

An Experiment on Behavior in Queues*

Mariagiovanna Baccara[†] SangMok Lee[‡] Brian Rogers[§] Chen Wei[¶]

July 1, 2025

Abstract

In a dynamic allocation setup, we experimentally study agents' choices under the First-In-First-Out (FIFO) and Last-In-First-Out (LIFO) queuing protocols. We find that agents are nearly rational under FIFO but tend to be overselective under LIFO. Within the model that anchors our experiment, we show that the magnitude of such an excessively selective bias reduces the welfare performance gap between FIFO and LIFO. As strategic complementarities under LIFO can reinforce overselective behavior, we show that such bias persists in a supplemental treatment where subjects interact with non-strategic robots.

Keywords: Dynamic Matching, Priority Protocols, Queuing.

*We gratefully acknowledge financial support from a Weidenbaum Center Impact Grant.

[†]Washington University in St. Louis, mbaccara@wustl.edu

[‡]Washington University in St. Louis, sangmoklee@wustl.edu

[§]Washington University in St. Louis, brogers@wustl.edu

[¶]Fannie Mae, cwei.phd21@gmail.com

1 Introduction

1.1 Overview

We study a dynamic allocation setting where agents match with items over time. Examples of these environments are widespread, including the assignment of public housing, medical appointments, deceased donors’ organs, etc. The theoretical literature has compared alternative queuing protocols, assessing their relative efficiency performances. In particular, much attention has been given to the First-In-First-Out (FIFO) and the Last-In-First-Out (LIFO) protocols. The FIFO protocol is the most common, as priority is often tied to arrival times. For example, a first-come-first-served rule is used by restaurants to allocate diners to tables, or by universities to assign course slots to students. However, implementing LIFO is more challenging, as it is subject to manipulation through exit-and-reenter and is considered “unfair,” since individuals who just arrived are served first (see Kahneman et al. (1986), Margaria (2024), and references therein). Nonetheless, LIFO naturally arises in some applications where agents not only choose but are also chosen to be matched. For example, in child adoption, adoptive parents tend to prefer younger children (see Baccara et al. (2014)). In such cases, birth mothers of later-born children have priority in the choice of adoptive parents, resulting in a setting equivalent to a LIFO protocol. Similar considerations apply to perishable goods and to environments where the desirability of a match decreases with the time spent on the market (e.g., unemployed workers or houses on the real estate market).

Since an agent’s priority deteriorates over time under LIFO, but improves under FIFO, agents are expected to be less selective under LIFO than under FIFO. This difference in queuing behavior implies mixed efficiency comparisons depending on the theoretical model under consideration. While a more selective behavior under FIFO yields the benefits of increased market thickness, it also imposes well-known negative queuing externalities when waiting costs are present, and could result in the LIFO protocol being socially more desirable.

The relative performances of these protocols in real life also depend on how well the individuals’ behavior approximates the canonical rationality assumption underlying the theoretical models. To explore this subject, in this paper, we design an experiment to study agents’ behavior in FIFO and LIFO queues. Exploiting the model that anchors the experiment, we also evaluate the efficiency implications of the behavioral patterns observed in the experiment.

In particular, we build a dynamic allocation model with no waiting costs, which is well-suited for experimental investigation. At every period, one agent and one item arrive at

the market. Since agents have a lifespan of two periods, any agent on the market is either “young” or “old.” Items’ values are common across agents and drawn independently from the same uniform distribution. An item can be offered to an agent only if it is *compatible* with her, and compatibility is independent across items and agents. At any period, an arriving item is offered to agents according to either a FIFO or a LIFO protocol, conditional on compatibility. If the item is accepted by an agent, the agent and the item leave the market matched. If the item remains unmatched, it is discarded at the end of its arrival period. Also, any agent still unmatched at the end of her two-period lifespan leaves the market.

First, we assume that agents behave rationally and study the efficiency properties of the rational equilibrium under the FIFO and the LIFO protocols. In a rational equilibrium, a young agent accepts an item if and only if its value exceeds the agent’s continuation value from waiting. We show that, under FIFO, the rational equilibrium threshold is higher than the socially optimal one, while under LIFO, it is lower. Also, in our setting, the rational equilibrium under FIFO is more efficient than the one under LIFO.

Next, we consider *behavioral equilibria*, in which young agents still follow a threshold strategy, but the threshold is allowed to be different from their continuation payoff from waiting—that is, agents are allowed to be either underselective or overselective with respect to the rational benchmark. We introduce the notion of *Minimal Queuing Rationality* (or *MQ-rationality*). Specifically, an agent is MQ-rational if (i) she recognizes that her future prospect is worse than always being offered an item, so she accepts any item valued more than the items’ average, and (ii) she understands her prospects deteriorate more under LIFO than FIFO, making her relatively less selective under LIFO. In a result that ties directly to our experimental findings, we show that if agents behave rationally under FIFO, and as long as they are MQ-rational, the efficiency gap between the FIFO and LIFO protocols first decreases and then increases again as agents become increasingly overselective under LIFO.

We design a laboratory experiment featuring reduced-form settings that are strategically equivalent to the FIFO and LIFO protocols of our model. We find that the subjects by and large satisfy MQ-rationality, and their behavior under FIFO converges quickly to the rational benchmark. However, under LIFO, the subjects remain consistently overselective with respect to the rational benchmark. We show that, within our model, the magnitude of this bias reduces the efficiency gap between FIFO and LIFO, bringing the efficiency under LIFO closer to the efficiency generated by the (approximately rational) behavior under FIFO.

To explore the causes of the overselective behavior under LIFO, we conduct a supplementary experiment and some additional analysis. First, note that the LIFO setup displays

strategic complementarities. Specifically, if a subject expects others to be overselective, the best response is to be overselective herself. To offset this strategic effect, in a supplemental treatment, subjects play the reduced-form LIFO game against a non-strategic robot. We find that overselectivity persists in this setting, though to a lesser extent. To explain the residual overselectivity, we apply some well-known behavioral models and we find that the Quantal Response Equilibrium framework, combined with mimicking behavior, provides the best, although not conclusive, fit for our findings.

To our knowledge, this is the first attempt to empirically evaluate the rationality of agents in FIFO and LIFO queues. The main takeaway of our analysis is that, in an experiment anchored by a simple theoretical model, agents are approximately rational under FIFO but overselective under LIFO with respect to the rational benchmark. The goal of our model is to provide a tractable queuing setup delivering robust and testable results. Specifically, in our model, a single parameter measures agents’ rational and behavioral decisions, allowing us to assess and compare the rationality of agents across alternative queuing protocols. We stress that the aim of our theoretical analysis is not to draw a general conclusion on the welfare implications of our experimental results. In fact, it is already well known that welfare rankings among queuing protocols are highly dependent on the environment. Therefore, the overselectivity we uncover under LIFO can have different welfare implications depending on the model at hand. As an illustration, we evaluate the welfare implications of the overselective bias in the model anchoring our experiment. If we had introduced substantial waiting costs into our model, the theoretical efficiency ranking between the protocols would have reversed because, as mentioned above, queuing externalities would arise under FIFO. Then, a moderate overselective bias under LIFO, which still improves the protocol’s performance, could potentially *increase* the welfare gap between FIFO and LIFO.

1.2 Related Literature

Dynamic matching problems, in particular dynamic allocation ones in which agents match with items over time, have received significant attention in the theoretical queueing and economics literature. For a survey, see Baccara and Yariv (2021) and references therein.

Starting from Naor (1969), many papers in this literature have explored the trade-offs associated with alternative priority protocols. In settings with waiting costs, agents under FIFO tend to wait inefficiently long because they do not internalize the negative externalities on others when they decide to stay in the queue rather than leave. In contrast, LIFO rules out such externalities, often achieving socially superior outcomes (see also Hassin (1985);

Hassin and Haviv (2003); Su and Zenios (2004), and Baccara et al. (2020)).

Generally, welfare rankings across protocols vary depending on the specific model under consideration. Notably, Bloch and Cantala (2017) study a dynamic allocation model with a fixed-length queue, and they show that FIFO can be optimal among protocols assigning some priority to agents who arrived first (i.e., mixtures between FIFO and uniform random priority, hence excluding LIFO). The negative externality pointed out by Naor (1969) does not exist in Bloch and Cantala (2017) because the number of agents in the queue remains the same, regardless of each agent’s decision. More recent work on the optimal design of queuing protocols includes Ashlagi et al. (2025) and Che and Tercieux (forthcoming).

The experimental and empirical literature on queuing behavior is sparse, and to our knowledge, it has not addressed behavioral patterns across protocols. Conte et al. (2014) consider an experiment where agents under time pressure select a FIFO queue among multiple ones differing in length, server speed, and entry fee. Dold and Khadjavi (2017) estimate the willingness to pay for a more favorable queue slot under FIFO. Kremer and Debo (2012) experimentally study queue-joining behavior when service quality is uncertain and, therefore, herding may arise.¹ On the empirical side, Batt and Terwiesch (2015) and Chan et al. (2017) document the impact of delays on queue abandonment and outcomes in medical emergency departments.

2 The Model

2.1 Setup

We consider a discrete-time, infinite-horizon, overlapping-generation matching market between agents and items. At the beginning of each period, $t \in \{1, 2, \dots\}$, one agent and one item arrive on the market. While each item must leave at the end of its arrival period, an agent can stay on the market for up to two periods before leaving. Hence, we refer to agents in their first and second periods on the market as “young” and “old,” respectively. When two agents are present simultaneously, they are ranked according to their arrival times by a first-in-first-out (FIFO) or a last-in-first-out (LIFO) protocol.

