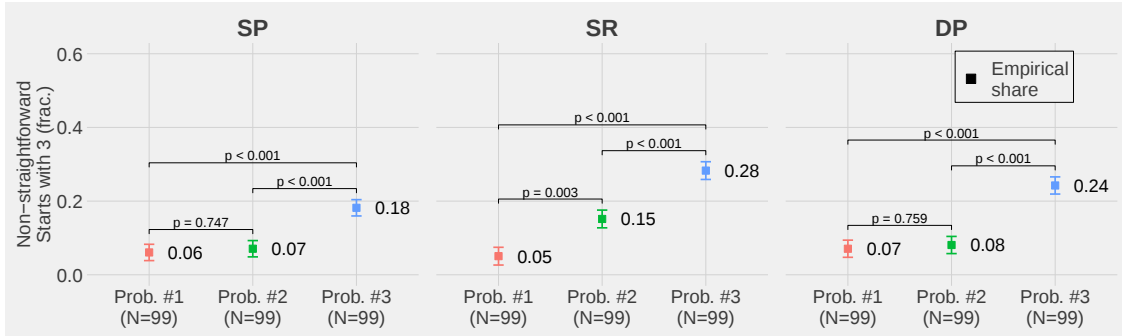
(b) Results (2 on Top)



(c) Theoretical Predictions (3 on Top): Classical & EBRD Preferences



(d) Results (3 on Top)



Notes: Fraction of NSF ROLs where the 2nd (3rd) highest value school is the one ranked first (weak-student problems only). Panels (a) and (c) present the indicated NSF strategy predictions by problem type under Classical (downward triangles) and EBRD (upward triangles) preferences. Panels (b) and (d) present the empirical shares. Error bars: standard errors from a regression of NSF behavior on problem type, clustering at the individual level. p -values: Wald tests.

DR and, moreover, implies that NSF behavior in our DR data typically occurs in situations that subjects often face.

Second, we do not find that observed NSF behavior increases with the total number of offers received in a round. In fact, NSF shares monotonically *decrease*, from 26 percent in one-offer rounds to 15 percent in five-offer rounds (Table B.1 in Online Appendix B).

Moreover, if subjects followed a strategy consistent with a ROL such as 21345, we would expect a higher fraction of observed NSF behavior conditional on receiving an offer from the highest value school compared to not receiving it. However, we do not find that in the data: out of 621 observations in which subjects received an offer from the highest value school and 1,409 in which they did not, respectively, 98 (16 percent) and 240 (17 percent) are classified as NSF.

3.5 Alternative Measures

As discussed in Section 2, in addition to the share of all NSF behavior, we had two additional NSF measures that allow for more direct comparisons across treatments. The first, costly NSF, counts only payoff-relevant NSF behavior, and allows for direct comparisons across all four treatments. The second, dynamically observable NSF (DO NSF), counts only NSF behavior that would have been observable in a dynamic implementation, and allows for direct comparisons between SP and DP, and between SR and DR.

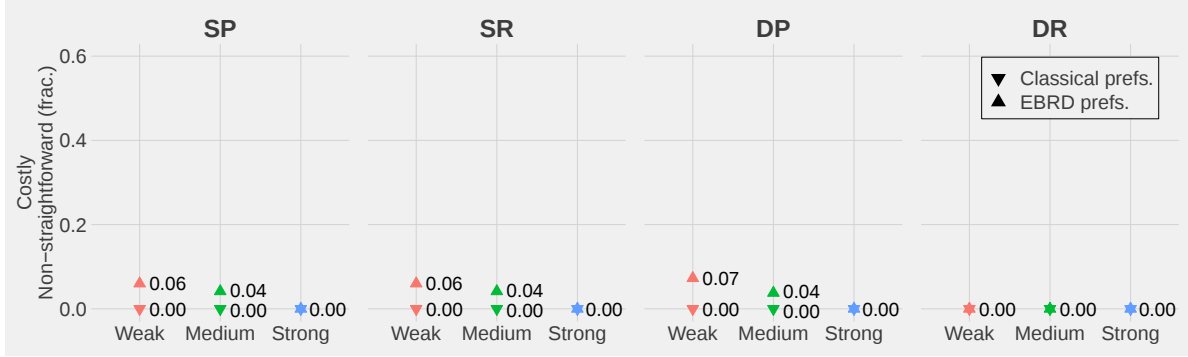
Panels (a) and (c) in Figure 6 show our predictions (see also Table 2 on page 22 and its accompanying discussion). As panel (a) shows, we predict a small share of costly NSF, with little variation across and within treatments (0–7 percent). Panel (c) shows that for DO NSF, in SP we predict significant variation—the same variation predicted for NSF—across problem types, and in SR we predict small variation—the same variation predicted for costly NSF—across problem types.

Empirically, panel (b) shows a slightly higher-than-predicted level of costly NSF behavior (6–14 percent), with no systematic differences across treatments and problem types. Panel (d) shows that for SP, when moving from weak- to medium- to strong-student problem types, DO NSF drops from 31 to 27 to 15 percent, consistent with the EBRD-predicted trend. In SR, the level of DO NSF is higher than predicted (11–17 percent), with no clear trend across problem types.

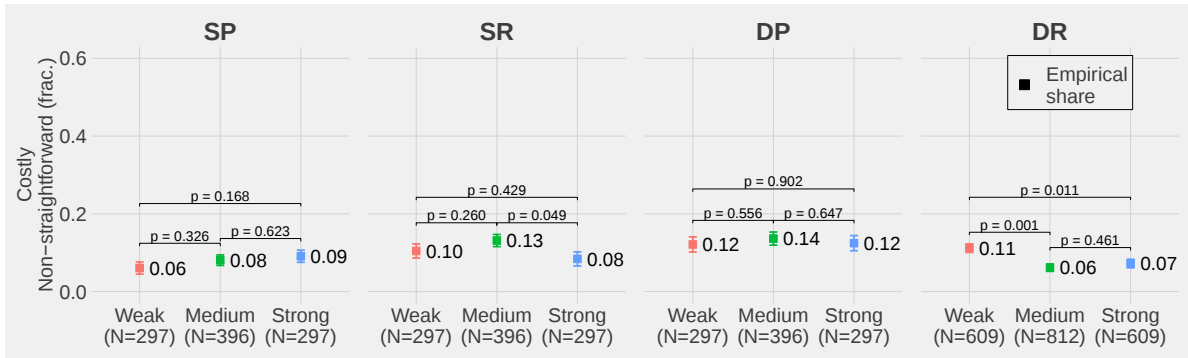
While we did not optimize our experiment to detect differences in these alternative measures, in principle we did collect enough observations to marginally detect a 4–6 percent predicted difference in them, both within each of the non-DR treatments and

Figure 6: Alternative Measures

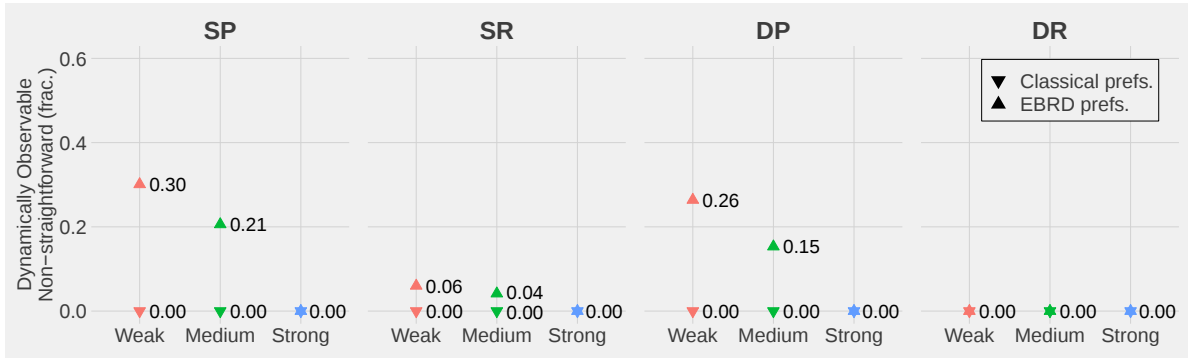
(a) Theoretical Predictions (Costly NSF): Classical & EBRD Preferences



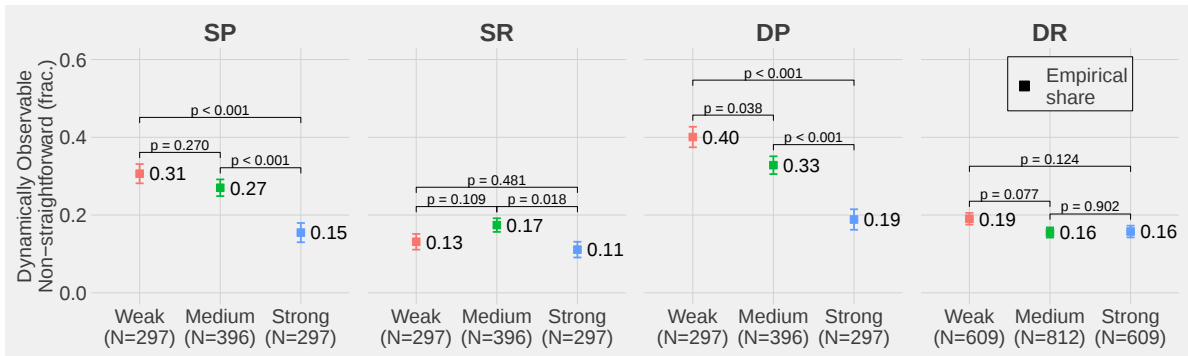
(b) Results (Costly NSF)



(c) Theoretical Predictions (Dynamically Observable NSF): Classical & EBRD Preferences



(d) Results (Dynamically Observable NSF)



Notes: Fraction of costly and dynamically observable NSF. Panels (a) and (c) present predictions for the share of costly NSF and dynamically observable NSF by problem type under Classical (downward triangles) and EBRD (upward triangles) preferences. Panels (b) and (d) present the empirical shares. Error bars: standard errors from a regression of NSF behavior on problem type, clustering at the individual level. p -values: Wald tests.

between DR and non-DR treatments, under ideal conditions—i.e., if the data were perfectly described by the model. However, as discussed above, we find higher NSF shares than our model predicts. These higher NSF shares muddy the picture for both alternative measures.

Start with costly NSF. We find that NSF behavior where the model does not predict it—in DR and in strong-student rounds—is costly 36–63 percent of the time, whereas in weak- and medium-student rounds in the non-DR treatments it is costly only 18–42 percent of the time (and always less, within each non-DR treatment, than in strong-student rounds). Therefore, the empirical *NSF-to-costly-NSF ratio* differs both between DR and non-DR treatments, and between problem types within the non-DR treatments, reflecting behavior that our model does not predict.

Moving to DO NSF, in SR we see a similarly high NSF-to-DO-NSF ratio in strong- relative to medium- and weak-student problems, masking the (small) predicted difference in DO NSF. In contrast, in SP, this ratio is similarly high across problem types. Both of these findings simply reflect the empirical distribution of NSF ROLs, discussed in Section 3.3: since NSF ROLs typically involve downranking the highest value school, they are typically dynamically observable under DP, but much less so under DR.

To summarize, we examine two alternative measures, both allowing direct comparisons between DR and other treatments. However, for these measures, our model predicts small variation, making it difficult to detect such differences in our (inherently noisy) data. Digging further in, we find that this variation is masked by the presence and the type of NSF behavior in the treatment (DR) and problem type (strong student) in which the model predicts no NSF.

3.6 Alternative Explanations

Throughout this section, we have shown that the data fit our (no-degrees-of-freedom) EBRD model significantly better than they do classical preferences. But what about alternative models? How well can *they* explain the data?

Recall that we review three sets of empirical patterns, all consistent with our EBRD model: (i) variation in NSF shares within DA variants; (ii) variation in shares across variants; and (iii) prevalence of specific NSF types, both pooling across matching problems, and for specific problems. At the same time, as we note at the beginning of this section, we also find a little-varying “non-pattern”: an across-the-board “baseline” NSF share (of roughly 16–22 percent in Figure 3) that neither classical nor EBRD preferences can explain.

In summary, the EBRD model, while explaining a lot of the observed data, appears to be an incomplete explanation. Alternative explanations—noisy decision making, cognitive limitations, strategic confusion, misunderstanding, or other explanations—likely play a role too. However, no alternative model we are aware of can replace EBRD in explaining the three patterns above.

Trembling-hand models, with an exogenous probability of making errors, are inconsistent with the observed within-treatment variation (as long as the trembles are independent of the matching problem—as is typically assumed). They are also inconsistent with the observed cross-treatment variation: in DR, subjects typically make multiple Keep/Reject decisions, and are therefore predicted by trembling-hand models to make *more* errors than in the (single-ROL-decision) static treatments. Finally, they cannot easily explain the prevalence of specific NSF strategies (unless one imposes a specific distribution over trembles), and in particular of specific NSF strategies in specific matching problems.

Random utility models (RUMs)—which in our context boil down to decision-makers randomly making errors, with the probability of making an error decreasing with its cost—also cannot easily explain our findings. (In fact, as explained in Online Appendix C, when choosing parameters for our experiment, one of our goals was to tell apart EBRD from RUMs; we chose parameters such that EBRD and logit—a specific RUM—make different predictions.) First, they cannot generally explain the observed within-variant patterns: when moving from weak- to medium- and strong-student problems, the cost of submitting ROLs that flip/omit low-value schools, e.g., the ROL 12354, instead of 12345 becomes extremely low,²⁹ and therefore such ROLs’ RUM-predicted share increases. Online Appendix C shows that under logit, as we move from weak- to strong- and medium-student problems, predicted NSF share indeed *increases*—the opposite of what we find in the data (see Figure 3).³⁰ Second, these models are also inconsistent with the observed cross-treatment variation, for the same reason that trembling-hand models are. Third, they are also inconsistent with the specific ROL types that are prevalent in our data (see Figure 4): as discussed above, the NSF ROLs that logit predicts to be prevalent omit/flip low-value schools, but these are rarely seen in the data. Finally, this same feature also makes the predicted variation of ROL types across problems inconsistent with the findings presented in Figure 5.

²⁹Submitting 12354 costs 25 cents with probability 0.013–0.319 in weak-student problems, and with probability 0.0005–0.006 in strong- and medium-student problems.

³⁰This predicted pattern holds for the entire range of the (calibrationally reasonable) logit scale parameter we tested.

Recent models of cognitive limitations, strategic confusion, and misunderstanding—such as [Li \(2017\)](#), [Pycia and Troyan \(2021\)](#), [Börgers and Li \(2019\)](#) and [Gonczarowski et al. \(2022\)](#)—classify mechanisms, their variants, or their implementations (e.g., a variant’s description to participants) by different notions of how simple they are. Therefore, they do not generally explain variation in behavior *within* a specific implementation of a specific variant of a specific mechanism. Moreover, while some such models suggest that dynamic implementation increases straightforward behavior, none of the models we are aware of tells apart dynamic-proposing and dynamic-receiving DA—implying that these models cannot easily explain the observed variation across treatments either. Finally, these models are also silent on the specific types of errors people make in specific choice situations (within a specific mechanism) and thus cannot easily explain the prevalence of specific NSF types in general or within specific matching problems.³¹

4 Discussion

Our findings have several implications. First, we replicate previous findings that a substantial fraction of participants in simple allocation games choose seemingly dominated actions. Such findings are particularly striking in controlled lab settings, where said seemingly dominated actions are equivalent to choosing first-order stochastically dominated (FOSD) lotteries over sums of money. While in clear violation of the predictions of classical-preferences models, our analysis suggests that the inclusion of EBRD preferences can explain a substantial fraction of such behavior. Moreover, the EBRD model can explain specific *types* of such FOSD-violating behavior, both predicted by the model and observed in the data. Importantly, these predictions are made while adding no additional degrees of freedom relative to the no-EBRD model; rather, we merely replace the default assumption of no loss-aversion—i.e., $\lambda = 1$ for everybody—with an empirical population distribution of λ estimated in a previous study.

