

Evaluating Search Cost Models: Estimation and Prediction*

Adrian Düll[†] Heiko Karle[‡] Simon Martin[§] Heiner Schumacher[¶]

Version: July 23, 2026

Abstract

The classic search models assume that consumers adhere to a particular method of search (sequential or non-sequential) and that they know the true price distribution. In this paper, we evaluate how well the search cost estimates from classic models predict search outcomes – the amount of search and purchase prices – when these assumptions are violated. To this end, we conduct an online experiment in which we vary searchers’ information about the price distribution of a homogeneous good. For each treatment, we (i) estimate search costs, (ii) fit each model to the estimated search cost distribution to obtain in- and out-of-sample predictions about outcomes, and (iii) compare predicted and realized outcomes. We find that the prediction performance of each model is largely robust to violations of the informational assumption. Further, the prediction performance of the sequential and non-sequential search model are similar, despite the fact that the search environment strongly favors sequential search.

Keywords: Online Search, Search Costs, Search Cost Estimation

JEL Classification: C90, D12, D83

*We thank Atabek Atayev, Max Breitenlechner, Ben Casner, Jonas Dovern, Daniel Garcia, Martin Geiger, Jürgen Huber, Maarten Janssen, Johannes Johnen, José Luis Moraga-González, Wieland Müller, Georg Nöldeke, David Ronayne, Karl Schlag, Felix Schleef, Philipp Schmidt-Dengler, Anton Sobolev, Michelle Sovinsky, Frank Verboven, and Achim Zeileis as well as seminar audiences at University of Basel, University of Cologne, DFG Workshop “Consumer Preferences, Consumer Mistakes, and Firms’ Response”, EARIE 2025 in Valencia, IIOC 2025 in Philadelphia, MaCCI, Annual Meeting of the Committee for Industrial Economics at the University of Vienna, World Congress of the Econometric Society 2025 in Seoul, and ZEW Mannheim for valuable comments and suggestions. Financial support from a Methusalem grant of KU Leuven, Frankfurt School of Finance and Management, German Research Foundation (DFG project 462020252), and from the OeNB Anniversary Fund (Project Grant No. 19009) is gratefully acknowledged. This RCT was registered as AEARCTR-0012497.

[†]University of Innsbruck. E-Mail: adrian.duell@uibk.ac.at.

[‡]Frankfurt School of Finance & Management, CEPR, and CESifo. E-Mail: h.karle@fs.de.

[§]University of Vienna, CEPR, and CESifo. E-Mail: simon.martin@univie.ac.at.

[¶]University of Innsbruck and KU Leuven. E-Mail: heiner.schumacher@uibk.ac.at.

1 Introduction

In many markets, consumers have to invest time and effort to acquire product and price information before making a purchase (Stigler, 1961). This informational friction potentially creates price dispersion: Some firms may charge relatively low prices in order to serve those customers who compare many prices, while other firms may charge high prices to maximize profits from customers who search only few options. A large literature in industrial organization quantifies the extent of informational frictions in markets by estimating consumers' search costs from observed prices and – if available – search behavior.¹ Given an estimated model of consumer search and firm pricing, researchers can study how changes in the market setting would affect the market equilibrium and consumer welfare.

The identification of search costs builds on the conjecture that consumers acquire additional information until the marginal costs of further search equal its marginal expected gains. There are two classic paradigms that capture this trade-off: *sequential* and *non-sequential* (or *fixed sample size* or *simultaneous*) search. Under sequential search, the decision maker (DM) decides after each search whether to purchase the good from a previously sampled shop or to continue search. Under non-sequential search, the DM first chooses the number of shops she wants to search and, after completing the search spell, purchases the good with the lowest price (or the highest net utility) in her sample.

Both search paradigms have been used extensively in the theoretical and empirical literature on consumer search. Sequential search first has been examined by McCall (1970) and Weitzman (1979). It is used in the theoretical literature in the context of price search, e.g., by Stahl (1989), Janssen et al. (2005), and most recently by Mauring and Williams (2023). Several papers including Kim et al. (2010), Kim et al. (2017), and Moraga-González et al. (2023) estimate search costs based on the sequential search paradigm. Ursu et al. (2025) provide a general framework for estimating search costs when also non-price attributes are uncertain. Non-sequential search has been introduced by Stigler (1961). In the theoretical literature, it is used, for example, by Burdett and Judd (1983), Janssen and Moraga-González (2004), and Atayev (2022). Much of the empirical literature builds on the non-sequential search model, e.g., Moraga-González and Wildenbeest (2008), De Los Santos et al. (2012), Honka (2014), Ghose et al. (2019), Lin and Wildenbeest (2020), as well as Fischer et al. (2024).

Using the classic search paradigms to estimate search costs is not innocuous though, for two reasons. First, an important assumption in both models is that the DM knows the distribution over potential outcomes. At each stage of the search process, this distribution defines the expected gains from (further) search and hence the continuation or completion of the search

¹ Additionally, a large literature in labor economics studies job search and the impact of informational frictions on wages and unemployment; see Le Barbanchon et al. (2024) for a recent review of this literature.

spell. However, when consumers are unfamiliar with the product category, the informational assumption may not be satisfied. This concern is well recognized in the literature (below we describe how previous studies have relaxed the assumption of full information).

Second, it is by no means clear that all consumers follow the same search paradigm or that they follow any of these paradigms at all. Some consumers may apply sequential, others non-sequential search. They may also follow a combination of sequential and non-sequential search (Morgan and Manning, 1985). The empirical literature on consumer search finds support for both search paradigms. Alternatively, some consumers may apply heuristics such as “search sequentially, but conduct no more than n searches” or “conduct n searches, regardless of the costs and (expected) gains from search.” Thus, it is not clear which empirical model researchers should use to estimate search costs. Observational data is typically not rich enough to jointly identify consumers’ expectations about the price distribution, their search method, and the corresponding search costs (provided such a value is well-defined at all).

Despite these concerns, the two classic search models may be good *as-if* models. To see whether this is the case in the context of price search, we examine the following questions in this paper: If a researcher has estimated a distribution over search costs using one of the classic models, then how well can one reconstruct the amount of search and the prices paid by consumers from this distribution? And how well can one predict the amount of search and purchase prices under alternative price distributions? Specifically, we are interested in the answers to these questions in settings for which we know that the informational and search protocol assumptions are violated.

In order to obtain appropriate data, we conduct an experiment in which subjects can search for the lowest price of a (hypothetical) product. The treatment variation is the price distribution of the product as well as subjects’ information about this distribution. All other aspects of the transaction (such as the utility from the product) are held constant and are made transparent to subjects. To examine the predictive performance of a sequential and a non-sequential search model in the context of price search, we apply three steps for each treatment:

- (i) We estimate search costs using an ordered-probit regression framework. The estimation yields us a fully specified search cost distribution.
- (ii) We use a given search model, parameters of the search environment, and the estimated search cost distribution to simulate predicted search outcomes (number of searches and purchase prices). In this manner, we obtain both in- and out-of-sample predictions.
- (iii) We compare the predicted and realized distributions over search outcomes in terms of means, medians, and two measures from statistics that quantify the similarity between distributions, namely *Hellinger distance* and *Kullback-Leibler divergence*.

This research design allows us to evaluate whether the classic search models have desirable as-if properties, even when we know that certain assumptions are not satisfied.²

In our online search experiment, we consider conditions where subjects receive information about the price distribution as well as conditions where no such information is provided. Our subjects are online workers on Amazon Mechanical Turk (AMT) and Prolific (PRO). They can search for the lowest price of a hypothetical homogeneous product in up to 100 online shops. Search creates time and hassle costs: To obtain the price quote from an online shop, they have to enter a 16-digit code. The price savings that subjects realize constitutes their payoff in the experiment. The main treatment variations are the price distribution and subjects' information about the price distribution. We consider settings with left-skewed, uniform, and right-skewed price distributions. For every price distribution there is both a treatment where the distribution is made salient to subjects and a treatment where subjects obtain no information about it. Through additional treatments, we identify the behavioral parameters in our search environment, that is, the level of context effects (the tendency that people become less sensitive to price variations of fixed size when the price level or range of prices increases) and the extent to which search costs are convex in the number of searches (e.g., due to fatigue).

We observe fairly heterogeneous search outcomes in our experiment. There is both a large fraction of subjects who do not search more than two shops and a substantial share of subjects who search more than ten shops. Accordingly, the dispersion of purchase prices is large. Neither the pure sequential nor the pure non-sequential search model is fully consistent with subjects' search behavior. All treatments exhibit a substantial degree of recall, which is inconsistent with sequential search under constant search costs. Additionally, several treatments show evidence of price-dependent search, which is inconsistent with non-sequential search.

Our first main result is that the prediction performance of both the sequential and non-sequential search model is hardly affected when the underlying assumption on consumers' information is violated. One may conjecture that the estimated search cost distribution from the classic models captures search outcomes better in settings where subjects know the true price distribution. Indeed, we find for the sequential search model that the point estimates of the prediction errors (regarding the number of searches and purchase prices) are on average smaller when subjects are informed about the price distribution than when this is not the case. However, the difference between the prediction errors is small relative to their variability. For the non-sequential search model we even find that the point estimates of the prediction errors are smaller for treatments where subjects are not informed about the price distribution.

The second main result is that the prediction performance of the non-sequential search

²We focus on the sequential and non-sequential search model in the main part of this paper since in most empirical search models, decision-makers are assumed to know the outcome distribution. However, in an extension, we also evaluate the predictive performance of a search without priors model based on [Parakhonyak \(2014\)](#).

model is fairly close to that of the sequential search model, despite the fact that the experimental protocol is tailored to the sequential search framework. Again, the differences in prediction errors are small relative to the variability of the prediction error point estimates. The non-sequential search model even predicts the median of search outcomes slightly better than the sequential search model. We obtain this pattern both for in- and out-of-sample predictions. Thus, despite the stark simplification assumed by the non-sequential search model, it seems to describe aggregate search outcomes fairly well, regardless of whether consumers have information about the price distribution or not. This result suggests that using the non-sequential search model in empirical work – as it is frequently done in the literature – is an appropriate choice even in settings where consumers can continue their search at any stage and are not constrained by a predetermined number of searches.

In order to provide benchmark values for the predictive performance of the two models, we conduct the following analyses. First, we perform a simulation study in which the true data-generating process is fully known. We vary the fraction of simulated subjects who search sequentially and non-sequentially, respectively, and apply the same estimation and evaluation procedures to the simulated dataset as for the experimental data. Second, we compare the prediction performance of the classic search models to the prediction performance of models in the forecasting literature (for retail sales and inflation). Both analyses confirm that the two search models are robust as-if descriptions of search behavior. Our simulation study also helps to explain a pattern that recurs throughout the analysis: The models systematically predict purchase prices more accurately than the number of searches. We find that purchase prices are relatively insensitive to the assumed search protocol, while the number of searches directly reflects the stopping rule and is more sensitive to behavioral heterogeneity that the models do not capture.

Overall, we find that the sequential and non-sequential search model are good as-if models, even in settings when the underlying assumptions on the search protocol and consumers' information about the price distribution are violated. Both models capture heterogeneity in search outcomes: Lower search costs imply (in expectation) more searches and lower prices. This has implications both for empirical and theoretical research. In empirical research, neither the precise form of search nor the initial information of consumers is known very well. Our findings suggest that the key findings are expected to be robust even when these components cannot be captured precisely. Choosing the right search protocol and the level of consumer information is a challenge also in the theoretical literature. While it is clearly important to understand the implications of different search models, our findings indicate that as far as estimated search costs are concerned, most common methods tend to be quite robust.

Related Literature. This paper continues the literature that evaluates the validity of structural

models by using experiments. [Bajari and Hortacısu \(2005\)](#) consider models of first-price auctions and [Salz and Vespa \(2020\)](#) evaluate models of dynamic competition. In the context of search, [Brown et al. \(2011\)](#) examine how reservation wages change over the search spell and [Karle et al. \(2025\)](#) identify the impact of context effects on search cost estimates. Most recently, [Fong et al. \(2025\)](#) consider a search experiment in order to disentangle prominence effects and biased beliefs in online shopping behavior. The present paper continues this line of research and evaluates the performance of search models when their underlying assumptions are (partially) violated.

Further, the paper contributes to the search literature that considers a relaxation of the full-information assumption. Two approaches address the issue of information in the theoretical literature. First, consumers may have a prior about the potential distributions over outcomes and update their beliefs in Bayesian manner; see, e.g., [Rothschild \(1974\)](#), [Rosenfield and Shapiro \(1981\)](#), and [Bikhchandani and Sharma \(1996\)](#). These papers examine under which conditions the optimal search strategy is myopic, so that it can be characterized by a reservation value. They show that this is the case when searchers' beliefs are given by Dirichlet priors. The second approach considers search without priors: Consumers use their observations to non-parametrically estimate the distribution over outcomes, see [Chou and Talmain \(1993\)](#), [Parakhonyak \(2014\)](#), or [Schlag and Zapechelnyuk \(2021\)](#). The first two papers consider an updating rule that maximizes entropy and generates a piece-wise uniform belief about the distribution over outcomes. The empirical literature that estimates search costs without the full information assumption largely follows the parametric approach. [Koulayev \(2013\)](#), [De los Santos et al. \(2017\)](#), and [Hu et al. \(2019\)](#) assume Dirichlet priors or Dirichlet process priors; [Ursu et al. \(2020\)](#) and [Ursu et al. \(2024\)](#) use normally distributed priors. [Jindal and Aribarg \(2021\)](#) conduct a search experiment (with monetary search costs) in which they elicit subjects' beliefs about the price distribution after each search. They find that prior beliefs are fairly heterogeneous and that the assumption of Bayesian updating about the price distribution does not significantly bias the search cost estimates as long as researchers take this heterogeneity into account. Our paper is the first that evaluates the predictive performance of different search models in settings with varying degrees of consumer information about prices. We focus on search models that consumers are informed about the price distribution. In an extension, we also consider a model of search without priors that is based on [Parakhonyak \(2014\)](#).

Finally, our paper contributes to the literature that tests the validity of the sequential and non-sequential search paradigm in the context of consumer search.³ Sequential search with constant search costs implies that the DM has a constant reservation value and trades with the

³See [Honka et al. \(2019\)](#) for a comprehensive overview of the literature on search cost estimation and its connection to the literature on consideration sets ([Abaluck and Adams-Prassl, 2021](#), [Amano et al., 2022](#)).

last sampled shop or searches all available shops. These predictions are frequently violated in experimental data, see, e.g., [Brown et al. \(2011\)](#) and [Casner \(2021\)](#). Within the empirical literature, [Chen et al. \(2007\)](#), [De Los Santos et al. \(2012\)](#), and [Honka and Chintagunta \(2017\)](#) test the two models against each other. [Chen et al. \(2007\)](#) apply a non-parametric likelihood ratio test to discriminate between the two search models (using price data from a price comparison website). They do not find significant differences between them in terms of fit. [De Los Santos et al. \(2012\)](#) examine web browsing data and find (for the online book market) that consumers recall shops and that there is no significant positive correlation between prices and the decision to continue search. Both observations are consistent with non-sequential search, but inconsistent with the sequential search model. [Honka and Chintagunta \(2017\)](#) observe consideration sets and purchases (in the U.S. auto insurance market). They find that the set of observed prices and choices is consistent with non-sequential search and that applying an empirical model that is based on sequential search would generate biased estimation results. In contrast, [Bronnenberg et al. \(2016\)](#) find that chosen options on average are discovered late in the search spell, which is consistent with sequential, but inconsistent with non-sequential search. Overall, the evidence from observational data generates mixed results on the validity of the two search paradigms. In the present paper, we therefore evaluate the prediction performance of both models.

The rest of the paper is organized as follows. In Section 2, we formally describe the sequential and non-sequential search model. In Section 3, we present our experimental design and procedures. In Section 4, we explain how we estimate search costs in our setting and how we evaluate the prediction performance of the two search models. In Section 5, we present the experimental results, the evaluation of the search models, and extensions of our analysis. Section 6 concludes. Appendix A contains all mathematical proofs and the instructions to the experiment. Online Appendix B contains further empirical analyses.

2 Search Models

In this section, we formally introduce the sequential and non-sequential search model. We allow for context effects and convex search costs. For each model, we state a condition for optimal search that will be used in the search cost estimation.

We consider a decision-maker who has unit demand for a homogeneous good and can search a (large) finite number N of firms that offer this good at varying prices. Each firm chooses its price p according to the distribution function F with support $[a, b] \subset [0, \infty)$ and piecewise continuously differentiable density f . The DM can search one firm at a time. Let $c(n)$ denote the time and hassle costs of n searches. We assume that $c(0) = 0$ and that c is

weakly convex. If the DM purchases the good at price p after searching n firms, her (decision) utility is given by the function

$$V^{du}(p, F, n) = u - v(p, F) - c(n) + wn, \quad (1)$$

where u is her utility from the product and w is a monetary payoff for each search. In a standard search setting, w equals zero. In the experiment, we will have treatments with positive values w in order to identify the shape of the search cost function. The function v defines how the benefits (price savings) from search are perceived relative to the physical costs of search. It may capture the standard case $v(p, F) = p$ as well as preferences with diminishing sensitivity or relative thinking as in [Karle et al. \(2025\)](#).

Sequential Search. We first assume that the DM searches optimally according to the random sequential search paradigm. After each search, the DM chooses between purchasing the good at the lowest price discovered so far and conducting one more search. By a simple induction argument we can show that this implies that the DM follows a reservation price policy⁴: Suppose the DM has searched $n - 1$ firms and the smallest price discovered so far equals $\hat{p}$. Then there is a value r_n so that it is optimal for the DM to stop search and to purchase the good at price $\hat{p}$ if $\hat{p} < r_n$, and to conduct the n 'th search if $\hat{p} \geq r_n$. This reservation price is defined by the indifference condition

$$c(n) - c(n - 1) = \int_a^{r_n} (v(r_n, F) - v(p, F)) f(p) dp + w. \quad (2)$$

The left-hand side of this equation represents the cost of the n 'th search and the right-hand side the benefits from search, i.e., the expected gains from learning one more price.

Non-sequential Search. Next, we assume that the DM searches optimally according to the non-sequential search paradigm. She chooses (and commits to) the number n of price quotes she wishes to obtain. After completing $n \geq 1$ searches, she trades with the firm in her sample that offers the product at the lowest price. The distribution of this price is given by

$$F(p; n) = 1 - (1 - F(p))^n. \quad (3)$$

Therefore, the expected expenses weighted by function v are equal to

$$\mathbb{E}(v; n) = \int_a^b v(p, F) n (1 - F(p))^{n-1} f(p) dp. \quad (4)$$

⁴In Appendix [A.1](#), we present the proof of this statement.

The DM chooses the number $n \leq N$ of searches that maximizes her expected payoff

$$u - \mathbb{E}(v; n) - c(n) + wn. \quad (5)$$

Note that the expected expenses $\mathbb{E}(v; n)$ decrease, while search costs $c(n)$ increase in the number of searches n . The DM searches n shops only if

$$-\mathbb{E}(v; n) - c(n) + wn \geq -\mathbb{E}(v; n') - c(n') + wn' \quad (6)$$

for any number $n' \leq N$ of searches.

Empirical Predictions. The two search models make different empirical predictions about the relationship between the DM's purchase decision and the sequence of observed prices as well as about the relationship between observed prices and the DM's decision to continue search. [De Los Santos et al. \(2012\)](#), [Honka and Chintagunta \(2017\)](#), and [Karle et al. \(2025\)](#) use such predictions to examine whether consumer search is consistent (or inconsistent) with sequential and non-sequential search. We follow this approach and consider three tests on whether recall behavior and the price-dependency of search are in line with a given search model. To formulate these tests, we assume that there is an infinite population of DMs who can search a (large) finite number of firms. Further, we assume that there is a positive fraction of DMs in this population who search exactly one firm and a positive fraction of DMs who search at least two firms. The first test considers recall.

Test 1 (Recall). *Consider the identity of the shop from which a DM purchases the product.*

- (a) *If all DMs search sequentially and the costs per search are constant in the number of searches, every DM buys the product from the last sampled shop or searches all shops.*
- (b) *If all DMs search non-sequentially, then among those who search n shops the probability of purchase from a given sampled shop is $\frac{1}{n}$.*

If marginal search costs are constant in the number of searches, then under sequential search, a DM will not recall a previously sampled shop, unless she has searched all shops. This prediction has been discussed repeatedly in the literature; see, for example, [Brown et al. \(2011\)](#) and [De Los Santos et al. \(2012\)](#). However, recall of previously sampled shops is possible under sequential search if marginal search costs are convex in the number of searches. We obtain a stronger test for non-sequential search. Independent of the search cost function, recall should take place with probability $\frac{n-1}{n}$ if a DM searches n shops. The next test considers the link between observed prices and the decision to continue the search spell.

Test 2 (Price-Dependence I). *Consider the correlation between the price at the first sampled shop and the decision to continue search.*

- (a) *If all DMs search sequentially, this correlation is positive.*
- (b) *If all DMs search non-sequentially, this correlation is zero.*

The sequential search model implies a reservation price policy. Therefore, it would generate a positive correlation between prices and the decision to continue search. In contrast, under non-sequential search, the number of searched shops is chosen ex-ante so that there is no such correlation. This is also reflected in the last test.