Upon arrival, each item’s value θ is drawn independently across items from the uniform distribution over $[0, 1]$, and it becomes publicly known. Any agent matched with the item

¹Moreover, Wang and Zhou (2018) and Rosokha and Wei (2024) examine how server behavior can be influenced by the queue configuration. Implementing dedicated queues for individual servers, as opposed to a single shared queue, can improve service times.

receives the payoff θ . When an agent leaves the market unmatched – that is, without having obtained any item – she receives a zero payoff. In addition, an item is compatible with an agent with probability $p \in (0, 1)$, independently across items and agents. An agent can be offered an item only if the item is compatible with her. We assume that there is no discounting or waiting cost to stay on the market.

At every period, once an item of value θ enters the market, the first-ranked agent chooses whether to accept the item or not, conditional on compatibility. If she accepts, the agent and the item leave the market. If the item is not compatible with the first-ranked agent, or if it is compatible but the agent rejects it, the item is offered to the second-ranked agent (if there is any), conditional on compatibility. If the match occurs, the second-ranked agent and the item leave the market. Otherwise, the item is discarded. In addition, at the end of each period, if an old agent is still present and is unmatched, she leaves. Therefore, one item and either one or two agents leave the market at each period, either because they match or because their lifespan ends.

For simplicity, we assume that an agent indifferent between accepting and rejecting an item always accepts it. Our analysis remains unchanged under any other tie-breaking rule.

2.2 Rational and Behavioral Equilibrium

We study stationary Markov perfect equilibria, in which agents' strategies only depend on their age (young or old) and their rank in the queue. Since an old agent prefers accepting any compatible item to leaving the market unmatched, an equilibrium is fully described by young agents' behavior. While the rational equilibrium requires the minimum acceptable item's value to coincide with an agent's continuation payoff, we also consider behavioral equilibria, in which agents' thresholds can differ from it.

2.2.1 First-In-First-Out

In the FIFO protocol, an old agent is always ranked first, so a young agent's continuation payoff from waiting is independent of her ranking upon arrival. Therefore, an equilibrium is identified by the minimum acceptable value for any young agent, regardless of her initial ranking.

At the beginning of each period, there are either two agents (one old and one young, ranked first and second, respectively) or only one (young) agent on the market. Let $\{v_o^F, v_{y2}^F, v_{y1}^F\}$ denote the expected payoffs of the agents in these scenarios at the beginning of a period,

before the item's value is realized. Suppose that a young agent uses a threshold strategy $\hat{\theta} \in [0, 1]$. If there are two agents, one old and one young, their expected payoffs are:

$$\begin{aligned} v_o^F &= \frac{p}{2}, \\ v_{y2}^F &= p(1 - \hat{\theta}) \left[(1 - p) \frac{1 + \hat{\theta}}{2} + p v_o^F \right] + [1 - p(1 - \hat{\theta})] v_o^F. \end{aligned} \quad (1)$$

Since the old agent is ranked first and accepts any compatible item, his expected payoff is simply $v_o^F = \frac{p}{2}$. The young agent is ranked second, and the item is both compatible with her and acceptable with probability $p(1 - \hat{\theta})$. The item can be offered to the young agent only if it is incompatible with the old agent, which occurs with probability $1 - p$, yielding an expected value $\frac{1 + \hat{\theta}}{2}$ for the young agent. Otherwise, the young agent remains unmatched, with continuation payoff v_o^F .

If a young agent enters the market when no old agent is present, she is ranked first and her expected payoff is:

$$v_{y1}^F = p(1 - \hat{\theta}) \frac{1 + \hat{\theta}}{2} + [1 - p(1 - \hat{\theta})] v_o^F. \quad (2)$$

In words, if the item is compatible and has an acceptable type, which occurs with probability $p(1 - \theta)$, the agent matches immediately and obtains the expected payoff of $\frac{1 + \hat{\theta}}{2}$. Otherwise, the young agent stays on the market with continuation payoff v_o^F . We denote the rational Markov perfect equilibrium threshold and expected payoffs under FIFO by $(\bar{\theta}^F, \bar{v}_o^F, \bar{v}_{y2}^F, \bar{v}_{y1}^F)$, where $\bar{\theta}^F = \bar{v}_o^F = \frac{p}{2}$.

Next, we consider a behavioral equilibrium where a young agent's decision can differ from the rational choice. Namely, we have

$$\tilde{\theta}^F = \alpha v_o^F = \alpha \frac{p}{2}, \quad (3)$$

where $0 < \alpha < \frac{2}{p}$ to guarantee $0 < \tilde{\theta}^F < 1$. While $\alpha = 1$ corresponds to the rational benchmark, we refer to agents with $\alpha > 1$ and $\alpha < 1$ as *overselective* and *undersensitive* under FIFO, respectively. For any $\alpha \in (0, \frac{2}{p})$, we denote the unique solution of the equilibrium conditions (1), (2), and (3) by $(\tilde{\theta}^F, \tilde{v}_o^F, \tilde{v}_{y2}^F, \tilde{v}_{y1}^F)$.

2.2.2 Last-In-First-Out

In the LIFO protocol, a young agent is always ranked first. If all young agents play a threshold strategy $\hat{\theta} \in [0, 1]$, an old (second-ranked) agent's expected payoff at the beginning of the period is

$$v_o^L = p \left[(1-p) \frac{1}{2} + p \hat{\theta} \frac{\hat{\theta}}{2} \right]. \quad (4)$$

In particular, the old agent can obtain an item if the item is compatible with her (probability p) and either incompatible with the young agent (probability $1-p$), yielding the expected value $\frac{1}{2}$, or compatible but unacceptable to the young agent (probability $p\hat{\theta}$), yielding the expected value $\frac{\hat{\theta}}{2}$.

A young agent's expected payoff is

$$v_y^L = p(1-\hat{\theta}) \frac{1+\hat{\theta}}{2} + [1-p(1-\hat{\theta})] v_o^L. \quad (5)$$

Specifically, since the young agent is ranked first, she obtains the item if it is compatible and acceptable to her (probability $p(1-\hat{\theta})$), yielding the expected payoff $\frac{1+\hat{\theta}}{2}$. Otherwise, her continuation value is v_o^L .

Let $(\bar{\theta}^L, \bar{v}_y^L, \bar{v}_o^L)$ be the rational Markov perfect equilibrium threshold and the expected continuation payoffs under LIFO. It is easy to check that $\bar{\theta}^L = \bar{v}_o^L = \frac{1-\sqrt{1-p^3(1-p)}}{p^2}$.²

Again, a behavioral equilibrium allows agents to use thresholds different from their continuation values. Namely, we consider

$$\tilde{\theta}^L = \beta v_o^L \quad (6)$$

for some $\beta > 0$. While $\beta = 1$ corresponds to the rational benchmark, we refer to agents with $\beta > 1$ and $\beta < 1$ as *overselective* and *underselective* under LIFO, respectively. We denote the unique solution of the equilibrium conditions (4), (5), and (6) by $(\tilde{\theta}^L, \tilde{v}_y^L, \tilde{v}_o^L)$. In particular, after some algebraic steps, we obtain

$$\tilde{\theta}^L = \frac{1 - \sqrt{1 - \beta^2 p^3 (1-p)}}{\beta p^2}. \quad (7)$$

²The rational-equilibrium threshold $\bar{\theta}^L = \bar{v}_o^L$ is a single-peaked function of p and converges to 0 as p approaches either 0 or 1. When p is very small, agents are arbitrarily unlikely to match with the next-period item. When $p = 1$, agents accept any item since they expect their successors to do the same, ruling out the possibility of receiving the next item.

It is easy to verify that $0 < \beta < \frac{2}{p}$ guarantees $0 < \tilde{\theta}^L < 1$.

The behavioral equilibrium threshold $\tilde{\theta}^L$ increases in β . Naturally, as β increases, for the same continuation value v_o^L , an agent becomes more selective. In addition, since each agent's continuation value v_o^L increases as other agents become more selective, there is an indirect effect driven by strategic complementarity. More formally, if an agent is overselective ($\beta > 1$), but others are not and play the rational equilibrium strategy $\bar{\theta}^L$, then the agent's threshold would be $\beta \bar{v}_o^L < \tilde{\theta}^L$. The opposite applies if an agent is underselective ($\beta < 1$) but the others are not. The following lemma summarizes this strategic effect.

Lemma 1. *Under LIFO, the behavioral equilibrium threshold $\tilde{\theta}^L$ is such that $\tilde{\theta}^L < \beta \bar{v}_o^L$ for $\beta < 1$, and $\tilde{\theta}^L > \beta \bar{v}_o^L$ for $\beta > 1$.*

2.3 Equilibrium and Efficiency Comparisons

Next, we compare the rational equilibrium under FIFO and LIFO, and study how behavioral patterns influence this comparison.