At the same time, while our DR treatment is predicted to be EBRD-proof, a non-negligible fraction of participants still choose dominated actions—actions that remain

³¹[Meisner’s \(2022\)](#) model of report-dependent utility can be consistent with some of our within-treatment findings in the static treatments. As noted in that paper, telling apart EBRD and report-dependence is possible only under specific conditions, (e.g., when the admission probability of low-value schools is higher than that of high-value schools). Our experiment is not designed to tell the two models apart, and does not have those features. In addition, report-dependent utility is not well defined in dynamic mechanisms, and in particular does not tell apart dynamic student proposing and receiving—and hence cannot explain our findings within and between our two dynamic treatments.

dominated at *any* degree of loss aversion. We also find similar fractions of participants choosing such actions in “strong student” problems under our other DA variants, where the (parameter-constrained) EBRD model *also* predicts no such behavior. These findings underscore the need for additional explanations, for example, along the lines of [Li \(2017\)](#) and [Gonczarowski et al. \(2022\)](#).

Another important prediction of the EBRD model that our experiment confirms is that while the fraction of such non-straightforward (NSF) behavior varies significantly across both DA variants and matching problems, the fraction of *costly* NSF behavior does not vary much across either. From a theoretical point of view, the intuition behind this prediction is simple: for a fixed distribution of loss aversion λ , the EBRD model predicts more FOSD violations the lower the probabilities of acceptance at higher value schools. However, the lower those probabilities, the lower the chance that such violations will end up costly.

From a policymaker’s point of view, one potential reaction is that reducing NSF behavior is of limited value if it is mostly bottom-line-outcome irrelevant. However, such reactions may overlook two nuances. First, in our data, NSF behavior does not only *decrease* under DR; it also *changes*, and could have distributional consequences. The (reduced) NSF behavior that we find under DR—and that the EBRD model cannot explain—may be consistent with, e.g., (potentially overoptimistic) subjects rejecting low-value offers and expecting to get better ones; whereas EBRD-consistent NSF in the non-DR variants is consistent with (potentially over-pessimistic) loss-averse subjects shying away from applying to high-value schools.

Second, informal conversations with policymakers suggest that at least some of them strongly view prevalent NSF behavior as a problem to be minimized, regardless of its actual cost. For example, “leveling the playing field” and other equity arguments in favor of DA become more complicated when NSF behavior is prevalent. From that perspective, reducing NSF behavior is a goal in itself.

Finally, for policymakers, the question of implementation feasibility looms large—a question that we do not address in this paper. By definition, dynamic implementations of a matching mechanism require it to run for longer periods of time, during which participants are repeatedly asked to make decisions. For investigations like ours to have real-world impact, progress will have to be made on whether and how to implement such variants.

References

- Atila Abdulkadiroğlu, Yeon-Koo Che, and Yosuke Yasuda. Expanding “choice” in school choice. *American Economic Journal: Microeconomics*, 7(1):1–42, 2015.
- Georgy Artemov, Yeon-Koo Che, and Yinghua He. Strategic ‘mistakes’: Implications for market design research. *NBER Working Paper*, 2017.
- Eduardo M Azevedo and Jacob D Leshno. A supply and demand framework for two-sided matching markets. *Journal of Political Economy*, 124(5):1235–1268, 2016.
- Abhijit Banerjee, Esther Duflo, Amy Finkelstein, Lawrence F Katz, Benjamin A Olken, and Anja Sautmann. In praise of moderation: Suggestions for the scope and use of pre-analysis plans for rcts in economics. Working Paper 26993, National Bureau of Economic Research, 2020.
- Tilman Börgers and Jiangtao Li. Strategically simple mechanisms. *Econometrica*, 87(6): 2003–2035, 2019.
- Daniel L Chen, Martin Schonger, and Chris Wickens. otree—an open-source platform for laboratory, online, and field experiments. *Journal of Behavioral and Experimental Finance*, 9:88–97, 2016.
- Simon Dato, Andreas Grunewald, Daniel Müller, and Philipp Strack. Expectation-based loss aversion and strategic interaction. *Games and Economic Behavior*, 104:681–705, 2017.
- Bnaya Dreyfuss, Ori Heffetz, and Matthew Rabin. Expectations-based loss aversion may help explain seemingly dominated choices in strategy-proof mechanisms. *American Economic Journal: Microeconomics*, 14(4):515–55, November 2022.
- Federico Echenique, Alistair J Wilson, and Leeat Yariv. Clearinghouses for two-sided matching: An experimental study. *Quantitative Economics*, 7(2):449–482, 2016.
- David Gale and Lloyd S Shapley. College admissions and the stability of marriage. *The American Mathematical Monthly*, 69(1):9–15, 1962.
- Yannai A. Gonczarowski, Ori Heffetz, and Clayton Thomas. Strategyproofness-exposing mechanism descriptions, 2022. URL <https://arxiv.org/abs/2209.13148>.

- Julien Grenet, YingHua He, and Dorothea Kübler. Preference discovery in university admissions: The case for dynamic multioffer mechanisms. *Journal of Political Economy*, 130(6):1427–1476, 2022.
- Rustamdjan Hakimov and Dorothea Kübler. Experiments on centralized school choice and college admissions: a survey. *Experimental Economics*, 24(2):434–488, 2021.
- Avinatan Hassidim, Assaf Romm, and Ran I Shorrer. The limits of incentives in economic matching procedures. *Management Science*, 67(2):951–963, 2021.
- Alexander Karpov. A necessary and sufficient condition for uniqueness consistency in the stable marriage matching problem. *Economics Letters*, 178:63–65, 2019.
- Flip Klijn, Joana Pais, and Marc Vorsatz. Static versus dynamic deferred acceptance in school choice: Theory and experiment. *Games and Economic Behavior*, 113:147–163, 2019.
- Botond Kőszegi and Matthew Rabin. A Model of Reference-Dependent Preferences. *The Quarterly Journal of Economics*, 121(4):1133–1165, 2006.
- Botond Kőszegi and Matthew Rabin. Reference-dependent risk attitudes. *American Economic Review*, 97(4):1047–1073, 2007.
- Botond Kőszegi and Matthew Rabin. Reference-dependent consumption plans. *American Economic Review*, 99(3):909–36, 2009.
- Shengwu Li. Obviously strategy-proof mechanisms. *American Economic Review*, 107(11): 3257–87, 2017.
- Vincent Meisner. Report-dependent utility and strategy-proofness. *Management Science*, 2022.
- Vincent Meisner and Jonas von Wangenheim. School choice and loss aversion. *Working paper*, 2021.
- Marek Pycia and Peter Troyan. A theory of simplicity in games and mechanism design. *Working paper*, 2021.
- Alex Rees-Jones and Samuel Skowronek. An experimental investigation of preference misrepresentation in the residency match. *Proceedings of the National Academy of Sciences*, 115(45):11471–11476, 2018.

Alex Rees-Jones, Ran Shorrer, and Chloe J Tergiman. Correlation neglect in student-to-school matching. *NBER Working Paper No. 26734.*, 2020.

Antonio Rosato. Loss aversion in sequential auctions: Endogenous interdependence, informational externalities and the” afternoon effect”. *Working paper*, 2014.

Alvin E. Roth and Marilda A. Oliveira Sotomayor. *Two-Sided Matching: A Study in Game-Theoretic Modeling and Analysis*. Econometric Society Monographs. 1990.

Ran I Shorrer and Sándor Sóvágó. Dominated choices in a strategically simple college admissions environment: The effect of admission selectivity. *Working paper*, 2022.

A Large-Market DA Variants as Games Against Nature

A.1 Static Variants

There are only two periods in the static SP and SR: at $t = 1$, the student submits a ROL, and $t = 2$ is the terminal period in which she learns her final match. In both variants, the set of strategies $\mathcal{L}$ is the set of possible ROLs, i.e., all permutations of all sets in the power set of S . Under our large-markets assumption, both variants are outcome equivalent: the student is matched with the highest ranked school s_j such that $\rho_j \geq T_j$.

A.2 DP

In DP, the student applies to single schools in odd periods, and Nature sends a response in the following even periods until an application is accepted. Specifically, at $t = 1$, the student applies to some school $s_j \in S$, and at $t = 2$ learns the outcome (acceptance/rejection). In case of acceptance, the game ends at $t = 2$. In case of rejection, at $t = 3$ the student applies to a different school that has not previously rejected her (i.e., some $s_{j'} \in S \setminus \{s_j\}$) and learns the outcome at $t = 4$. The game continues until an application is accepted.³² In our large-market framework, an application to s_j is accepted iff $\rho_j \geq T_j$. Notice that any set of on-path equivalent strategies in DP can also be represented by a ROL, which determines the order of (on-path) applications.

A.3 DR

In DR, the student receives offers during odd periods and responds during even periods. Nature's actions are determined by a set of (known) functions $\phi_j(\cdot) : [0, 1] \rightarrow 2\mathbb{N} + 1$ which map priority scores at $s_j \in S$ to (odd) time periods in which the student receives an offer from s_j . The student receives an offer from s_j during the matching process iff $\rho_j \geq T_j$. Otherwise, we have $\phi_j(\rho_j) = \infty$, interpreted as not receiving an offer from s_j during the matching process.³³ We denote the set of offers the student receives in each period by $\zeta_t = \{s_j : \phi_j(\rho_j) = t\}$. We assume $\phi_{n+1}(\cdot) = 1$, i.e, the outside option is always offered on the first period.

At $t = 1$, the student receives a set of offers ζ_1 , then at $t = 2$, she must keep exactly one offer from ζ_1 . Denote the kept offer by $\kappa_2 \in \zeta_1$. Then at $t = 3$, she might receive offers

³²Since the outside option is in S , an application will always eventually be accepted.

³³While these functions depend on capacities and other students' actions, it is always the case that all else equal, a higher priority score at a school is associated with an offer that arrives at a (weakly) earlier period.

from a disjoint set of schools, $\zeta_3 \subseteq S \setminus \zeta_1$, and at $t = 4$ she must keep exactly one offer from $\{\kappa_2\} \cup \zeta_3$ (i.e., the currently held offer and the new set of offers). The game continues until no more offers are sent, i.e., until the terminal period, in which $\zeta_t = \emptyset$.

A strategy l in DR determines the offer the student retains from any $\{\kappa_t\} \cup \zeta_{t+1}$, given any history of prior offers and decisions. Notice that the strategy space in DR contains, but is not limited to, all possible ROLs, as not all strategies conform to a ROL.³⁴

B Proofs

B.1 Proof of Proposition 1

Proof. We prove that DR is EBRD-strategyproof. The proof for the other variants is a matter of constructing a simple counter-example, e.g., Figure 2.

The proof is structured as follows: we first show (Lemma 1) that without uncertainty, SF behavior must be optimal.³⁵ Next, we show that given a stream of offers, news utility is maximized for a strategy that induces a path of beliefs about consumption which deviates minimally from the shortest possible path (Lemma 3). We then use this result to show that, fixing SF behavior in all continuation histories, SF behavior in the current period maximizes news and consumption utility and is hence uniquely credible. Since there is always a positive probability of reaching an action period with no uncertainty, by backward induction, a PPE must dictate SF behavior.

Let Γ be the game induced by a matching market governed by DR, summarized from student i 's point of view by $\langle G, \mathbf{T}, \mathbf{m}, \text{DR} \rangle$. We slightly abuse notation by treating schools as representing their utility values, and so $s_k > s_l$ means that schools k and l are such that $m_k > m_l$. Similarly, for a set of schools A , $\max A$ represents the school with the highest consumption value. Denote the subgame induced by history h by Γ^h .

Define a *sure-thing period* as an action period that is terminal with certainty (i.e., the student knows that she is below the threshold at all schools that have not yet sent an offer, if there are any).

We first show that an optimal strategy must prescribe SF behavior in sure-thing periods:

³⁴In particular, treating the available offers $\{\kappa_t\} \cup \zeta_{t+1}$ as the state variable, a ROL cannot describe non-Markovian strategies.

³⁵While this lemma sounds trivial, it is not. It states that *regardless* of beliefs at the beginning of the period, SF behavior is optimal. This result is tightly related to some of our modelling assumptions. For example, under horizontal differentiation (which would imply multidimensional consumption utility), the loss from giving up on the currently held offer might loom larger than the added utility from choosing a new offer with a higher consumption utility.

Lemma 1. *Let z be a terminal history that contains a sure-thing period, $\bar{t}_z - 1$. Then a PPE prescribes SF behavior in the sure-thing period, i.e., $\kappa_{\bar{t}_z-1} = \max(\zeta_{\bar{t}_z-2} \cup \{\kappa_{\bar{t}_z-3}\})$.*

Proof of Lemma 1. Denote the currently held offer by s_k . First, assume that the student only receives an offer from one school, denoted s_j .

We consider the two cases: $s_j > s_k$ (the new offer has a higher value) and $s_k > s_j$ (lower value), and show that SF behavior (i.e., picking the higher value offer) is uniquely optimal in both.

Assume $s_j > s_k$. Assume by contradiction that NSF behavior in this period is optimal. Then the student enters the period believing she will consume m_k with certainty. Following through and rejecting s_j yields m_k consumption utils (and no news utility). Deviating and keeping s_j instead of s_k yields

$$m_j + (m_j - m_k) > m_k,$$

where the first and second expressions on the LHS correspond to consumption and news utility, respectively. This is a profitable deviation, a contradiction. Therefore, a strategy dictating NSF behavior in this period is not credible.

An SF strategy prescribes keeping s_j , so the student enters the period believing that she will consume m_j with certainty. Following through yields m_j consumption utils and no news utility.³⁶ Deviating and rejecting s_j yields

$$m_k - \lambda(m_j - m_k) < m_j,$$

where the last term on the LHS is news utility from a decrease in beliefs about consumption. This implies that in this case, all elements $\mathcal{L}_{z_{\bar{t}_z-1}}^*$ must prescribe SF behavior.

Assume $s_j < s_k$. By the same argument (switching the indices) we conclude that in this case as well, all elements $\mathcal{L}_{z_{\bar{t}_z-1}}^*$ must prescribe SF behavior.

In similar vain, using symmetric arguments it is easy to show that when there are multiple offers in the sure-thing period, taking the highest value one is always better and so SF behavior is optimal in this case as well. QED

Having shown this, Proposition 1 follows by backward induction from the following lemma:

³⁶Notice that while there is news utility from receiving the offer, we focus on the decision after the student learns about the offer.

Lemma 2. Let h be a history in which the student acts, and fix behavior in the remainder of the game to be SF (i.e., assume that starting next period, all credible strategies prescribe SF behavior), then a credible strategy in the subgame Γ^h must be SF.

In other words, if the student knows that, no matter what she does now, her behavior in all future periods is SF, then SF behavior is optimal also in the current period.

To prove Lemma 2 we introduce *excess deviation* and show that minimizing it maximizes news utility.

Definition 5. Let $W \equiv (W_t)_{t=0}^T$ be a sequence of beliefs (CDFs) over consumption. Similarly, let W_t^p denote the consumption level at percentile p of W_t , and let $W^p = (W_t^p)_{t=0}^T$ denote the sequence of consumption levels at percentile p . The **excess deviation at p** , denoted $\hat{W}^p$, is defined as half the total distance traveled along the sequence W^p in excess of the shortest route from W_0^p to W_T^p : $\hat{W}^p \equiv \frac{\sum_{t=0}^{T-1} |W_{t+1}^p - W_t^p| - |W_T^p - W_0^p|}{2}$. Similarly, the **total excess deviation** of a belief sequence W , denoted $\hat{W}$, is the integral of excess deviations over percentiles: $\hat{W} \equiv \int_0^1 \hat{W}^p dp$.