Test 3 (Price-Dependence II). *Consider the correlation between the price at the current shop and the decision to continue search.*

- (a) *If all DMs search sequentially, this correlation is positive.*
- (b) *If all DMs search non-sequentially, this correlation is zero.*

3 Experimental Design and Procedures

We design an online search experiment in which we vary both the price distribution as well as subjects' information about the price distribution. We will use the experimental data to estimate search costs and evaluate the extent to which search outcomes (number of searches and purchase prices) can be recovered from these estimates. In the following, we describe the design of our baseline treatment, treatment variations, and experimental procedures.

Experimental Design Baseline Treatment. We invite subjects (on Amazon Mechanical Turk and Prolific) to an online experiment which consists of two parts. The first part is a survey in which we elicit demographic information (age, gender, education) as well as measures of cognitive ability, risk preferences, trust, and online labor supply on the platform. After the survey, we describe the second part, which consists of an experimental search task. Screenshots of the instructions and the search task are shown in Appendix A.5.

In the search task, subjects have to buy a fictitious product that we call “product A.” They can search for the lowest price of this product in up to 100 shops. The price at each shop (in USD) is drawn from a uniform distribution on the interval $[a, b] = [3.00, 6.00]$. This distribution is shown to subjects, both in the instructions and on the screen where they conduct price search. If they purchase product A at price p , their payoff from the search task equals

$b - p$. Upon entering the search screen they can also push a button to indicate that they do not want to search at all. In this case, their payoff from the search task is zero.

To see the price of an online shop, subjects have to enter a 16-digit code. We disable the copy-and-paste option so that subjects have to write it down on some device or take a picture before entering it on the next screen. The code varies between shops and subjects. After observing the price at a shop, subjects see an overview with all prices discovered so far. They can then purchase product A at any previously sampled shop (by clicking on the corresponding price) or continue the search spell. Hence, recall is essentially for free. If a subject wishes to purchase product A at a price that is not the smallest one discovered so far, a warning message appears with the offer to purchase it at the lowest available price (this offer can be declined). In the first part of the experiment, subjects have to enter an example code so that they are informed about the physical costs of search.

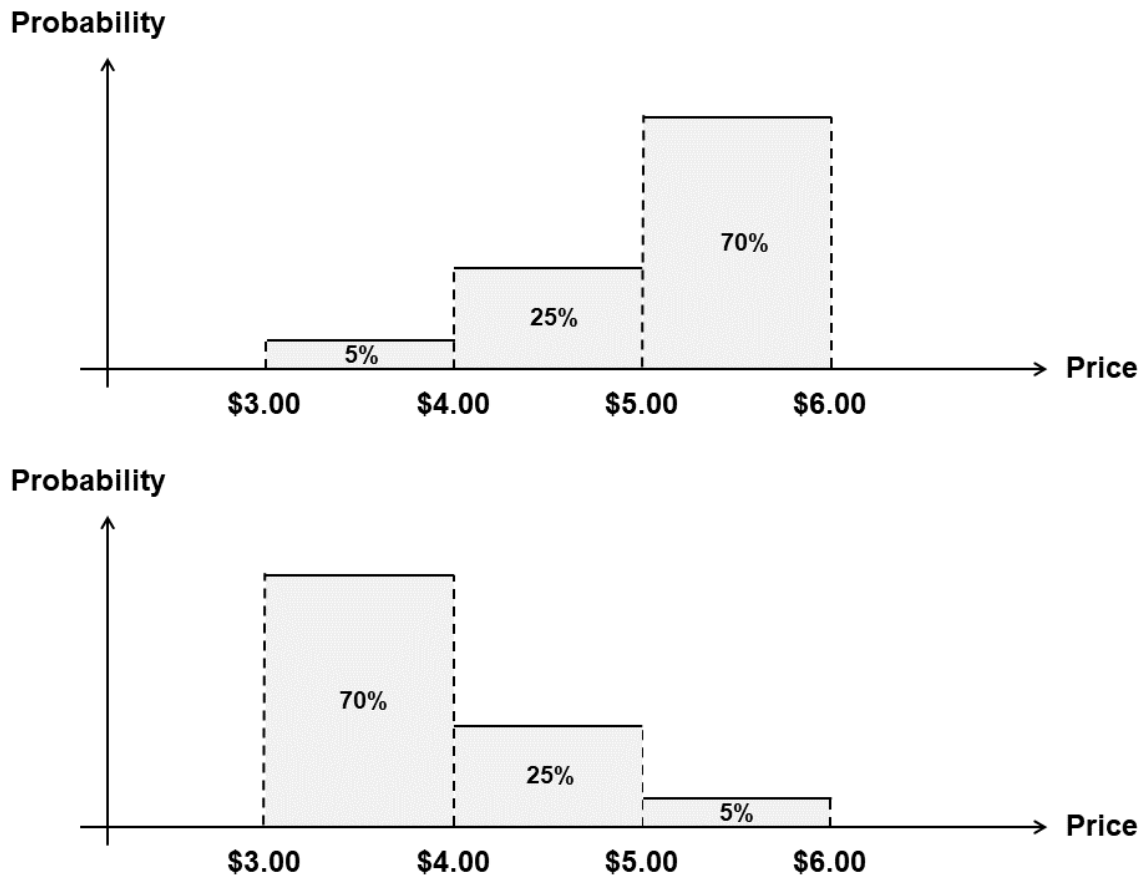


Figure 1: Left-skewed price distribution in treatment $S\ 1.0 - l$ (upper graph) and right-skewed price distribution in treatment $S\ 1.0 - r$ (lower graph). These graphs were also shown in the experimental instructions as well as on the search screen during the experiment.

Treatment Variation. We implement twelve different treatments that we classify in five treatment types: *information*, *no-information*, *piece rate*, *scale*, and *search cost treatments*. In the following, we explain each type of treatment and the objective behind the treatment variation.

We implement three *information treatments* in which we vary the price distribution. The first information treatment is the baseline treatment and denoted by $S\ 1.0$. The other two information treatments are identical to the baseline treatment (including the support of the price distribution), except that the price distributions in these treatments are either skewed to the left (more probability mass on high prices) or skewed to the right (more probability mass on low prices). Figure 1 displays the two distributions. We call these two treatments $S\ 1.0 - l$ and $S\ 1.0 - r$, respectively.

Next, we have three *no-information treatments* which are identical to the information treatments $S\ 1.0$, $S\ 1.0 - l$, and $S\ 1.0 - r$, except that subjects do not receive any information about the price distribution, neither in the instructions nor on the search screen. Subjects are informed that the highest price at each shop is $b = 6.00$ USD and we explicitly state in the instructions that no further information about the price distribution is provided. We call these treatments $S\ 1.0 - \text{noinfo}$, $S\ 1.0 - l.\text{noinfo}$, and $S\ 1.0 - r.\text{noinfo}$, respectively.

Further, we implement three *piece rate treatments*. They are identical to the baseline treatment, except that the price interval is $[0.2a, 0.2b] = [0.60, 1.20]$ and that subjects earn a positive amount w for each shop that they search, in addition to the realized price savings. We implement piece rate treatments with $w = 0.04$ USD, $w = 0.08$ USD, and $w = 0.12$ USD, and we denote these treatments by $S\ 0.2 - pr4$, $S\ 0.2 - pr8$, and $S\ 0.2 - pr12$, respectively. The instructions to the experiment have the same structure as in the baseline treatment; only on the last screen we announce that, in addition to the gains from search, subjects would earn the amount w for each completed search.

With the piece rate treatments we can identify the shape of the search cost function c . Note that search incentives vanish as searchers sequentially find smaller prices. Therefore, we would need a lot of observations (to get sufficiently many unlucky searchers who find low prices only at later searches) if we rely on search incentives alone to identify the level of convexity of c , especially when $c(n)$ increases slowly in the number of searches n .

Next, we have one *scale treatment*. It is identical to the baseline treatment except that the price interval is $[5.0a, 5.0b] = [15.00, 30.00]$ USD instead of $[a, b]$ USD. We denote this treatment by $S\ 5.0$. The potential gains from search are larger in the scale treatment than in the baseline treatment. With the scale treatment we can identify the shape of function v , that is, subjects' perception of the gains from search relative to the physical cost of search.

Finally, we have two *search cost treatments* that we denote by $S\ 1.0 - sc4$ and $S\ 1.0 - sc16$. They are identical to the baseline treatment, but may vary in the length of the code that subjects

have to enter in order to see a price quote. In treatment $S\ 1.0 - sc4$, subjects have to enter a 4-digit code, while in treatment $S\ 1.0 - sc16$, they have to insert a 16-digit code (as in the baseline treatment). With the search cost treatments, we can verify that our experimental task captures time- and hassle costs of price search.

Table 1: Demographic Variables and Labor Supply

Sample, Treatment	N	Share Females	Age	Average hourly Earnings	Average hours per Week
<i>Information and No-Information Treatments</i>					
AMT $S\ 1.0$	213	0.441	34.88 (10.09)	12.62 (14.10)	35.90 (18.29)
AMT $S\ 1.0 - \text{noinfo}$	210	0.419	35.66 (10.32)	11.41 (13.40)	34.49 (20.62)
AMT $S\ 1.0 - l$	202	0.431	33.42 (8.75)	11.17 (12.70)	38.00 (22.75)
AMT $S\ 1.0 - l.\text{noinfo}$	192	0.385	34.31 (9.33)	11.41 (13.28)	38.01 (21.06)
AMT $S\ 1.0 - r$	199	0.342	34.89 (9.59)	11.29 (14.02)	39.88 (25.53)
AMT $S\ 1.0 - r.\text{noinfo}$	209	0.340	35.04 (9.39)	10.76 (13.39)	38.24 (21.48)
PRO $S\ 1.0 - l$	115	0.591	39.84 (13.63)	8.57 (4.65)	11.00 (8.74)
PRO $S\ 1.0 - l.\text{noinfo}$	134	0.530	38.90 (12.12)	8.77 (5.97)	10.82 (8.52)
PRO $S\ 1.0 - r$	152	0.605	37.55 (12.24)	8.14 (3.90)	11.33 (8.27)
PRO $S\ 1.0 - r.\text{noinfo}$	161	0.590	39.43 (12.55)	7.62 (4.06)	10.98 (8.82)
<i>Piece Rate, Scale, and Search Cost Treatments</i>					
AMT $S\ 0.2 - pr4$	216	0.407	34.55 (10.04)	11.71 (12.74)	34.60 (18.08)
AMT $S\ 0.2 - pr8$	202	0.396	34.99 (9.81)	13.44 (14.20)	33.38 (18.57)
AMT $S\ 0.2 - pr12$	194	0.418	35.37 (10.64)	12.57 (13.08)	35.35 (18.03)
AMT $S\ 5.0$	190	0.416	34.57 (10.34)	11.94 (12.88)	35.26 (22.54)
PRO $S\ 1.0 - sc4$	188	0.479	39.12 (12.54)	11.44 (6.07)	16.56 (14.33)
PRO $S\ 1.0 - sc16$	150	0.507	37.77 (10.56)	11.21 (5.83)	18.72 (15.42)

Notes: Standard deviation in parentheses. Age in years, average hourly earnings in USD.

Experimental Procedures. In the first part of the experiment, we measure subjects' cognitive ability, willingness to take risk, and trust.⁵ Further, we measure online labor supply by asking subjects about their expected hourly earnings and about how many hours per week they spend working on the online platform (Amazon Mechanical Turk or Prolific) on average. Before subjects can commence price search, they have to answer a comprehension question.⁶ We programmed the experimental software with oTree (Chen et al., 2016) and ran the experiment with online workers from Amazon Mechanical Turk (AMT) as well as with subjects from Prolific (PRO). Specifically, we ran the treatments $S0.2 - pr4$, $S0.2 - pr8$, $S0.2 - pr12$, $S5.0$, $S1.0$, and $S1.0 - noinfo$ on AMT in November 2023, the treatments $S1.0 - l$, $S1.0 - r$, $S1.0 - l.noinfo$, and $S1.0 - r.noinfo$ with both AMT workers and Prolific subjects in June 2024, and the treatments $S1.0 - sc4$ and $S1.0 - sc16$ on Prolific in April 2026.

Before starting the experiment, we pre-registered it on the AEA RCT Registry (registry number AEARCTR-0012497) and obtained IRB approval from the Board for Ethical Questions in Science of the University of Innsbruck. In the pre-registration, we indicate that we exclude subjects from the final sample who (i) do not search at least one shop and also do not indicate that they do not want to search at all (by pressing the corresponding button on the search screen) as well as subjects who (ii) conduct at least one search, but who do not purchase the product at the lowest observed price.⁷

Table 1 provides an overview of all treatment samples. Each of our experimental samples is balanced in terms of demographic variables, labor supply, as well as education and personal characteristics (cognitive ability, willingness to take risks, and trust); see Appendix A.6. In particular, this holds for each pair of information and no-information treatment with the same price distribution.⁸

⁵We measure cognitive ability through a cognitive reflection test. We elicit the general willingness to take risk as in Dohmen et al. (2011) by asking the question: *Do you generally avoid taking risks or are you a risk-taker?* This question has to be answered on a scale between zero (not at all willing to take risks) and ten (very willing to take risks). To elicit trust, we ask the standard question from the world values survey: *Generally speaking, would you say that you need to be very careful in dealing with people or that most people can be trusted?* The response to this question has to be provided on a scale between zero (people cannot be trusted at all) and ten (people can fully be trusted).

⁶Specifically, we ask the following question: *Suppose that after searching two times for the lowest price, a shopper buys Product A for USD [70 percent of the maximal price in the treatment]. What will be the shopper's bonus?* Subjects have to choose the answer from a list of four possible options. If they indicate the wrong answer, we explain to them how to find the correct value and they get a second chance to provide the correct response. All subjects are admitted to the search task.

⁷For organizational reasons, the experiment on Prolific differed from the one on AMT in two details: First, Prolific subjects earn 2.00 USD for the completion of the first part of the experiment and an additional bonus of 0.10 USD for completing the search task (regardless of how many shops they search). Second, Prolific subjects have to complete the search task right after finishing the first part of the experiment, while AMT workers have three days for searching (since almost all AMT workers start searching right away, it is unlikely that this difference matters for our results).

⁸The only exception is that the CRT score differs between treatments in the November 2023 AMT sample.

4 Search Cost Estimation and Evaluation

4.1 Search Cost Estimation

We derive the empirical search models from the search paradigms of Section 2. For this, we have to specify the shape of the function v and the shape of the search cost function c . To parametrize context effects, we assume

$$v(p, F) = \frac{p}{(\max\{1, \Delta_F\})^\rho}, \quad (7)$$

where $\Delta_F = b - a$. This function captures relative thinking of degree ρ . A similar function has been used by [Somerville \(2022\)](#) to quantify the extent of relative thinking. Relative thinking of degree $\rho = 0$ implies $v(p, F) = p$ so that the DM takes absolute price savings fully into account (as in most empirical search models), while $\rho = 1$ implies scale-independent search behavior: This means that the DM exerts the same search effort, regardless of whether all prices (and hence price savings) are scaled up by some factor $z > 1$ or not. Based on the results from [Karle et al. \(2025\)](#) we expect that our estimate of ρ is close to one. To parametrize the search cost function, we follow [DellaVigna and Pope \(2018\)](#) and assume the functional form

$$c(n) = \frac{kn^{1+\kappa}}{1+\kappa}, \quad (8)$$

where $\kappa \geq 0$. The value $\kappa = 0$ implies that marginal search costs are constant in the number of searches, as it is frequently assumed in theoretical search models; $\kappa > 0$ captures convex search costs. One mechanism that may generate convex search costs is fatigue ([Ursu et al., 2023](#)). However, we do not expect that fatigue plays an important role in our setup. There is no time pressure in the experimental task and the task difficulty remains constant over time. Further, the time needed to identify a price quote (around one minute) is fairly small compared to the average time subjects work on AMT and Prolific, see Table 1. Thus, a significant level of convexity would be difficult to reconcile with the labor supply on these platforms. We therefore expect that the estimated value of κ is close to zero.

Sequential Search. We first outline how we estimate search costs assuming optimal sequential search. Suppose we know that the DM searches n shops and that her reservation value at the n 'th shop equals r_n . For given parameter values ρ, κ and a uniform distribution F on the interval $[a, b]$ with $\Delta_F \geq 1$ we can derive her search cost parameter k from equation (2) as

$$k^{\text{seq}}(r_n, \rho, \kappa, n) = \frac{1 + \kappa}{n^{1+\kappa} - (n - 1)^{1+\kappa}} \left(\frac{1}{\Delta_F^\rho} \frac{(r_n - a)^2}{2(b - a)} + w \right). \quad (9)$$

In Appendix A.2, we generalize this term to the left- and right-skewed price distributions in our experiment. Since we do not observe the reservation price r_n directly, we have to derive it from the discovered prices. For this, we assume that DM i 's search cost parameter k_i is drawn from a lognormal distribution G_i , truncated at t . Let x_i be the observable characteristics of DM i . The log of DM i 's search cost parameter k_i is then given by the equation

$$\ln k_i = x_i\beta + \sigma\varepsilon_i, \quad (10)$$

where ε_i follows a standard normal distribution Φ , β is a vector of parameters affecting the mean of the underlying normal distribution, and σ is the standard deviation of the underlying normal distribution, see Appendix A.3 for computational details (in a robustness check, we also consider a more flexible specification of the search cost distribution). For each DM i with number of searches $n_i \in \{2, \dots, 99\}$, we observe the two smallest prices p_i^1, p_i^2 discovered before the DM purchased the product. The reservation price r_n must be located in the interval $[p_i^1, p_i^2]$. For given parameters ρ and κ , we obtain the likelihood contribution

$$\begin{aligned} P_i &= \Pr(k^{\text{seq}}(p_i^1, \rho, \kappa, n_i) \leq k_i < k^{\text{seq}}(p_i^2, \rho, \kappa, n_i)) \\ &= \left[\Phi\left(\frac{\ln k^{\text{seq}}(p_i^2, \rho, \kappa, n_i) - x_i\beta}{\sigma}\right) - \Phi\left(\frac{\ln k^{\text{seq}}(p_i^1, \rho, \kappa, n_i) - x_i\beta}{\sigma}\right) \right] \Phi\left(\frac{\ln t - x_i\beta}{\sigma}\right)^{-1}. \end{aligned} \quad (11)$$

For censored observations with $n_i = 1$ or $n_i = N$, we adjust the likelihood contribution accordingly; see Appendix A.2. Adding up the logarithm of these terms for all subjects yields the log-likelihood function for maximum likelihood estimation. To ensure that all observed outcomes can be rationalized through search costs drawn from G , we set the truncation point $t = 3 \max_i k^{\text{seq}}(p_i^1, \rho, \kappa, n_i)$ during the estimation. We use the treatments with variation in scales to identify context effects ρ , and those with variation in the piece rates to identify κ . Given these, we obtain the treatment-specific parameters of the respective search cost distributions.

Non-sequential Search. Next, we describe how we estimate search costs assuming optimal non-sequential search. Consider a DM who plans to search $n \geq 1$ shops. For given parameter values ρ , κ and a uniform distribution F on the interval $[a, b]$ with $\Delta_F \geq 1$, her expected expenses can be calculated from equation (4) as

$$\mathbb{E}^{[n]}(v) = \frac{1}{\Delta_F^\rho} \left(a + \frac{b-a}{n+1} \right). \quad (12)$$

In Appendix A.2, we derive the corresponding term for left- and right-skewed price distributions, respectively. From the fact that the DM weakly prefers searching n shops to searching

$n + 1$ shops we obtain a lower bound on the search cost parameter k from equation (6):

$$k_-^{\text{nsq}}(\rho, \kappa, n) = \frac{1 + \kappa}{(n + 1)^{1+\kappa} - n^{1+\kappa}} \left(\mathbb{E}^{[n]}(v) - \mathbb{E}^{[n+1]}(v) + w \right). \quad (13)$$

Similarly, we obtain an upper bound on k from the fact that the DM weakly prefers searching n shops to searching $n - 1$ shops:

$$k_+^{\text{nsq}}(\rho, \kappa, n) = \frac{1 + \kappa}{n^{1+\kappa} - (n - 1)^{1+\kappa}} \left(\mathbb{E}^{[n-1]}(v) - \mathbb{E}^{[n]}(v) + w \right). \quad (14)$$

For given parameters ρ and κ , the true value of the search cost parameter k_i for DM i must be in the interval between $k_-^{\text{nsq}}(\rho, \kappa, n_i)$ and $k_+^{\text{nsq}}(\rho, \kappa, n_i)$. Assuming that the log of search costs is normally distributed, we obtain the likelihood contribution

$$\begin{aligned} P_i &= \Pr(k_-^{\text{nsq}}(\rho, \kappa, n_i) \leq k_i < k_+^{\text{nsq}}(\rho, \kappa, n_i)) \\ &= \left[\Phi \left(\frac{\ln k_+^{\text{nsq}}(\rho, \kappa, n_i) - x_i \beta}{\sigma} \right) - \Phi \left(\frac{\ln k_-^{\text{nsq}}(\rho, \kappa, n_i) - x_i \beta}{\sigma} \right) \right] \Phi \left(\frac{\ln t - x_i \beta}{\sigma} \right)^{-1}. \end{aligned} \quad (15)$$

For censored observations we adjust the likelihood contribution accordingly; see Appendix A.2. As for the sequential search model, we estimate ρ from the treatments with scale variation, and subsequently the treatment-specific search cost distribution parameters. However, the parameter κ is not identified under non-sequential search due to the interaction of the piece rate w and the number of searches n . Thus, we impute the value of κ that we obtain from the sequential search model also to the non-sequential search setting.