We start with some preliminaries. For $h = F, L$, given a threshold $\hat{\theta} \in (0, 1)$ used by all agents, we denote by $W^h(\hat{\theta})$ the expected time an agent waits on the market, which is equivalent to the probability that the agent is not matched upon her arrival. Since the probability of matching on arrival under LIFO is $1 - W^L(\hat{\theta}) = p(1 - \hat{\theta})$, we have

$$W^L(\hat{\theta}) = 1 - p(1 - \hat{\theta}). \quad (8)$$

Under FIFO, $W^F(\hat{\theta})$ has to satisfy

$$W^F(\hat{\theta}) = (1 - W^F(\hat{\theta}))[1 - p(1 - \hat{\theta})] + W^F(\hat{\theta})[1 - (1 - p)p(1 - \hat{\theta})]. \quad (9)$$

This is because a young agent is ranked first on arrival if and only if her immediate predecessor is matched on arrival (which occurs with probability $1 - W^h(\hat{\theta})$). Moreover, the young agent stays on the market into the second period with probabilities $1 - p(1 - \hat{\theta})$ and $1 - (1 - p)p(1 - \hat{\theta})$ if she is ranked first and second on arrival, respectively. Solving (9) for $W^F(\hat{\theta})$, we obtain

$$W^F(\hat{\theta}) = 1 - \frac{(1 - p)p(1 - \hat{\theta})}{1 - p^2(1 - \hat{\theta})}. \quad (10)$$

We measure the efficiency of a queuing protocol using its average match value. More formally, given a threshold $\hat{\theta} \in (0, 1)$ used by all agents, the *allocation efficiency* of protocol

$h = F, L$ is given by

$$E^h(\hat{\theta}) = \frac{pW^h(\hat{\theta})}{2} + (1 - pW^h(\hat{\theta}))\frac{p(1 - \hat{\theta}^2)}{2}, \quad (11)$$

where $W^h(\theta)$ is determined by (8) for LIFO and by (10) for FIFO.³ The first term of (11) represents the scenario in which an old agent is present on the market and is compatible with an item (with probability $pW^h(\hat{\theta})$). Then, the item is matched to some agent for sure, yielding an expected efficiency gain of $\frac{1}{2}$. Otherwise, only a young agent can obtain the item, conditional on the item being both compatible and acceptable to her with probability $p(1 - \hat{\theta})$. Then, the expected efficiency gain is $\frac{1+\hat{\theta}}{2}$, resulting in the second term of (11).

2.3.1 Allocation Efficiency in the Rational Equilibrium

Setting $\alpha = \beta = 1$, we now compare the efficiency in the rational equilibrium with the socially optimal (i.e., efficiency maximizing) stationary mechanisms under FIFO and LIFO.⁴ Our first result describes an optimal mechanism.

Proposition 1 (Socially Optimal Protocol). *1. A socially optimal stationary allocation mechanism is a FIFO protocol with*

$$\theta^{F-opt} = \frac{-(1 - p^2) + \sqrt{1 - 2p^2 + p^3}}{p^2}; \quad (12)$$

2. We have $0 < \theta^{F-opt} < \bar{\theta}^F$.

To see why Part (1) of Proposition 1 holds, note that a social planner prefers to match an item with an old agent rather than a young agent subject to compatibility. This is because, while the surplus of these two matches is the same, the young agent could still realize some positive surplus by staying on the market through the next period. Since the planner prioritizes old agents (i.e., uses FIFO), an optimal mechanism is identified by the threshold θ^{F-opt} that maximizes the allocation efficiency (11) under FIFO, resulting in (12).

To see the intuition of Part (2) of Proposition 1, suppose a young agent under FIFO is offered an item of value $\bar{\theta}^F$, and is therefore indifferent between accepting or rejecting

³Alternatively, the same allocation efficiency (11) can be derived from the agents' average payoffs: given $\hat{\theta} \in (0, 1)$, we have $E^L(\hat{\theta}) = v_y^L$ under LIFO, and $E^F(\hat{\theta}) = (1 - W^F(\hat{\theta}))v_{y1}^F + W^F(\hat{\theta})v_{y2}^F$ under FIFO.

⁴A stationary mechanism determines an item's allocation only based on the item's value and the presence of an old agent on the market in that period.

it. Rejection strictly increases the probability of the next agent remaining unmatched, since, conditional on compatibility with both agents, the next item will not be offered to her. Therefore, a social planner strictly prefers that match to occur, implying the rational threshold $\bar{\theta}^F$ being strictly above the optimal one θ^{F-opt} —that is, rational agents under FIFO are too selective compared to the social optimum.

While Proposition 1 guarantees that a social planner prefers a FIFO protocol to a LIFO one in our setting, the next result identifies the most efficient threshold under LIFO.

Proposition 2 (Socially Optimal LIFO). *The optimal threshold θ^{L-opt} under LIFO is such that $\theta^{L-opt} > \bar{\theta}^L$.*

To understand Proposition 2, suppose a young agent is offered an item of value $\bar{\theta}^L$ under LIFO, and is therefore indifferent between accepting or rejecting it. If an old agent is present and the item is compatible with her, then by allocating the item to the old agent, she leaves the market matched rather than unmatched, generating a strictly positive surplus. Moreover, the continued presence of the young agent in the next period generates no externalities on future agents, as the agent will be ranked second and, therefore, may obtain the next item only when the item would be wasted otherwise. Thus, a social planner strictly prefers the young agent to reject, implying that the equilibrium threshold $\bar{\theta}^L$ is below the optimal one θ^{L-opt} , i.e. rational agents under LIFO are too accommodating compared to the social optimum.

The next result compares the rational equilibrium under FIFO and LIFO, and their respective allocation efficiencies.

Proposition 3 (Efficiency Comparison under Rationality). *In the rational equilibrium,*

1. *Agents are more selective under FIFO than under LIFO—that is, $\bar{\theta}^L < \bar{\theta}^F$;*
2. *The expected time on the market under FIFO is longer than under LIFO—that is, $W^F(\bar{\theta}^F) > W^L(\bar{\theta}^L)$;*
3. *The allocation efficiency is higher under FIFO than under LIFO—that is, $E^F(\bar{\theta}^F) > E^L(\bar{\theta}^L)$.*

Part (1) of Proposition 3 holds because young agents' positions deteriorate more under LIFO than under FIFO, making them relatively less selective. Therefore, on average, fewer old agents are present on the market under LIFO, as highlighted in Part (2). Part (3) of

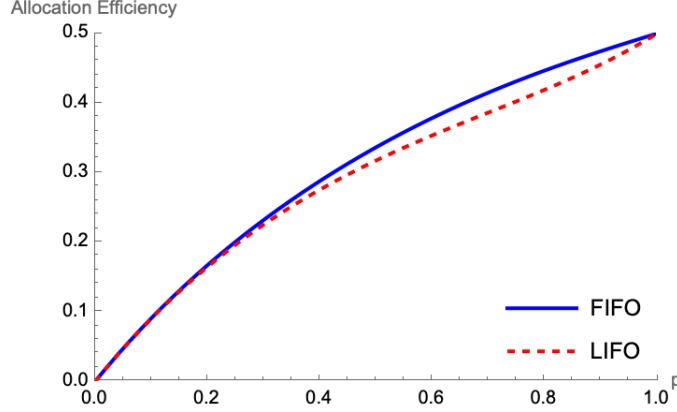


Figure 1: Rational Equilibrium Allocation Efficiency under FIFO and LIFO

Proposition 3 shows that, the FIFO rational equilibrium is more efficient than the LIFO one in our setting, as illustrated in Figure 1.

To gain intuition about Part (3), it is useful to compare the expected payoffs of a young agent arriving as ranked first under FIFO, as ranked second under FIFO, and under LIFO (where she is always ranked first), respectively, conditional on *all agents using the same strategy* $\hat{\theta} \in (0, 1)$. The expected payoff of a young agent who is ranked first under FIFO is v_{y1}^F , as described in (2).

Suppose now that the young agent is ranked second in FIFO (obtaining v_{y2}^F) rather than first. The lower rank affects the agent's payoff only if the arriving item is compatible with and acceptable to the agent, and also compatible with the (first-ranked) old agent. This event occurs with probability $p^2(1 - \hat{\theta})$, and in this case, the young agent receives a continuation value of $\frac{p}{2}$, rather than $\frac{1+\hat{\theta}}{2}$, implying

$$v_{y1}^F - v_{y2}^F = p^2(1 - \hat{\theta}) \left(\frac{1 + \hat{\theta}}{2} - \frac{p}{2} \right).$$

Finally, consider a young agent under LIFO, rather than ranked first under FIFO. Since the agent is ranked first in both protocols, the change in protocol affects her payoff only if the arriving item is either incompatible or unacceptable, which occurs with probability $1 - p(1 - \hat{\theta})$. Conditional on this scenario, her payoff changes only if next period's item is compatible with both her and the next agent, which occurs with probability p^2 . Then, our agent, who is now old, surely obtains the next item under FIFO (with expected value $\frac{1}{2}$), while under LIFO she obtains it only if the item's value is unacceptable to the next agent

(probability $\hat{\theta}$), with expected value $\frac{\hat{\theta}}{2}$. Hence, we obtain

$$v_{y1}^F - v_y^L = \left[1 - p(1 - \hat{\theta})\right] p^2 \left(\frac{1}{2} - \hat{\theta}\frac{\hat{\theta}}{2}\right).$$

It is easy to verify that $v_{y1}^F - v_{y2}^F < v_{y1}^F - v_y^L$, which implies $v_{y2}^F > v_y^L$ —that is, conditional on all agents using the same threshold $\hat{\theta}$, being second upon arrival under FIFO is preferable to entering a market under LIFO. In particular, this holds for $\hat{\theta} = \bar{\theta}^L$. Furthermore, if a young second-ranked agent under FIFO chooses the rational threshold $\bar{\theta}^F = \frac{p}{2}$ instead of $\bar{\theta}^L$, her expected payoff v_{y2}^F becomes even higher. Finally, the expected payoff of an arriving agent under FIFO, which is a weighted average of the values of being ranked first or second upon arrival, must be even higher than the expected payoff conditional on being ranked second, yielding Part (3) of Proposition 3.