Intuitively, given a prior belief W_0 and a posterior belief W_T , excess deviation of percentile p , $\hat{W}^p$, is the total distance traveled away from W_T^p , i.e., the sum of downward movements in all periods with $W_t^p < W_T^p$ plus the sum of upwards movements in all periods with $W_t^p > W_T^p$.³⁷ Similarly, $\hat{W}$ integrates deviations over all percentiles. As shown below, it measures the *excess disappointment* in the process of gradually updating from a given prior to a given posterior.

The *per-period* excess deviation in period t is defined as follows:

$$e(X_t^p, X_{t-1}^p, X_T^p) = \begin{cases} 0 & X_T^p > X_t^p > X_{t-1}^p \text{ (up \& towards)} \\ X_{t-1}^p - X_t^p & X_T^p > X_{t-1}^p > X_t^p \text{ (down \& away)} \\ 0 & X_{t-1}^p > X_t^p > X_T^p \text{ (down \& towards)} \\ X_t^p - X_{t-1}^p & X_t^p > X_{t-1}^p > X_T^p \text{ (up \& away)} \\ X_t^p - X_T^p & X_t^p > X_T^p > X_{t-1}^p \text{ (overshoot upwards)} \\ X_T^p - X_t^p & X_{t-1}^p > X_T^p > X_t^p \text{ (overshoot downwards)} \end{cases}$$

$e(\cdot) \geq 0$ measures movements away from W_T^p as well as those overshooting W_T^p in a given period. Clearly, $\sum_{t=1}^T e(X_t^p, X_{t-1}^p, X_T^p) = \hat{W}^p$.

The following lemma relates excess deviation to total news utility:

³⁷For example, if $W^p = (0, 3, 1)$ then $\hat{W}^p = 2$ because at $t = 2$ the sequence “unnecessarily” deviated two units above the terminal value.

Lemma 3. Let X and Y be two sequences of beliefs such that $X_0 = Y_0 = A$ and $X_T = Y_T = B$. Then, news utility from the stream X is greater than that from the stream Y if and only if total excess deviation under Y is greater than under X :

$$\sum_{t=0}^{T-1} N(X_{t+1} | X_t) > \sum_{t=0}^{T-1} N(Y_{t+1} | Y_t) \iff \hat{Y} > \hat{X}.$$

Proof of Lemma 3. Let $W^p \in \{X^p, Y^p\}$. Since $W_0^p = A^p$ and $W_T^p = B^p$ for all p , then

$$\sum_{t=0}^{T-1} \mu(W_{t+1}^p - W_t^p) = \mu(B^p - A^p) - (\lambda - 1) \hat{W}^p. \quad 38$$

Thus, the following holds:

$$\begin{aligned} & \sum_{t=0}^{T-1} N(X_{t+1} | X_t) - \sum_{t=0}^{T-1} N(Y_{t+1} | Y_t) = \\ & \int_0^1 \left(\sum_{t=0}^{T-1} \mu(X_{t+1}^p - X_t^p) - \sum_{t=0}^{T-1} \mu(Y_{t+1}^p - Y_t^p) \right) dp = \\ & \int_0^1 \left(\mu(B^p - A^p) - (\lambda - 1) \hat{X}^p - \mu(B^p - A^p) + (\lambda - 1) \hat{Y}^p \right) dp = \\ & (\lambda - 1) \int_0^1 (\hat{Y}^p - \hat{X}^p) dp = (\lambda - 1)(\hat{Y} - \hat{X}) > 0 \text{ iff } \hat{Y} > \hat{X}. \end{aligned}$$

QED

Proof of Lemma 2. Consider WLOG a history h in which the student faces a decision between two offers from schools $s_j > s_k$. Abusing notation, denote the SF and NSF strategies by SF and NSF, respectively. The belief distributions induced by these strategies are then denoted by $F_{\text{SF}|h}$ and $F_{\text{NSF}|h}$.

Suppose that in the remainder of the game, the student follows an SF strategy.³⁹ We call two histories that contain the same sequence of offers, but diverge on whether the student chose SF or NSF at h *twin histories*.⁴⁰ Then, for any twin histories h' and h'' future straightforward behavior implies that (i) realized consumption utility would be at least as

³⁸For example, if $A^p = 0$, $B^p = 1$, $W^p = (0, 3, 1)$ then: $\sum_{t=0}^{T-1} \mu(W_{t+1}^p - W_t^p) = 3 - \lambda \cdot 2 = 1 - (\lambda - 1) \cdot 2$.

³⁹Therefore, NSF chooses the lower value s_k in history h , but prescribes SF behavior in all subsequent periods.

⁴⁰While behavior under both histories is straightforward, observed behavior might diverge between the two histories, since the student starts each continuation history holding a different offer.

high as the currently held offer, i.e., at least m_k under NSF, and at least m_j under SF, and (ii) for all $s > s_j$ the probability of matching with s is the same under SF and NSF. Taken together, (i) and (ii) imply:

$$\begin{aligned} F_{\text{SF}|h'}^p &> F_{\text{NSF}|h''}^p \text{ for } p \in [0, \bar{p}_{h'}) \\ F_{\text{SF}|h'}^p &= F_{\text{NSF}|h''}^p \text{ for } p \in [\bar{p}_{h'}, 1] \end{aligned} \quad (5)$$

where $\bar{p}_{h'} = \bar{p}_{h''}$ is the probability of *not* receiving an offer from some $s > s_j$.⁴¹

We now show that, no matter if SF or NSF was initially planned, for any stream of offers, total excess deviation is always higher under NSF than under SF, which, by Lemma 3 implies that news utility is higher under SF.

Let z and $\dot{z}$ be a pair of terminal twin histories that contain h , with z and $\dot{z}$ representing SF- and NSF-behaving history, respectively. Start by looking at the case where the student receives (in some future period) an offer from $s_d > s_j$ (below we show that this is WLOG). Therefore, the student's initial and terminal beliefs in the game Γ^h are the same, regardless of whether she plays SF or NSF: she starts the game with some plan, and ends the game with the degenerate belief $F_T^p = m_d$ for all p .⁴² Since these are twin histories, we have $\bar{p}_{z_{t-1}} = \bar{p}_{\dot{z}_{t-1}}$ and $\bar{p}_{z_t} = \bar{p}_{\dot{z}_t}$ for all t . We introduce the following shorthand notation: Let $F_t^p := F_{\text{SF}|z_t}^p$ and similarly $\dot{F}_t^p := F_{\text{SF}|\dot{z}_t}^p$, and let $e_t^p := e(F_t^p, F_{t-1}^p, F_{t_z}^p)$ and similarly $\dot{e}_t^p := e(\dot{F}_t^p, \dot{F}_{t-1}^p, \dot{F}_{t_z}^p)$.

Start with the first period in Γ^h , denoted t_0 . Notice that this is the only period where the initial plan matters. Assume that the initial plan is NSF (i.e., keeping s_k). Following through implies no belief movement, and hence $e_{t_0}^p = 0$ for all p . Deviating to SF implies an “up and towards” movement for $p \in [0, \bar{p}_h]$ (and no movement elsewhere), which again implies $e_{t_0}^p = 0$ for all p . Now assume that the initial plan is SF. Following through implies no movement and hence $e_{t_0}^p = 0$ for all p . Deviating to NSF (keeping s_k) implies a “down and away” movement and hence $\dot{e}_{t_0}^p > 0$ for $p \in [0, \bar{p}_h]$, and no movement elsewhere. We conclude that regardless of the initial plan, $e_{t_0}^p = 0 \leq \dot{e}_{t_0}^p$ for all p .

Next, let z_{t-1} and z_t be two consecutive subhistories of z and similarly let $\dot{z}_{t-1}$ and $\dot{z}_t$ be their twin subhistories of $\dot{z}$.

Assume $\bar{p}_{z_{t-1}} \leq \bar{p}_{z_t}$, corresponding to a reduction in the probability of the student receiving an offer from some $s > s_j$. For $p \in [0, \bar{p}_{z_{t-1}})$, $F_t^p = F_{t-1}^p$ and hence $e_t^p = 0 \leq \dot{e}_t^p$. For

⁴¹Notice that an update to the belief $G(\cdot|h')$ will possibly change $\bar{p}_{h'}$, but (5) will still hold.

⁴²This holds regardless of whether the student enters the period planning to play SF or NSF.

$p \in [\bar{p}_{z_{t-1}}, \bar{p}_{z_t}]$, by (5) $F_{t-1}^p = \dot{F}_{t-1}^p$ and $F_t^p = m_j > \dot{F}_t^p$, implying

$$e_t^p = F_{t-1}^p - m_j < \dot{F}_{t-1}^p - \dot{F}_t^p = \dot{e}_t^p.$$

For $p \in (\bar{p}_{z_t}, 1]$, we have $F_t^p = \dot{F}_t^p$ and $F_{t-1}^p = \dot{F}_{t-1}^p$ which implies $e_t^p = \dot{e}_t^p$. We conclude that in this case, $e_t^p \leq \dot{e}_t^p$ for all p .

Now assume $\bar{p}_{z_{t-1}} \geq \bar{p}_{z_t}$, corresponding to a (weak) increase in the probability of the student receiving an offer from some $s > s_j$. As before, for $p \in [0, \bar{p}_{z_t})$ we have $F_t^p = F_{t-1}^p$, i.e., $e_t^p = 0 \leq \dot{e}_t^p$. For $p \in [\bar{p}_{z_t}, \bar{p}_{z_{t-1}}]$ we have $F_t^p = \dot{F}_t^p$. We have $e_t^p > 0$ iff $F_t^p > m_d$ (where m_d is the utility of the terminal belief under both histories), since in this case beliefs overshoot upwards under both strategies. However, notice that $e_t^p = \dot{e}_t^p = F_t^p - m_d$. Last, as before, for $p \in (\bar{p}_{z_{t-1}}, 1]$ we have $e_t^p = \dot{e}_t^p$. We conclude that in this case as well, $e_t^p \leq \dot{e}_t^p$ for all p .

The above implies that for any pair of twin terminal histories such that the student receives an offer from $s_d > s_j$, total excess deviation is greater under NSF, and therefore by Lemma 3 news utility is weakly higher under SF.

However, this trivially implies that when the student does not receive an offer from some school $s_d > s_j$, news utility is also higher under SF. To see why, add an auxiliary period in which the student receives and accepts a second offer from s_j (i.e., in addition to the offer received in h). We can then use the analysis above to show that in this modified game, news utility is weakly higher under SF than under NSF, which in turn implies that the same holds in the original game as well.

We have shown that news utility is weakly higher under SF in the game Γ^h , for any terminal history, independent of the initial plan entering the game. Since expected consumption utility is also higher under SF, we conclude that the SF strategy is credible in the subgame Γ^h . Moreover, whenever there is uncertainty over final consumption, these inequalities are strict, implying that if Γ^h has some uncertainty, then *any* credible strategy in Γ^h is SF. QED

Lemmas 1 and 2 imply that whenever there is a positive probability of reaching a sure-thing period, only an SF strategy is credible. However, in any history, either the student is in a sure-thing period, or there is a positive probability of reaching one. Therefore, any element in $\mathcal{L}^*$ must prescribe SF behavior. QED

B.2 Proof of Proposition 2

Proof. This result is based on and extends Meisner and von Wangenheim’s (2021) Proposition 2. While they use the static Kőszegi and Rabin (2007) whereas we use Kőszegi and Rabin, 2009, for the static mechanisms these models essentially coincide. Taking into account news utility from ROL submission (which equals the expected value of the induced lottery and is absent from their framework), and noticing that in their parameterization $\Lambda = (\lambda - 1)$,⁴³ the two models are identical up to the scale of the loss-aversion parameter. Following the same steps in the proof of the first part of their Proposition 2, we get that if $\lambda < 1 + \frac{2}{1-q_1} \equiv \underline{\lambda}$, the SF ROL is strictly optimal in the static SP and SR.

The second part of our Proposition 2 is an extension that relies on monotonic and independent probabilities of admission. Ported into our notation, their proof shows the following condition for the suboptimality of the SF ROL:

$$\frac{1}{\lambda - 1} \geq -q_1 + \epsilon + (1 - q_1),$$

with q_1 the probability of exceeding the threshold at s_1 , and $\epsilon := Pr(\rho_1 > T_1 \wedge \rho_2 > T_2)$ is the probability of being above the threshold at the two highest value schools.

Meisner and von Wangenheim (2021) get an upper bound by taking the lower bound of ϵ (zero). Instead, we use Assumption 1 to tighten the bound. In particular, under Assumption 1 we have $\epsilon = q_1 \cdot q_2 > q_1^2$. Therefore, we get that if $\lambda > 1 + \frac{2}{(1-q_1)^2} \equiv \bar{\lambda}$, the SF ROL is strictly suboptimal in SP and SR. QED

⁴³Up to the normalization of η ; see footnote 10.

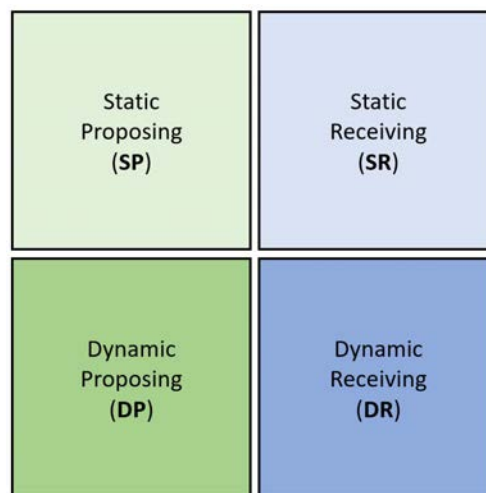
Online Appendix for “Deferred Acceptance with News Utility”

A Experiment Instructions and Interface

The experiment consists of the following parts:

1. Consent form
2. General instructions
3. Specific instructions (for each treatment)
4. Tutorial
5. Attention question
6. Ten order-randomized matching problems
7. Demographic questions
8. Feedback

We provide screenshots for the General instructions, Specific instructions and an example for the interface in a matching problem. These screenshots are color-coded (only for illustration here) to highlight the differences between treatments. We will use the following color scheme:



Color-coding for screenshots

All pages shown here will have a color-coded symbol at the bottom right corner, illustrating how different parts of the instructions appeared on different treatments.

A.1 General Instructions

This is a study about decision making. Based on your decisions, you might earn a considerable amount of money **in addition** to the \$4 participation payment. In this study, we simulate a procedure to allocate students to schools. The procedure, payment rules, and student allocation method will be described in the next few screens. Everything will be exactly as specified in the instructions.

Next >



Figure A.1: General instructions - page 1

Imagine that there are five different schools, and you are one of many candidates considering applying this year. These schools differ in size, geographic location, topics taught, and quality of instruction. Each school has a dollar **value** that reflects all these qualities. The *value* of the school you are matched with will be added to your payment.

< Back

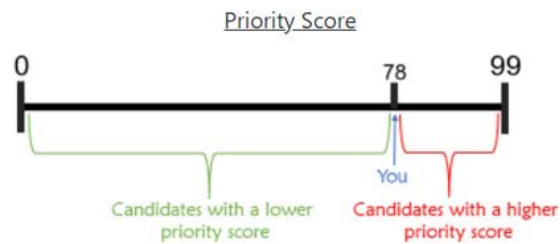
Next >



Figure A.2: General instructions - page 2

School seats are limited, and each school can accept only a small share of the candidates applying.

Each school gives you, the candidate, a **priority score** which is chosen at random between 0 and 99 (every round number between 0 and 99 is equally likely to be chosen). *Priority scores* represent your rank in a school relative to the rest of the candidates. For example, if your *priority score* at some school is 78, it means that you are ranked higher than 78 percent of the candidates, and lower than or equal to the other 22 percent.