4.2 Evaluation of Search Cost Estimates

Each estimation procedure generates a fully specified distribution over search costs. This distribution can be inserted back into the search model to obtain predictions about the amount of search and realized purchase prices. This allows us to evaluate the validity of the search cost estimates by comparing predicted and realized outcomes in the experiment.

We proceed as follows: Let $\hat{G}$ be any given estimated distribution over the cost parameter k . For any given value k and price distribution F , we calculate (or numerically approximate) the probability distribution over the number of searches n and purchase price p , for each of the search models. In Appendix A.4, we derive the analytical expressions assuming $\kappa = 0$ and $w = 0$. Integrating over $\hat{G}$ we derive the predicted distribution over the number of searches, $\hat{G}^{[1]}(n)$ for $n \in \{0, 1, \dots, 100\}$ and the predicted distribution over purchase prices $\hat{G}^{[2]}(p)$ for

$p \in \{3.00, 3.01, \dots, 6.00\}$.⁹ Let $\hat{g}^{[1]}$ and $\hat{g}^{[2]}$ be the corresponding densities. Accordingly, define by $G^{[1]}$ and $G^{[2]}$ the realized distributions over the number of searches and purchase prices, respectively, with densities $g^{[1]}$ and $g^{[2]}$.

First, we consider the difference in means (and medians) between realized and predicted distribution for each information and no-information treatment. We call this difference *prediction error in means (medians)*.¹⁰ To quantify the variability of the point estimate of the prediction error in means, we use the standard error, as defined by

$$SE = \sqrt{\frac{\sigma_{G^{[j]}}^2}{l_{G^{[j]}}} + \frac{\sigma_{\hat{G}^{[j]}}^2}{l_{\hat{G}^{[j]}}}} \quad (16)$$

for $j \in \{1, 2\}$, where $\sigma_{G^{[j]}}$ is the standard deviation of the realized values of variable j (number of searches or purchase prices) and $l_{G^{[j]}}$ is the number of observations in the considered treatment. Further, $\sigma_{\hat{G}^{[j]}}$ is the standard deviation in a sample of 1,000 random draws from the predicted distribution $\hat{G}^{[j]}$ (thus, the value $l_{\hat{G}^{[j]}}$ equals 1,000).

Next, we evaluate how close the realized and predicted distribution are to each other. There exist various methods to measure the similarity of probability distributions. In our context, the most convenient way to quantify how close two distributions are to each other is the Hellinger distance (HD). For the discrete distributions $G^{[1]}$ and $\hat{G}^{[1]}$, it is defined as

$$H(G^{[1]}, \hat{G}^{[1]}) = \frac{1}{\sqrt{2}} \sqrt{\sum_{n \in \{0, 1, \dots, 100\}} \left(\sqrt{g^{[1]}(n)} - \sqrt{\hat{g}^{[1]}(n)} \right)^2}, \quad (17)$$

and for the distributions $G^{[2]}$ and $\hat{G}^{[2]}$, it is given by

$$H(G^{[2]}, \hat{G}^{[2]}) = \frac{1}{\sqrt{2}} \sqrt{\sum_{p \in \{3.00, 3.01, \dots, 6.00\}} \left(\sqrt{g^{[2]}(p)} - \sqrt{\hat{g}^{[2]}(p)} \right)^2}. \quad (18)$$

The Hellinger distance takes on values between zero and one, where the value zero indicates that the two distributions are identical and the value one is attained when there is no overlap in the support of the two distributions. In addition to the Hellinger distance, we report the Kullback-Leibler (KL) divergence. For the discrete distributions $G^{[1]}$ and $\hat{G}^{[1]}$, it is defined as

$$D_{KL}(G^{[1]} \parallel \hat{G}^{[1]}) = \sum_{n \in \{0, 1, \dots, 100\}} g^{[1]}(n) \log \left(\frac{g^{[1]}(n)}{\hat{g}^{[1]}(n)} \right). \quad (19)$$

⁹We numerically approximate the integral over $\hat{G}$ by obtaining 1000 simulation draws from the estimated search cost distribution.

¹⁰The literature on prediction accuracy and “Wisdom of the Crowd” also frequently uses the term *mean absolute error (MAE)* for this value; see, for example, [König-Kersting et al. \(2025\)](#).

This value can be interpreted as the amount of additional “surprise” if one expects the distribution to be $\hat{G}^{[1]}$ when the true distribution is given by $G^{[1]}$. For the distributions $G^{[2]}$ and $\hat{G}^{[2]}$, the KL divergence is given by

$$D_{KL}(G^{[2]} \parallel \hat{G}^{[2]}) = \sum_{p \in \{3.00, 3.01, \dots, 6.00\}} g^{[2]}(p) \log \left(\frac{g^{[2]}(p)}{\hat{g}^{[2]}(p)} \right). \quad (20)$$

The KL divergence is not a metric like the Hellinger distance. It satisfies neither symmetry nor the triangle inequality.

5 Results

We present our results as follows. In Subsection 5.1, we provide the descriptive statistics of search behavior and outcomes in our experiment. In Subsection 5.2, we examine the extent to which subjects’ search behavior is consistent with the search models. In Subsection 5.3, we present the results from the search cost estimation for each model. In Subsection 5.4 we evaluate how well the search cost estimates capture the outcomes from search in-sample and out-of-sample. In Subsection 5.5, we briefly discuss two extensions: the estimation and evaluation of a search without priors model and the evaluation of search models according to an alternative outcome variable. In Subsection 5.6, we evaluate the prediction performance of all models and compare it to two benchmarks.

5.1 Descriptive Statistics

Table 2 provides an overview of subjects’ search behavior and search outcomes in our experiment. It shows the share of searchers (subjects who search at least one shop), the mean and median number of searches, as well as the mean and median purchase price (the purchase price is set to b for subjects who do not search at all). In the last column, Table 2 shows the p -values of two-sided t-tests which examine whether there is a statistically significant difference in the number of searches (or the realized purchase price) between a given information treatment and the corresponding no-information treatment with the same price distribution.

We make five important observations from Table 2. First, the share of searchers is fairly large, between 91.7 and 100.0 percent, so most subjects conduct at least one search. There are mostly no significant differences in the share of searchers between a given information treatment and its corresponding no-information treatment (t-test p -values > 0.050). Thus, whether or not subjects obtain information about the price distribution does not seem to matter for their decision to commence price search.

Table 2: Average Search Outcomes

Sample Treatment	Share Searchers	Mean and Median No. Searches		Mean and Median Purchase Price		Diff. p -value
<i>Information and No-Information Treatments</i>						
AMT S 1.0	0.972	10.49 (21.13)	2	3.76 (0.88)	3.40	0.344/0.705
AMT S 1.0 – noinfo	0.957	12.49 (22.12)	3	3.79 (0.89)	3.39	
AMT S 1.0 – l	0.980	20.56 (28.22)	10	4.07 (0.88)	3.92	0.849/0.225
AMT S 1.0 – l .noinfo	1.000	20.02 (27.73)	8	4.18 (0.91)	4.07	
AMT S 1.0 – r	0.975	13.78 (22.83)	4	3.41 (0.64)	3.14	0.970/0.483
AMT S 1.0 – r .noinfo	0.995	13.70 (22.61)	5	3.37 (0.52)	3.16	
PRO S 1.0 – l	0.922	6.35 (6.47)	4	4.43 (0.74)	4.29	0.132/0.574
PRO S 1.0 – l .noinfo	0.918	8.22 (11.87)	5	4.37 (0.82)	4.28	
PRO S 1.0 – r	0.954	2.97 (3.26)	2	3.60 (0.68)	3.37	0.000/0.028
PRO S 1.0 – r .noinfo	0.963	6.55 (7.78)	4	3.44 (0.62)	3.24	
<i>Piece Rate, Scale, and Search Cost Treatments</i>						
AMT S 0.2 – pr 4	0.917	16.68 (32.09)	2	0.81 (0.21)	0.72	–
AMT S 0.2 – pr 8	0.970	23.28 (37.75)	2	0.76 (0.18)	0.69	
AMT S 0.2 – pr 12	0.974	18.96 (34.17)	2	0.77 (0.19)	0.69	
AMT S 5.0	0.974	10.55 (20.09)	2	18.94 (4.73)	16.52	–
PRO S 1.0 – sc 4	0.968	7.46 (10.16)	4	3.62 (0.73)	3.33	0.000/0.001
PRO S 1.0 – sc 16	0.920	3.88 (4.24)	3	3.92 (0.90)	3.56	

Notes: Standard deviation in parentheses. The first (second) p -value in the last column originates from a two-sided t-test which compares the mean number of searches (purchase price) between an information treatment and the no-information treatment with the same price distribution.

Second, if we compare the information treatment AMT S 1.0 and the scale treatment AMT S 5.0, we observe that the amount of search is fairly similar in these two treatments, despite the fact that search incentives are five times larger in the latter than in the former treatment. This suggests that the context effect parameter ρ should be around one, meaning that subjects use relative price differences to balance the gains and losses from search rather than absolute price differences. The search cost estimation in Subsection 5.3 will confirm this conjecture. Thus, the extent of context effects in our experiment is consistent with those documented for AMT workers and Prolific subjects in Karle et al. (2025).

Third, the availability of information about the price distribution has mostly no effect on the average number of searches and purchase prices. For four out of five pairs of information and

no-information treatments we do not find significant differences in the number of searches or purchase prices. Only for the Prolific treatment pair $S\ 1.0 - r$ and $S\ 1.0 - r.\text{noinfo}$ we observe that subjects search more and on average realize lower purchase prices when they are not informed about the price distribution. The differences are significant at least on the 5-percent level. In all other treatments, subjects on average realize roughly the same gains from search, regardless of whether they know the price distribution or not.

Fourth, subjects react to the price distribution: They search more shops if the price distribution is left-skewed (more probability mass on high prices) than if it is right-skewed (more probability mass on low prices). Combining the data from the information and no-information treatments, this effect is significant for both the AMT (t-test p -value < 0.001) and the Prolific sample (t-test p -value < 0.001). There are also significant differences in the realized purchase prices, but this effect is of course (in part) driven by the different price distributions. We conclude that the amount of search and realized purchase prices are relatively unresponsive to the availability of price distribution information, but responsive to the price distribution.

Fifth, subjects also react to physical search costs: They search significantly more shops if the length of the code they have to enter at each shop is only 4 digits (instead of 16 digits). Accordingly, there are also significant differences in realized purchase prices. We therefore believe that the search task in our experiment captures the real effort costs of search.

5.2 Empirical Predictions: Results

We use the data from the information- and no-information treatments to examine the empirical predictions from Section 2. The detailed results can be found in Appendix A.7. There is a substantial amount of recall in our experiment: At least 29.7 percent of subjects recall a previously sampled shop (without having searched all shops) in a given treatment. This behavior is inconsistent with sequential search when search costs are constant in the number of searches. However, it is consistent with non-sequential search. Further, there is a significant level of price-dependency of search in six treatments: The probability of search spell continuation tends to increase in the price of the last sampled shop. This observation is inconsistent with non-sequential search. We find few systematic differences between information and no-information treatments. Overall, the important observation here is that none of the considered search models is fully consistent with subjects' search behavior in all treatments.

5.3 Search Cost Estimation: Results

We estimate search costs according to the procedures outlined in Section 4. Table 3 shows the estimation results for the sequential and Table 4 for the non-sequential search model.

Table 3: Search Cost Estimates – Sequential Search

	(1)	(2)	(3)	(4)	(5)
<i>Parameter Estimates</i>					
<i>S</i> 5.0	-0.67* (0.28)		-2.94*** (0.39)		
<i>S</i> 1.0	-2.35*** (0.23)		-2.42*** (0.40)		
<i>S</i> 1.0 – noinfo	-2.54*** (0.23)		-2.75*** (0.38)		
<i>S</i> 1.0 – <i>l</i>				-3.92*** (0.22)	-3.14*** (0.19)
<i>S</i> 1.0 – <i>l</i> .noinfo				-3.83*** (0.22)	-3.37*** (0.18)
<i>S</i> 1.0 – <i>r</i>				-4.43*** (0.22)	-2.81*** (0.18)
<i>S</i> 1.0 – <i>r</i> .noinfo				-4.42*** (0.21)	-4.21*** (0.16)
β		-1.81*** (0.16)			
σ	2.87*** (0.12)	2.47*** (0.12)	3.37*** (0.20)	2.65*** (0.09)	1.87*** (0.07)
κ		0.00 (0.03)			
ρ		1.00*** (0.00)			
Data	AMT	AMT	AMT	AMT	PRO
Treatments	<i>S</i> 5.0, <i>S</i> 1.0, <i>S</i> 1.0 – noinfo	<i>S</i> 5.0, <i>S</i> 1.0, <i>S</i> 1.0 – <i>pr</i> 4, <i>S</i> 1.0 – <i>pr</i> 8, <i>S</i> 1.0 – <i>pr</i> 12	<i>S</i> 5.0, <i>S</i> 1.0, <i>S</i> 1.0 – noinfo	<i>S</i> 1.0 – <i>l</i> , <i>S</i> 1.0 – <i>r</i> , <i>S</i> 1.0 – <i>l</i> .noinfo, <i>S</i> 1.0 – <i>r</i> .noinfo	<i>S</i> 1.0 – <i>l</i> , <i>S</i> 1.0 – <i>r</i> , <i>S</i> 1.0 – <i>l</i> .noinfo, <i>S</i> 1.0 – <i>r</i> .noinfo
ρ est.	fixed 0	est.	fixed at est.	fixed at est.	fixed at est.
κ est.	fixed 0	est.	fixed at est.	fixed at est.	fixed at est.
Observations	613	1015	613	802	562
Log.Lik.	-1230.89	-2451.85	-1205.23	-1613.49	-1061.22
AIC	2469.79	4911.70	2418.46	3236.99	2132.44
BIC	2487.46	4931.39	2436.13	3260.42	2154.10
<i>Implied Search Costs</i>					
SC <i>S</i> 5.0	2.093 (3.972)		0.161 (0.291)		
SC <i>S</i> 1.0	1.009 (2.709)		0.185 (0.310)		
SC <i>S</i> 1.0 – noinfo	0.915 (2.567)		0.169 (0.298)		
SC <i>S</i> 1.0 – <i>l</i>				0.130 (0.298)	0.149 (0.285)
SC <i>S</i> 1.0 – <i>l</i> .noinfo				0.136 (0.306)	0.127 (0.258)
SC <i>S</i> 1.0 – <i>r</i>				0.100 (0.258)	0.184 (0.322)
SC <i>S</i> 1.0 – <i>r</i> .noinfo				0.101 (0.259)	0.067 (0.172)

Notes: Ordered probit regressions with truncation for the sequential search model. In the top panel, standard errors of the estimated parameters are in parentheses. The bottom panel presents the search cost estimates for each treatment, derived from the parameter estimates shown in the top panel. The first value represents the mean of the implied search cost distribution, with the standard deviation of the distribution provided in parentheses. Significance at * $p < 0.1$, ** $p < 0.05$, and *** $p < 0.01$.

Table 4: Search Cost Estimates – Non-Sequential Search

	(1)	(2)	(3)	(4)	(5)
<i>Parameter Estimates</i>					
<i>S</i> 5.0	0.20 (0.28)		-1.77*** (0.46)		
<i>S</i> 1.0	-1.89*** (0.22)		-1.70*** (0.45)		
<i>S</i> 1.0 – noinfo	-2.23*** (0.21)		-2.26*** (0.41)		
<i>S</i> 1.0 – <i>l</i>				-3.80*** (0.22)	-3.07*** (0.16)
<i>S</i> 1.0 – <i>l</i> .noinfo				-3.63*** (0.23)	-3.27*** (0.15)
<i>S</i> 1.0 – <i>r</i>				-3.99*** (0.22)	-2.54*** (0.15)
<i>S</i> 1.0 – <i>r</i> .noinfo				-4.09*** (0.21)	-3.74*** (0.14)
β		-1.89*** (0.14)			
σ	2.79*** (0.12)	2.32*** (0.10)	3.40*** (0.22)	2.72*** (0.10)	1.67*** (0.06)
ρ		1.00*** (0.00)			
Data	AMT	AMT	AMT	AMT	PRO
Treatments	<i>S</i> 5.0, <i>S</i> 1.0, <i>S</i> 1.0 – noinfo	<i>S</i> 5.0, <i>S</i> 1.0, <i>S</i> 1.0 – <i>pr</i> 4, <i>S</i> 1.0 – <i>pr</i> 8, <i>S</i> 1.0 – <i>pr</i> 12	<i>S</i> 5.0, <i>S</i> 1.0, <i>S</i> 1.0 – noinfo	<i>S</i> 1.0 – <i>l</i> , <i>S</i> 1.0 – <i>r</i> , <i>S</i> 1.0 – <i>l</i> .noinfo, <i>S</i> 1.0 – <i>r</i> .noinfo	<i>S</i> 1.0 – <i>l</i> , <i>S</i> 1.0 – <i>r</i> , <i>S</i> 1.0 – <i>l</i> .noinfo, <i>S</i> 1.0 – <i>r</i> .noinfo
ρ est.	fixed 0	est.	fixed at est.	fixed at est.	fixed at est.
Observations	613	1015	613	802	562
Log.Lik.	-1865.28	-3121.28	-1831.91	-2729.22	-1550.10
AIC	3738.57	6248.55	3671.83	5468.45	3110.20
BIC	3756.24	6263.32	3689.50	5491.88	3131.86
<i>Implied Search Costs</i>					
SC <i>S</i> 5.0	2.841 (4.575)		0.215 (0.332)		
SC 1.0	1.227 (2.991)		0.219 (0.334)		
SC <i>S</i> 1.0 – noinfo	1.034 (2.723)		0.192 (0.316)		
SC <i>S</i> 1.0 – <i>l</i>				0.141 (0.314)	0.138 (0.257)
SC <i>S</i> 1.0 – <i>l</i> .noinfo				0.153 (0.327)	0.120 (0.233)
SC <i>S</i> 1.0 – <i>r</i>				0.129 (0.299)	0.202 (0.324)
SC <i>S</i> 1.0 – <i>r</i> .noinfo				0.123 (0.292)	0.081 (0.180)

Notes: Ordered probit regressions with truncation for the non-sequential search model. In the top panel, standard errors of the estimated parameters are in parentheses. The bottom panel presents the search cost estimates for each treatment, derived from the parameter estimates shown in the top panel. The first value represents the mean of the implied search cost distribution, with the standard deviation of the distribution provided in parentheses. Significance at * $p < 0.1$, ** $p < 0.05$, and *** $p < 0.01$.

Sequential Search. We first estimate search costs assuming that there are no context effects ($\rho = 0$) and that the marginal search costs are constant ($\kappa = 0$), see specification (1) in Table 3. We apply this model to the treatments $S5.0$, $S1.0$, and $S1.0 - \text{noinfo}$ from AMT. As expected, the estimated coefficients (which govern the respective search cost estimates for these treatments) are different between treatment $S5.0$ and the other two treatments. Specifically, the estimated mean search cost for treatment $S5.0$ is 2.09 USD per search, compared to 1.01 USD and 0.92 USD per search in treatment $S1.0$ and treatment $S1.0 - \text{noinfo}$, respectively. Recall that subjects exert roughly the same search effort in all these treatments. The difference in search costs between treatments is significant (p -values < 0.001), which indicates that we need to take context effects into account to obtain realistic search cost estimates.

We identify the level of context effects ρ from the variation in scales (i.e., the treatments $S1.0$ and $S5.0$) and possible non-linearities in marginal search costs κ from the piece-rate treatments $S0.2 - pr4$, $S0.2 - pr8$, and $S0.2 - pr12$, see specification (2) in Table 3. The estimated values of ρ and κ are 1.00 and 0.00, respectively. Hence, the level of context effects in our setting is comparable to that in Karle et al. (2025). Next, we re-estimate search costs for the treatments $S5.0$, $S1.0$, and $S1.0 - \text{noinfo}$ from AMT, using the estimated values $\rho = 1.00$ and $\kappa = 0.00$, see specification (3) in Table 3. The estimated mean search costs for all treatments are now relatively close to each other, between 0.16 USD and 0.19 USD per search. The differences in search costs between these treatments are not significant (p -values > 0.273).