While we focus on $p \in (0, 1)$, FIFO and LIFO are equally efficient for both $p = 0$ and $p = 1$. If $p = 0$, the efficiency is obviously zero under both protocols. If $p = 1$, $\bar{\theta}^L = 0$ is the unique equilibrium under LIFO, so the equilibrium efficiency under LIFO is $E^L(0) = \frac{1}{2}$.⁵ Under FIFO, it is easy to check that $\bar{\theta}^F = \frac{1}{2}$, which guarantees $E^F\left(\frac{1}{2}\right) = \frac{1}{2}$ as well. The equal efficiency of the two protocols at $p = 1$ is also intuitive. Under LIFO, a zero equilibrium threshold ensures that each item is always matched with the first-ranked (young) agent. Under FIFO, a strictly positive equilibrium threshold ensures that the steady-state probability of being matched upon arrival is zero, so that each item is always matched with the first-ranked (old) agent. Therefore, the allocation efficiency of both protocols must equal the average item value, which is $\frac{1}{2}$.

2.3.2 Allocation Efficiency in the Behavioral Equilibrium

Next, we describe how agents' behavioral patterns affect the previous results. As we allow for a wide range of behavioral equilibria in each protocol, a meaningful comparison between FIFO and LIFO requires a notion of some minimal rationality of agents' behavior in queues.

Definition 1 (MQ-Rationality). *An agent is “minimally queuing-rational” (MQ-rational) if her strategies θ^L, θ^F satisfy (i) $\theta^L, \theta^F \leq 1/2$, and (ii) $\theta^L \leq \theta^F$.⁶*

⁵For uniqueness, note that if agents played a threshold strategy $\hat{\theta} > 0$, (4) implies $v_o^L = \frac{\hat{\theta}^2}{2}$. By (6), we have $\hat{\theta} = \frac{\beta \hat{\theta}^2}{2}$. For any $\beta \in (0, 2)$, and in particular for $\beta = 1$, this yields $\hat{\theta} > 1$, which is not feasible.

⁶In equilibrium, MQ-rationality translates into conditions on α and β . Specifically, $\alpha < \frac{1}{p}$ and, if all agents are MQ-rational, then $\beta < \bar{\beta}$, where $\bar{\beta}$ is the unique positive solution of $\frac{1 - \sqrt{1 - \beta^2 p^3 (1 - p)}}{\beta p^2} = \frac{\alpha p}{2}$.

Requirement (i) of MQ-rationality guarantees that an agent always accepts any item of value greater than $\frac{1}{2}$, coinciding with the continuation value if the next period's item is both compatible and allocated to the agent for sure. Requirement (ii) amounts to an agent understanding that her future position deteriorates more under LIFO than under FIFO. It is easy to verify that Parts (1) and (2) of Proposition 3 continue to hold in a behavioral equilibrium if agents are MQ-rational.

However, if MQ-rational agents under LIFO become increasingly selective – that is, if $\tilde{\theta}^L$ moves toward $\tilde{\theta}^F$ or, equivalently, β increases – the difference in expected wait time between LIFO and FIFO decreases. In the next result, we illustrate how behavioral patterns under LIFO influence the protocol's equilibrium efficiency and how such efficiency compares with the FIFO protocol.⁷

Proposition 4 (Overselection under LIFO and Efficiency). *We have*

1. $E^L(\theta^{L-opt}) < E^F(\bar{\theta}^F)$, and
2. *If agents are MQ-rational, the allocation efficiency of the LIFO equilibrium is a single-peaked function of β , and is maximized at some $\beta^{L-opt} > 1$. Hence, the efficiency gap between the FIFO rational equilibrium ($\alpha = 1$) and the LIFO behavioral equilibrium decreases for $\beta \in [1, \beta^{L-opt})$ and increases for $\beta \in (\beta^{L-opt}, \bar{\beta}]$.*

Part (1) guarantees that the efficiency under LIFO is lower than the rational equilibrium efficiency of FIFO even at the socially optimal threshold. The intuition is almost identical to Part (3) of Proposition 3, substituting $\bar{\theta}^L$ with θ^{L-opt} . Part (2) of Proposition 4 follows from Proposition 2: efficiency is maximized at a threshold higher than $\bar{\theta}^L$, associated with $\beta^{L-opt} > 1$.

Proposition 4 summarizes the efficiency implication of agents departing from rational behavior under LIFO. In particular, if the agents are moderately overselective under LIFO ($1 < \beta < \beta^{L-opt}$), such departure is efficiency-improving, and it reduces the efficiency gap between LIFO and FIFO. However, if they become extremely overselective ($\beta > \beta^{L-opt}$), the efficiency gap starts to grow again and can ultimately surpass the one associated with rational equilibrium. Proposition 4 constitutes the underpinning of the efficiency implications of our experimental analysis, which we present in the next section.

⁷Since the allocation efficiencies E^L and E^F are continuous in the agents' threshold choices, Proposition 4 still holds under moderate deviations from MQ-rationality. While some subjects in our experiment violate MQ-rationality, such deviations are not significant.

For $p = 0.5$, Figure 4 shows the efficiency gap between FIFO under rational behavior ($\bar{\theta}^F = 0.25$, corresponding to $\alpha = 1$), LIFO under rational behavior ($\bar{\theta}^L = 0.127$, corresponding to $\beta = 1$), and LIFO at the social optimum ($\theta^{L-opt} = 0.184$, corresponding to $\beta^{L-opt} = 1.42$), obtained from Proposition 2.

3 Experimental Design

3.1 Reduced-Form Games

To replicate an infinite-horizon overlapping generation game in an experimental setting, we follow the approach of Lim et al. (1986), Aliprantis and Plott (1992), and Marimon and Sunder (1993). Specifically, we consider reduced-form static setups, which are strategically equivalent to the FIFO and LIFO environments presented in Section 2.

Reduced-Form FIFO A single player selects a minimum acceptable value $\tilde{\theta} \in (0, 1)$. Then, an item with value $\theta_1 \sim U[0, 1]$ arrives and is compatible with the player with probability $p \in (0, 1)$. If the item is compatible and $\theta_1 \geq \tilde{\theta}$, the player receives the payoff θ_1 , and the process ends. Otherwise, a second item arrives with value $\theta_2 \sim U[0, 1]$ and is compatible with probability $p \in (0, 1)$. If the second item is compatible, the player's payoff is θ_2 . Otherwise, the player obtains zero payoff.

This decision problem is equivalent to the FIFO protocol studied in Section 2.2.1, and the strategy for a player with a behavioral parameter $\alpha \in (0, \frac{2}{p})$ is $\tilde{\theta}^F = \alpha \frac{p}{2}$.

Reduced-Form LIFO Players A and B simultaneously select their minimum acceptable values $\tilde{\theta}_A \in (0, 1)$ and $\tilde{\theta}_B \in (0, 1)$, respectively. Then, each player is equally likely to be chosen as an 'active' player. The non-active player obtains zero payoff.

Suppose that player $i = A, B$ is chosen to be active. An item with value $\theta_1 \sim U[0, 1]$ arrives and is compatible with player i with probability $p \in (0, 1)$. If the item is compatible with player i and $\theta_1 \geq \tilde{\theta}_i$, player i obtains θ_1 and the game ends. Otherwise, a second item arrives with value $\theta_2 \sim U[0, 1]$, compatible with each player with probability $p \in (0, 1)$, with compatibility independent between players. Then, one of these cases follows:

1. If $\theta_2 \geq \tilde{\theta}_{-i}$ and the second item is not compatible with player $-i$ but is compatible with player i (with probability $p(1 - p)$), player i obtains θ_2 .

2. If $\theta_2 < \tilde{\theta}_{-i}$ and the second item is compatible with player i with probability p , then player i obtains θ_2 .
3. Otherwise, player i obtains zero payoff.

This setup is strategically equivalent to the LIFO protocol studied in Section 2.2.2. If the players of this game have a behavioral parameter $\beta \in (0, \frac{2}{p})$, there exists a unique behavioral equilibrium where both choose $\tilde{\theta}^L$, described in (7).

3.2 Experiment Description

The experiment was carried out at Washington University in Saint Louis, with 60 undergraduate students as participants. Each session comprises 30 rounds of the FIFO game and 30 rounds of the LIFO game. The subjects accumulated tokens during the 60 rounds and were paid in cash at the end with a conversion rate of \$1 for every 650 tokens. The average payment was \$24.6, in addition to a \$5 show-up fee. Here, we will briefly describe the experiment, while detailed instructions can be found in the Online Appendix.

Subjects are randomly paired at the start of each round in both FIFO and LIFO.⁸ Each item is represented by a jar that contains a random number of tokens between 1 and 1000, representing the jar's value.⁹ Each subject selects an integer between 1 and 1000 as the minimum acceptable jar value. Then, one subject in each pair is randomly chosen as 'active.'

In the *FIFO treatment*, a computer generates the first jar with a random value. If the value meets or exceeds the active subject's threshold, that subject obtains the jar with probability 50%, concluding the round. Otherwise, a second jar is generated, and the active subject has a 50% chance of obtaining the second jar regardless of its value. If she does not obtain it, she receives zero, and the round ends.