Please note: your *priority score* at each school is chosen randomly and separately, therefore, your *priority scores* at different schools are not connected to each other in any way.

Each school has an **acceptance threshold**. Schools are only willing to accept candidates whose *priority score* is greater than or equal to this *threshold*. For example, a school with a *threshold* of 75 is willing to accept you if your *priority score* is 80 and not willing to if your *priority score* is 70.

< Back

Next >



Figure A.3: General instructions - page 3

In this study you will participate in 10 rounds. Each round will simulate the matching process to 5 schools. In each round, you can be matched with at most 1 school.

At the beginning of each round, you will know the *value* of each school and its *threshold*, but you won't know your *priority score* until the end of that round. Note: since your *priority score* at each school is a random number between 0 and 99, a school's *threshold* conveys the probability it is willing to accept you (that is, the chance that your *priority score* is greater than or equal to its *threshold*). For example, if some school's *threshold* is 25, then there is a 75% chance that your *priority score* is greater than or equal to this *threshold*, meaning a 75% chance they will be willing to accept you.

In each round, before the matching process begins, you will be presented with a table similar to this one. This table will also be visible to you throughout the process.

School	Value	Threshold	Your Priority Score	Chance that Your Priority Score $\geq$ Threshold
Pine Peak	\$0.75	60	?	40%
Birch Hill	\$0.25	0	?	100%
Hickory Bridge	\$0.50	50	?	50%
Maplecrest	\$1.25	80	?	20%
Elm South	\$1.00	70	?	30%

In this example, Elm South's *threshold* is 70, which means that the chance that your *priority score* is greater than or equal to this *threshold* is 30%. If you are matched with Elm South at the end of the matching process, you will receive \$1.00 in addition to your \$4.00 participation payment.

< Back

Next >



Figure A.4: General instructions - page 4

A.2 Specific Instructions

The following four screenshots present the instruction page that differs between treatments. These pages also include color-coded symbols for specific lines (**left to the specific line**), indicating whether this line is unique for this treatment, or whether it is the same across treatments.

Description of the Matching Process

- Before the matching process starts, you will see a list of schools, their *thresholds*, and their associated *values*.
- Throughout the matching process we will send applications to schools on your behalf.
- If your *priority score* at a school is greater than or equal to that school's *threshold*, then that school will accept your application.

The matching process will proceed as follows:

- You will be prompted to rank any of the five schools in any order of your choice.
- At the first stage we will send an application to the top school, according to your ranking.
- If your *priority score* is lower than the school's *threshold*, your application will be rejected, and we will continue sending applications to schools according to the order you ranked them.
- This process will continue until an application is accepted, or until there are no more schools left on your list.
- Keep in mind that we will not send an additional application to a school that has previously rejected your application.
- Your final match will be the school that accepted your application. Your earnings for that round will be the *value* of that school.

In order to familiarize you with the matching process, you will now participate in a quick tutorial. Notice that **the sums of money in the tutorial are hypothetical**. During this short tutorial, we will walk you through the matching process step by step. *Explanations in this font and color* are a part of the tutorial. You will not see them during the real matching process.

When you are ready to start the tutorial press "Start Tutorial".

< Back
Start Tutorial
SP

Description of the Matching Process

- Before the matching process starts, you will see a list of schools, their *thresholds*, and their associated *values*.
- Throughout the matching process we will respond to offers from schools on your behalf.
- If your *priority score* at a school is greater than or equal to that school's *threshold*, then you will receive an offer from that school, at some point.

The matching process will proceed as follows:

- You will be prompted to rank any of the five schools in any order of your choice.
- At the first stage you might receive offers from some schools. We will keep the offer of the highest-ranked school among these, according to your ranking, and the rest will be rejected.
- At every subsequent stage you might receive additional offers from some schools. From among these offers and the offer you are currently holding, we will keep the offer of the highest-ranked school (according to your ranking) and reject the rest.
- This process will continue until there are no more offers.
- Keep in mind that a school whose offer has been rejected will not send you any more offers.
- Your final match will be the school whose offer you are holding at the end of the round. Your earnings for that round will be the *value* of that school.

In order to familiarize you with the matching process, you will now participate in a quick tutorial. Notice that **the sums of money in the tutorial are hypothetical**. During this short tutorial, we will walk you through the matching process step by step. *Explanations in this font and color* are a part of the tutorial. You will not see them during the real matching process.

When you are ready to start the tutorial press "Start Tutorial".

< Back
Start Tutorial
SR

Figure A.5: Specific instructions - SP and SR treatments

Description of the Matching Process

- Before the matching process starts, you will see a list of schools, their *thresholds*, and their associated *values*.
- Throughout the matching process you will send applications to schools.
- If your *priority score* at a school is greater than or equal to that school's *threshold*, then that school will accept your application.

The matching process will proceed as follows:

- At the first stage you can send an application to any school you choose.
- If your *priority score* is lower than the school's *threshold*, your application will be rejected, and you can continue sending applications to the remaining schools in any order you choose.
- This process will continue until an application is accepted, or until there are no more schools left. You can choose to stop the process at any point (in which case, you will not be matched with any school).
- Keep in mind that you cannot send an additional application to a school that has previously rejected your application.
- Your final match will be the school that accepted your application. Your earnings for that round will be the *value* of that school.

In order to familiarize you with the matching process, you will now participate in a quick tutorial. Notice that **the sums of money in the tutorial are hypothetical**. During this short tutorial, we will walk you through the matching process step by step. *Explanations in this font and color* are a part of the tutorial. You will not see them during the real matching process.

When you are ready to start the tutorial press "Start Tutorial".

< Back
Start Tutorial
DP

Description of the Matching Process

- Before the matching process starts, you will see a list of schools, their *thresholds*, and their associated *values*.
- Throughout the matching process you will respond to offers from schools.
- If your *priority score* at a school is greater than or equal to that school's *threshold*, then you will receive an offer from that school, at some point.

The matching process will proceed as follows:

- At the first stage you might receive offers from some schools. You can keep at most one of these offers and the rest will be rejected.
- At every subsequent stage you might receive additional offers from some schools. From among these offers and the offer you are currently holding, you can then keep at most one offer, and the rest will be rejected.
- This process will continue until there are no more offers.
- Keep in mind the that a school whose offer has been rejected will not send you any more offers.
- Your final match will be the school whose offer you are holding at the end of the round. Your earnings for that round will be the *value* of that school.

In order to familiarize you with the matching process, you will now participate in a quick tutorial. Notice that **the sums of money in the tutorial are hypothetical**. During this short tutorial, we will walk you through the matching process step by step. *Explanations in this font and color* are a part of the tutorial. You will not see them during the real matching process.

When you are ready to start the tutorial press "Start Tutorial".

< Back
Start Tutorial
DR

Figure A.6: Specific instructions - DP and DR treatments

A.3 Matching problem example

We provide here an example for a matching problem for each of the treatments. In each treatment the subject was introduced to the problem, and then interacted with a treatment interface, and then got learned the match outcome as well as the realization of priority scores in all schools.

For each participant we randomized the order of the matching problems. The order of appearance of schools was also randomized for each subject in each round. Last, we randomized (at the subject level) whether the Value column will be shown to the left or right of the Threshold column.

You will now participate in an additional matching process. The money you earn from holding a slot at a school at the end of this matching will be **added** to what you have earned so far in the study. The rules are identical.

Note
Schools' *values* and *thresholds* in this matching process are presented in the following table:

School	Value	Threshold	Chance that Your Priority Score $\geq$ Threshold
Fig Point	\$1.25	75	25%
Grapevine Pass	\$0.25	0	100%
Orange Terrace	\$1.00	25	75%
Citrus Park	\$0.75	20	80%
Juniper Square	\$0.50	15	85%

Press "Continue" to start the matching process

Continue

Round: 7/10



Figure A.7: matching problem example

Round: 7/10

Please rank any of the five schools in any order of your choice.
(click "Submit" when you are done):

Description of the Matching Process

Fig Point

Grapevine Pass

Orange Terrace

Citrus Park

Juniper Square

First-ranked School:

Choose your first-ranked school

Second-ranked School:

Third-ranked School:

Fourth-ranked School:

Fifth-ranked School:

Submit

School	Value	Threshold	Your Priority Score	Chance that Your Priority Score $\geq$ Threshold
Fig Point	\$1.25	75	?	25%
Grapevine Pass	\$0.25	0	?	100%
Orange Terrace	\$1.00	25	?	75%
Citrus Park	\$0.75	20	?	80%
Juniper Square	\$0.50	15	?	85%

SR

SP

The matching process has ended!

You ranked the schools in the following order: **Fig Point** , **Grapevine Pass**, **Orange Terrace**, **Citrus Park** and **Juniper Square**.

Your final match is: **Grapevine Pass (\$0.25)**

School	Value	Threshold	Your Priority Score
Fig Point	\$1.25	75	18 (Below threshold)
Grapevine Pass	\$0.25	0	3 (Above threshold)
Orange Terrace	\$1.00	25	23 (Below threshold)
Citrus Park	\$0.75	20	75 (Above threshold)
Juniper Square	\$0.50	15	94 (Above threshold)

Continue

Round: 7/10

SR

SP

Figure A.8: matching problem interface and results - SP&SR treatments

Please choose where to send your next application:

Description of the Matching Process

Fig Point

Grapevine Pass

Orange Terrace

Juniper Square

Citrus Park

Stop matching process (you will not match with any school)

School	Threshold	Value	Your Priority Score	Chance that Your Priority Score $\geq$ Threshold
Fig Point	75	\$1.25	?	Originally: 25%. Now: 0%
Grapevine Pass	0	\$0.25	?	Originally: 100%
Orange Terrace	25	\$1.00	?	Originally: 75%
Juniper Square	15	\$0.50	?	Originally: 85%
Citrus Park	20	\$0.75	?	Originally: 80%

Schools that have already rejected you: **Fig Point**

Round: 5/10

DP

The matching process has ended!

You sent applications to schools in the following order: **Fig Point** and **Grapevine Pass**.

Your final match is: **Grapevine Pass (\$0.25)**

School	Threshold	Value	Your Priority Score
Fig Point	75	\$1.25	40 (Below threshold)
Grapevine Pass	0	\$0.25	48 (Above threshold)
Orange Terrace	25	\$1.00	94 (Above threshold)
Juniper Square	15	\$0.50	61 (Above threshold)
Citrus Park	20	\$0.75	77 (Above threshold)

Continue

Round: 5/10

DP

Figure A.9: matching problem interface and results - DP treatment

New offer received!

Description of the Matching Process

The new offer is from:

Citrus Park

Remember: this means that your Priority Score at this school is higher than this school's Threshold.

If you want to switch to this offer, click on it.

Otherwise, you can choose to reject the new offer (and keep holding the offer you currently hold). **Reject**

Offer you currently hold: **Orange Terrace**

Offers you already rejected: **Grapevine Pass**

School	Threshold	Value	Your Priority Score	Chance that Your Priority Score $\geq$ Threshold
Citrus Park	20	\$0.75	?	Originally: 80%. Now: 100%
Orange Terrace	25	\$1.00	?	Originally: 75%. Now: 100%
Juniper Square	15	\$0.50	?	Originally: 85%
Fig Point	75	\$1.25	?	Originally: 25%
Grapevine Pass	0	\$0.25	?	Originally: 100%. Now: 100%

Round: 2/10

DR

The matching process has ended!

Your Priority Scores at **Orange Terrace**, **Grapevine Pass**, **Citrus Park** and **Juniper Square** were all higher than their Thresholds, and therefore you received offers from all of them during the matching process.

Your final match is: **Orange Terrace (\$1.00)**

School	Threshold	Value	Your Priority Score
Citrus Park	20	\$0.75	55 (Above threshold)
Orange Terrace	25	\$1.00	91 (Above threshold)
Juniper Square	15	\$0.50	33 (Above threshold)
Fig Point	75	\$1.25	29 (Below threshold)
Grapevine Pass	0	\$0.25	78 (Above threshold)

Continue

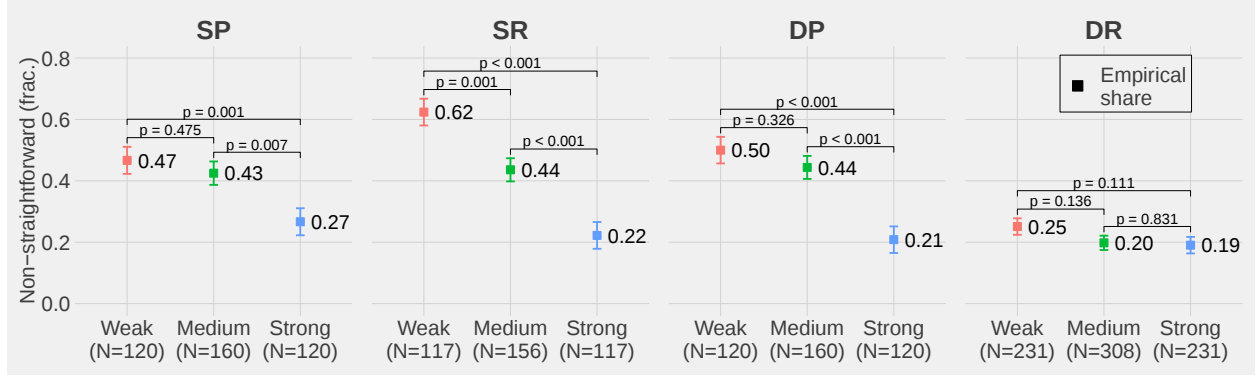
Round: 2/10

DR

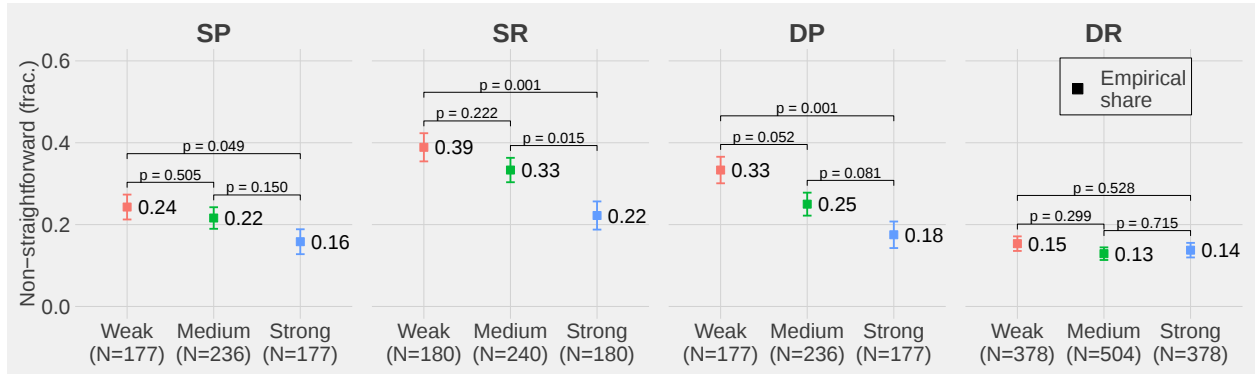
Figure A.10: matching problem interface and results - DR treatment

Figure B.11: Non-straightforward (NSF) behavior: Unpooled Samples

(a) Cornell BSL



(b) IPMM



Notes: Empirical NSF shares for Cornell BSL (Panel (a)) and IPMM (Panel (b)) Samples. Error bars: standard errors from a regression of NSF behavior on problem type, clustering at the individual level. p -values: Wald tests.