Using the estimated values $\rho = 1.00$ and $\kappa = 0.00$, we estimate search costs for the June 2024 AMT and Prolific samples, see specification (4) and (5), respectively, in Table 3. In the AMT sample, the estimated mean search costs are similar between treatments, varying from 0.10 USD to 0.14 USD. There are no significant differences between a given information treatment and the no-information treatment with the same price distribution (p -values > 0.760). In the Prolific sample, estimated mean search costs vary between 0.07 USD and 0.18 USD. There is no significant difference between the estimated search costs in PRO $S1.0 - l$ and PRO $S1.0 - l.\text{noinfo}$ (p -value = 0.373). However, the difference in estimated search costs is statistically significant for PRO $S1.0 - r$ and PRO $S1.0 - r.\text{noinfo}$ (p -value < 0.001). This is the only pair of information and no-information treatment in our samples where subjects' search behavior depends on the availability of price information.

Non-Sequential Search. We consider the same specifications for non-sequential search as for the sequential search model, see Table 4. The only difference is that we take $\kappa = 0.00$ as given, since the parameter κ is not identified under non-sequential search due to the interaction of the piece rate w and the number of searches n . We again obtain a level of context effects of $\rho = 1.00$. The search cost estimates are somewhat larger for non-sequential than for sequential search. In the November 2023 AMT sample, they vary between 0.19 USD and 0.22 USD;

the difference between the information and the no-information treatment is not statistically significant (p -value = 0.224). In the June 2024 AMT sample, search costs vary between 0.12 USD and 0.15 USD; the difference between a given information treatment and the no-information treatment with the same price distribution is again not significant (p -values > 0.561). Finally, in the June 2024 PRO sample, search costs vary between 0.08 USD and 0.20 USD. Here the difference in search costs is not significant between PRO $S 1.0 - l$ and PRO $S 1.0 - l.noinfo$ (p -value = 0.397), while the difference between PRO $S 1.0 - r$ and PRO $S 1.0 - r.noinfo$ is statistically significant (p -value < 0.001).

Physical Search Costs. Using the data from the search cost treatments, we verify that a change in subjects' physical search costs (the 4-digit and 16-digit treatments on Prolific) also impacts our estimation results. Using these treatments, we estimate search costs as for the other treatments, assuming $\rho = 1.00$ and $\kappa = 0.00$. Under the sequential search model, the estimated mean search costs are 0.08 USD when the code length is 4 digits and 0.14 USD when the code length is 16 digits. The difference is statistically significant (p -value < 0.001). Under the non-sequential search model these values are 0.03 USD and 0.04 USD, respectively. The difference is again significant (p -value < 0.001). Thus, under both models, reducing the real effort costs of search also reduces estimated search costs.

Direct Search Costs. In Appendix A.8, we compare these search cost estimates to subjects' opportunity costs of time (which we elicit through the first part of the experiment). It turns out that estimated search costs are roughly in line with subjects' opportunity costs of time.

Goodness of Fit Measures. We consider the goodness of fit of our empirical search models through the Akaike information criterion (AIC) and the Bayesian information criterion (BIC). These goodness of fit measures attempt to resolve the problem of model overfitting by introducing a penalty term for the number of parameters in the model.¹¹ All else equal, the preferred model is the one with the minimal AIC (or BIC) value. Tables 3 and 4 show for each model the AIC and BIC values. For a given specification and sample, the log-likelihood for the sequential search model is larger than that of the non-sequential search model (the number of parameters is identical for these models). Both AIC and BIC suggest that the sequential search model better fits the data.

Box-Cox Transformation. In our main specification, we model search costs as drawn from a truncated lognormal distribution, consistent with prior evidence that lognormality fits search costs well (Chen et al., 2007, Wildenbeest, 2011, Fischer et al., 2024). To assess the robustness

¹¹Let h be the number of estimated parameters in the model, ξ the number of data points, and $\hat{L}$ the maximized value of the likelihood function for the model. We then have $AIC = 2h - 2 \ln \hat{L}$ and $BIC = h \ln \xi - 2 \ln \hat{L}$.

of our estimates to this assumption, we relax lognormality and allow for more general distributions via a Box-Cox transformation with transformation parameter λ (Box and Cox, 1964, Birchall et al., 2024). The Box-Cox transformation nests both the lognormal and the normal distribution as special cases. We estimate the parameter λ and to the compare estimation results under the Box-Cox transformation to those under the lognormal restriction.

Online Appendix B.1 presents the details of the procedure as well as the estimation results. We are interested in whether the following outcomes are robust: Do we again get a level of context effects ρ close to 1 and a convexity parameter κ close to 0? Are estimated search costs similar in treatments from the same experimental sample? Are estimated search costs similar in treatments with and without price information? And is the model fit again better for the sequential than for the non-sequential search model? Reassuringly, the answer to all these questions is affirmative. The Box-Cox transformation delivers results that are fairly consistent with our baseline findings under the lognormal specification. For both search protocols, the search cost estimates from the same experimental sample as well as estimates from treatments with and without information about the price distribution are similar to each other. We therefore conclude that our results are not driven by the distributional assumption in our main specification. The relative pattern of predictive performance across models and outcome variables also remains qualitatively unchanged under the Box-Cox transformation.

5.4 Evaluation of Search Cost Estimates: Results

How well do the search cost estimates from the classic search models capture subjects' search outcomes in-sample and out-of-sample? In this subsection, we examine the prediction performance of the search models by following the procedure outlined in Subsection 4.2. For a given treatment, we obtain the distribution $\hat{G}$ over the search cost parameter k from the search cost estimation. Given this distribution, we derive, for each treatment, a predicted distribution $\hat{G}^{[1]}$ over the number of searches and a predicted distribution $\hat{G}^{[2]}$ over purchase prices. Then, for any treatment combination, we compute the absolute distance in means and medians between the predicted and realized distributions, $\hat{G}^{[1]}/G^{[1]}$ and $\hat{G}^{[2]}/G^{[2]}$, as well as the corresponding values of Hellinger distance and Kullback-Leibler divergence.

Table 5 presents the results in aggregated form. For in-sample predictions, we average how well the prediction from a given treatment captures behavior in the same treatment. For out-of-sample predictions, we average how well the prediction from a given treatment captures outcomes in other treatments from the same experimental sample. The detailed results can be found in Online Appendix B.2. Figure A1 to Figure A4 in the Appendix display in-sample predictions and outcomes graphically for each treatment. They are helpful for putting the numbers from Table 5 into context.

Table 5: Predicted versus Observed Distribution over Outcomes

	Prediction Error		Difference predicted and realized distribution	
	Means (SE)	Medians	HD	KL
Number of Searches				
<i>sequential search (in-sample)</i>				
information treatments	1.84 (1.32)	1.40	0.30	0.27
no-information treatments	2.52 (1.48)	1.70	0.32	0.33
all treatments	2.18 (1.40)	1.55	0.31	0.30
<i>non-sequential search (in-sample)</i>				
information treatments	2.25 (1.31)	0.40	0.30	0.30
no-information treatments	1.56 (1.48)	0.50	0.33	0.36
all treatments	1.91 (1.40)	0.45	0.32	0.33
<i>sequential search (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	2.83 (1.41)	1.85	0.32	0.33
any treatment → any treatment	3.43 (1.38)	1.87	0.32	0.33
<i>non-sequential search (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	2.64 (1.40)	1.35	0.35	0.38
any treatment → any treatment	3.67 (1.36)	1.48	0.34	0.36
Purchase Prices				
<i>sequential search (in-sample)</i>				
information treatments	0.05 (0.06)	0.07	0.33	0.33
no-information treatments	0.09 (0.06)	0.10	0.32	0.32
all treatments	0.07 (0.06)	0.09	0.32	0.32
<i>non-sequential search (in-sample)</i>				
information treatments	0.16 (0.06)	0.11	0.36	0.39
no-information treatments	0.09 (0.06)	0.05	0.33	0.31
all treatments	0.13 (0.06)	0.08	0.34	0.35
<i>sequential search (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	0.09 (0.06)	0.09	0.33	0.32
any treatment → any treatment	0.13 (0.06)	0.14	0.33	0.33
<i>non-sequential search (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	0.15 (0.06)	0.10	0.36	0.39
any treatment → any treatment	0.13 (0.06)	0.11	0.35	0.38

Notes: Average differences between predicted and observed distribution, as outlined in Subsection 4.2. The detailed values for each treatment combination are presented in Table B3 to Table B6 in the Online Appendix.

In-Sample Predictions. We first examine whether the sequential and the non-sequential search model capture search outcomes better when subjects have full information about the price distribution compared to the case when they do not have any information about the price distribution. One may hypothesize that the prediction performance of these models is higher when subjects are informed about the price distribution.

Let us first consider the sequential search model. The average prediction error in means in the information treatments is 1.84 for number of searches and 0.05 for purchase prices. The prediction error is slightly higher in the no-information treatments: Here the average prediction error in means is 2.52 for number of searches and 0.09 for purchase prices. We observe a similar pattern for the prediction error in medians, Hellinger distance, and Kullback-Leibler divergence.

Given the variability of the point estimates of the prediction errors in means, the difference in prediction performance seems small. The value in parentheses next to the prediction error in means is the average standard error (SE) for the prediction error. For every pair of information and no-information treatment, we find that the prediction error in means for the latter treatment is well within the 95-percent confidence interval of the prediction error in means for the former treatment.

Next, we consider the non-sequential search model. Here the pattern for the prediction performance is reversed. In the information treatments, the average prediction error in means is 2.25 for number of searches and 0.16 for purchase prices. In contrast, in the no-information treatments, the average prediction error in means is only 1.56 for number of searches and 0.09 for purchase prices. In terms of prediction error in medians, Hellinger distance, and Kullback-Leibler divergence, we mostly see a similar pattern. Again, the difference in prediction performance is modest as the prediction error in means of any information treatment is within the 95-percent confidence interval of the prediction error in means of the corresponding no-information treatment (and vice versa).

Interestingly, if we take the average prediction performance over both information and no-information treatments, the sequential and the non-sequential search model perform similarly well, both for number of searches and purchase prices. In terms of prediction error in medians, the non-sequential search model is even slightly better than the sequential search model. This result is remarkable for (at least) two reasons. First, we chose the experimental design so that the search task is as close as possible to the sequential search model. Second, as discussed in Subsection 5.3, the model fit in the regression is strictly better for the sequential than for the non-sequential search model.

Out-of-Sample Predictions. We further examine how well the search cost estimates from one treatment predict the search outcomes in other treatments. We do this in two steps.

First, we use the distribution over the search cost parameter from a given information (no-information) treatment to derive the prediction for the no-information (information) treatment with the same price distribution. The values for “inf./no-inf. $\rightarrow$ no-inf./inf.” in Table 5 show the corresponding prediction error in means and medians as well as the distribution distance measures. Second, we take any combination of information and/or no-information treatments from the November 2023 AMT sample, any combination of information and/or no-information treatments from the June 2024 AMT sample, and any combination of information and/or no-information treatments from the June 2024 PRO sample to compare out-of-sample predictions and realized outcomes. In the tables, the corresponding prediction errors are displayed in the lines labeled “any treatment $\rightarrow$ any treatment.”

As one may expect, the precision of the out-of-sample predictions is worse than the precision of the in-sample predictions; one can see that from comparing the out-of-sample predictions to the in-sample predictions that take all treatments into account. However, the differences in prediction quality are modest, given the overall variability of prediction errors. This holds for all search models and all outcome variables.

Consistent with the in-sample predictions, we observe that the sequential search model produces slightly more precise out-of-sample predictions than the non-sequential search model in terms of prediction error in means, Hellinger distance, and Kullback-Leibler divergence. Taking all treatments together, the average prediction error in means for the number of searches is 3.43 for the sequential search model and 3.67 for the non-sequential search model; the average prediction error in means for purchase prices is 0.13 USD for both the sequential search and the non-sequential search model. In terms of prediction error in medians, the non-sequential search model is slightly better than the sequential search model. Again, the differences in the prediction performance between the two models are fairly modest.

5.5 Search without Priors and Restricted Number of Searches

We briefly consider two extensions of our analysis. First, we apply our method to a search model where consumers do not have priors about the price distribution, but estimate the price distribution based on their price observations. We examine how well the search cost estimates from this model predict search outcomes in the information and no-information treatments. Second, we evaluate the prediction performance of our search models according to another outcome variable, the *restricted number of searches*.

Search without Priors. We consider a model of search without priors, building on Parakhonyak (2014). In this model, the DM does not know the price distribution F and has no prior about

it. Instead, she estimates the price distribution from observed prices. Let

$$p^{[n]} = (p^{1,n}, p^{2,n}, \dots, p^{n,n}) \quad (21)$$

be the ordered set of observed prices after n searches, where $p^{1,n}$ is the lowest and $p^{n,n}$ the highest price observed so far. We set $p^{[0]} = \emptyset$. The DM knows the upper bound of the support of the price distribution b and assumes that the lower bound equals θb with $\theta < 1$. If $\theta = \frac{a}{b}$, the DM has the correct belief about the support of the price distribution F . Throughout, we assume that $\theta \leq \frac{a}{b}$ so that there are no surprises for the DM. After n searches, the DM applies the quantile preserving estimation procedure from [Chou and Talmain \(1993\)](#) and [Parakhonyak \(2014\)](#) to form beliefs about the price distribution:

- (i) she assigns equal probability mass $\frac{1}{1+n}$ to each interval $[\theta b, p^{1,n}]$, $[p^{1,n}, p^{2,n}]$, ..., $[p^{n,n}, b]$ when these prices are different; she distributes the probability mass uniformly within each interval.

- (ii) she assigns probability mass $\frac{\xi-1}{n+1}$ to a price that appears ξ times in the sample.

Denote by $\hat{F}^{[n]}(p \mid p^{[n]})$ the DM's belief about the price distribution F after n searches if the set of observed prices is given by $p^{[n]}$. Accordingly, $\hat{F}^{[0]}(p \mid p^{[0]})$ is the uniform distribution on the interval $[\theta b, b]$. The estimator $\hat{F}^{[n]}$ is consistent, i.e., for $n \rightarrow \infty$ it converges uniformly to the true price distribution; see [Chou and Talmain \(1993\)](#). Under the proposed estimation procedure the DM conducts the n 'th search if

$$c(n) - c(n-1) \leq \int_{\theta b}^{p^{1,n-1}} (v(p^{1,n-1}, \hat{F}^{[n-1]}) - v(p, \hat{F}^{[n-1]})) \frac{1}{p^{1,n-1} - \theta b} \frac{1}{n} dp + w. \quad (22)$$

Thus, the DM's reservation price r_n^{swp} at the n 'th search is implicitly defined by

$$c(n) - c(n-1) = \int_{\theta b}^{r_n^{\text{swp}}} (v(r_n^{\text{swp}}, \hat{F}^{[n-1]}) - v(p, \hat{F}^{[n-1]})) \frac{1}{r_n^{\text{swp}} - \theta b} \frac{1}{n} dp + w. \quad (23)$$

Note that the reservation price r_n^{swp} depends on the set of observed prices $p^{[n-1]}$ only through the function v that captures context effects. If this function is independent of $p^{[n-1]}$, then also r_n^{swp} is independent of $p^{[n-1]}$.

We estimate search costs based on the search without priors model and evaluate it according to the same measures as the classic search models in the previous subsection, see Online Appendix [B.3](#) for details. One may conjecture that the search without priors model performs better in the no-information than in the information treatments. However, this is not what we find. In the information treatments, the average prediction error in means is 2.09 for number of

searches and 0.05 for purchase prices. In contrast, in the no-information treatments, the average prediction error in means is 2.57 for number of searches and 0.07 for purchase prices. The other measures mostly confirm that the prediction performance of the search without priors model is, if anything, slightly worse for the no-information treatments.

The prediction performance of the search without priors model is similar to that of the classic search models according to most indicators. For in-sample predictions, the average prediction error in means is 2.33 for number of searches and 0.06 for purchase prices (these values were 2.18 and 0.07 for the sequential search model; 1.91 and 0.13 for the non-sequential search model). For out-of-sample predictions, the average prediction error in means is 4.30 for number of searches and 0.15 for purchase prices (these values were 3.43 and 0.13 for the sequential search model; 3.67 and 0.13 for the non-sequential search model). Thus, the search without priors model performs similarly as the classic search models, but there does not seem to be a strong reason for using this model instead of the others.

Restricted Number of Searches. In empirical applications of search models, the number of shops is typically much smaller than 100. It therefore may be important to know how many individuals search only one or two shops (or do not search at all) and how many individuals search more than that. To capture this information in succinct manner, we compute for each observation the *restricted number of searches*. This value equals 0, 1, 2, 3, or “4 or more”, where the last value is assigned if a subject searches four or more shops (thus, all subjects who search four or more shops receive the same value; we do not drop any observation). We then apply the same steps as outlined in Subsection 4.2 to calculate the prediction error in means and medians as well as the Hellinger distance and Kullback-Leibler divergence with respect to this variable.

Online Appendix B.4 shows the detailed results. We find that the prediction performance of the two classic search models for the restricted number of searches is slightly better in the information than in the no information treatments. Again, the differences in prediction error in means are small relative to the variability of the point estimates. Further, the non-sequential search model exhibits smaller prediction errors in means and medians than the sequential search model. However, the latter model better captures the distribution over outcomes according to the Hellinger distance and Kullback-Leibler divergence. Overall, both models perform similarly well in predicting the distribution over the restricted number of searches.

5.6 How well do search models predict the outcomes from search?

There does not exist an objective benchmark to which we can easily compare the predictive performance of search models. In order to provide some perspective on the prediction quality

of our search cost estimates, we perform the following three steps. First, we consider the *relative prediction error in means (medians)* in- and out-of sample for all search models and outcome variables. This value is defined as the prediction error divided by the realized value.¹² It allows us to compare the prediction performance across models and domains. Second, we conduct a simulation study to examine the relative prediction error of the sequential and non-sequential search model when the true data-generating process is known. Third, we compare the relative prediction error of our models to relative prediction errors taken from the literature on predictions of sales and macroeconomic variables (inflation).

Table 6: Relative Prediction Error

	Number of Searches		Rest. Number of Searches		Purchase Prices	
	Mean	Median	Mean	Median	Mean	Median
In-Sample Predictions						
<i>sequential search</i>	0.209	0.310	0.101	0.208	0.018	0.023
<i>non-sequential search</i>	0.165	0.120	0.048	0.083	0.032	0.022
<i>search without priors</i>	0.189	0.140	0.036	0.050	0.017	0.013
Out-of-Sample Predictions						
<i>sequential search</i>	0.389	0.353	0.120	0.234	0.032	0.036
<i>non-sequential search</i>	0.371	0.350	0.098	0.202	0.034	0.030
<i>search without priors</i>	0.416	0.419	0.126	0.226	0.039	0.034

Notes: Differences between predicted and observed values (mean and median), relative to the the realized value.

Table 6 shows the relative prediction error in means and medians for all models and outcome variables. The predictions of the number of searches tend to be quite noisy with relative prediction errors up to around 40 percent. The predictions of the restricted number of searches are substantially less noisy, and the predictions of purchase prices are quite accurate, with relative prediction errors of less than 4 percent. There is no search model that strictly dominates the others (or that is strictly dominated by another model). The non-sequential search model performs fairly well, despite the fact that the search setting favors the sequential search model.

To put these findings into context, we conduct a simulation study in which the true data-generating process is known, see Online Appendix B.5 for all details. We proceed as follows:

¹²Here is an example: The average number of searches in treatment AMT S 1.0 is 10.49, the prediction error in means of the sequential search model for this treatment is 1.48. Hence, the relative prediction error in means of the sequential search model for the number of searches in treatment AMT S 1.0 is $\frac{1.48}{10.49} \approx 0.141$.

First, we draw the search cost parameter k for 10,000 simulated subjects from a truncated lognormal distribution that implies an average search cost parameter of $\mathbb{E}(k) = 0.216$ (close to our experimental estimates). Throughout, we assume $\kappa = 0.00$ and $\rho = 1.00$. Second, we specify their search method: The fraction α of subjects searches non-sequentially and the fraction $1 - \alpha$ searches sequentially. We vary α from zero to 100 percent in steps of 20 percentage points (so that we have six specifications that span the full range from purely sequential to purely non-sequential search). Third, we simulate (for each specification) the search outcomes of the 10,000 subjects when prices are uniformly distributed on the interval $[3.00, 6.00]$. Fourth, we apply (for each specification) the estimation and evaluation procedure from Section 4 to the simulated data, assuming either that all subjects search sequentially or that all subjects search non-sequentially. We thus can assess the predictive performance of the two search models under scenarios in which the true data-generating process either coincides with or (partly) differs from these models.

Table 7: Simulated Data – Relative Prediction Error

	Number of Searches		Rest. Number of Searches		Purchase Prices	
	Mean	Median	Mean	Median	Mean	Median
True model: Sequential						
<i>sequential search</i>	0.048	0.000	0.008	0.000	0.000	0.001
<i>non-sequential search</i>	0.021	0.000	0.025	0.000	0.086	0.101
True model: Non-sequential						
<i>sequential search</i>	0.030	0.333	0.130	0.333	0.038	0.052
<i>non-sequential search</i>	0.008	0.000	0.001	0.000	0.002	0.001

Notes: Differences between predicted and observed values (mean and median), relative to the the realized value from the respective simulation specification.