In the *LIFO treatment*, a computer generates the first jar with a random value. If the value meets or exceeds the active subject's threshold, the active subject obtains the jar with a 50% chance, concluding the round. Otherwise, a second jar is generated. If the second jar's value meets or exceeds the inactive subject's threshold, the active subject has a 25% chance of obtaining the jar. If the second jar's value is strictly below the inactive subject's threshold, the active subject has a 50% chance of obtaining it. If the active subject fails to obtain the second jar, she receives zero, ending the round.

⁸The pairing is strategically irrelevant under FIFO but ensures a similar structure between treatments.

⁹We set the minimal jar value at 1 instead of 0 to make each value's probability $\frac{1}{1000}$, which is simpler for the subjects to grasp. In what follows, we round each jar's average value to 500.

4 Experimental Results

First, we assess whether the observed choices satisfy MQ-rationality (Definition 1). MQ-rationality requires a subject’s threshold choices under LIFO to be lower than under FIFO, and both to be at most 500.

Result 1. *Most subjects satisfy MQ-rational, i.e., their choices in both treatments are lower than 500, and the choices under LIFO are lower than under FIFO. The subjects that do not satisfy it, do not deviate significantly from it.*

Table 1 shows the threshold choices in each treatment. Since subjects may gradually learn, we present results from three datasets: all 30 rounds, the last 15 rounds, and the last 5 rounds. For each dataset, we calculate the average threshold chosen by each subject under FIFO and LIFO, respectively, and present the summary statistics.

Rational Equilibrium Thresholds		Experimental Data		
		All rounds	Last 15	Last 5
FIFO	250	270.6 (23.1)	258.5 (26.8)	266.8 (28.2)
LIFO	127	218.3 (22.9)	197.8 (24.7)	193.5 (27.1)

Table 1: Summary Statistics of the Subjects’ Average Choices

The last three columns of Table 1 summarize the experimental data. The second, third, and fourth column of the table show the mean and standard deviation (in parentheses) of the average thresholds for the 60 subjects across all 30, last 15, and last 5 rounds, respectively. The aggregate data are consistent with MQ-rationality.

We also observe that individual subjects tend to be MQ-rational. Figure 2 show each subject’s average choices (each dot representing one subject), averaged across all 30, last 15, or last 5 rounds under LIFO on the x -axis, and FIFO on the y -axis. MQ-rationality requires the FIFO threshold choices to be below 500 (i.e., below the red horizontal line), and the LIFO choices to be not higher than the FIFO ones (i.e., on the left side of the blue 45° line). Hence, the dots in the bottom-left triangle of each plot represent subjects that satisfy MQ-rationality on average.

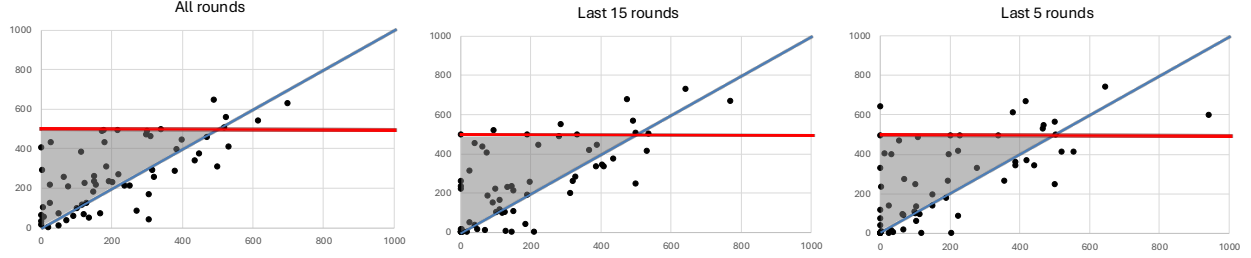


Figure 2: Individual subjects' average choices under LIFO (on the x -axis) and FIFO (on the y -axis).

Specifically, the subjects' threshold choices do not deviate enough from the bottom-left triangle to reject the MQ-rationality hypothesis. Let $\{\hat{\theta}_i^F, \hat{\theta}_i^L\}_{i=1}^{60}$ be the observed threshold choices under FIFO and LIFO. Assuming that the chosen thresholds are independent samples of the equilibrium thresholds $(\tilde{\theta}^F, \tilde{\theta}^L)$, the 95% confidence interval for the FIFO threshold $\tilde{\theta}^F$ is [224.4, 316.8], [204.8, 312.2], or [210.4, 323.1] for the data from all 30, last 15, and last 5 rounds, respectively. Similarly, the 95% confidence interval for the threshold difference, $\tilde{\theta}^F - \tilde{\theta}^L$, is [15.3, 89.3], [19.2, 102.2], or [26.2, 120.3] for the data from all 30, last 15, and last 5 rounds. One-sided tests reject the hypotheses $\tilde{\theta}^F \geq 500$ and $\tilde{\theta}^L \geq \tilde{\theta}^F$.

Next, we compare the subjects' choices to the rational equilibria. As shown in the first column of Table 1, full rationality ($\alpha = \beta = 1$) entails the equilibrium threshold under FIFO and LIFO of $\bar{\theta}^F = 250$ and $\bar{\theta}^L = 127$, respectively.

Result 2. (i) *The subjects' behavior is close to rational in FIFO, but overselective in LIFO;* (ii) *The observed difference between the FIFO and LIFO thresholds is smaller than in the rational equilibrium.*

Figure 3 shows the evolution of the subjects' threshold choices over all rounds. For each round of each treatment, the figure shows the mean of the subjects' choices (solid lines) and the 95% confidence interval (shaded areas).

Panel (a) depicts the choices under FIFO, and the rational threshold of 250. The threshold choices are generally aligned with the rational benchmark, starting higher in the early rounds but converging to 250 within 5-10 rounds. Panel (b) shows the choices under LIFO, and the rational equilibrium threshold of 127. Although the mean of the chosen thresholds tend to decrease slightly over time, they remain significantly higher than the rational equilibrium of 127. Therefore, from Results 1 and 2, we conclude that the observed difference between the FIFO and LIFO thresholds is positive, but smaller than what the rational equilibrium predicts.

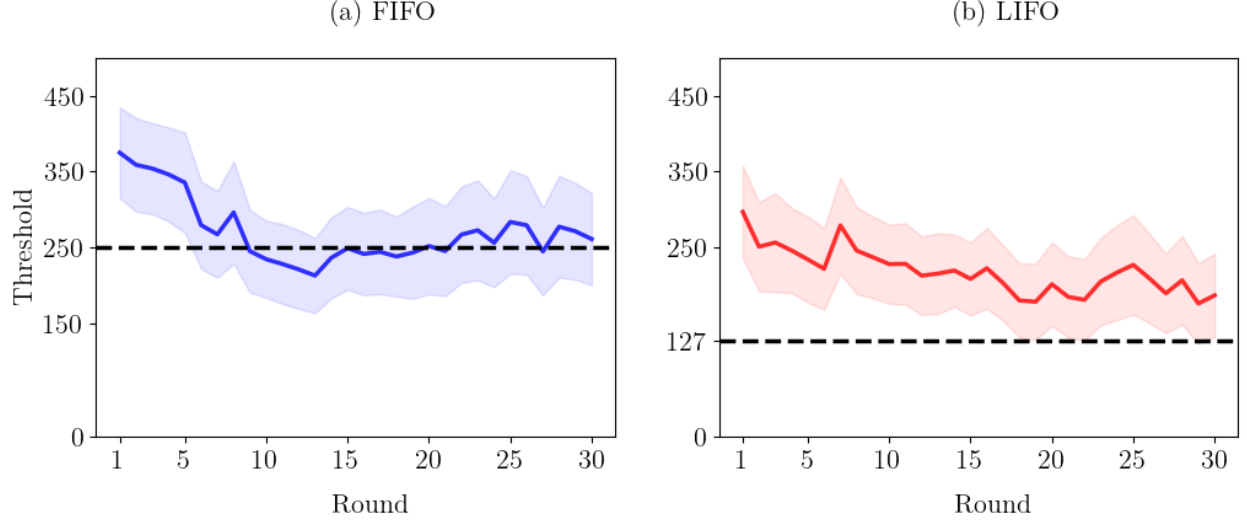


Figure 3: Rational vs. Observed Choices under FIFO [Panel (a)] and LIFO [Panel (b)]

Lastly, we estimate the behavioral parameters α and β using the observed threshold choices.

Result 3. (i) *The estimated behavioral parameters from the threshold choices from all, last 15, and last 5 rounds, are $\hat{\alpha} \in [1.03, 1.08]$ and $\hat{\beta} \in [1.49, 1.67]$.* (ii) *Such estimates imply a reduction in the allocation efficiency gap between the FIFO and LIFO protocols compared to the rational equilibrium.*

The estimate of the behavioral parameter $\hat{\alpha}$ is derived from $\hat{\theta}^F = \hat{\alpha}^{\frac{p}{2}}$ with $p = \frac{1}{2}$. Here, $\hat{\theta}^F$ represents the threshold choices 270.6, 258.5, or 266.8 in Table 1, averaged across all, the last 15, or the last 5 rounds, respectively. Similarly, the estimate of the parameter $\hat{\beta}$ is obtained from (7) and Table 1.