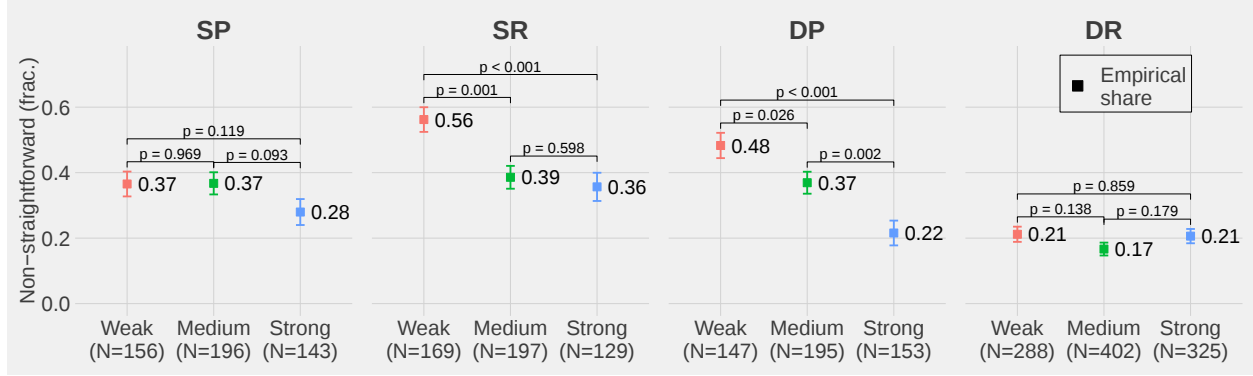
B Additional Analyses

B.1 Unpooled Samples

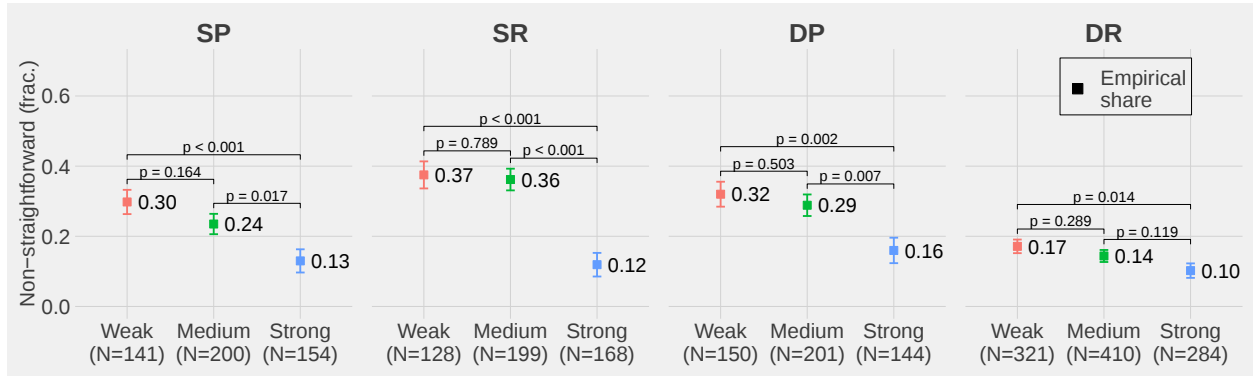
B.2 First vs. Last Five Rounds

Figure B.12: Non-straightforward (NSF) behavior: First vs. Last Five Rounds

(a) Rounds 1–5



(b) Rounds 6–10



Notes: Empirical NSF shares for first (panel (a)) and last (panel (b)) five rounds. Error bars: standard errors from a regression of NSF behavior on problem type, clustering at the individual level. p -values: Wald tests.

B.3 Individual Behavior

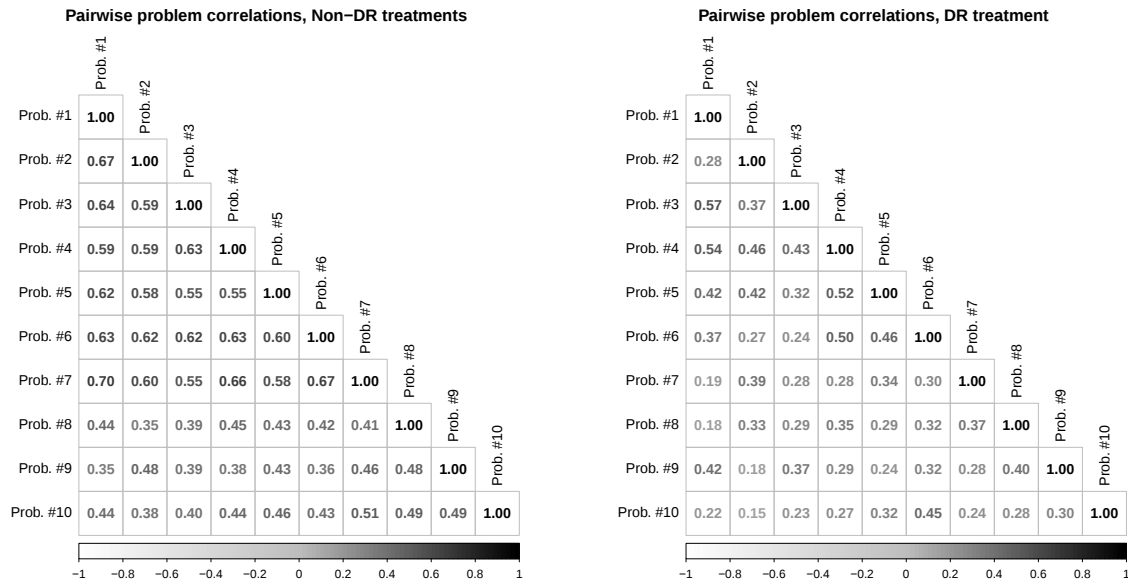
B.4 Behavior in DR

Table B.1: NSF by number of received offers in DR

# offers	# obs.	# NSF behavior obs.	% NSF behavior obs.
1	95	25	26%
2	294	51	17%
3	592	94	16%
4	743	122	16%
5	306	46	15%
1–5	2,030	338	17%

Notes: Distribution of DR observations by the number of offers the subject received ($N = 2,030$).

Figure B.13: Pairwise correlations in NSF across the ten problems



Notes: Each cell shows the correlation in NSF behavior (as a dummy variable) for a given pair of problems.

B.5 Regressions

Table B.2: The effect of interface and experience on NSF

Dependent Variable:	NSF				
<i>Variables</i>					
(Intercept)	0.17 (0.02)	0.10 (0.02)	0.05 (0.02)	0.01 (0.02)	0.03 (0.03)
SP	0.11 (0.04)	0.11 (0.04)	0.11 (0.04)	0.11 (0.04)	0.11 (0.04)
SR	0.19 (0.04)	0.19 (0.04)	0.19 (0.04)	0.19 (0.04)	0.19 (0.04)
DP	0.14 (0.04)	0.14 (0.04)	0.14 (0.04)	0.14 (0.04)	
Weak		0.13 (0.02)	0.13 (0.02)	0.13 (0.02)	0.11 (0.02)
Medium		0.08 (0.01)	0.08 (0.01)	0.08 (0.01)	0.06 (0.01)
English			0.12 (0.03)	0.12 (0.03)	0.11 (0.03)
First 5 rounds				0.09 (0.01)	0.09 (0.01)
Values first					-0.02 (0.03)
<i>Fit statistics</i>					
Observations	5,000	5,000	5,000	5,000	4,010
R ²	0.03	0.05	0.06	0.07	0.08
Adjusted R ²	0.03	0.04	0.06	0.07	0.07

Notes: DR is treated as baseline, and is omitted. Standard errors (clustered at the subject level) in parenthesis. “Values first” is a dummy variable denoting the order of “Value” and “Threshold” columns in the user interface (see Figure 1). It was not recorded in DP due to a data-logging bug.

Table B.3: NSF behavior by treatment

Dependent Variable: Matching Problems:	NSF				
	Weak	Medium	Strong	Weak and Medium	All
(Intercept)	0.19 (0.02)	0.16 (0.02)	0.16 (0.02)	0.17 (0.02)	0.17 (0.02)
SP	0.14 (0.05)	0.15 (0.04)	0.04 (0.04)	0.14 (0.04)	0.11 (0.04)
SR	0.29 (0.05)	0.22 (0.04)	0.06 (0.04)	0.25 (0.04)	0.19 (0.04)
DP	0.21 (0.05)	0.17 (0.04)	0.03 (0.04)	0.19 (0.04)	0.14 (0.04)
<i>Fit statistics</i>					
Observations	1,500	2,000	1,500	3,500	5,000
R ²	0.06	0.04	0.00	0.05	0.03
Adjusted R ²	0.06	0.04	0.00	0.05	0.03
$SP = SR = DP = 0$					
<i>p</i> -value	0.00	0.00	0.32	0.00	0.00

Notes: DR is treated as baseline, and is omitted. Standard errors (clustered at the subject level) in parenthesis.

Table B.4: NSF: difference-in-difference specification

Dependent Variable: Matching Problems:	NSF				
	Weak	Medium	Strong	Weak and Medium	All
(Intercept)	0.33 (0.04)	0.30 (0.04)	0.20 (0.03)	0.31 (0.04)	0.28 (0.03)
Dynamic	0.07 (0.06)	0.03 (0.06)	-0.01 (0.05)	0.04 (0.06)	0.03 (0.05)
Receiving	0.15 (0.06)	0.07 (0.06)	0.02 (0.05)	0.11 (0.06)	0.08 (0.05)
Dynamic × Receiving	-0.36 (0.08)	-0.25 (0.07)	-0.05 (0.06)	-0.29 (0.07)	-0.22 (0.06)
<i>Fit statistics</i>					
Observations	1,500	2,000	1,500	3,500	5,000
R ²	0.06	0.04	0.00	0.05	0.03
Adjusted R ²	0.06	0.04	0.00	0.05	0.03

Notes: Standard errors (clustered at the subject level) in parenthesis.

Table B.5: Costly NSF behavior by treatment

Dependent Variable: Matching Problems:	Costly NSF				
	Weak	Medium	Strong	Weak and Medium	All
(Intercept)	0.11 (0.01)	0.06 (0.01)	0.07 (0.01)	0.08 (0.01)	0.08 (0.01)
SP	-0.05 (0.02)	0.02 (0.02)	0.02 (0.02)	-0.01 (0.02)	-0.00 (0.02)
SR	-0.01 (0.02)	0.07 (0.02)	0.01 (0.03)	0.04 (0.02)	0.03 (0.02)
DP	0.01 (0.03)	0.07 (0.03)	0.05 (0.03)	0.05 (0.02)	0.05 (0.02)
<i>Fit statistics</i>					
Observations	1,500	2,000	1,500	3,500	5,000
R ²	0.00	0.01	0.00	0.01	0.00
Adjusted R ²	0.00	0.01	0.00	0.01	0.00
$SP = SR = DP = 0$					
<i>p</i> -value	0.05	0.01	0.35	0.03	0.07

Notes: DR is treated as baseline, and is omitted. Standard errors (clustered at the subject level) in parenthesis.

Table B.6: Costly NSF: difference-in-difference specification

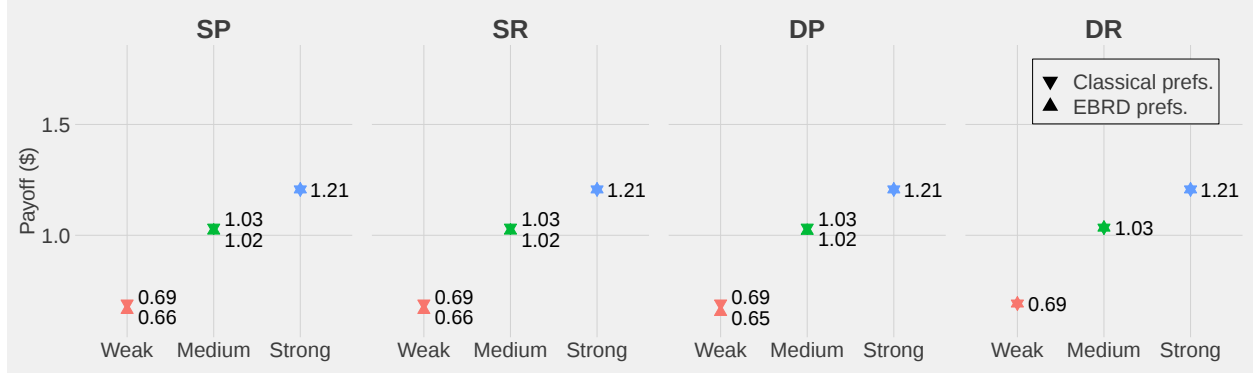
Dependent Variable: Matching Problems:	Costly NSF				
	Weak	Medium	Strong	Weak and Medium	All
(Intercept)	0.06 (0.02)	0.08 (0.02)	0.09 (0.02)	0.07 (0.01)	0.08 (0.01)
Dynamic	0.06 (0.03)	0.06 (0.03)	0.03 (0.03)	0.06 (0.03)	0.05 (0.02)
Receiving	0.04 (0.02)	0.05 (0.03)	-0.01 (0.03)	0.05 (0.02)	0.03 (0.02)
Dynamic × Receiving	-0.05 (0.04)	-0.13 (0.04)	-0.05 (0.04)	-0.09 (0.03)	-0.08 (0.03)
<i>Fit statistics</i>					
Observations	1,500	2,000	1,500	3,500	5,000
R ²	0.00	0.01	0.00	0.01	0.00
Adjusted R ²	0.00	0.01	0.00	0.01	0.00

Notes: Standard errors (clustered at the subject level) in parenthesis.

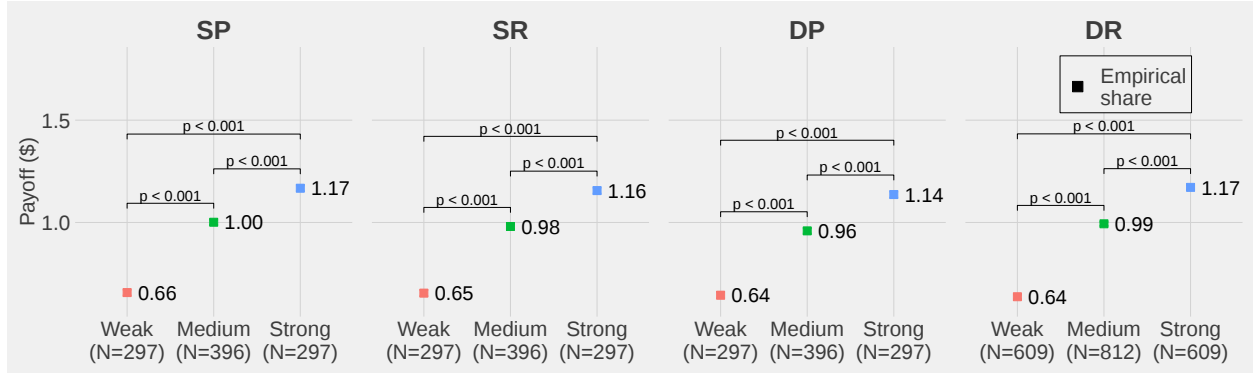
B.6 Total Earnings

Figure B.14: Total Payoff (\$)

(a) Theoretical Predictions: Classical & EBRD Preferences



(b) Results



Notes: Panel (a): Predicted payoff by treatment and problem type under classical preferences (downwards triangle) and EBRD preferences (upwards triangle). Panel (b): average (empirically observed) payoff, excluding participation fee. In order to compare earnings across samples, we use the same exchange rate of \$4 = 1NIS as was used to translate school values. Error bars: standard errors from a regression of earnings on problem type, clustering at the individual level. p -values: Wald tests.

C Parameters for the Experiment and Logit Predictions

When choosing the parameters for our experiment, our goal was to construct matching problems that generate high-powered tests for the EBRD model's predictions, both within and across-treatments. In addition, as explained in Online Appendix D, we originally planned to use the data from the static treatments to estimate a discrete-choice logit model and test EBRD against a $\lambda = 1$ benchmark. Therefore, when choosing our parameters we maximized the separation between the predictions of classical (reference-independent) random-utility logit and an EBRD random-utility model.