Table 7 shows the relative prediction errors for the two extreme specifications, $\alpha = 0.00$ and $\alpha = 1.00$.¹³ The first case is represented in the top panel: All subjects search sequentially so that the sequential search model is correct (and the non-sequential search model misspecified). The second case is shown in the bottom panel: All subjects search non-sequentially so that the roles of the two search models are reversed. When correctly specified, the relative prediction errors of the two models are between 0.8 and 4.8 percent for the number of searches, and be-

¹³The full breakdown of results for intermediate specifications is in Online Appendix B.5.

tween 0.0 and 0.2 percent for purchase prices. Note that positive values for relative prediction errors are expected since the estimation targets the full distribution of outcomes subject to a parametric search cost distribution (and not the mean directly). Hence, a small residual gap between predicted and realized means is inherent to the approach. When models are misspecified, prediction errors remain modest: These are between 2.1 and 3.0 percent for the number of searches, and between 3.8 and 8.6 percent for purchase prices.

We compare the relative prediction errors for in-sample predictions from the experimental data (Table 6) to the relative prediction errors from the simulations (Table 7), focusing on means. For purchase prices, the relative prediction errors from the experimental data (1.8 percent for sequential and 3.2 percent for non-sequential search) fall squarely within the range of the relative prediction errors from the simulation when the search model is misspecified. For the number of searches, however, the relative prediction errors from the experimental data (20.9 percent for sequential and 18.9 from non-sequential search) exceed those from the simulations considerably. This indicates that the prediction errors from the experimental data are not primarily driven by protocol misspecification, but by additional sources (such as behavioral heterogeneity and deviations from fully optimal search) that the simulation, by design, does not capture.

The simulations help rationalize why the models predict purchase prices more accurately than the number of searches, a pattern visible from Table 6 and throughout the analysis. Three factors contribute to this pattern. First, the true mean number of searches varies by as much as 41.7 percent across specifications, whereas the true mean purchase price varies by only about 3.8 percent. There is simply more room for a misspecified model to generate errors in the former than in the latter dimension. Second, the bounded support of the price distribution mechanically limits the scope for large errors on purchase prices. Third, the assumed protocol governs the stopping rule and hence the distribution over the number of searches directly. In contrast, purchase prices are relatively insensitive to the assumed search protocol.

Finally, we compare the relative prediction errors from the experimental data to those obtained in the forecasting literature. Table 8 displays the relative prediction errors for several examples. The forecasting literature is fairly diverse and reports different measures of prediction accuracy. We selected examples for which we can compute the relative prediction error based on the data provided in the paper. The prediction of retail sales based on neural networks (Alon et al., 2001) is fairly precise with a relative prediction error of 1.5 to 2.8 percent, while the prediction of inflation year-on-year is more noisy with a relative prediction error of 25.9 percent. Compared to these values, the predictions of the search models hold up well, especially in the domain of purchase prices (where our relative prediction errors are between 1.7 and 3.9 percent).

Table 8: Examples for Sales and Macro Variables – Relative Prediction Error

Variable	Method	Study/Data	Relative Prediction Error
Market shares	expert prediction market	Matzler et al. (2013)	0.075 to 0.575
Individual product sales	statistical and AI forecasting models	Aras et al. (2017)	0.139 to 0.513
Aggregate retail sales	neural network forecasting models	Alon et al. (2001)	0.015 to 0.028
Inventory	stat. forecasting with expert judgment	Trapero et al. (2013)	0.059
US Inflation, one year ahead	macroeconomic model	Federal Reserve economic data ¹⁴	0.259

6 Conclusion

In this paper, we examined the extent to which the classic search models – sequential and non-sequential search – are good as-if models. How well do the search cost estimates from these models capture the outcomes from search, especially in settings where the underlying assumptions on consumers’ information about prices or their search method are violated? To address this question, we conducted an online search experiment in which we manipulated the price distribution as well as subjects’ knowledge about the price distribution. For each treatment, we estimated search costs, fitted the models to the estimated search cost distribution to obtain predictions about the amount of search and purchase prices, and then compared the predicted and realized distributions over search outcomes.

Our main results are two-fold. First, we found that the predictive performance of the sequential search model is slightly higher in settings in which consumers know the price distribution than in settings where this is not the case. For the non-sequential search models, it is the other way around. In both cases, however, the differences in predictive performance between settings with and without price information are modest. Second, we found that the predictive performance of the non-sequential search model is close to that of the sequential search model, despite the fact that our search environment strongly favors sequential search. Overall, the predictive performance of the classic search models (as summarized in Table 6)

¹⁴Own calculations. Time horizon: 01/1983 to 08/2024. Data for the one-year-ahead expected inflation series was obtained from <https://fred.stlouisfed.org/series/EXPINF1YR>. Data for the realized year-over-year inflation rates was obtained from <https://fred.stlouisfed.org/series/MEDCPIM159SFRBCLE>.

indicates that these models are decent as-if models even when their assumptions on consumer information and the search protocol are violated.

Despite the stylized nature of the search environment in our experiment, our results have the following managerial implication. The consumers' search method affects their consideration sets and hence choices. The present research shows that the non-sequential search model is a fairly good as-if description of aggregate search outcomes (even though it might not fully capture individual search behavior). One implication of non-sequential search is that the consumers' decision to continue the search spell is, on average, not affected by their current price observation. This makes it easier for firms to model their consumers' consideration sets. For this, they only need to gather data on what options consumers see before they choose (e.g. through online surveys, clickstream, or consumer panel data), as the decision to see another option is independent of the features of the discovered option. Specifically, this also holds if consumers can search options sequentially. This observation is related to current research that supports the consideration set approach for search markets.

In this project, we focused on price search: Consumers know the product they wish to purchase and only search for the lowest price. The trade-off they face is the trade-off between spending more time on search and paying a higher price. However, in many settings, consumers search not only for lower prices, but also for higher match values from differentiated products. Further, they may have initial information about the different options that they could search. The relevant search paradigm is then directed sequential search as considered in [Weitzman \(1979\)](#). The open question is to what extent our conclusions carry over to such a setting. Our conjecture is that the search cost estimates from such a model yield reasonable predictions about search outcomes as long as the model captures the heterogeneity in outcomes through varying degrees of search costs. To verify this conjecture, researchers can update our experimental design so that subjects have to choose between different options that can be searched (unlike in our setting where, at each stage, they can search only one option). We believe that the results outlined in this paper are useful for examining more complex search settings.

References

- Abaluck, Jason and Abi Adams-Prassl (2021) “What do consumers consider before they choose? Identification from asymmetric demand responses,” *Quarterly Journal of Economics*, 136 (3), 1611–1663. 5
- Alon, Ilan, Min Qi, and Robert Sadowski (2001) “Forecasting aggregate retail sales: A comparison of artificial neural networks and traditional methods,” *Journal of Retailing and Consumer Services*, 8 (3), 147–156. 33, 34
- Amano, Tomomichi, Andrew Rhodes, and Stephan Seiler (2022) “Flexible demand estimation with search data,” CEPR Discussion Paper No. 16933, CEPR Press, Paris & London. 5
- Aras, Serkan, İpek Deveci Kocakoç, and Cigdem Polat (2017) “Comparative study on retail sales forecasting between single and combination methods,” *Journal of Business Economics and Management*, 18 (5), 803–832. 34
- Atayev, Atabek (2022) “Uncertain product availability in search markets,” *Journal of Economic Theory*, 204, 105524. 1
- Bajari, Patrick and Ali Hortaçsu (2005) “Are structural estimates of auction models reasonable? Evidence from experimental data,” *Journal of Political Economy*, 113 (4), 703–741. 5
- Bikhchandani, Sushil and Sunil Sharma (1996) “Optimal search with learning,” *Journal of Economic Dynamics and Control*, 20 (1–3), 333–359. 5
- Birchall, Cameron, Debashrita Mohapatra, and Frank Verboven (2024) “Estimating substitution patterns and demand curvature in discrete-choice models of product differentiation,” *Review of Economics and Statistics*, forthcoming. 25, 2
- Box, George and David Cox (1964) “An analysis of transformations,” *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 26 (2), 211–243. 25, 2
- Bronnenberg, Bart, Jun Kim, and Carl Mela (2016) “Zooming in on choice: How do consumers search for cameras online?” *Marketing Science*, 35 (5), 693–712. 6
- Brown, Meta, Christophe Flinn, and Andrew Schotter (2011) “Real-Time search in the laboratory and the market,” *American Economic Review*, 101 (2), 948–974. 5, 6, 8
- Burdett, Kenneth and Kenneth Judd (1983) “Equilibrium price dispersion,” *Econometrica*, 51 (4), 955–969. 1

- Casner, Ben (2021) “Learning while shopping: An experimental investigation into the effect of learning on consumer search,” *Experimental Economics*, 24, 232–260. 6
- Chen, Daniel, Martin Schonger, and Chris Wickens (2016) “oTree – An open-source platform for laboratory, online, and field experiments,” *Journal of Behavioral and Experimental Finance*, 9, 88–97. 13
- Chen, Xiaohong, Han Hong, and Matthew Shum (2007) “Nonparametric likelihood ratio model selection tests between parametric likelihood and moment condition models,” *Journal of Econometrics*, 141 (1), 109–140. 6, 24
- Chou, Chien-fu and Gabriel Talmain (1993) “Nonparametric search,” *Journal of Economic Dynamics and Control*, 17 (5–6), 771–784. 5, 29
- De Los Santos, Babur, Ali Hortaçsu, and Matthijs Wildenbeest (2012) “Testing models of consumer search using data on web browsing and purchasing behavior,” *American Economic Review*, 102 (6), 2955–2980. 1, 6, 8
- DellaVigna, Stefano and Devin Pope (2018) “What motivates effort? Evidence and expert forecast,” *Review of Economic Studies*, 85 (2), 1029–1069. 14
- Dohmen, Thomas, Armin Falk, David Huffman, Uwe Sunde, Jürgen Schupp, and Gert Wagner (2011) “Individual risk attitudes: Measurement, determinants, and behavioral consequences,” *Journal of the European Economic Association*, 9 (2), 522–550. 13
- Fischer, Kai, Simon Martin, and Philipp Schmidt-Dengler (2024) “Indirect taxation in consumer search markets: The case of retail fuel,” CEPR Discussion Paper No. 19095, CEPR Press, Paris & London. 1, 24
- Fong, Jessica, Olivia Natan, and Ranmit Pantle (2025) “Consumer inferences from Product rankings: The role of beliefs in search behavior,” Working Paper, University of Michigan. 5
- Ghose, Anindya, Panagiotis Ipeirotis, and Beibei Li (2019) “Modeling consumer footprints on search engines: An interplay with social media,” *Management Science*, 65 (3), 1363–1385. 1
- Honka, Elisabeth (2014) “Quantifying search and switching costs in the US auto insurance industry,” *RAND Journal of Economics*, 45 (4), 847–884. 1
- Honka, Elisabeth and Pradeep Chintagunta (2017) “Simultaneous or sequential? Search strategies in the U.S. auto insurance industry,” *Marketing Science*, 36 (1), 21–42. 6, 8

- Honka, Elisabeth, Ali Hortaçsu, and Matthijs Wildenbeest (2019) “Empirical search and consideration sets,” in *Handbook of the Economics of Marketing*, 1, 193–257: Elsevier. 5
- Hu, Mantian, Chu Dang, and Pradeep Chintagunta (2019) “Search and learning at a daily deals website,” *Marketing Science*, 38 (4), 609–642. 5
- Janssen, Maarten and José Luis Moraga-González (2004) “Strategic pricing, consumer search and the number of firms,” *Review of Economic Studies*, 71 (4), 1089–1118. 1
- Janssen, Maarten, José Luis Moraga-González, and Matthijs Wildenbeest (2005) “Truly costly sequential search and oligopolistic pricing,” *International Journal of Industrial Organization*, 23, 451–466. 1
- Jindal, Pranav and Anocha Aribarg (2021) “The importance of price beliefs in consumer search,” *Journal of Marketing Research*, 58 (2), 321–342. 5
- Karle, Heiko, Florian Kerzenmacher, Heiner Schumacher, and Frank Verboven (2025) “Search costs and context effects,” *American Economic Journal: Microeconomics*, 17 (2), 127–161. 5, 7, 8, 14, 19, 23
- Kim, Jun, Paulo Albuquerque, and Bart Bronnenberg (2010) “Online demand under limited consumer search,” *Marketing Science*, 29 (6), 1001–1023. 1
- (2017) “The probit choice model under sequential search with an application to online retailing,” *Management Science*, 63 (11), 3911–3929. 1
- König-Kersting, Christian, Yana Litovsky, Robert Böhm, Igor Grossmann, Jürgen Huber, and Michael Kirchler (2025) “A many-designs study crowdsourcing 516 aggregation algorithms to increase the Wisdom of the Crowd,” Working Paper, University of Innsbruck. 17
- Koulayev, Sergei (2013) “Search with Dirichlet priors: Estimation and implications for consumer demand,” *Journal of Business and Economic Statistics*, 31 (2), 226–239. 5
- Le Barbanchon, Thomas, Johannes Schmieder, and Andrea Weber (2024) “Job search, unemployment insurance, and active labor market policies,” NBER Working Paper No. 32720, NBER. 1
- Lin, Haizhen and Matthijs Wildenbeest (2020) “Nonparametric estimation of search costs for differentiated products: Evidence from Medigap,” *Journal of Business and Economic Statistics*, 38 (4), 754–770. 1

- Matzler, Kurt, Christopher Grabher, Jürgen Huber, and Johann Füller (2013) “Predicting new product success with prediction markets in online communities,” *R&D Management*, 43 (5), 420–432. 34
- Mauring, Eeva and Cole Williams (2023) “Coasian Dynamics in Sequential Search,” CEPR Discussion Paper No. 17907, CEPR Press, Paris & London. 1
- McCall, J. J. (1970) “Economics of information and job search,” *Quarterly Journal of Economics*, 84 (1), 113–126. 1
- Moraga-González, José Luis, Zsolt Sándor, and Matthijs Wildenbeest (2023) “Consumer search and prices in the automobile market,” *Review of Economic Studies*, 90 (3), 1394–1440. 1
- Moraga-González, José Luis and Matthijs Wildenbeest (2008) “Maximum likelihood estimation of search costs,” *European Economic Review*, 52 (5), 820–848. 1
- Morgan, Peter and Richard Manning (1985) “Optimal search,” *Econometrica*, 53 (4), 923–944. 2
- Parakhonyak, Alexei (2014) “Oligopolistic competition and search without priors,” *Economic Journal*, 124 (576), 594–606. 3, 5, 28, 29
- Rosenfield, Donald and Roy Shapiro (1981) “Optimal adaptive price search,” *Journal of Economic Theory*, 25 (1), 1–20. 5
- Rothschild, Michael (1974) “Searching for the lowest price when the distribution of prices is unknown,” *Journal of Political Economy*, 82 (4), 689–711. 5
- Salz, Tobias and Emanuel Vespa (2020) “Estimating dynamic games of oligopolistic competition: an experimental investigation,” *RAND Journal of Economics*, 51 (2), 447–469. 5
- De los Santos, Babur, Ali Hortaçsu, and Matthijs Wildenbeest (2017) “Search with learning for differentiated products: Evidence from e-commerce,” *Journal of Business and Economic Statistics*, 35 (4), 626–641. 5
- Schlag, Karl and Andriy Zapechelnjuk (2021) “Robust sequential search,” *Theoretical Economics*, 16 (4), 1431–1470. 5
- Somerville, Jason (2022) “Range-dependent attribute weighting in consumer choice: An experimental test,” *Econometrica*, 90 (2), 799–830. 14

- Stahl, Dale (1989) “Oligopolistic pricing with sequential consumer search,” *American Economic Review*, 79 (4), 700–712. [1](#)
- Stigler, George (1961) “The economics of information,” *Journal of Political Economy*, 69 (3), 213–225. [1](#)
- Trapero, Juan, Diego Pedregal, Robert Fildes, and Nikolaos Kourentzes (2013) “Analysis of judgmental adjustments in the presence of promotions,” *International Journal of Forecasting*, 29 (2), 234–243. [34](#)
- Ursu, Raluca, Tülin Erdem, Qingliang Wang, and Qianyun Zhang (2024) “Prior information and consumer search: Evidence from eye-tracking,” *Management Science*, 23 (1), 165–213. [5](#)
- Ursu, Raluca, Stephan Seiler, and Elisabeth Honka (2025) “The sequential search model: A framework for empirical research,” *Quantitative Marketing and Economics*, 23 (1), 165–213. [1](#)
- Ursu, Raluca, Qingliang Wang, and Pradeep Chintagunta (2020) “Search duration,” *Marketing Science*, 39 (5), 849–871. [5](#)
- Ursu, Raluca, Qianyun Zhang, and Elisabeth Honka (2023) “Search gaps and consumer fatigue,” *Marketing Science*, 42 (1), 110–136. [14](#)
- Weitzman, Martin (1979) “Optimal search for the best alternative,” *Econometrica*, 47 (3), 641–654. [1](#), [35](#)
- Wildenbeest, Matthijs (2011) “An empirical model of search with vertically differentiated products,” *RAND Journal of Economics*, 42 (4), 729–757. [24](#)

A Appendix

A.1 Optimal Sequential Search with Convex Search Costs

We prove that equation (2) characterizes optimal sequential search when the cost function c is weakly convex. The proof proceeds by induction. Let $p_N^{\min}$ be the smallest price discovered after $N - 1$ searches and let p_N be the price discovered at the N 'th search. Suppose the DM has searched $N - 1$ shops so far. Let r_N be defined by equation (2). By definition, it is then optimal for the DM to search the N 'th shop if $p_N^{\min} \geq r_N$, and to stop search otherwise. Next, suppose the DM has searched $N - 2$ shops so far and let r_{N-1} be given by equation (2). Since $v(p, F)$ is strictly increasing in p , we must have $r_{N-1} \leq r_N$. Obviously, it is optimal to search shop $N - 1$ if $p_{N-1}^{\min} \geq r_{N-1}$. Thus, assume that $p_{N-1}^{\min} < r_{N-1}$. We show that it is then optimal for the DM not to search shop $N - 1$. Note that $r_{N-1} \leq r_N$ implies $p_{N-1}^{\min} < r_N$. If the DM searches shop $N - 1$ and discovers a price $p_{N-1} > r_N$, then, by the DM's optimal behavior at the N 'th shop, it is optimal for her to stop search because $p_N^{\min} = p_{N-1}^{\min} < r_N$. If the DM searches shop $N - 1$ and discovers a price $p_{N-1} \leq r_N$, we have $\min\{p_{N-1}, p_{N-1}^{\min}\} < r_N$. Again, by the DM's optimal behavior at the N 'th shop, it is optimal for her to stop search. Thus, in any case, searching shop $N - 1$ would be the last search if $p_{N-1}^{\min} < r_{N-1}$. By the definition of r_{N-1} , it therefore cannot be optimal to search shop $N - 1$ if $p_{N-1}^{\min} < r_{N-1}$, which completes the proof.

A.2 Derivation of Search Costs from Search Behavior

Search cost parameter k for sequential search and skewed distributions. We generalize equation (9) to left- and right-skewed distributions as we have them in the experiment. Suppose the distribution F is piece-wise constant and has support $[a, b]$. Let there be three intervals, $l = 1, 2, 3$, that partition $[a, b]$. The first interval is given by $[a, d_1) = [d_0, d_1)$, the second interval by $[d_1, d_2)$, and the third interval by $[d_2, b] = [d_2, d_3]$. The density in interval l is given by $f^{[l]}$. If the reservation price is located in the first interval, $l^* = 1$, the search cost parameter k equals

$$k^{\text{seq}}(r_n, \rho, \kappa, n) = \frac{1 + \kappa}{n^{1+\kappa} - (n-1)^{1+\kappa}} \left(\frac{f^{[1]}(r_n - a)^2}{2\Delta_F^\rho} + w \right), \quad (24)$$

and if the reservation price is located in interval $l^* \in \{2, 3\}$, parameter k is given by

$$k^{\text{seq}}(r_n, \rho, \kappa, n) = \frac{1 + \kappa}{n^{1+\kappa} - (n-1)^{1+\kappa}} \times \left(\sum_{l=1}^{l^*-1} \frac{f^{[l]}(d_l - d_{l-1})(r_n - \frac{1}{2}(d_l + d_{l-1}))}{\Delta_F^\rho} + \frac{f^{[l^*]}(r_n - d_{l^*-1})^2}{2\Delta_F^\rho} + w \right). \quad (25)$$

Likelihood contributions for sequential search and censored observations. We define the likelihood contributions for censored observations in the sequential search model. If subject i searches $n_i = 0$ times, her likelihood contribution is given by

$$P_i = 1 - \Phi\left(\frac{\ln k^{\text{seq}}(b, \rho, \kappa, 1) - x_i\beta}{\sigma}\right). \quad (26)$$

Here we set $n_i = 1$ so that this expression is well-defined. If she searches $n_i = 1$ times, her likelihood contribution equals

$$P_i = \Phi\left(\frac{\ln k^{\text{seq}}(b, \rho, \kappa, 1) - x_i\beta}{\sigma}\right) - \Phi\left(\frac{\ln k^{\text{seq}}(p_i^1, \rho, \kappa, 1) - x_i\beta}{\sigma}\right), \quad (27)$$

and if she searches all shops, $n_i = N$, her likelihood contribution is given by

$$P_i = \Phi\left(\frac{\ln k^{\text{seq}}(b, \rho, \kappa, N) - x_i\beta}{\sigma}\right). \quad (28)$$

Search cost parameter k for non-sequential search and skewed distributions. We generalize equation (12) to left- and right-skewed distributions as we have them in the experiment. Note that with our parametrization of relative thinking we can rewrite equation (4) as

$$\mathbb{E}^{[n]}(v) = \frac{1}{\Delta_F^\rho} \left(a + \int_a^b (1 - F(p))^n dp \right). \quad (29)$$

Let F be the piece-wise constant distribution as defined above. We then have $F(p) = (p-a)f^{[1]}$ for $p \in [a, d_1]$; $F(p) = (d_1 - a)f^{[1]} + (p - d_1)f^{[2]}$ for $p \in [d_1, d_2]$; and $F(p) = (d_1 - a)f^{[1]} + (d_2 - d_1)f^{[2]} + (p - d_2)f^{[3]}$ for $p \in [d_2, b]$. With this, we can write equation (29) as

$$\begin{aligned} \mathbb{E}^{[n]}(v) = & \frac{1}{\Delta_F^\rho} \left(a + \int_a^{d_1} (1 + af^{[1]} - pf^{[1]})^n dp + \int_{d_1}^{d_2} (1 + af^{[1]} - d_1f^{[1]} + d_1f^{[2]} - pf^{[2]})^n dp \right. \\ & \left. + \int_{d_2}^b (1 + af^{[1]} - d_1f^{[1]} + d_1f^{[2]} - d_2f^{[2]} + d_2f^{[3]} - pf^{[3]})^n dp \right), \end{aligned} \quad (30)$$

which can be simplified to

$$\begin{aligned} \mathbb{E}^{[n]}(v) = & \frac{1}{\Delta_F^\rho} \left(a + \frac{1}{f^{[1]} n + 1} \left[1 - (1 - (d_1 - a)f^{[1]})^{n+1} \right] \right. \\ & + \frac{1}{f^{[2]} n + 1} \left[(1 - (d_1 - a)f^{[1]})^{n+1} - (1 - (d_1 - a)f^{[1]} - (d_2 - d_1)f^{[2]})^{n+1} \right] \\ & \left. + \frac{1}{f^{[3]} n + 1} (1 - (d_1 - a)f^{[1]} - (d_2 - d_1)f^{[2]})^{n+1} \right). \end{aligned} \quad (31)$$

Using this expression in the equations (13) and (14) yields the search cost parameter k for non-sequential search and skewed distributions.