The behavioral patterns described in Result 3 draw the efficiency implications of our experimental results. Recall from Proposition 2 that $\theta^{L-opt} > \bar{\theta}^L$. According to Part (2) of Proposition 4, the overselective behavior under LIFO has the potential to enhance efficiency, if it is not too severe.

Hence, the observed thresholds under LIFO in Table 1 are excessively high, leading to a decrease in the efficiency relative to the maximal one. However, even the highest mean threshold $\hat{\theta}^L = 0.218$ observed from all 30 rounds satisfies $E^L(\hat{\theta}^L) > E^L(\bar{\theta}^L)$. The allocation efficiency under the observed behavior is lower than the optimal level but still higher than what rationality would entail. Therefore, the behavioral pattern we document *reduces the efficiency gap between the FIFO and LIFO protocols with respect to the rational equilibrium*

and helps mitigate the efficiency loss caused by the LIFO protocol.

In Figure 4, for $p = 0.5$, we show the effect of the observed behavior in LIFO ($\hat{\theta}^L = 0.218$, corresponding to $\hat{\beta} = 1.67$), obtained in Result 3 on the allocation efficiency gap between FIFO under rational behavior ($\bar{\theta}^F = 0.25$, corresponding to $\alpha = 1$) and LIFO.

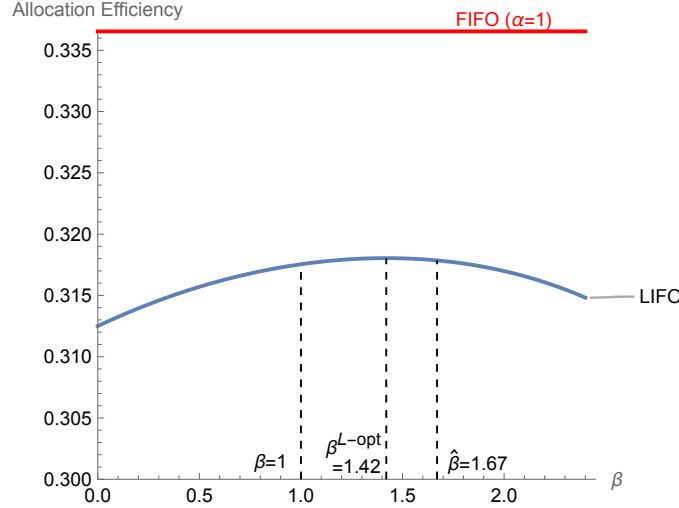


Figure 4: Allocation Efficiency of FIFO under Rational Equilibrium ($\alpha = 1$) and LIFO under Rational Equilibrium ($\beta = 1$), Optimal LIFO ($\beta^{L-opt} = 1.42$), and LIFO under Observed Behavior in Experiment ($\hat{\beta} = 1.67$), for $p = 0.5$

As already mentioned above, it is worth noting that, while agents' overselection behavior under LIFO is the key takeaway from our experiment, its welfare implication is a consequence of this behavior in the context of our model. The welfare impact of overselection under LIFO varies across models, including cases where the welfare ranking between FIFO and LIFO may be reversed. For example, consider a setting similar to ours, but where agents must pay a positive waiting cost if they enter old age unmatched. In this setting, a LIFO protocol generates an efficiency benefit relative to FIFO, since equilibrium waits are longer under FIFO. If waiting costs are high enough, Part (3) of Proposition 3 reverses, and LIFO becomes more efficient than FIFO under rational behavior. However, since Proposition 2 still holds, a moderate overselective bias under LIFO would continue to yield an efficiency gain. Then, the efficiency gap between the two protocols would increase, not decrease, relative to the rational benchmark.

5 Discussion

This section explores potential explanations for the overselective behavior observed under the LIFO treatment. Such behavior conflicts with standard economic theories. If our subjects experienced time discounting, waiting costs, or risk aversion – all omitted in our model – these considerations would have implied $\hat{\alpha} < 1$ and $\hat{\beta} < 1$, the opposite of what we observe in the experimental data. Moreover, subjects’ continuation values under LIFO are more complex to compute than under FIFO, as they depend on one’s beliefs about another agent’s choices. Therefore, aversion to complex or ambiguous consequences from waiting – also omitted in our model – would have implied a more cautious behavior under LIFO than under FIFO, or $\hat{\beta} < \hat{\alpha}$, again in contradiction with our results.

In the next section, we begin with the possibility of the overselective bias being driven by strategic complementarity.

5.1 Human-to-Robot LIFO Treatment

While choices under FIFO are non-strategic, a subject’s optimal choice under LIFO depends on the continuation payoff, which increases in the opponent’s threshold choice, as illustrated in Lemma 1. If subjects form beliefs about their opponents’ choices based on observed high thresholds from earlier rounds, they may fail to adjust their choices downward, thereby preventing convergence to the rational equilibrium.

To address and quantify this effect, we conducted an additional LIFO experiment, in which 47 subjects played the reduced-form LIFO game against robots with known strategies. In this additional treatment, a subject selects an integer between 1 and 1000 for each of 7 possible robot choices $Z \in \{1, 100, 200, 300, 400, 500, 1000\}$. For each robot’s choice and subject’s response, we apply the LIFO treatment described in Section 3.2. The rational best response to a robot’s choice can be derived by replacing $\hat{\theta}$ with $z := \frac{Z}{1000} \in [0, 1]$ in (4) to obtain $\bar{\theta}(z) = \frac{1}{2}[p(1 - p) + p^2 z^2]$, and multiplying it by 1000.

The human-to-robot LIFO treatment was repeated for 5 rounds. Figure 5 shows that the subjects’ average choices still consistently exceeded the rational best responses to the robot’s choices.

We estimate a subject’s behavioral parameter $\hat{\beta}$ using (6), as the ratio of their threshold choice $\hat{\theta}$ to the rational best response $\bar{\theta}(z)$ for each robot’s choice z . After averaging this estimation across all rounds and subjects, the resulting estimate $\hat{\beta}$ falls within the range $[1.31, 1.56]$, still exceeding the rationality benchmark of $\beta = 1$, similarly to the range

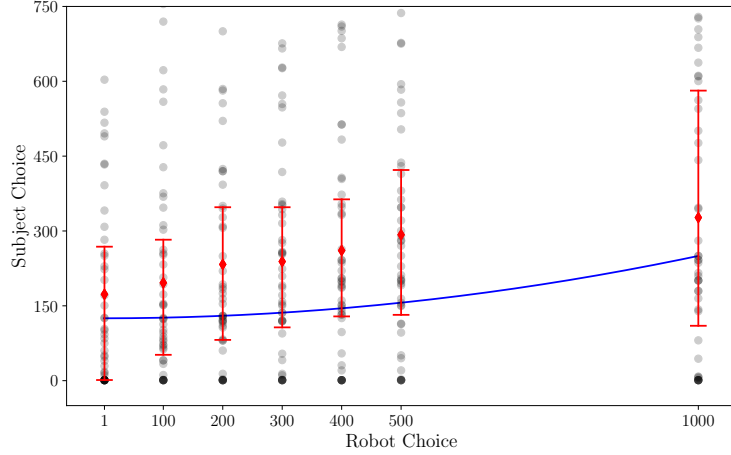


Figure 5: Subjects’ Choices in the Human-to-Robot LIFO Treatment. Each gray dot represents a subject’s average threshold choice across 5 rounds. Red diamonds and segments represent observed averages and 25-to-75 percentile ranges, respectively. The blue curve is the rational best response.

[1.49, 1.67] found in Result 3 of the human-to-human LIFO treatment.¹⁰ To summarize, *without strategic complementarities in the human-to-robot LIFO treatment, subjects still exhibit an overselective bias similar to the one found in the original LIFO treatment.* The rest of this section explores to what extent some well-known behavioral frameworks can explain this residual bias.

5.2 Alternative Behavioral Theories

In this section, we employ well-known behavioral models to examine how well they fit the residual overselective bias observed in the human-to-robot LIFO experiment.¹¹

5.2.1 Quantal Response Equilibrium

In the Quantal Response Equilibrium (QRE) framework, introduced by McKelvey and Palfrey (1995), players do not necessarily choose an optimal action; instead, they choose actions yielding higher expected payoffs with greater probabilities. In the absence of strategic interaction as in the human-to-robot LIFO treatment, QRE simplifies to a logit model.

¹⁰Specifically, the estimated values of $\hat{\beta}$ are 1.38, 1.43, 1.55, 1.47, 1.49, 1.56, or 1.31, corresponding to subjects’ responses to robot’s choices of 1, 100, 200, 300, 400, 500, or 1000, respectively.

¹¹In addition to the ones considered in this section, other behavioral models can potentially generate excessively optimistic expectations in the LIFO game. However, we show in the Online Appendix that they fail to explain our results quantitatively.

Specifically, for each robot's choice z , a threshold $\tilde{\theta}_Q$ is selected with probability

$$Pr(\tilde{\theta}_Q|z, \lambda_Q) = \frac{\exp[\lambda_Q U(\tilde{\theta}_Q; z)]}{\sum_{\theta=1}^{1000} \exp[\lambda_Q U(\theta; z)]},$$

where $U(\theta; z)$ is the subject's expected payoff for a given choice θ , and $\lambda_Q > 0$ measures the decision-making precision. As λ_Q diverges, the choice converges to the rational best response to z , and when $\lambda_Q = 0$, the choice is uniformly random.