As we explain in Online Appendix D, we opted for a more transparent test that simply

compares the EBRD model with our previously estimated distribution of λ to classical preferences (with no noise). To rule out logit as a potential explanation of our findings (as we discuss in Section 3.6 in the main text), we show in Supplementary Material Appendix B that a classical ($\lambda = 1$) random-utility model predicts *higher* NSF shares in medium- and strong-student problems compared to weak-student problems, for a wide range of parameters, whereas in our data we find the opposite (a finding consistent with the EBRD-model’s predictions).

C.1 Search Algorithms

In this section we explain the procedures we used in choosing the five monetary values and thresholds (equivalently, admission probabilities) in each of our ten matching problems (the chosen parameters are listed in Table 2).

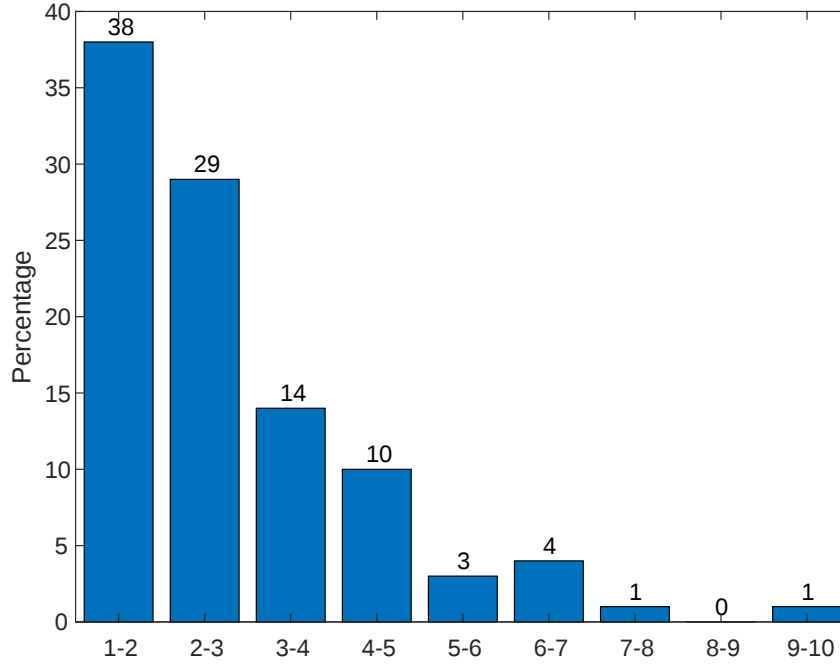
We used six different search algorithms, each optimizing a different objective. In each of the six algorithms, we first randomly sample 1,000 candidate matching problems, where each problem is composed of a threshold and a value for each of the five schools. Then, for each problem we calculate the EBRD-predicted behavior under the static variants (see Section 2) for every loss-aversion parameter λ over a grid $[1,10]$ (grid step = 0.1). We then aggregate behavior to generate aggregated predictions assuming a population distribution of λ based on Dreyfuss et al. (2022). Finally, we choose the optimal problem(s) (out of the 1,000 candidate problems), according to various criteria we explain below.

A set of parameters for each matching problem contains five school values and five admission probabilities. We used the following sampling restrictions:

- School values are drawn uniformly (without replacement) from the set $\{0.25, 0.5, 0.75, 1, 1.25, 1.5\}$.
- The probability of admission to the lowest value school is one.
- The rest of the probabilities of admission are drawn uniformly (without replacement) from $\{0.05, 0.1, \dots, 0.9, 0.95\}$, except in algorithms 3-5 (in which the probability for the highest value school is fixed at 0.2, 0.25, and 0.3, respectively).
- Probabilities are assigned to schools’ values so that a higher value school will always have a lower admission probability.

In algorithms 1–5, our predictions (for both classical preferences and EBRD) include errors drawn from an Extreme Value Type I distribution added to each alternative (i.e.,

Figure C.15



Notes: Population distribution of loss aversion, based on [Dreyfuss et al. \(2022\)](#).

each ROL). This allows us to directly calculate predicted ROL shares, for each λ value. We use a scaling parameter $x = \frac{1}{300}$ that effectively controls the variance of the error term.⁴⁴ We chose this value because it seemed to generate predictions consistent with data from previous work. Supplementary Material Appendix B shows the predictions for a range of scaling-parameters values. For EBRD predictions, we integrated the predicted shares using a distribution based on past estimates ([Dreyfuss et al., 2022](#)). The distribution we used is presented in Figure C.15.

A detailed description of our algorithm is in Supplementary Materials Appendix B. We briefly explain the objective behind each algorithm here: Algorithm 1 (matching problems #1–2) maximizes, for each ROL, the absolute difference between the predicted share under our λ distribution (with logit noise) and under $\lambda = 1$ (with logit noise that has the same scaling parameter). Algorithm 2 (matching problem #3) maximizes this distance for ROLs

⁴⁴Denote the utility of subject i from a strategy l by $u(\lambda_i, l)$. Then the probability of this subject choosing ROL l is $\pi_{\lambda_i}(l) = \frac{\exp\left(\frac{u(\lambda_i, l)}{x}\right)}{\sum_{k \in \mathcal{L}} \exp\left(\frac{u(\lambda_i, k)}{x}\right)}$, where x is a scaling parameter that sets the variance of the noise term. The aggregate probability of play under our distribution of loss aversion, $\pi_a(l)$, is calculated by integrating π_{λ_i} over λ using our population distribution distribution of λ .

that start with 3 (effectively maximizing the share of predicted ROLs that start with 3 under our λ distribution). Algorithms 3–5 (matching problems #4–7) create intermediate level of predicted NSF by maximizing this distance, fixing the probability of admission to the highest value school to 0.2, 0.25, and 0.3. Algorithms 6 (matching problems #8–10) *minimizes* the predicted NSF under our λ distribution (without noise) .

C.2 Logit vs. EBRD

Supplementary Material Appendix B shows that the choice probability of ROLs under $\lambda = 1$ and a range of scaling parameters, for the ten problems. It shows that in all of them, predicted NSF shares are higher in medium- and strong-student problems (problems #1–3 and #4–7, respectively) than in weak-student problems (#8–10).

D Main Text vs. Preregistration

We preregistered each experiment before running it. Both preregistrations contain our main between- and within-treatment predictions, and we collected the preregistered number of observations. However, the main-text analysis sometimes deviates from what we preregistered, both in content and in emphasis. In doing so, we follow the guidelines in [Banerjee et al. \(2020\)](#), who warn against strict adherence to pre-specified plans or the discounting of non-prespecified work. Instead, our main goal is to understand and explain our findings as best we can, and we therefore omit some of the preregistered analyses, and, in addition, include analyses that have not been preregistered. For completeness, in this appendix we highlight and explain these differences.

D.1 Preregistration #1: Cornell BSL

We submitted this preregistration before collecting any data (except small-scale pilot sessions). Below we highlight the main differences between our main-text analysis and this preregistration. We note that the updated analysis plan in preregistration #2 (discussed below) addresses many of the issues we are about to discuss, and is therefore much closer to our main-text analysis.

D.1.1 Between-treatments analysis:

In the preregistration, we emphasize costly and DO NSF as the more relevant measures for cross-treatment comparisons. In doing so, we overlooked the fact that we optimized our experiment to detect differences in NSF, and not in DO or costly NSF. We provide those comparisons in Section 3.5 in the main text, and highlight the fact that indeed, the EBRD-predicted differences in costly NSF (between DR and the other three treatments) and DO NSF (between DR and SR) are small. As we explain there, these smaller predicted differences, coupled with the presence of noise and NSF behavior not easily explained by loss aversion or classical preferences, make it impossible to detect cross-treatment differences in these measures.

Moreover, as we explain in our (non-preregistered) analysis in Section 3.4 in the main text, our findings suggest that using plain NSF behavior to compare between DR and the other three treatments likely does not suffer from the mechanical bias that initially worried us, and NSF likely provides relevant cross-treatment comparisons.

In addition, the preregistration includes a diff-in-diff specification that is meant to separately identify the effect of dynamic vs. static implementations relative to the effect of proposing vs. receiving implementation. We include this analysis in Online Appendix B. Table B.6's column 4 there shows a result consistent with our preregistered hypothesis: Using observations from the seven weak- and medium-student problems (as preregistered) where costly NSF is the dependent variable, the coefficient of interest on Dynamic×Receiving is large and statistically strong with the predicted negative sign (-0.09 ; $SE = 0.03$).⁴⁵ However, as discussed above, the main-text Figure 6 panel (b) implies a more nuanced and complicated relationship between the theory and the results for costly NSF (see the discussion above, and our discussion of Figure 6 in the main text). We therefore decided to place this diff-in-diff specification in the Online Appendix—in spite of it providing strong results, supportive of our hypotheses!—and include the more informative Figure 6 in the main text instead.

D.1.2 Within-treatment analysis:

The proposed analysis in the preregistration mentions the use of alternative measures (costly NSF, DO NSF, and total payoff), in addition to NSF, when comparing across problem types within each treatment. We do so in Section 3.5 in the main text. As explained above,

⁴⁵As expected, using the same specification with NSF as the dependent variable, the coefficient of interest is even larger, with the same sign (-0.29 ; $SE = 0.07$).

given the patterns of unexplained NSF behavior we observe, these measures are less informative also within treatments.

D.1.3 Using NSF types:

We originally planned to estimate the distribution of loss aversion and test the model fit by estimating a mixed logit model with Maximum Simulated Likelihood. To help with identification and precision, we chose the parameters of the experiment to create variation in the EBRD-predicted behavior across problems (based on a previously estimated population distribution of loss aversion). Moreover, we chose those parameters to make the predicted variation as distinct as possible from the variation predicted by (classical-preferences) logit. See Online Appendix C for further details on the process of choosing the parameters for the experiment. In the end, we opted for what we view as a cleaner and more transparent test of the theory against classical preferences: a zero-degrees-of-freedom test that compares the model’s predictions (based on the same previously estimated distribution, with no noise in decision making) to our empirical findings.

D.2 Preregistration #2: IPMM

We preregistered this experiment after preliminary analysis of the data from the Cornell BSL pool ($N = 196$), and prior to collecting any data from the IPMM ($N = 304$). This preregistration is much closer to our main-text analysis than our original preregistration. In particular, in this preregistration, we explicitly mention the low-power issue when using measures such as costly and DO NSF, and hence we place much greater emphasis on NSF. In addition, while this preregistration still mentions the possibility of estimating the discrete-choice mixed-logit model, we also explicitly discuss our ROL-types analysis, which we discuss in Section 3.3 in the main text.

Supplemental Material for “Deferred
Acceptance with News Utility” (NOT FOR
PUBLICATION)

A Alternative Measures

We consider in this analysis three different measures: NSF, costly NSF, and dynamically observable NSF (DO NSF). Online Appendix B considers total earnings as an additional relevant measure. Figure A.1 illustrates how measures are related to each other across treatments.

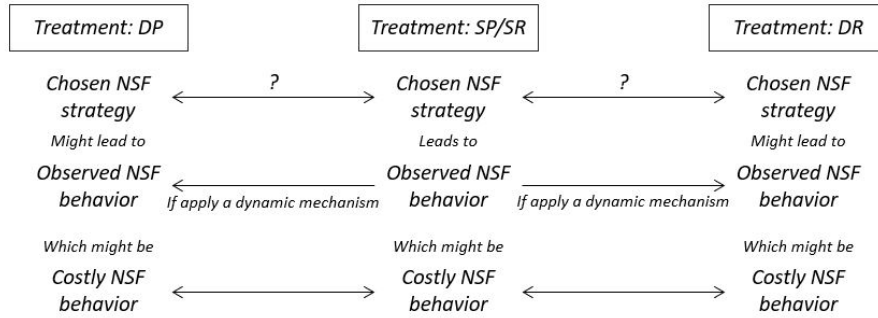


Figure A.1: *Measures across treatments*

Costly NSF behavior only considers outcomes and provides a crude measure across treatments. Total earnings extend this measure by considering *how* costly a given NSF behavior was. Dynamically observable (DO) NSF behavior takes the chosen strategy in a static treatment and asks what would have been the observed behavior under a dynamic implementation. For example, the ROL 21345 in the SP treatment would be seen under the DP treatment as a sequence of applications (conditional on rejections) of school 2, then school 1, then school 3, etc. We now explain how we created predictions for each of these measures.

A.1 Ex-ante predictions

Calculating these alternative measures ex-post is rather immediate: we use the realized priority scores to determine the highest-value school that was attainable and whether a ROL, implemented as a strategy, would have appeared SF-consistent under a dynamic implementation with identical realized priority scores.

Generating ex-ante predictions is slightly more complicated, so we provide some details here. First, we predict no NSF behavior in the DR treatment, which immediately results in a prediction of zero percent costly NSF behavior. To calculate the aggregate predicted share of costly NSF for the other three treatments, we first calculate the probability that each NSF ROL (or strategy defined by a ROL under the DP treatment) will be costly, i.e.,

for each ROL, we calculate the probability that the student’s final match has a lower value than the highest-value attainable school.

The algorithm we applied to generate the ex-ante predictions for a given ROL l is as follows:

1. For the highest-ranked school, calculate the probability of matching under l (in this case, the probability of exceeding the threshold) multiplied by the probability of exceeding the threshold in at least one higher value school (if one exists).
2. For the next highest-ranked school, calculate the probability of matching under l , multiplied by the probability of exceeding the threshold in at least one higher value school *that was not ranked higher on the ROL l* (if one exists).
3. Repeat step 2 until the ROL is exhausted.
4. Sum the probabilities to get the probability of ROL l being costly.

We then calculated for each $\lambda \in [1, 10]$ (with 0.1 intervals) the optimal ROL and used our previously estimated distribution of loss-aversion [Dreyfuss et al. \(2022\)](#) (see Online Appendix C) to create aggregate predictions. In particular, we get the probability of observing each ROL and multiply it by the probability of this ROL being costly.

For dynamically observable (DO) NSF behavior, the following two claims hold:

Claim 1. *Under SP, all predicted NSF ROLs are also dynamically observable.*

Claim 2. *Under SR, all predicted dynamically observable NSF ROLs are costly.*

We refer the curious reader to [Meisner and von Wangenheim \(2021\)](#), Proposition 1. The claims follow directly from their characterization.