Likelihood contributions for non-sequential search and censored observations. We define the likelihood contributions for censored observations in the non-sequential search model. If subject i searches $n_i = 0$ times, her likelihood contribution is given by

$$P_i = 1 - \Phi\left(\frac{\ln k_-^{\text{nseq}}(\rho, \kappa, 0) - x_i\beta}{\sigma}\right), \quad (32)$$

and if she searches all shops, $n_i = N$, her likelihood contribution is assumed to be

$$P_i = \Phi\left(\frac{\ln k_+^{\text{nseq}}(\rho, \kappa, N) - x_i\beta}{\sigma}\right). \quad (33)$$

Here we cannot distinguish between subjects who wish to search exactly N shops and subjects who would search more than N shops. This distinction is, however, immaterial for our results.

A.3 Computational Details for Truncated Lognormal Distribution

Given that search costs are lognormally distributed and truncated at t , the distribution of the search cost parameter k_i of subject i is given by

$$G_i(k_i) = \Phi\left(\frac{\ln k_i - x_i\beta}{\sigma}\right) \Phi\left(\frac{\ln t - x_i\beta}{\sigma}\right)^{-1} \quad (34)$$

and the associated density function equals

$$g_i(k_i) = \frac{1}{k_i\sigma} \phi\left(\frac{\ln k_i - x_i\beta}{\sigma}\right) \Phi\left(\frac{\ln t - x_i\beta}{\sigma}\right)^{-1}, \quad (35)$$

where Φ and ϕ denote the CDF and the PDF of the standard normal distribution, respectively. The expected value of k_i is obtained through

$$\mathbb{E}(k_i) = \frac{1}{\sigma} \int_0^t \phi\left(\frac{\ln k_i - x_i\beta}{\sigma}\right) dk_i \times \Phi\left(\frac{\ln t - x_i\beta}{\sigma}\right)^{-1} \quad (36)$$

and the median through

$$\text{med}(k_i) = \exp\left(x_i\beta + \sqrt{2}\sigma \text{erf}^{-1}\left(\Phi\left(\frac{\ln t - x_i\beta}{\sigma}\right) - 1\right)\right) = \exp\left(x_i\beta + \sigma\Phi^{-1}\left(\frac{1}{2}\Phi\left(\frac{\ln t - x_i\beta}{\sigma}\right)\right)\right) \quad (37)$$

where erf^{-1} denotes the inverse of the Gauss error function.

A.4 Computations for the Evaluation of Search Cost Estimates

We derive for each search model the distribution over the number of searches and purchase prices for given search cost parameter k . Throughout we assume that $\kappa = 0$ and $w = 0$. For convenience, we write r for the reservation price (instead of r_n). Further, for any skewed (piece-wise constant) price distribution we define $F^{[1]}(r) = (r - a)f^{[1]}$ for $r \in [a, d_1]$; $F^{[2]}(r) = (d_1 - a)f^{[1]} + (r - d_1)f^{[2]}$ for $r \in [d_1, d_2]$; and $F^{[3]}(r) = (d_1 - a)f^{[1]} + (d_2 - d_1)f^{[2]} + (r - d_2)f^{[3]}$ for $r \in [d_2, b]$.

A.4.1 Computations for the Sequential Search Model

Derivation of the distribution over the number of searches for given search cost parameter k , uniform price distribution. We derive the distribution over the number of searches from equality (9). We obtain an upper bound $\bar{k}$ for the search cost parameter when we set the reservation price equal to the highest possible price, $r = b$:

$$\bar{k} = \frac{1}{\Delta_F^\rho} \frac{b - a}{2}. \quad (38)$$

It is optimal for the DM to not conduct any search if $k \geq \bar{k}$ and to conduct at least one search if $k < \bar{k}$. From equality (9) we get that, if $k < \bar{k}$, then the reservation price is given by

$$r = a + \sqrt{2k\Delta_F^\rho(b - a)}. \quad (39)$$

At a given reservation price $r < b$, the probability that the number of searches equals $n < 100$ is given by $\Pr(n) = F(r)(1 - F(r))^{n-1}$ and the probability that the number of searches is $n = 100$ is given by $\Pr(n = 100) = (1 - F(r))^{99}$. Using equation (39) and the fact that F is the uniform distribution, we obtain

$$\Pr(n; k) = \sqrt{\frac{2k\Delta_F^\rho}{b - a}} \left(1 - \sqrt{\frac{2k\Delta_F^\rho}{b - a}} \right)^{n-1} \quad (40)$$

for $n < 100$ and

$$\Pr(n = 100; k) = \left(1 - \sqrt{\frac{2k\Delta_F^\rho}{b - a}} \right)^{99}. \quad (41)$$

This constitutes the distribution over the number of searches for given cost parameter k .

Derivation of the distribution over the number of searches for given search cost parameter k , skewed price distribution. We derive the distribution over the number of searches for a skewed price distribution from equations (24) and (25). To this end, we derive three boundary values

from these equations:

$$\bar{k}_1 = \frac{1}{\Delta_F^\rho} \frac{f^{[1]}(d_1 - d_0)^2}{2}, \quad (42)$$

$$\bar{k}_2 = \frac{1}{\Delta_F^\rho} \left[f^{[1]}(d_1 - d_0) \left(d_2 - \frac{d_1 + d_0}{2} \right) + \frac{f^{[2]}(d_2 - d_1)^2}{2} \right], \quad (43)$$

$$\bar{k}_3 = \frac{1}{\Delta_F^\rho} \left[f^{[1]}(d_1 - d_0) \left(d_3 - \frac{d_1 + d_0}{2} \right) + f^{[2]}(d_2 - d_1) \left(d_3 - \frac{d_2 + d_1}{2} \right) + \frac{f^{[3]}(d_3 - d_2)^2}{2} \right]. \quad (44)$$

If $k < \bar{k}_1$, we have $r \in [a, d_1)$ and the DM's reservation price is defined by equality (24); if $\bar{k}_1 \leq k < \bar{k}_2$, we have $r \in [d_1, d_2)$ and the reservation price is defined by equality (25) for $l^* = 2$; if $\bar{k}_2 \leq k < \bar{k}_3$, we have $r \in [d_2, b)$ and the reservation price is defined by equality (25) for $l^* = 3$; and if $k \geq \bar{k}_3$, it is optimal for the DM not to conduct any search. We derive the reservation price for the three former cases. If $k < \bar{k}_1$, the reservation price equals

$$r = a + \sqrt{\frac{2k\Delta_F^\rho}{f^{[1]}}}. \quad (45)$$

If $\bar{k}_1 \leq k < \bar{k}_2$, it is given by

$$r = \frac{d_1 f^{[2]} - (d_1 - d_0) f^{[1]} + \sqrt{(d_1 - d_0)^2 f^{[1]}(f^{[1]} - f^{[2]}) + 2f^{[2]}k\Delta_F^\rho}}{f^{[2]}}. \quad (46)$$

If $\bar{k}_2 \leq k < \bar{k}_3$, the reservation price is given by

$$r = \frac{d_2 f^{[3]} - (d_1 - d_0) f^{[1]} - (d_2 - d_1) f^{[2]} + \sqrt{((d_1 - d_0) f^{[1]} + (d_2 - d_1) f^{[2]} - d_2 f^{[3]})^2 + f^{[3]} A}}{f^{[3]}}. \quad (47)$$

with

$$A = (d_1 - d_0)(d_1 + d_0) f^{[1]} + (d_2 - d_1)(d_2 + d_1) f^{[2]} - d_2^2 f^{[3]} + 2k\Delta_F^\rho. \quad (48)$$

Given the reservation price r , we can derive the probability distribution over the number of searches: If $k < \bar{k}_1$, we have

$$\Pr(n; k) = F^{[1]}(r)(1 - F^{[1]}(r))^{n-1} \quad \text{and} \quad \Pr(n = 100; k) = (1 - F^{[1]}(r))^{99}. \quad (49)$$

If $\bar{k}_1 \leq k < \bar{k}_2$, we have

$$\Pr(n; k) = F^{[2]}(r)(1 - F^{[2]}(r))^{n-1} \quad \text{and} \quad \Pr(n = 100; k) = (1 - F^{[2]}(r))^{99}, \quad (50)$$

and if $\bar{k}_2 \leq k < \bar{k}_3$, we have

$$\Pr(n; k) = F^{[3]}(r)(1 - F^{[3]}(r))^{n-1} \quad \text{and} \quad \Pr(n = 100; k) = (1 - F^{[3]}(r))^{99}. \quad (51)$$

Derivation of the distribution over purchase prices for given search cost parameter k , uniform price distribution. We derive the distribution over purchase prices. The search cost parameter $k < \bar{k}$ implies the reservation price $r_k < b$, see equation (39). Given that the price distribution is uniform on the interval $[a, b]$, the density over purchase prices equals

$$f(p; k) = \frac{1}{r_k - a} \quad (52)$$

for $p \in [a, r]$ and zero otherwise.

Derivation of the distribution over purchase prices for given search cost parameter k , skewed price distribution. The cost parameter k implies the reservation price r_k as defined in the equations (45), (47), and (48). If $r_k \in [a, d_1)$, the density over purchase prices is given by

$$f(p; k) = \begin{cases} \frac{1}{r_k - d_0} & \text{if } p \in [a, r_k] \\ 0 & \text{else} \end{cases}. \quad (53)$$

If $r_k \in [d_1, d_2)$, it equals

$$f(p; k) = \begin{cases} \frac{f^{[1]}}{F^{[2]}(r_k)} & \text{if } p \in [a, d_1) \\ \frac{f^{[2]}}{F^{[2]}(r_k)} & \text{if } p \in [d_1, r_k] \\ 0 & \text{else} \end{cases} \quad (54)$$

and if $r_k \in [d_2, b]$, it is given by

$$f(p; k) = \begin{cases} \frac{f^{[1]}}{F^{[3]}(r_k)} & \text{if } p \in [a, d_1) \\ \frac{f^{[2]}}{F^{[3]}(r_k)} & \text{if } p \in [d_1, d_2) \\ \frac{f^{[3]}}{F^{[3]}(r_k)} & \text{if } p \in [d_2, r_k] \\ 0 & \text{else} \end{cases}. \quad (55)$$

A.4.2 Computations for the Non-Sequential Search Model

Derivation of the distribution over the number of searches for given search cost parameter k , uniform price distribution. We derive the distribution over the number of searches for the non-sequential search model when the price distribution is uniform on the interval $[a, b]$. From inequality (6) we obtain a cost bound k_n with the following interpretation: If $k > k_n$, the optimal number of searches for the DM is $n - 1$ or lower; if $k \leq k_n$, the optimal number of searches is n or higher. At a uniform price distribution, the expected value of expenses after n searches

$\mathbb{E}(v; n)$ is given by equation (12). With this, we obtain the cost bound

$$k_n = \frac{b-a}{\Delta_F^\rho} \frac{1}{n(n+1)}. \quad (56)$$

The set of cost bounds $k_1, k_2, \dots, k_{100}$ characterizes the distribution over the number of searches.

Derivation of the distribution over the number of searches for given search cost parameter k , skewed price distribution. We derive the distribution over the number of searches for the non-sequential search model when the price distribution is skewed. As for the case of a uniform price distribution, we derive cost bounds k_n . With a skewed price distribution, the expected value of expenses after n searches $\mathbb{E}(v; n)$ is given by equation (31). With this, we obtain the cost bound

$$\begin{aligned} k_n = & \frac{1}{\Delta_F^\rho} \left[\frac{1}{f^{[1]}} \frac{1}{n} [1 - (1 - (d_1 - a)f^{[1]})^n] - \frac{1}{f^{[1]}} \frac{1}{n+1} [1 - (1 - (d_1 - a)f^{[1]})^{n+1}] \right. \\ & + \frac{1}{f^{[2]}} \frac{1}{n} [(1 - (d_1 - a)f^{[1]})^n - (1 - (d_1 - a)f^{[1]} - (d_2 - d_1)f^{[2]})^n] \\ & - \frac{1}{f^{[2]}} \frac{1}{n+1} [(1 - (d_1 - a)f^{[1]})^{n+1} - (1 - (d_1 - a)f^{[1]} - (d_2 - d_1)f^{[2]})^{n+1}] \\ & + \frac{1}{f^{[3]}} \frac{1}{n} [1 - (d_1 - a)f^{[1]} - (d_2 - d_1)f^{[2]}]^n \\ & \left. - \frac{1}{f^{[3]}} \frac{1}{n+1} [1 - (d_1 - a)f^{[1]} - (d_2 - d_1)f^{[2]}]^{n+1} \right]. \quad (57) \end{aligned}$$

The set of cost bounds $k_1, k_2, \dots, k_{100}$ characterizes for each given cost parameter k the distribution over the number of searches.

Derivation of the distribution over purchase prices for given search cost parameter k , uniform price distribution. We derive the distribution over purchase prices for the non-sequential search model when the price distribution is uniform on the interval $[a, b]$. The search cost parameter k implies the number of searches n_k according to the bounds defined in (56). For a given price distribution F , the distribution over purchase prices after n_k searches is given by

$$F(p; n_k) = 1 - (1 - F(p))^{n_k}. \quad (58)$$

Hence, for the uniform distribution, the density over purchase prices when the search cost parameter equals k is given by

$$f(p; k) = \frac{n_k}{b-a} \left(\frac{b-p}{b-a} \right)^{n_k-1}. \quad (59)$$

for $p \in [a, b]$ and zero otherwise.

Derivation of the distribution over purchase prices for given search cost parameter k , skewed price distribution. The search cost parameter k implies the number of searches n_k according to the bounds defined in (57). The density over the purchase price when the search cost parameter equals k is given by

$$f(p; k) = \begin{cases} n_k f^{[1]}(1 - F^{[1]}(p))^{n_k-1} & \text{if } p \in [a, d_1) \\ n_k f^{[2]}(1 - F^{[2]}(p))^{n_k-1} & \text{if } p \in [d_1, d_2) \\ n_k f^{[3]}(1 - F^{[3]}(p))^{n_k-1} & \text{if } p \in [d_2, b] \\ 0 & \text{else} \end{cases} . \quad (60)$$

A.5 Instructions and Screenshots

This appendix shows the instructions to the baseline treatment *S* 1.0.

Instructions Shopping Task Part 1/5

The next part of the study is about buying a fictitious product at the lowest possible price. We call it "Product A".

Your budget for Product A is \$6.00. Your earnings from this part of the study increase in the price savings that you realize. That is, if you buy Product A at price P , then your earnings will be $\$6.00 - P$.

The lower the price at which you purchase Product A, the higher will be your earnings.

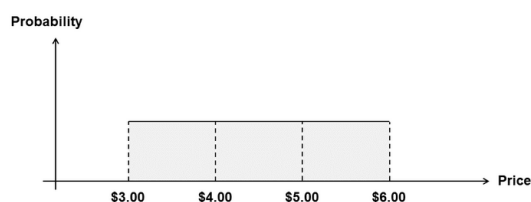
The earnings from this part of the study will be paid as a bonus in MTurk (additional to the \$1 participation fee).

You can search for the lowest price for product A in up to 100 online shops. On the next page we will explain how this works.

[Next](#)

Instructions Shopping Task Part 2/5

You will get access to up to 100 online shops that offer Product A. The price of each online shop ranges from \$3.00 to \$6.00. The following graph shows how likely it is that the price of a shop is in a certain range.



Therefore:

- Prices range from \$3.00 to \$6.00.
- All prices are equally probable.

You do not have to search immediately. You can do this anytime within the next 3 days. You can also continue your search after a break. You can use the link to this study to access the online shopping task anytime.

[Back](#)[Next](#)

Instructions Shopping Task Part 3/5

To find out the price of an online shop, a 16-digit code must be entered. This code will be given to you as soon as you click on the "Show Code" button. It needs to be entered into an input field, which appears as soon as you click on the "Show Input Field" button.

Note that the code cannot be entered by copy and paste. Thus, you have to record it somehow (for example by writing it down). After entering the code, the price will be displayed.

Example Shop:

To illustrate this procedure, please search the shop below. Once you successfully searched the shop click on the "Next" button to proceed.

Show Code

Back

Instructions Shopping Task Part 4/5

Once you learn the price of Product A at an online shop, it appears on a list of the prices that you have already seen.

You can then buy the product by clicking on the price of an online shop on this list **or** you can continue searching for the lowest price.

As soon as you purchase Product A, you will reach the final page with the completion code to copy in your HIT.

If you do not want to search for the lowest price at all, you can press the red "I do not want to search" button (your bonus then will be zero).

If you visit some online shops but do not buy Product A from any of them, we cannot pay you a bonus.

Back

Next

Instructions Shopping Task Part 5/5

Bonus Payment:

Your bonus from this part of the study is determined by the price savings you realize:

If you buy Product A at price P in one of the online shops, we pay you a bonus of $\$6.00 - P$.

Payment Scheme:

Your Bonus Payment = $\$6.00 - \text{Price at which you buy Product A}$

Comprehension Question:

To make sure you understand the payment structure, please answer the following question:

Suppose that after searching 2 times for the lowest price, a shopper buys Product A for \$4.20. What will be the shopper's bonus?

Open this select menu ▾

Submit

Back

These two figures show the screen of the search task in the baseline treatment after zero and three searches, respectively.

Shopping Task

Number of shops searched: 0

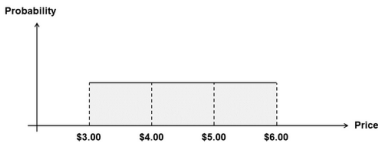
Remember: Prices at each shop range from \$3.00 to \$6.00.

Shop Number 1

Show Code

I do not want to search

Keep in mind how likely the different prices are:



Therefore:

- Prices range from \$3.00 to \$6.00.
- All prices are equally probable.

Prices seen so far

Click on a price to buy Product A for the specific price and end the experiment; the lowest price is always highlighted in green:

Shopping Task

Number of shops searched: 3

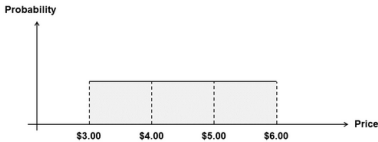
Remember: Prices at each shop range from \$3.00 to \$6.00.

Shop Number 3

You successfully searched the shop! The price at this shop is \$4.29.

Next Shop

Keep in mind how likely the different prices are:



Therefore:

- Prices range from \$3.00 to \$6.00.
- All prices are equally probable.