The maximum likelihood estimate of the parameter λ_Q is $\hat{\lambda}_Q = 9.59$.¹² Using the estimated $\hat{\lambda}_Q$, we compute the expected value of threshold choices in response to any robot's choice z and present as black dashed line in Figure 6. The QRE predictions are generally higher than the averages of the choices made by the subjects as observed.

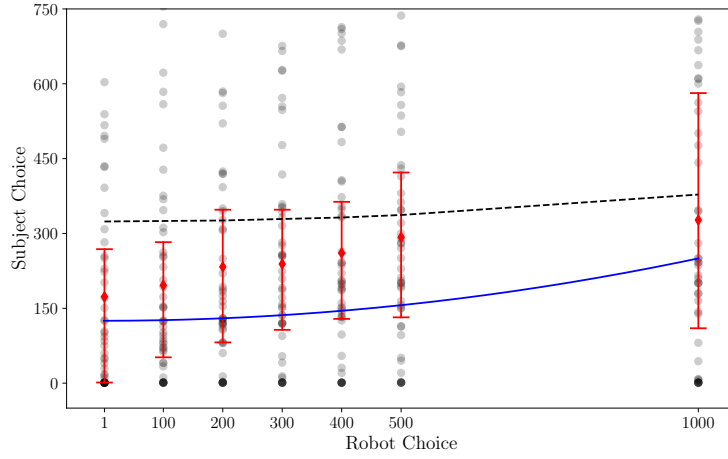


Figure 6: QRE Model Results. The black dashed line corresponds to the choices predicted by the QRE model. A gray dot represents a subject's average threshold choice across 5 rounds. Red diamonds and segments represent observed averages and 25-to-75 percentile ranges, respectively. The blue curve is the rational best response.

5.2.2 Anchoring on the Mean, Mimicking, and Unified Model

Next, we consider two additional well-known behavioral frameworks, and then we combine them with the QRE model in a unified analysis.

¹²Since the human-to-robot LIFO treatment had 47 subjects, each making choices over 5 rounds in response to 7 choices by the robot, our data set includes $N = 47 \times 5 \times 7$ observations, $\{z_i, \hat{\theta}_i\}_{i=1}^N$ where $\hat{\theta}_i$ is a threshold chosen by a subject as a response to a robot's choice z_i . Hence, the log-likelihood function is $\log L(\lambda_Q) = \sum_{i=1}^N \log Pr(\hat{\theta}_i|z_i, \lambda_Q)$.

First, in the *Anchoring on the Mean* model, subjects choose a threshold by departing from the mean value of a jar and adjusting toward the rational best response.¹³ Therefore, we have

$$\tilde{\theta}_A(z) = \lambda_A \times 0.5 + (1 - \lambda_A) \times \bar{\theta}(z) + \epsilon,$$

where $\bar{\theta}(z)$ represents the rational best response to z , and $\epsilon \sim N(0, \sigma^2)$.

Second, when a subject is uncertain about how to respond to a robot's choice, *mimicking* the robot's choice and adjusting toward the rational best response could also be appealing. Similarly to before, this pattern can be formalized as

$$\tilde{\theta}_M(z) = \lambda_M \times z + (1 - \lambda_M) \times \bar{\theta}(z) + \epsilon,$$

where $\epsilon \sim N(0, \sigma^2)$.

Finally, we consider a unified framework combining Quantal Response Equilibrium (Q), Anchoring (A), and Mimicking (M), in the following weighted average:

$$Pr(\tilde{\theta}; z) = \sum_{m \in \{Q, A, M\}} q_m Pr_m(\tilde{\theta}; z, \lambda_m),$$

with $q_m \geq 0$ for $m = Q, A, M$, and $q_Q + q_A + q_M = 1$. For each model, we discretize the threshold prediction.¹⁴

We report the maximum likelihood estimates of up to three weight parameters (q_Q , q_A , and q_M), up to three behavioral parameters (λ_Q , λ_A , and λ_M), and σ in Table 2. QRE explains most of the threshold choices, because the estimated weight of QRE in the unified model is $q_Q = 0.88$. However, the QRE-only model is statistically rejected by the likelihood-ratio test and must be combined with mimicking to explain the subjects' choices.¹⁵

¹³Anchoring is a well-documented bias in the behavioral management literature. See for example Schweitzer and Cachon (2000); Bostian et al. (2008); Katok and Wu (2009); Bolton et al. (2012); Becker-Peth et al. (2013); Moritz et al. (2013).

¹⁴Specifically, we partition the interval $[1, 1000]$ into 40 sub-intervals $[1, 25]$, $[26, 50]$, etc. and assign the probability of a random threshold falling into each sub-interval to its midpoint.

¹⁵The degree of freedom for the LR test of QRE+Anchor or QRE+Mimic against the unified model equals 1, as we only restrict $q_M = 0$ or $q_A = 0$, respectively. On the other hand, the degree of freedom for the LR test of QRE-only is 2.

	QRE	QRE+Anchor	QRE+Mimic	QRE+Anchor+Mimic
q_Q		1.00	0.89	0.88
q_A		0.00		0.01
q_M			0.11	0.11
λ_Q	9.59	9.59	10.02	9.89
λ_A		(0.04)		(0.00)
λ_M			0.97	0.97
σ		(17.16)	1.34	0.21
LogLikelihood	-5741	-5740	-5554	-5553
LR test against full unified model	$p < 0.001$	$p < 0.001$	$p = 0.094$	-

Table 2: Estimation of the Behavioral Models. Estimates in parentheses are unreliable since the estimated q_A is close to 0.

6 Conclusion

We study a one-sided dynamic matching setup, where agents and items arrive sequentially, and there are no wait costs. The items' types are random, and items are wasted if they are not matched with an agent upon arrival. Agents are offered items according to either FIFO or LIFO priority protocol, subject to compatibility, and they leave the market after two periods if still unmatched. We experimentally study the subjects' implementation of threshold strategies under the two priority protocols. While agents behave approximately rationally under FIFO, they adopt thresholds significantly higher than the rational ones under LIFO. In the context of our model, the magnitude of this overselective bias reduces the efficiency gap between FIFO and LIFO. After experimentally demonstrating that this bias persists in an environment where strategic effects are absent, we use well-known behavioral frameworks to explain the overselective bias under LIFO. Further research could be beneficial to understand the causes of this departure from rationality.

References

- ALIPRANTIS, C. D. AND C. R. PLOTT (1992): “Competitive equilibria in overlapping generations experiments,” *Economic Theory*, 2, 389–426.
- ASHLAGI, I., F. MONACHOU, AND A. NIKZAD (2025): “Optimal allocation via waitlists: Simplicity through information design,” *Review of Economic Studies*, 92, 40–68.
- BACCARA, M., A. COLLARD-WEXLER, L. FELLI, AND L. YARIV (2014): “Child-adoption Matching: Preferences for Gender and Race,” *American Economic Journal: Applied Economics*, 6, 133–58.
- BACCARA, M., S. LEE, AND L. YARIV (2020): “Optimal dynamic matching,” *Theoretical Economics*, 15, 1221–1278.
- BACCARA, M. AND L. YARIV (2021): “Dynamic matching,” *Online and Matching-Based Market Design*, 1221–1278.
- BATT, R. J. AND C. TERWIESCH (2015): “Waiting patiently: An empirical study of queue abandonment in an emergency department,” *Management Science*, 61, 39–59.
- BECKER-PETH, M., E. KATOK, AND U. W. THONEMANN (2013): “Designing Buy-back Contracts for Irrational But Predictable Newsvendors,” *Management Science*, 59, 1800–1816.
- BLOCH, F. AND D. CANTALA (2017): “Dynamic Assignment of Objects to Queuing Agents,” *American Economic Journal: Microeconomics*, 9, 88–122.
- BOLTON, G. E., A. OCKENFELS, AND U. W. THONEMANN (2012): “Managers and Students as Newsvendors,” *Management Science*, 58, 2225–2233.
- BOSTIAN, A. A., C. A. HOLT, AND A. M. SMITH (2008): “Newsvendor “Pull-to-Center” Effect: Adaptive Learning in a Laboratory Experiment,” *Manufacturing & Service Operations Management*, 10, 590–608.
- CHAN, C. W., V. F. FARIAS, AND G. J. ESCOBAR (2017): “The impact of delays on service times in the intensive care unit,” *Management Science*, 63, 2049–2072.
- CHE, Y.-K. AND O. TERCIEUX (forthcoming): “Optimal queue design,” *Journal of Political Economy*.