B Experimental parameters: Algorithms and Figures

B.1 Algorithms

Algorithm 1 (matching problems #1–2) In this algorithm, the criteria is based on NSF behavior which cannot be explained by classical preferences ($\lambda = 1$) with errors. We wanted to capture how different the predicted ROL distribution under $\lambda = 1$ is from the one predicted assuming our λ distribution. To do that, we use the *sum of absolute*

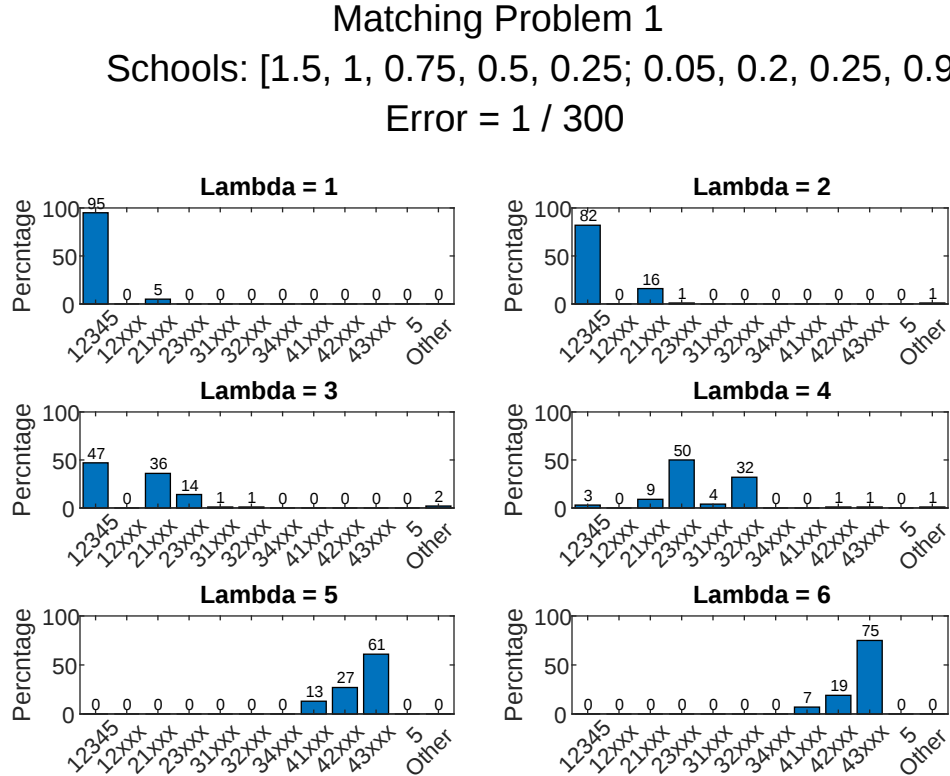
deviations across ROLs to measure the distance between the predictions under $\lambda = 1$ and the aggregate predictions under our λ distribution.

Denote ROL l 's probability of play (equivalently, predicted share) under $\lambda = 1$ by $\pi_1(l)$. Similarly, denote the aggregate probability of play (/predicted share) under our distribution of loss aversion by $\pi_a(l)$. Our measure is given by:

$$\sum_{l \in \mathcal{L}} |\pi_1(l) - \pi_a(l)|$$

where $\mathcal{L}$ contains all ROLs that induce different payoff distributions (i.e., the ROLs 1235 and 12354 are treated as one ROL since they induce identical payoff distributions). Figure B.2 shows an example of predicted distributions, for six different loss-aversion parameters.

Figure B.2



Notes: Predicted distribution of ROLs in matching problem #1 for error=1/300 and various λ 's, binning together non-12345 ROLs are binned together by the two top-ranked schools, with "xxx" representing the rest of the ROL.

As Figure B.2 shows, with our scaling parameter, under $\lambda = 1$, the predicted share of 12345 is 0.95, and the remaining 0.05 is predicted to be 21xxx (e.g., 21345). Naturally, as λ increases, the predicted NSF share increases. As mentioned above, we chose a scale

parameter $\frac{1}{300}$ to create predictions that seemed calibrationally reasonable, given prior empirical evidence.

We then created, for each model ($\lambda = 1$ and λ distributed according to our distribution), predictions for each of the randomly drawn 1000 matching problems. Finally, we chose the two matching problems that had the highest predicted sum of absolute deviations.⁴⁶

Algorithm 2 (matching problem #3) This algorithm is identical to Algorithm 1, with one important difference: when calculating the predicted sum of absolute deviations, we sum only over ROLs that rank school 3 on top (e.g., 32145). Denote $\mathcal{L}_3 \subset \mathcal{L}$ as the set of ROLs for which 3 is ranked on top; then, our measure is given by:

$$\sum_{l \in \mathcal{L}_3} |\pi_1(l) - \pi_a(l)|.$$

Since these ROLs are rarely predicted for $\lambda = 1$,⁴⁷ this effectively maximizes the predicted share of ROLs that rank school 3 on top under EBRD. This created another problem-specific EBRD prediction, which we used as an additional test of the model (see Section 3.3).

Algorithms 3–5 (matching problems #4–7) These algorithms are identical to Algorithm 1, with one important difference: We kept the probability of admission to the highest value school fixed at 0.2, 0.25, or 0.3, for each Algorithm 3, 4, 5, respectively. The criterion is the same.

The reason we wanted to include problems with a higher probability of admission to the highest value school was twofold: First, we wanted to include problems such that under DR, subjects will often face the highest value school (while maintaining a moderately high predicted NSF share under the non-DR treatments). In other words, we wanted to give subjects opportunities to exhibit behavior consistent with ROLs such as 21345 under DR. With $q_1 \in [0.2, 0.3]$, the chance of receiving an offer from school 1 in DR is non-negligible, yet NSF is still predicted (in the non-DR treatments) for subjects with $\lambda > \bar{\lambda} \in [4.13, 5.08]$ (see Proposition 2). Second, as an additional test of the theory, we

⁴⁶The two matching problems selected using this algorithm had lowest possible probability of admission to the highest value school (0.05), consistent with Proposition 2. This proposition shows that the probability of admission to the highest value school, q_1 , is directly related to the threshold of λ for which SF is no longer the optimal strategy, and consequently a low q_1 creates a markedly different predicted ROL distribution under $\lambda = 1$ relative to the predictions under our loss-aversion distribution.

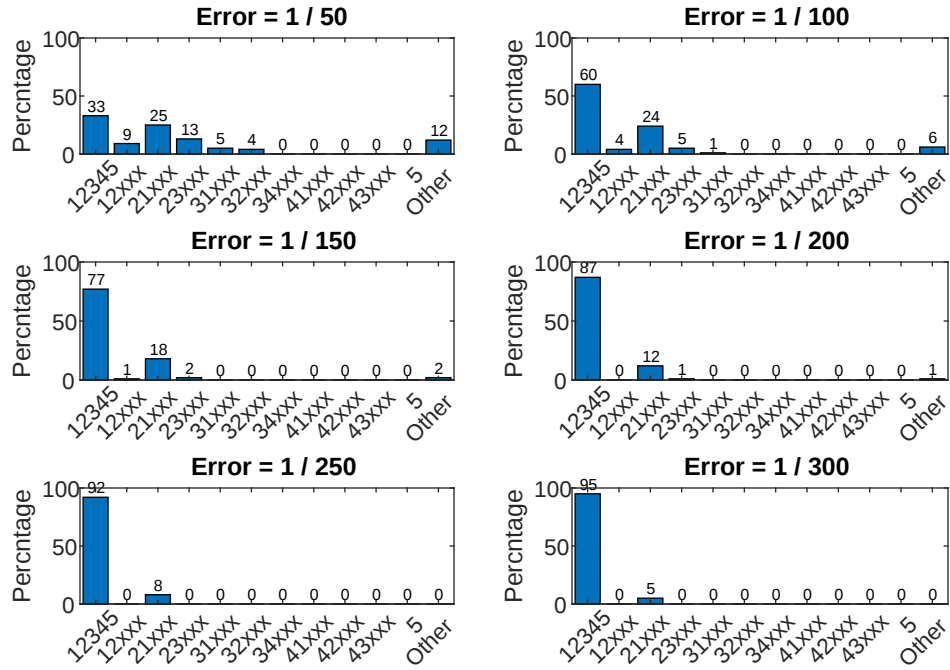
⁴⁷This is true not only under our specific scaling parameter: below we show that these ROLs are rarely predicted for a wide range of scaling parameters.

wanted to create an intermediate level of predicted NSF share that will be lower than the one predicted in problems #1–3, yet nonnegligible. Effectively, this generated three main levels of predicted NSF, corresponding to our three different problem types composing our main within-treatment test of the model (Section 3.2).

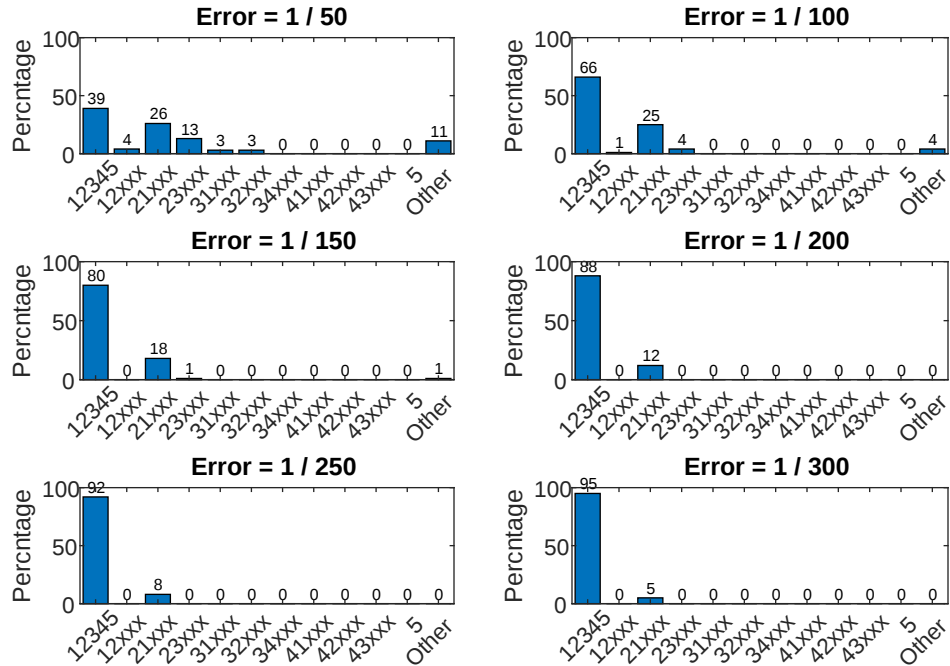
Algorithm 6 (matching problems #8–10) This algorithm chose problems with the *lowest* EBRD-predicted NSF share (without added noise). As expected, the chosen matching problems have high (0.6–0.65) admission probability to the highest value school, which resulted in a zero-predicted NSF share under our distribution of λ .

B.2 Logit Figures

Matching Problem 1 $\Lambda = 1$

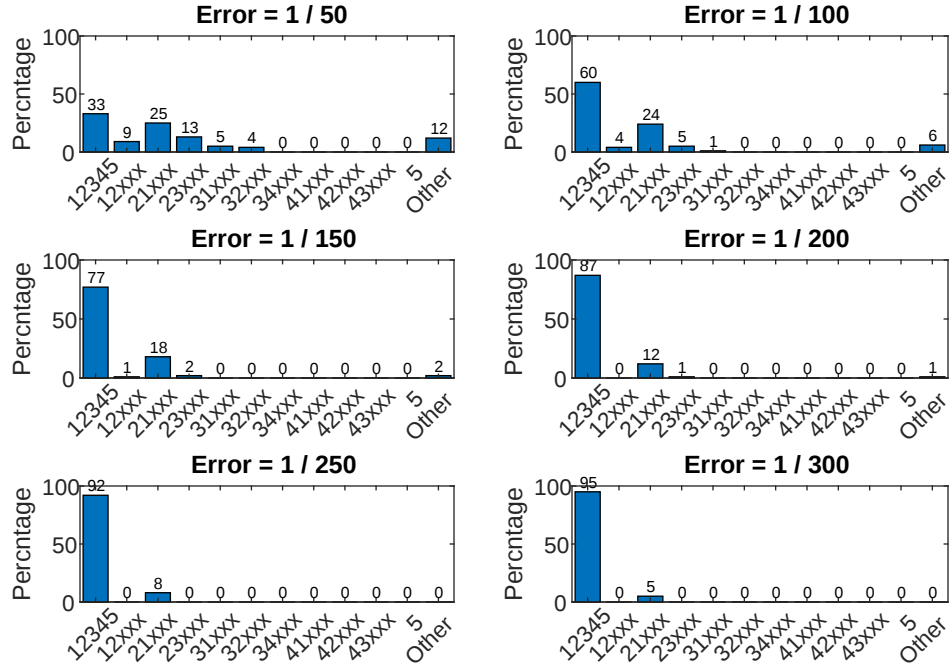


Matching Problem 2 $\Lambda = 1$

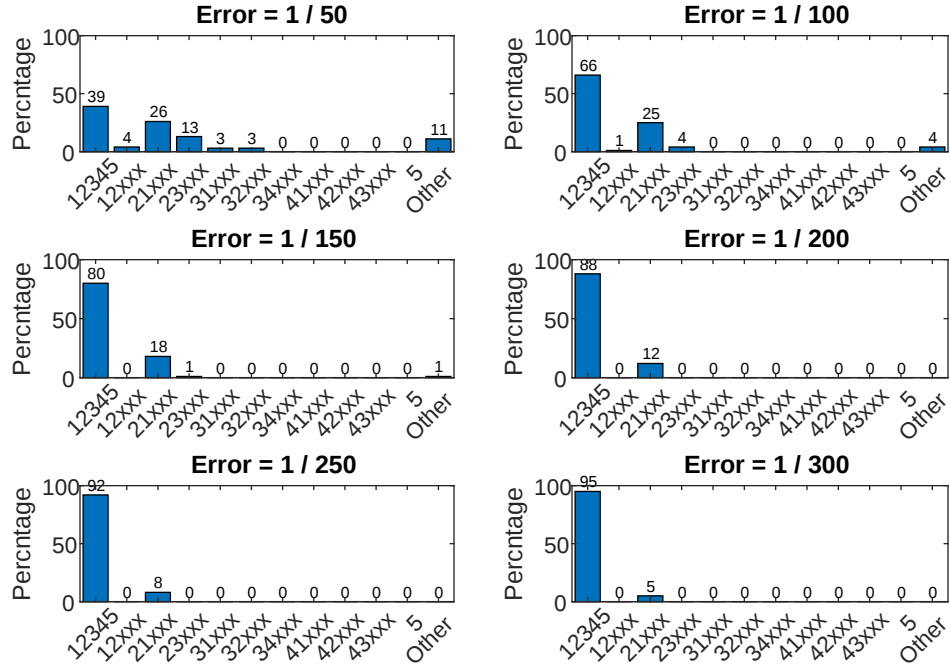


Notes: Predicted distribution of ROLs in matching problem #1 and #2 for $\lambda = 1$ and various error terms, binning together non-12345 ROLs are binned together by the two top-ranked schools, with “xxx” representing the rest of the ROL.