Prices seen so far

Click on a price to buy Product A for the specific price and end the experiment; the lowest price is always highlighted in green:

Buy for lowest price so far (\$3.92)

\$5.71

\$3.92

\$4.29

A.6 Balancing Tables

Table A1: Balancing Tables – Demographic Variables and Labor Supply

Sample, Treatment	<i>N</i>	Share Females	Age	Average hourly Earnings	Average hours per Week
AMT <i>S</i> 0.2 – <i>pr</i> 4	216	0.407	34.55 (10.04)	11.71 (12.74)	34.60 (18.08)
AMT <i>S</i> 0.2 – <i>pr</i> 8	202	0.396	34.99 (9.81)	13.44 (14.20)	33.38 (18.57)
AMT <i>S</i> 0.2 – <i>pr</i> 12	194	0.418	35.37 (10.64)	12.57 (13.08)	35.35 (18.03)
AMT <i>S</i> 5.0	190	0.416	34.57 (10.34)	11.94 (12.88)	35.26 (22.54)
AMT <i>S</i> 1.0	213	0.441	34.88 (10.09)	12.62 (14.10)	35.90 (18.29)
AMT <i>S</i> 1.0 – <i>noinfo</i>	210	0.419	35.66 (10.32)	11.41 (13.40)	34.49 (20.62)
ANOVA <i>p</i> -value		0.965	0.859	0.682	0.833
AMT <i>S</i> 1.0 – <i>l</i>	202	0.431	33.42 (8.75)	11.17 (12.70)	38.00 (22.75)
AMT <i>S</i> 1.0 – <i>l.noinfo</i>	192	0.385	34.31 (9.33)	11.41 (13.28)	38.01 (21.06)
AMT <i>S</i> 1.0 – <i>r</i>	199	0.342	34.89 (9.59)	11.29 (14.02)	39.88 (25.53)
AMT <i>S</i> 1.0 – <i>r.noinfo</i>	209	0.340	35.04 (9.39)	10.76 (13.39)	38.24 (21.48)
ANOVA <i>p</i> -value		0.186	0.279	0.965	0.814
PRO <i>S</i> 1.0 – <i>l</i>	115	0.591	39.84 (13.63)	8.57 (4.65)	11.00 (8.74)
PRO <i>S</i> 1.0 – <i>l.noinfo</i>	134	0.530	38.90 (12.12)	8.77 (5.97)	10.82 (8.52)
PRO <i>S</i> 1.0 – <i>r</i>	152	0.605	37.55 (12.24)	8.14 (3.90)	11.33 (8.27)
PRO <i>S</i> 1.0 – <i>r.noinfo</i>	161	0.590	39.43 (12.55)	7.62 (4.06)	10.98 (8.82)
ANOVA <i>p</i> -value		0.593	0.444	0.158	0.966
PRO <i>S</i> 1.0 – <i>sc</i> 4	188	0.479	39.12 (12.54)	11.44 (6.07)	16.56 (14.33)
PRO <i>S</i> 1.0 – <i>sc</i> 16	150	0.507	37.77 (10.56)	11.21 (5.83)	18.72 (15.42)
ANOVA <i>p</i> -value		0.611	0.293	0.727	0.184

Notes: Standard deviation in parentheses. Age in years, average hourly earnings in USD. Treatments are grouped together according to the points in time when we conducted the corresponding experiments (we ran the treatments *S* 0.2 – *pr*4, *S* 0.2 – *pr*8, *S* 0.2 – *pr*12, *S* 5.0, *S* 1.0, and *S* 1.0 – *noinfo* in November 2023, the treatments *S* 1.0 – *l*, *S* 1.0 – *r*, *S* 1.0 – *l.noinfo*, and *S* 1.0 – *r.noinfo* with both AMT workers and Prolific subjects in June 2024, and the treatments *S* 1.0 – *sc*4 and *S* 1.0 – *sc*16 in April 2026).

Table A2: Balancing Tables – Personal Characteristics

Sample, Treatment	<i>N</i>	Education Score	CRT Score	Willingness to take risks	Trust
AMT <i>S</i> 0.2 – <i>pr</i> 4	216	3.03 (0.58)	1.26 (0.94)	7.82 (1.89)	7.84 (1.98)
AMT <i>S</i> 0.2 – <i>pr</i> 8	202	3.14 (0.55)	1.46 (0.99)	7.65 (2.15)	7.75 (1.93)
AMT <i>S</i> 0.2 – <i>pr</i> 12	194	3.10 (0.59)	1.25 (0.95)	7.68 (2.10)	7.86 (1.93)
AMT <i>S</i> 5.0	190	3.12 (0.60)	1.36 (0.97)	7.54 (2.01)	7.62 (1.91)
AMT <i>S</i> 1.0	213	3.15 (0.60)	1.30 (0.99)	7.61 (2.12)	7.68 (2.01)
AMT <i>S</i> 1.0 – <i>noinfo</i>	210	3.13 (0.63)	1.50 (0.94)	7.71 (1.88)	7.65 (1.87)
ANOVA <i>p</i> -value		0.308	0.036	0.793	0.737
AMT <i>S</i> 1.0 – <i>l</i>	202	3.14 (0.58)	1.51 (1.04)	7.80 (1.91)	7.87 (1.83)
AMT <i>S</i> 1.0 – <i>l.noinfo</i>	192	3.08 (0.58)	1.40 (1.03)	7.86 (2.04)	7.68 (1.90)
AMT <i>S</i> 1.0 – <i>r</i>	199	3.22 (0.52)	1.45 (1.09)	7.79 (1.96)	7.80 (1.90)
AMT <i>S</i> 1.0 – <i>r.noinfo</i>	209	3.15 (0.51)	1.55 (1.11)	7.69 (2.30)	7.80 (1.91)
ANOVA <i>p</i> -value		0.102	0.498	0.879	0.809
PRO <i>S</i> 1.0 – <i>l</i>	115	2.82 (0.78)	1.56 (1.16)	4.66 (2.50)	4.47 (2.33)
PRO <i>S</i> 1.0 – <i>l.noinfo</i>	134	2.72 (0.74)	1.71 (1.16)	5.13 (2.44)	4.41 (2.30)
PRO <i>S</i> 1.0 – <i>r</i>	152	2.86 (0.75)	1.41 (1.11)	4.90 (2.54)	4.47 (2.30)
PRO <i>S</i> 1.0 – <i>r.noinfo</i>	161	2.84 (0.74)	1.50 (1.12)	4.97 (2.55)	4.66 (2.27)
ANOVA <i>p</i> -value		0.436	0.158	0.533	0.803
PRO <i>S</i> 1.0 – <i>sc</i> 4	188	2.72 (0.75)	1.49 (1.20)	5.49 (2.63)	4.72 (2.26)
PRO <i>S</i> 1.0 – <i>sc</i> 16	150	2.88 (0.81)	1.43 (1.18)	5.60 (2.58)	5.03 (2.37)
ANOVA <i>p</i> -value		0.067	0.639	0.712	0.231

Notes: Standard deviation in parentheses. Education is indicated on a scale from 0 to 4 (0 = No degree, 1 = Some high school, 2 = High school degree, 3 = Bachelor's degree, 4 = Master's degree or higher); CRT score is on a scale from 0 to 3 and shows the number of correct answers in the CRT test; willingness to take risk is on a scale from 0 (not willing to take risk at all) to 10 (very willing to take risk); trust is indicates on a scale from 0 (people can't be trusted at all) to 10 (people can be fully trusted).

A.7 Empirical Predictions: Detailed Results

We implement the three tests. Table A3 shows the results on recall (Test 1a and Test 1b) and Table A4 displays the results on the price-dependency of search (Test 2 and Test 3). For Test 1a we find that the share of subjects who purchase from the last sampled shop or search all shops (among those who search at least one shop) varies between 34.1 and 70.3 percent. Thus, there is a substantial share of subjects who recall previously searched shops in all treatments, which is inconsistent with sequential search when search cost are constant in the number of searches. In the Prolific samples, the share of subjects who recall previously sampled shops is significantly larger in the no-information treatment than in the information treatments. However, this is not the case in the AMT samples.

To implement Test 1b, we consider only those subjects who search at least two, but not all shops. For subjects who non-sequentially search n shops, the probability of trade with the last sampled shop is $\frac{1}{n}$. From this, we derive the predicted average probability of trading with the last sampled shop and we compare this value to the actual share of subjects who traded with the last sampled shop. Next, to implement Test 2, we compare the average price observed in the first shop between subjects who only search once and subjects who search multiple times. Under sequential search and search without priors, the average price observed in the first shop should be larger for subjects who search multiple times; under non-sequential search this average price should be the same for both groups. Finally, to implement Test 3, we derive through a linear regression on all individual searches by how much the probability of continued search increases if the price increases by one USD at the current shop. Under sequential search and search without priors, there should be a positive association between the price at the current shop and the probability of continued search; under non-sequential search, there should be no such association.

Taking Test 1b, Test 2, and Test 3 together generates the following result: Search is consistent with sequential search and inconsistent with non-sequential search in the treatments AMT S 1.0, AMT S 1.0 – noinfo, PRO S 1.0 – l , and PRO S 1.0 – r . In contrast, search is consistent with non-sequential search and inconsistent with sequential search in the treatments AMT S 1.0 – r , AMT S 1.0 – r .noinfo, and PRO S 1.0 – l .noinfo. For each of the remaining treatments there is at least one test that rejects sequential search as well as one test that rejects non-sequential search. There do not seem to be robust, systematic differences between treatments with and without information about the price distribution. Overall, we conclude that none of the three search models is fully consistent with subjects' search behavior.

Table A3: Sequential versus Non-Sequential Search – Recall

<i>Test 1a (Recall I)</i>	<i>N</i>	Share purchase last sampled shop or search all shops	difference <i>p</i> -value
AMT S 1.0	207	0.599	0.340
AMT S 1.0 – noinfo	201	0.552	
AMT S 1.0 – <i>l</i>	198	0.444	0.273
AMT S 1.0 – <i>l</i> .noinfo	192	0.500	
AMT S 1.0 – <i>r</i>	194	0.479	0.946
AMT S 1.0 – <i>r</i> .noinfo	208	0.476	
PRO S 1.0 – <i>l</i>	106	0.557	0.001
PRO S 1.0 – <i>l</i> .noinfo	123	0.341	
PRO S 1.0 – <i>r</i>	145	0.703	0.000
PRO S 1.0 – <i>r</i> .noinfo	155	0.439	
<i>Test 1b (Recall II)</i>	<i>N</i>	predicted vs. actual av. share purchase last sampled shop	difference <i>p</i> -value
AMT S 1.0	115	0.201 vs. 0.278	0.063
AMT S 1.0 – noinfo	124	0.184 vs. 0.274	0.008
AMT S 1.0 – <i>l</i>	134	0.132 vs. 0.179	0.143
AMT S 1.0 – <i>l</i> .noinfo	119	0.121 vs. 0.193	0.041
AMT S 1.0 – <i>r</i>	122	0.158 vs. 0.172	0.635
AMT S 1.0 – <i>r</i> .noinfo	133	0.169 vs. 0.180	0.679
PRO S 1.0 – <i>l</i>	90	0.205 vs. 0.478	0.000
PRO S 1.0 – <i>l</i> .noinfo	105	0.176 vs. 0.229	0.165
PRO S 1.0 – <i>r</i>	81	0.287 vs. 0.469	0.001
PRO S 1.0 – <i>r</i> .noinfo	114	0.198 vs. 0.237	0.307

Notes: Test 1a – For each treatment we indicate the share of subjects who purchase from the last sampled shop or who search all shops among those subjects who search at least one shop; *p*-values from two-sided t-tests which compare these shares between a treatment with information about the price distribution and the corresponding no-information treatment. Test 1b – Among all searchers who search at least two but not all shops we compare the (according to the non-sequential search paradigm) predicted share and the actual share of subjects who purchase from the last sampled shop; *p*-values from two-sided t-tests.

Table A4: Sequential versus Non-Sequential Search – Price-Dependence

<i>Test 2 (Price-Dependence I)</i>	<i>N</i>	one search vs. multiple searches	difference <i>p</i> -value
AMT <i>S</i> 1.0	207	4.25 vs. 4.52	0.037
AMT <i>S</i> 1.0 – noinfo	201	4.28 vs. 4.61	0.014
AMT <i>S</i> 1.0 – <i>l</i>	198	5.01 vs. 5.09	0.503
AMT <i>S</i> 1.0 – <i>l</i> .noinfo	192	5.06 vs. 5.21	0.140
AMT <i>S</i> 1.0 – <i>r</i>	194	3.75 vs. 3.77	0.896
AMT <i>S</i> 1.0 – <i>r</i> .noinfo	208	3.80 vs. 3.91	0.219
PRO <i>S</i> 1.0 – <i>l</i>	106	4.56 vs. 5.19	0.001
PRO <i>S</i> 1.0 – <i>l</i> .noinfo	123	5.10 vs. 5.15	0.778
PRO <i>S</i> 1.0 – <i>r</i>	145	3.69 vs. 3.99	0.003
PRO <i>S</i> 1.0 – <i>r</i> .noinfo	155	3.60 vs. 3.89	0.007
<i>Test 3 (Price-Dependence II)</i>	<i>N</i>	change in prob. of continued search	<i>p</i> -value
AMT <i>S</i> 1.0	2235	2.10 %	0.008
AMT <i>S</i> 1.0 – noinfo	2622	1.50 %	0.020
AMT <i>S</i> 1.0 – <i>l</i>	4154	1.48 %	0.018
AMT <i>S</i> 1.0 – <i>l</i> .noinfo	3845	1.01 %	0.125
AMT <i>S</i> 1.0 – <i>r</i>	2743	-0.93 %	0.269
AMT <i>S</i> 1.0 – <i>r</i> .noinfo	2863	1.01 %	0.172
PRO <i>S</i> 1.0 – <i>l</i>	730	14.82 %	0.000
PRO <i>S</i> 1.0 – <i>l</i> .noinfo	1102	1.29 %	0.384
PRO <i>S</i> 1.0 – <i>r</i>	452	14.50 %	0.000
PRO <i>S</i> 1.0 – <i>r</i> .noinfo	1054	3.32 %	0.038

Notes: Test 2 – Among all individuals who search at least one shop we compare the average price at the first shop observed by subjects who search exactly one shop and subjects who search more than one shop; *p*-values from two-sided t-tests. Test 3 – For all individual price observations we derive the average probability of continued search when the price at the current shop increases by one USD; *p*-values from OLS regressions; standard errors are clustered at the subject level.

A.8 Direct Search Costs

We briefly compare the search cost estimates from the two models to subjects' *direct search costs*. For each individual, we record the earnings for one hour of activity (on AMT or Prolific) as well as the average time she needs to obtain a price quote. An individual's direct search costs are defined as

$$\text{direct search costs} = \text{average hourly earnings} \times \frac{\text{mean search duration in seconds}}{3600}. \quad (61)$$

The average direct search costs are 0.25 USD (sd = 0.59) in the November 2023 AMT sample, 0.18 USD (sd = 0.25) in the June 2024 AMT sample, and 0.14 USD (sd = 0.15) in the June 2024 PRO sample. In Table A5 below we compare, for each information and no-information treatment, the search cost estimates to direct search costs. Overall, direct search costs are roughly in line with estimated search costs. They would differ by a factor of around five or more if we would not take context effects into account.

Table A5: Estimated Mean Search Costs and Direct Search Costs

Sample Treatment	Estimated Search Costs (sequential)	Estimated Search Costs (non-sequential)	Direct Search Costs
AMT S 1.0	0.185 (0.310)	0.219 (0.334)	0.228 (0.298)
AMT S 1.0 – noinfo	0.169 (0.298)	0.192 (0.316)	0.207 (0.312)
AMT S 1.0 – <i>l</i>	0.130 (0.298)	0.141 (0.314)	0.191 (0.273)
AMT S 1.0 – <i>l</i> .noinfo	0.136 (0.306)	0.153 (0.327)	0.187 (0.223)
AMT S 1.0 – <i>r</i>	0.100 (0.258)	0.129 (0.299)	0.197 (0.292)
AMT S 1.0 – <i>r</i> .noinfo	0.101 (0.259)	0.123 (0.292)	0.163 (0.198)
PRO S 1.0 – <i>l</i>	0.149 (0.285)	0.138 (0.257)	0.148 (0.126)
PRO S 1.0 – <i>l</i> .noinfo	0.127 (0.258)	0.120 (0.233)	0.154 (0.198)
PRO S 1.0 – <i>r</i>	0.184 (0.322)	0.202 (0.324)	0.144 (0.129)
PRO S 1.0 – <i>r</i> .noinfo	0.067 (0.172)	0.081 (0.180)	0.123 (0.128)

Notes: Standard errors in parentheses. The search cost estimates for each model originate from the regressions in Table 3 and Table 4, respectively. The last column shows the average direct search costs.