- CONTE, A., M. SCARSINI, AND O. SÜRÜCÜ (2014): “An Experimental Investigation into Queueing Behavior,” Tech. rep., Jena Economic Research Papers.
- DOLD, M. AND M. KHADJAVI (2017): “Jumping the queue: An experiment on procedural preferences,” *Games and Economic Behavior*, 102, 127–137.
- HASSIN, R. (1985): “On the Optimality of First Come Last Served Queues,” *Econometrica*, 53, 201–202.
- HASSIN, R. AND M. HAVIV (2003): *To Queue or not to Queue: Equilibrium behavior in queueing systems*, vol. 59, Springer Science & Business Media.
- KAHNEMAN, D., J. L. KNETSCH, AND R. H. THALER (1986): “Fairness and the Assumptions of Economics,” *Journal of Business*, S285–S300.
- KATOK, E. AND D. Y. WU (2009): “Contracting in Supply Chains: A Laboratory Investigation,” *Management Science*, 55, 1953–1968.
- KREMER, M. AND L. DEBO (2012): “Herding in a queue: A laboratory experiment,” *Chicago Booth Research Paper*.
- LIM, S. S., E. PRESCOTT, AND S. SUNDER (1986): “Stationary Solution to the Overlapping Generations Model of Fiat Money,” Tech. rep., Carnegie Mellon University Working Paper.
- MARGARIA, C. (2024): “Queueing to Learn,” *Theoretical Economics*, *forthcoming*.
- MARIMON, R. AND S. SUNDER (1993): “Indeterminacy of Equilibria in a Hyperinflationary World: Experimental Evidence,” *Econometrica*, 1073–1107.
- McKELVEY, R. D. AND T. R. PALFREY (1995): “Quantal response equilibria for normal form games,” *Games and Economic Behavior*, 10, 6–38.
- MORITZ, B. B., A. V. HILL, AND K. L. DONOHUE (2013): “Individual differences in the newsvendor problem: Behavior and cognitive reflection,” *Journal of Operations Management*, 31, 72–85, behavioral Operations.
- NAOR, P. (1969): “The regulation of Queue Size by Levying Tolls,” *Econometrica*, 15–24.
- ROSOKHA, Y. AND C. WEI (2024): “Cooperation in queueing systems,” *Management Science*, 70, 7597–7616.

SCHWEITZER, M. E. AND G. P. CACHON (2000): “Decision Bias in the Newsvendor Problem with a Known Demand Distribution: Experimental Evidence,” *Management Science*, 46, 404–420.

SU, X. AND S. ZENIOS (2004): “Patient choice in kidney allocation: The role of the queueing discipline,” *Manufacturing & Service Operations Management*, 6, 280–301.

WANG, J. AND Y.-P. ZHOU (2018): “Impact of queue configuration on service time: Evidence from a supermarket,” *Management Science*, 64, 3055–3075.

7 Appendix

Proof of Lemma 1 When $\beta = 1$, we have $\bar{\theta}^L = \bar{v}_o^L$. For the case of $\beta < 1$, let $x = \beta^2$ and $y = p^3(1 - p)$. Then, $\tilde{\theta}^L < \beta \bar{v}_o^L = \beta \bar{\theta}^L$ if and only if

$$\frac{1 - \sqrt{1 - xy}}{x} < 1 - \sqrt{1 - y}.$$

To verify the last inequality and complete the proof, we observe that the two sides of the last inequality are equal at $x = 1$ and

$$\begin{aligned} \frac{d\left(\frac{1 - \sqrt{1 - xy}}{x}\right)}{dx} > 0 &\iff \frac{xy}{2}(1 - xy)^{-1/2} > 1 - \sqrt{1 - xy} \\ &\iff 1 - \frac{xy}{2} - \sqrt{1 - xy} > 0 \\ &\iff 0 < xy < 1, \end{aligned}$$

which always holds. A similar argument can prove the $\beta > 1$ case. ■

Proof of Proposition 1 For Part (1), we substitute $W^h(\theta)$ with (10) in (11). It is easy to check that the threshold in (12) is optimal. Part (2) follows from easy algebraic steps, using (12) and $\bar{\theta}^F = \frac{p}{2}$. ■

Proof of Proposition 2 From (8) and (11), we obtain

$$\theta^{L-opt} = \frac{1 - p + p^2 - \sqrt{(1 - p + p^2)^2 - 3p^3(1 - p)}}{3p^2}. \quad (13)$$

Recall that $\bar{\theta}^L = \frac{1-\sqrt{1-p^3(1-p)}}{p^2}$. Let $x = p^2$ and $y = p(1-p)$. Then, $\bar{\theta}^L < \theta^{L-opt}$ if and only if

$$\begin{aligned}
1 - \sqrt{1-xy} &< \frac{(1-y) - \sqrt{(1-y)^2 - 3xy}}{3} \\
\iff (2+y) &< 3\sqrt{1-xy} - \sqrt{(1-y)^2 - 3xy} \\
\iff \sqrt{(1-xy)((1-y)^2 - 3xy)} &< (1-y) - 2xy \\
\iff -(1-y)^2 - 3(1-xy) &< 4xy - 4(1-y) \\
\iff (1-y)(3+y) &< 3+xy,
\end{aligned}$$

which always holds since $(1-y)(3+y) = 3 - 2y - y^2 < 3$. ■

Proof of Proposition 3 Part (1) For $h = F, L$, we have $\bar{\theta}^h = \bar{v}_o^h$. Since $\bar{v}_o^F = \frac{p}{2}$, and $\bar{v}_o^L$ is given by (4), it is easy to check that $\bar{v}_o^F > \bar{v}_o^L$ for any $\hat{\theta} \in (0, 1)$.

Part (2) By comparing (8) for LIFO and (10) for FIFO, we observe that $W^L(\hat{\theta}) < W^F(\hat{\theta})$ for any $\hat{\theta} \in (0, 1)$. Moreover, $W^L(\hat{\theta}) = 1 - p(1 - \hat{\theta})$ strictly increases in $\hat{\theta}$. Thus, Part (2) follows from Part (1), as $\bar{\theta}^L < \bar{\theta}^F$.

Part (3) First, we show that, for any $\hat{\theta} \in (0, 1)$, $v_{y2}^F(\hat{\theta}) > v_y^L(\hat{\theta})$, the expected utility of an agent who arrives into a FIFO market when an old agent is present is higher than the expected utility of an agent who arrives into a LIFO market.

From (1), we have

$$v_{y2}^F(\hat{\theta}) = p(1 - \hat{\theta}) \left[(1-p) \frac{1+\hat{\theta}}{2} + p \frac{p}{2} \right] + [1 - p(1 - \hat{\theta})] \frac{p}{2}.$$

From (4) and (5),

$$v_y^L(\hat{\theta}) = p(1 - \hat{\theta}) \frac{1+\hat{\theta}}{2} + [1 - p(1 - \hat{\theta})] \left[(1-p) \frac{p}{2} + (p\hat{\theta}) \frac{p\hat{\theta}}{2} \right].$$

Then,

$$\begin{aligned} v_{y2}^F(\hat{\theta}) - v_y^L(\hat{\theta}) &= p(1 - \hat{\theta}) \left[\frac{p^2 - p(1 + \hat{\theta})}{2} \right] + [1 - p(1 - \hat{\theta})] \frac{p^2(1 - \hat{\theta}^2)}{2} \\ &= p(1 - \hat{\theta}) \frac{p^2 - p^2(1 - \hat{\theta}^2)}{2} > 0. \end{aligned} \quad (14)$$

Finally,

$$E^L(\bar{\theta}^L) = v_y^L(\bar{\theta}^L) < v_{y2}^F(\bar{\theta}^L) < v_{y2}^F(\bar{\theta}^F) < E^F(\bar{\theta}^F).$$

The equality holds because an equal number of agents and items arrive over time, and the type of each matched item is equal to the utility of the associated agent. The first inequality follows from (14). The subsequent inequalities hold because $v_{y2}^F(\hat{\theta})$ increases as $\hat{\theta}$ approaches $\bar{\theta}^F = \frac{p}{2}$, and $v_{y2}^F(\bar{\theta}^F)$ must be lower than the average utility across all agents under FIFO who can be ranked first or second upon arrival. ■

Proof of Proposition 4 We can obtain $E^L(\hat{\theta}) = v_y^L(\hat{\theta})$ from (8) and (11) as a cubic function of $\hat{\theta}$. Then,

$$\frac{dE^L(\hat{\theta})}{d\hat{\theta}} = \frac{p}{2} \underbrace{\left[-2\hat{\theta} + p(1 - p + p\hat{\theta}^2) + (1 - p(1 - \hat{\theta}))2p\hat{\theta} \right]}_{:=h(\hat{\theta})}.$$

Note that $h(\hat{\theta})$ has a positive quadratic coefficient, that $h(0) = p(1 - p) > 0$, and that $h(\frac{1}{2}) = p(2 - \frac{5p}{4}) - 1 < 0$. Consequently, $E^L(\hat{\theta})$ is a single-peaked function over the interval $\hat{\theta} \in [0, \frac{1}{2}]$, which we focus on by MQ-rationality.

Part (1) Since $h(\frac{p}{2}) = p^3(\frac{3p}{4} - 1) < 0$, we have $\bar{\theta}^F = \frac{p}{2} > \theta^{L-opt}$. Then,

$$E^L(\theta^{L-opt}) = v_y^L(\theta^{L-opt}) < v_{y2}^F(\theta^{L-opt}) < v_{y2}^F(\bar{\theta}^F) < E^F(\bar{\theta}^F).$$

The equality and inequalities hold similarly to the last part of the proof of Proposition 3. The equality holds because an equal number of agents and items arrive over time, and the type of each matched item is equal to the utility of the associated agent. The first inequality follows from (14). The subsequent inequalities hold because $v_{y2}^F(\hat{\theta})$ increases as $\hat{\theta}$ approaches an agent's optimal threshold $\bar{\theta}^F = \frac{p}{2}$, and the expected utility of an agent ranked second

upon arrival in FIFO must be lower than the expected utility across all agents under FIFO who can be ranked either first or second upon arrival.

Part (2) Given that the behavioral equilibrium $\tilde{\theta}^L$ strictly increases in β by (7), $E^L(\hat{\theta})$ is a single-peaked function of β . Furthermore, $\theta^{L-opt} > \bar{\theta}^L$ from Proposition 2, so the maximal efficiency is achieved at some $\beta^{L-opt} > 1$. ■