Matching Problem 1 $\Lambda = 1$

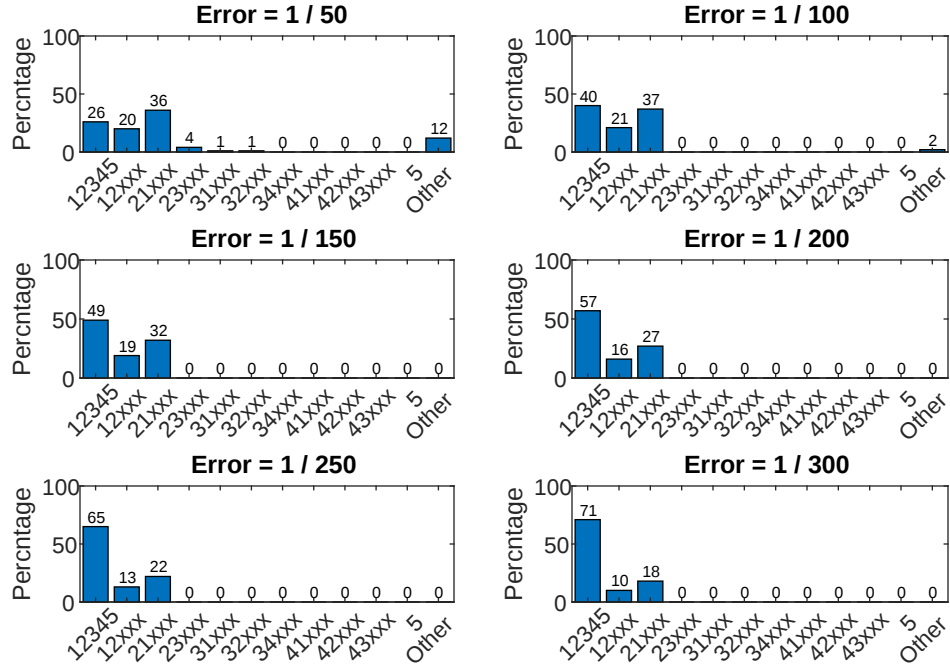


Matching Problem 2 $\Lambda = 1$

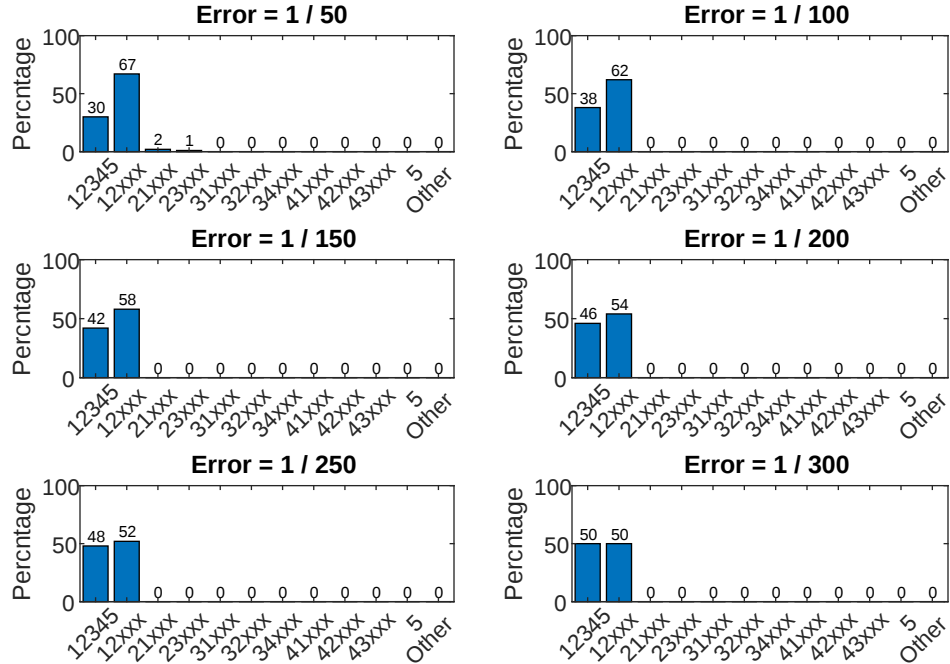


Notes: Predicted distribution of ROLs in matching problem #1 and #2 for $\lambda = 1$ and various error terms, binning together non-12345 ROLs are binned together by the two top-ranked schools, with “xxx” representing the rest of the ROL.

Matching Problem 3 Lambda = 1

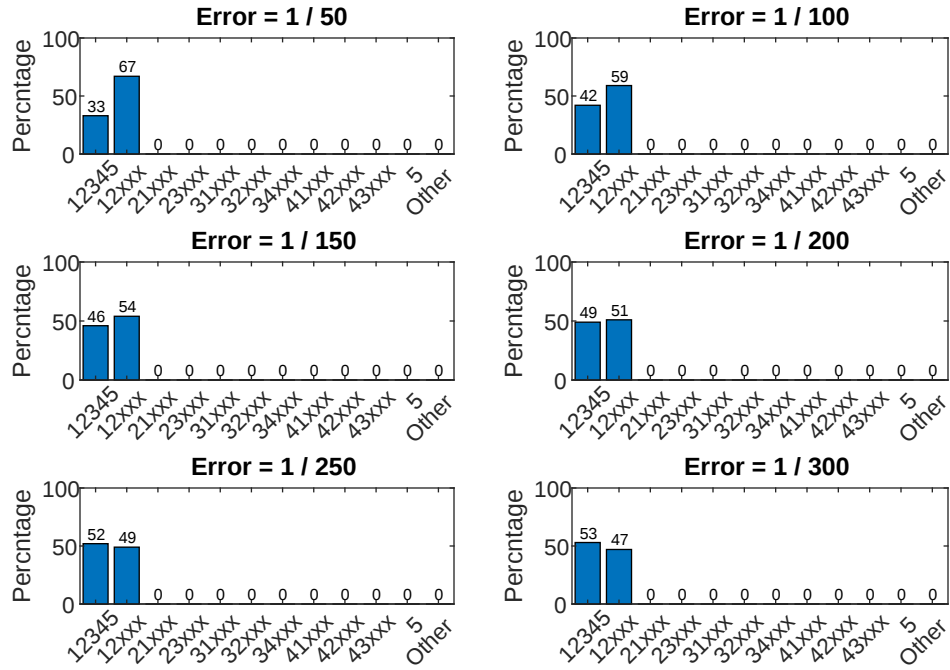


Matching Problem 4 Lambda = 1

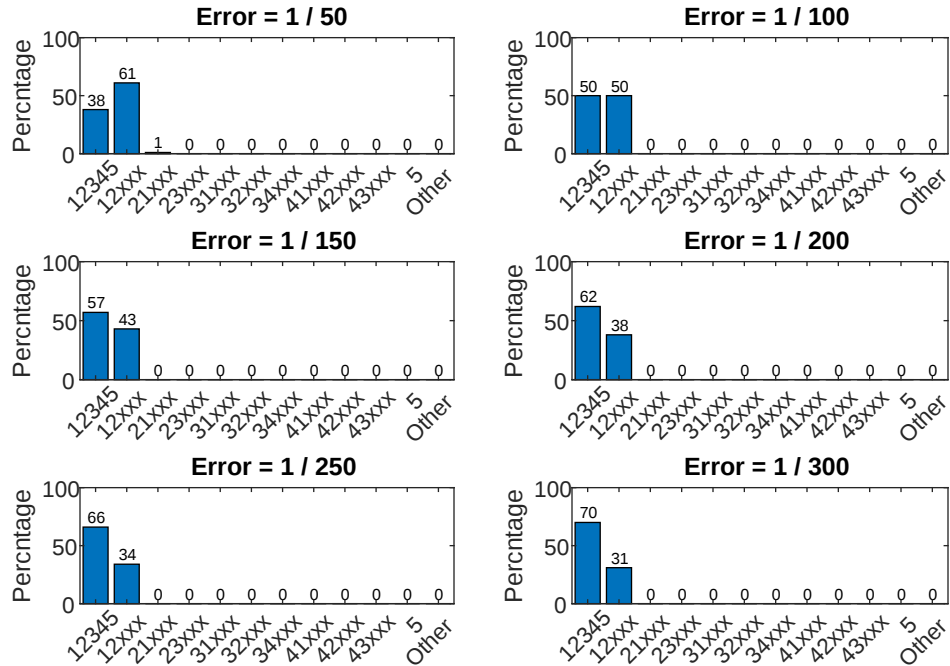


Notes: Predicted distribution of ROLs in matching problem #3 and #4 for $\lambda = 1$ and various error terms, binning together non-12345 ROLs are binned together by the two top-ranked schools, with “xxx” representing the rest of the ROL.

Matching Problem 5 Lambda = 1

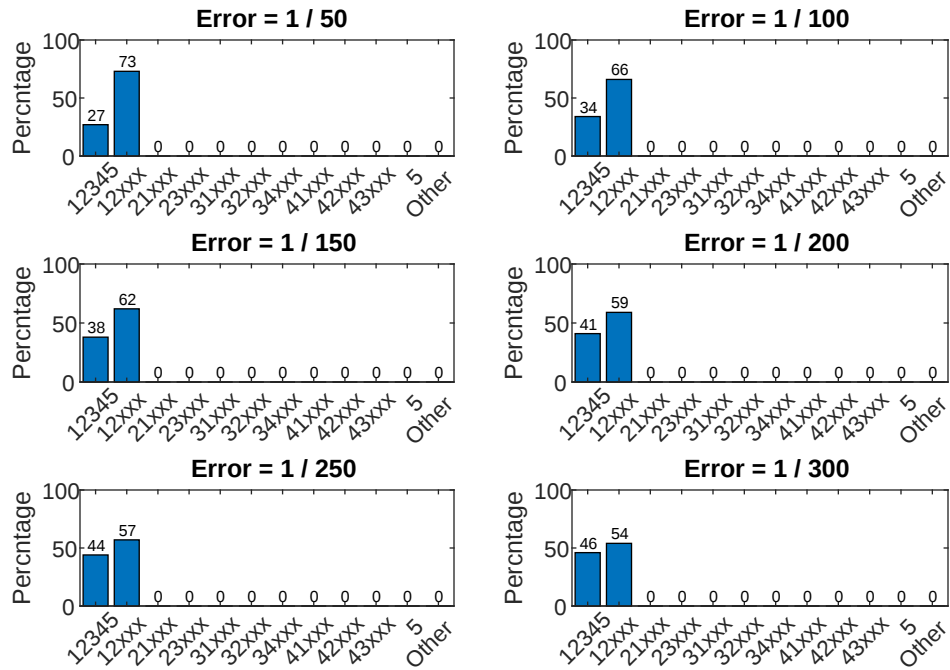


Matching Problem 6 Lambda = 1

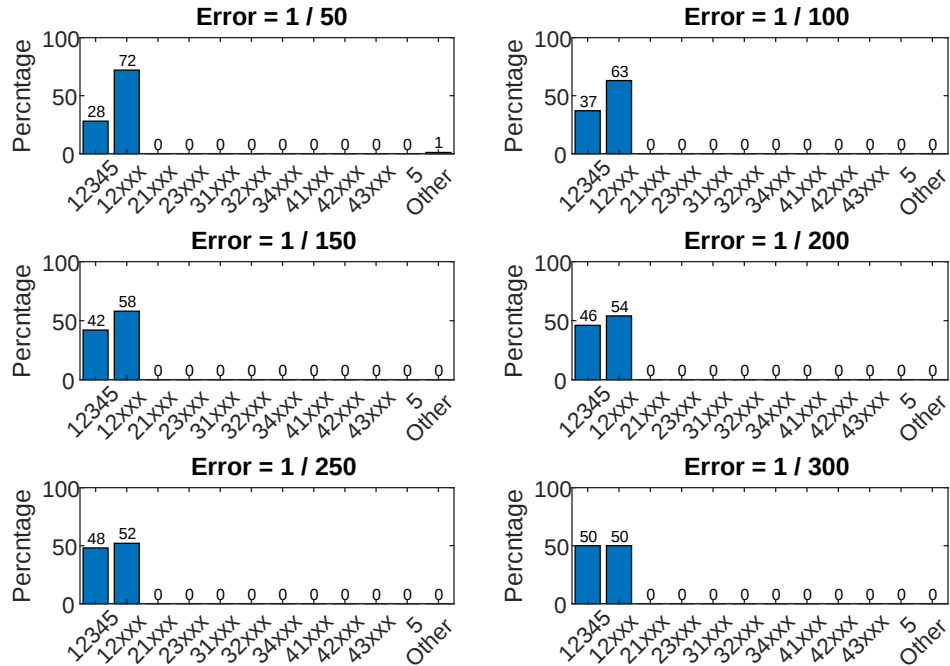


Notes: Predicted distribution of ROLs in matching problem #5 and #6 for $\lambda = 1$ and various error terms, binning together non-12345 ROLs are binned together by the two top-ranked schools, with “xxx” representing the rest of the ROL.

Matching Problem 7 $\Lambda = 1$

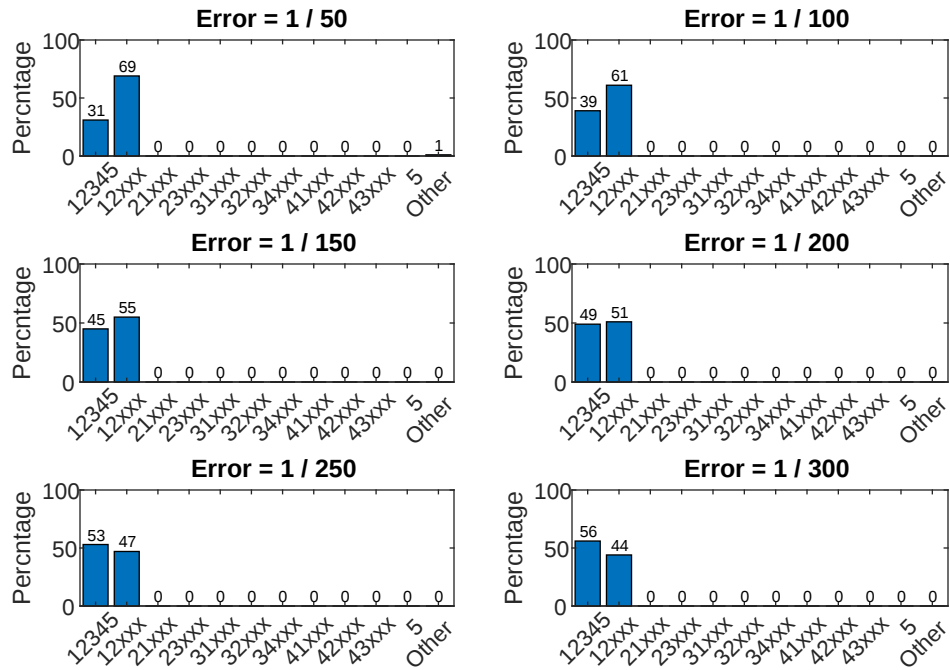


Matching Problem 8 $\Lambda = 1$

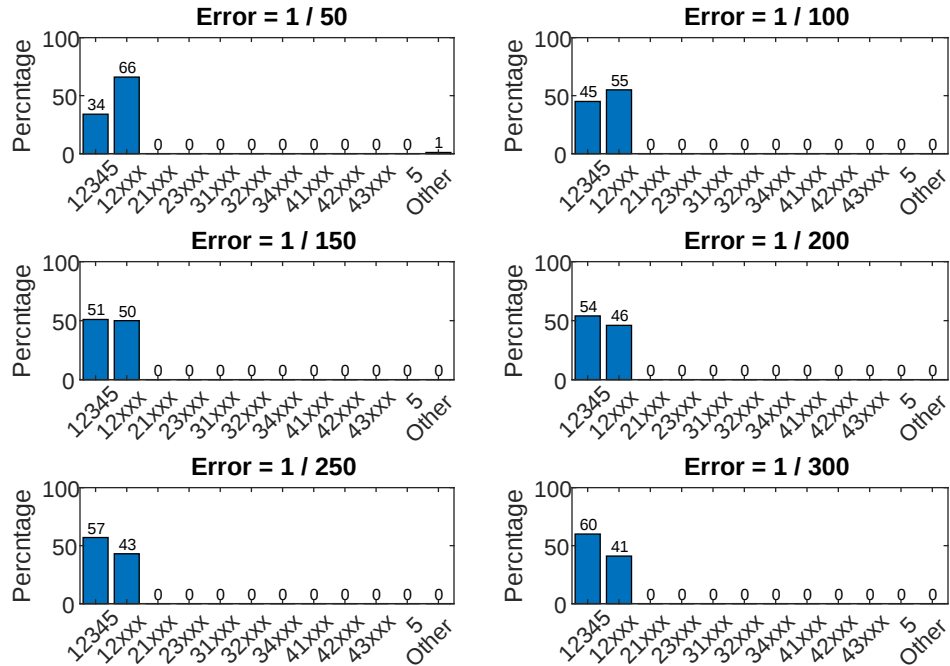


Notes: Predicted distribution of ROLs in matching problem #7 and #8 for $\lambda = 1$ and various error terms, binning together non-12345 ROLs are binned together by the two top-ranked schools, with “xxx” representing the rest of the ROL.

Matching Problem 9 Lambda = 1



Matching Problem 10 Lambda = 1



Notes: Predicted distribution of ROLs in matching problem #9 and #10 for $\lambda = 1$ and various error terms, binning together non-12345 ROLs are binned together by the two top-ranked schools, with “xxx” representing the rest of the ROL.