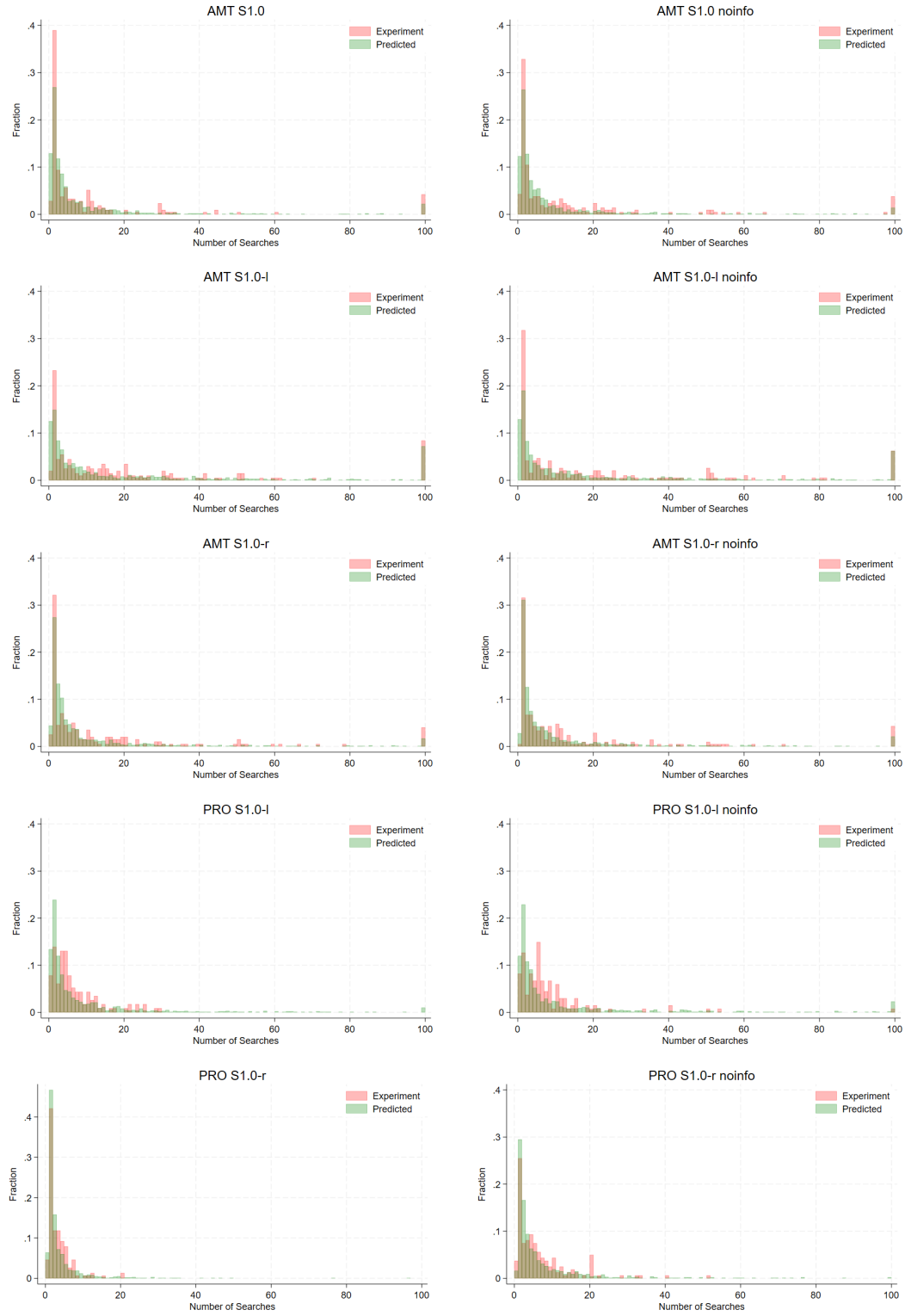


Figure A1: Sequential Search – Observed vs. predicted number of searches

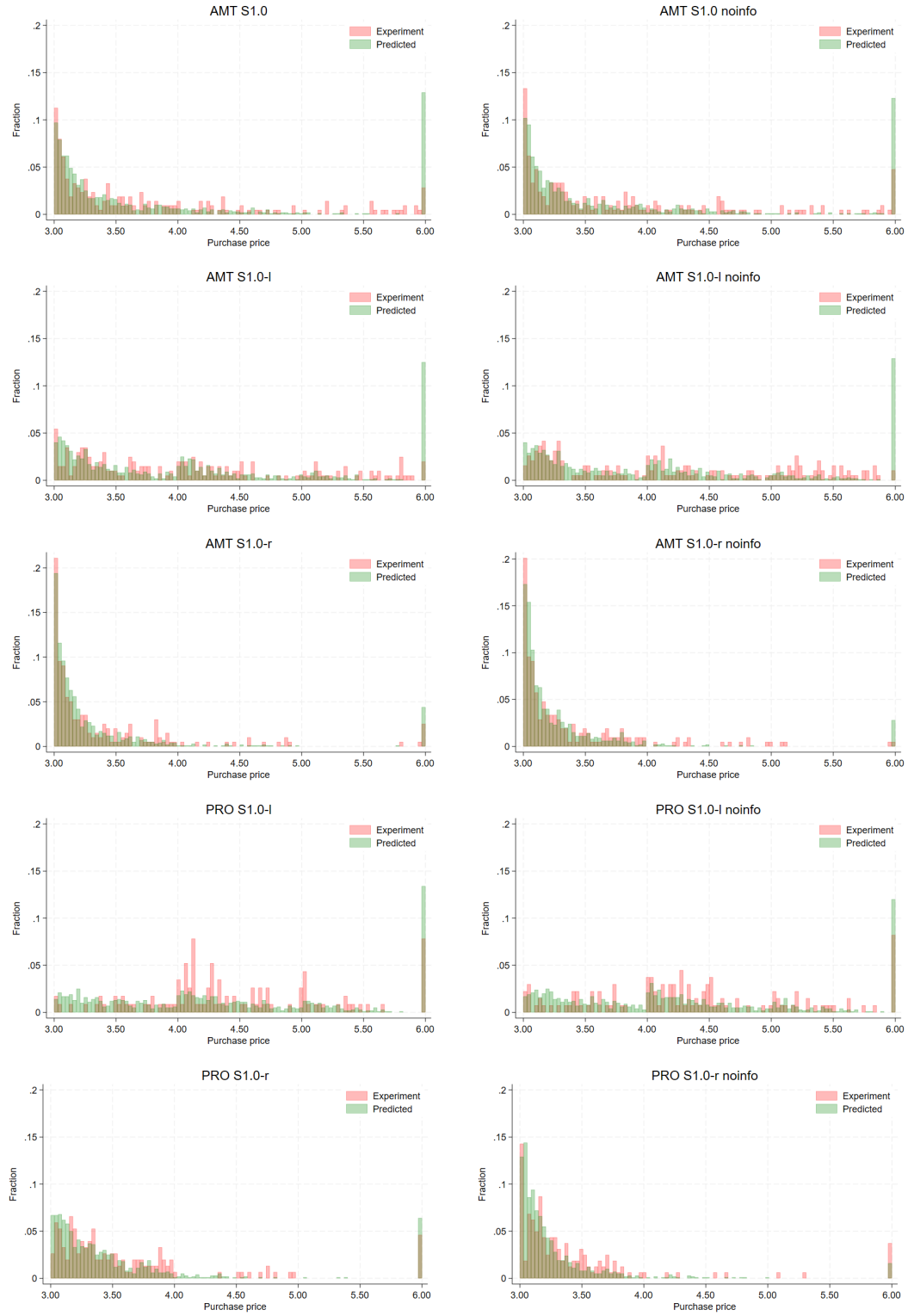


Figure A2: Sequential Search – Observed vs. predicted purchase prices

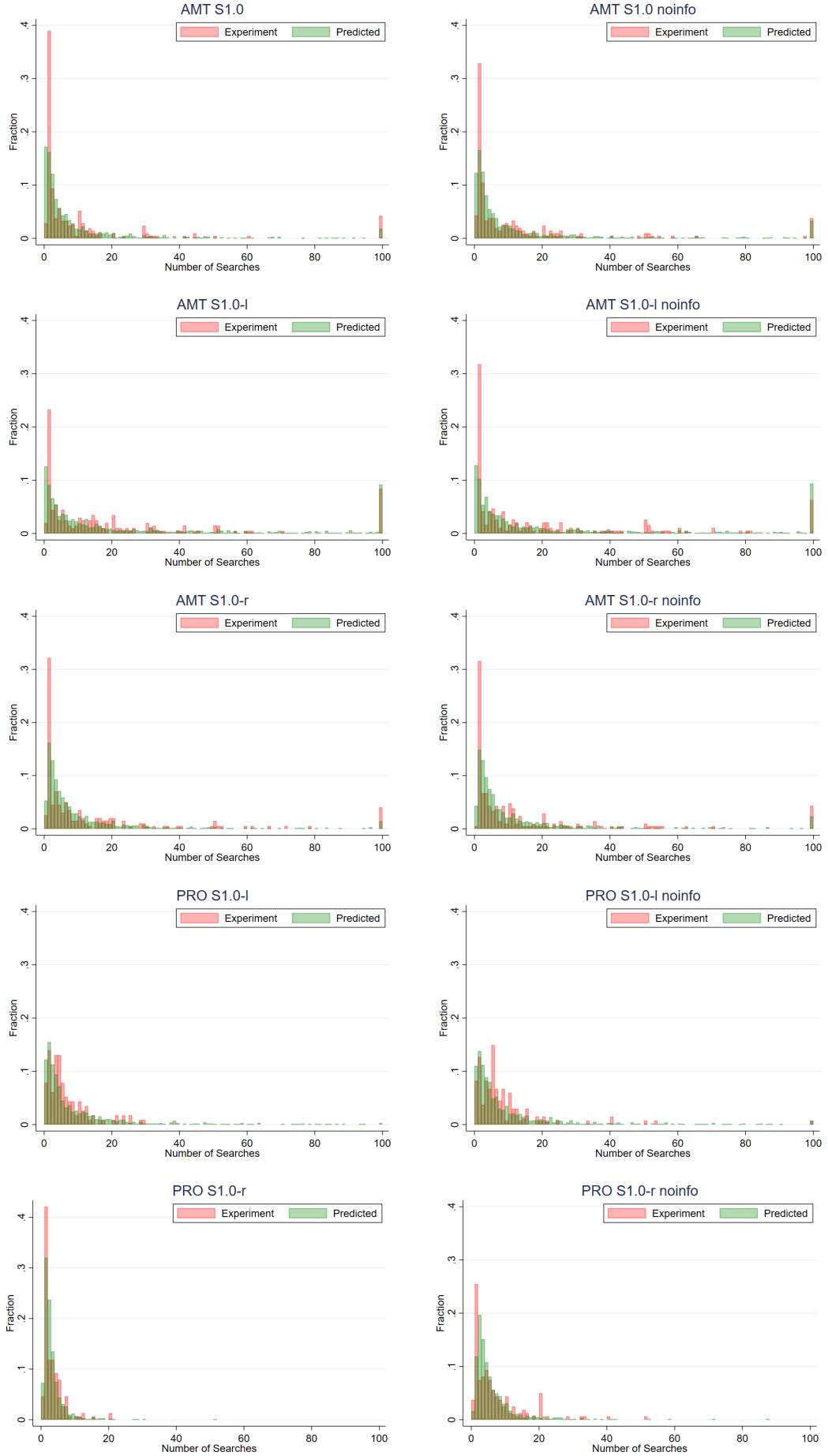


Figure A3: Non-sequential Search – Observed vs. predicted distributions number of searches

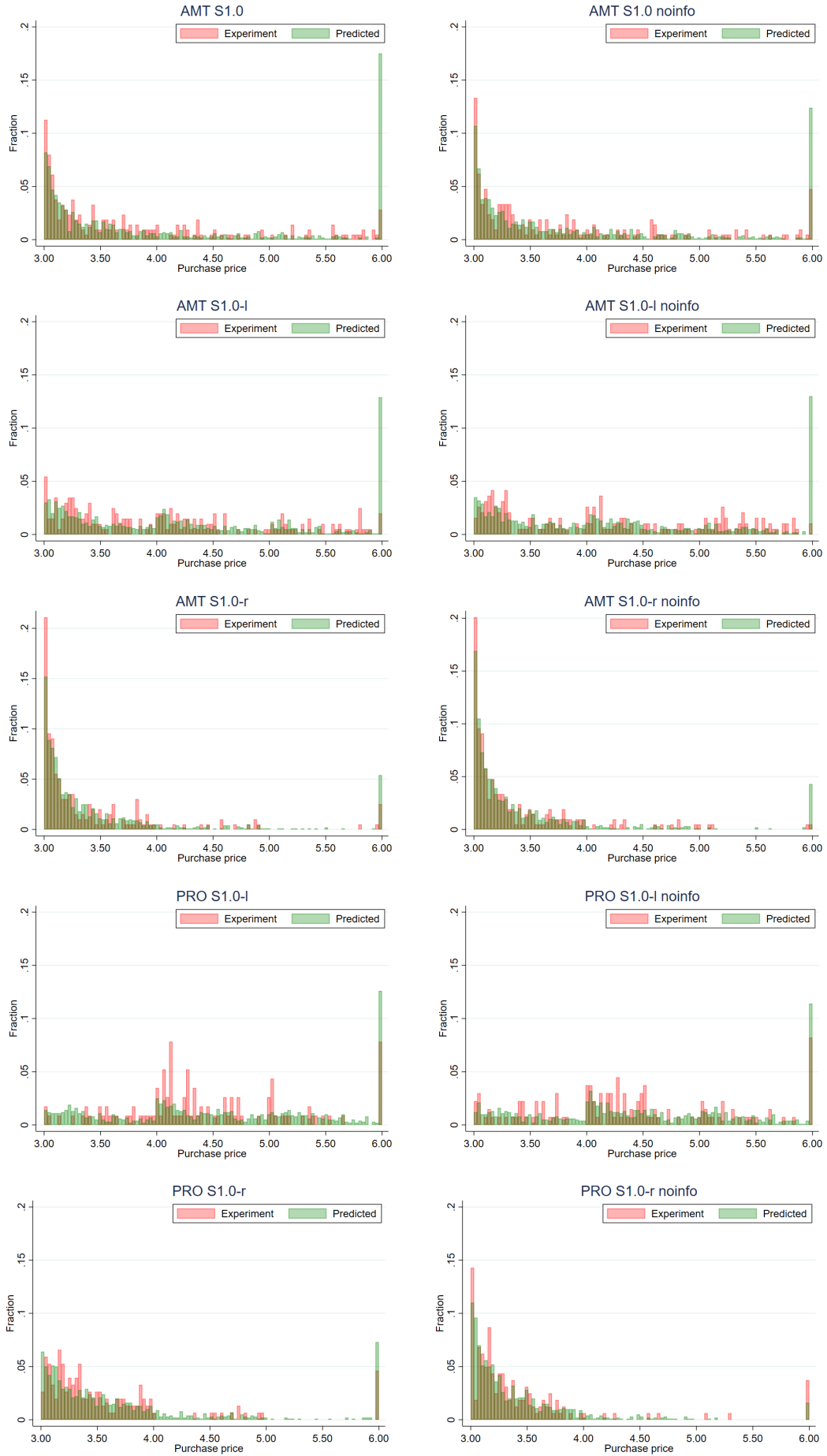


Figure A4: Non-sequential Search – Observed vs. predicted purchase prices

B Online Appendix

Online Appendix [B.1](#) explains the Box-Cox transformation and presents the corresponding estimation results. Online Appendix [B.2](#) contains the tables for the evaluation of search cost estimates. These generate the numbers in Table [5](#). Online Appendix [B.3](#) introduces the search without priors model and shows the search cost estimates as well as the evaluation tables for this model. Online Appendix [B.4](#) shows all results for the variable *restricted number of searches*. Online Appendix [B.5](#) presents our simulation study.

B.1 Box-Cox Transformation

To investigate the robustness of our estimation results with respect to the implicit distributional assumption, we allow for more general search cost distributions through a Box-Cox transformation (Box and Cox, 1964, Birchall et al., 2024). It transforms a non-normal dependent variable into a normally distributed variable. Using the transformation parameter λ , it turns variable k into the transformed variable $g(k)$ where

$$g(k) = \frac{k^\lambda - 1}{\lambda} \text{ if } \lambda \neq 0 \text{ and } g(x) = \ln k \text{ if } \lambda = 0. \quad (62)$$

This specification nests our special case of a lognormal distribution when $\lambda = 0$, and it smoothly converges to a normal distribution as λ approaches 1. Our framework allows estimating λ and comparing estimation outcomes with those obtained under the assumption of lognormality (i.e., holding $\lambda = 0$ fixed instead of estimating it).

We are interested in whether the estimation results with Box-Cox transformation are similar to those obtained under the assumption of lognormality. We consider the following questions: First, do we again get a level of context effects ρ close to 1 and a convexity parameter κ close to 0? Second, are estimated search costs similar in treatments from the same experimental sample? Third, are estimated search costs similar in treatments with and without information about the price distribution? Fourth, is the model fit again better for the sequential than for the non-sequential search model?

For the sequential search model the baseline estimation results are shown in Table 3 and the results with Box-Cox transformation are shown in Table B1 (where the first column is omitted for clarity, so columns (2) to (5) in Table 3 correspond to columns (1) to (4) in Table B1). With Box-Cox transformation, we obtain the context effect parameter $\rho = 0.94$ and the convexity parameter $\kappa = 0.00$. The transformation parameter λ is estimated to be 0.40, which differs from a strictly lognormal distribution, but still matches the lognormal distribution fairly well for small search costs. As in our baseline estimation results, we find that the search cost estimates from the same experimental sample are similar to each other, and also that they are mostly similar in treatments with and without information about the price distribution.

Next, we consider the non-sequential search model. The baseline estimation results are shown in Table 4 and the results with Box-Cox transformation are shown in Table B2. With Box-Cox transformation, we obtain the context effect parameter $\rho = 0.98$. Again, search cost estimates from the same experimental sample as well as estimates from treatments with and without information about the price distribution are similar to each other. As in our baseline estimation results, the model fit is better for the sequential than for the non-sequential search model. We conclude that our estimation results are robust under a Box-Cox transformation.

Table B1: Search Cost Estimates – Sequential Search with Box-Cox Transformation

	(1)	(2)	(3)	(4)
<i>Parameter Estimates</i>				
$S\ 1.0$		-1.60*** (0.05)		
$S\ 5.0$		-1.59*** (0.05)		
$S\ 1.0 - \text{noinfo}$		-1.63*** (0.05)		
$S\ 1.0 - l$			-1.85*** (0.04)	-1.69*** (0.05)
$S\ 1.0 - r$			-1.82*** (0.04)	-1.51*** (0.05)
$S\ 1.0 - l.\text{noinfo}$			-1.80*** (0.04)	-1.74*** (0.05)
$S\ 1.0 - r.\text{noinfo}$			-1.85*** (0.04)	-1.82*** (0.04)
β	-1.40*** (0.04)			
σ	0.51*** (0.03)	0.65*** (0.02)	0.53*** (0.02)	0.52*** (0.02)
κ	0.00 (0.03)			
ρ	0.94*** (0.04)			
λ	0.40*** (0.02)			
Data	AMT	AMT	AMT	PRO
Treatments	$S\ 5.0, S\ 1.0,$ $S\ 1.0 - pr4,$ $S\ 1.0 - pr8,$ $S\ 1.0 - pr12$	$S\ 5.0, S\ 1.0,$ $S\ 1.0 - \text{noinfo}$	$S\ 1.0 - l,$ $S\ 1.0 - r,$ $S\ 1.0 - l.\text{noinfo},$ $S\ 1.0 - r.\text{noinfo}$	$S\ 1.0 - l,$ $S\ 1.0 - r,$ $S\ 1.0 - l.\text{noinfo},$ $S\ 1.0 - r.\text{noinfo}$
ρ est.	est	fixed at est.	fixed at est.	fixed at est.
Observations	1015	613	802	562
Log.Lik.	-2265.24	-1289.31	-1769.66	-1158.29
AIC	4540.49	2586.62	3549.32	2326.57
BIC	4565.10	2604.30	3572.76	2348.23
<i>Implied Search Costs</i>				
SC $S\ 1.0$		0.163 (0.198)		
SC $S\ 5.0$		0.166 (0.200)		
SC $S\ 1.0 - \text{noinfo}$		0.156 (0.192)		
SC $S\ 1.0 - l$			0.085 (0.110)	0.111 (0.129)
SC $S\ 1.0 - r$			0.089 (0.113)	0.155 (0.161)
SC $S\ 1.0 - l.\text{noinfo}$			0.093 (0.116)	0.102 (0.122)
SC $S\ 1.0 - r.\text{noinfo}$			0.086 (0.110)	0.088 (0.110)

Notes: Ordered probit regressions with truncation for the sequential search model with Box-Cox transformation with estimated parameter λ . In the top panel, standard errors of the estimated parameters are in parentheses. The bottom panel presents the search cost estimates for each treatment, derived from the parameter estimates shown in the top panel. The first value represents the mean of the implied search cost distribution, with the standard deviation of the distribution provided in parentheses. Significance at * $p < 0.1$, ** $p < 0.05$, and *** $p < 0.01$.

Table B2: Search Cost Estimates – Non-sequential Search with Box-Cox Transformation

	(1)	(2)	(3)	(4)
<i>Parameter Estimates</i>				
$S\ 1.0$		-1.40*** (0.03)		
$S\ 5.0$		-1.40*** (0.04)		
$S\ 1.0 - \text{noinfo}$		-1.44*** (0.03)		
$S\ 1.0 - l$			-1.65*** (0.03)	-1.53*** (0.04)
$S\ 1.0 - r$			-1.57*** (0.03)	-1.34*** (0.04)
$S\ 1.0 - l.\text{noinfo}$			-1.61*** (0.03)	-1.56*** (0.03)
$S\ 1.0 - r.\text{noinfo}$			-1.61*** (0.03)	-1.56*** (0.03)
β	-1.30*** (0.04)			
σ	0.40*** (0.03)	0.49*** (0.02)	0.40*** (0.01)	0.40*** (0.01)
ρ	0.98*** (0.03)			
λ	0.49*** (0.03)			
<i>Data</i>				
Treatments	AMT $S\ 5.0, S\ 1.0,$ $S\ 1.0 - pr4,$ $S\ 1.0 - pr8,$ $S\ 1.0 - pr12$	AMT $S\ 5.0, S\ 1.0,$ $S\ 1.0 - \text{noinfo}$	AMT $S\ 1.0 - l,$ $S\ 1.0 - r,$ $S\ 1.0 - l.\text{noinfo},$ $S\ 1.0 - r.\text{noinfo}$	PRO $S\ 1.0 - l,$ $S\ 1.0 - r,$ $S\ 1.0 - l.\text{noinfo},$ $S\ 1.0 - r.\text{noinfo}$
ρ est.	est	fixed at est.	fixed at est.	fixed at est.
Observations	1015	613	802	562
Log.Lik.	-2902.34	-1847.12	-2851.85	-1702.20
AIC	5812.68	3702.24	5713.70	3414.41
BIC	5832.37	3719.91	5737.13	3436.06
<i>Implied Search Costs</i>				
SC $S\ 1.0$		0.167 (0.173)		
SC $S\ 5.0$		0.168 (0.173)		
SC $S\ 1.0 - \text{noinfo}$		0.155 (0.165)		
SC $S\ 1.0 - l$			0.084 (0.096)	0.108 (0.112)
SC $S\ 1.0 - r$			0.099 (0.107)	0.158 (0.143)
SC $S\ 1.0 - l.\text{noinfo}$			0.091 (0.101)	0.101 (0.108)
SC $S\ 1.0 - r.\text{noinfo}$			0.092 (0.102)	0.101 (0.108)

Notes: Ordered probit regressions with truncation for the non-sequential search model with Box-Cox transformation with estimated parameter λ . In the top panel, standard errors of the estimated parameters are in parentheses. The bottom panel presents the search cost estimates for each treatment, derived from the parameter estimates shown in the top panel. The first value represents the mean of the implied search cost distribution, with the standard deviation of the distribution provided in parentheses. Significance at * $p < 0.1$, ** $p < 0.05$, and *** $p < 0.01$.

B.2 Tables from Evaluation

Table B3: Observed vs. Predicted Number of Searches (Sequential Search)

$G^{[l]}$ $\hat{G}^{[l]}$	AMT $S^{1.0}$	AMT $S^{1.0}$ noinfo	AMT $S^{1.0-l}$	AMT $S^{1.0-l}$ noinfo	AMT $S^{1.0-r}$	AMT $S^{1.0-r}$ noinfo	PRO $S^{1.0-l}$	PRO $S^{1.0-l}$ noinfo	PRO $S^{1.0-r}$	PRO $S^{1.0-r}$ noinfo
AMT $S^{1.0}$	1.48	3.21	6.82	3.97	7.35	6.81	8.73	8.11	3.53	0.31
	(1.56)	(1.64)	(2.14)	(2.18)	(1.68)	(1.64)	(1.03)	(1.36)	(0.55)	(0.79)
	0.00	1.00	8.00	4.50	2.00	3.00	1.00	2.00	0.00	2.00
	0.30	0.28	0.39	0.44	0.30	0.29	0.42	0.41	0.27	0.30
AMT $S^{1.0}$ noinfo	0.24	0.24	0.52	0.54	0.30	0.27	0.53	0.56	0.19	0.35
	0.18	3.66	2.50	3.96	6.73	7.13	11.87	8.60	4.68	0.83
	(1.60)	(1.62)	(2.19)	(2.19)	(1.68)	(1.62)	(1.14)	(1.37)	(0.56)	(0.78)
	0.00	1.00	7.00	4.50	2.00	3.00	1.00	2.00	0.00	2.00
AMT $S^{1.0-l}$	0.31	0.28	0.38	0.45	0.28	0.30	0.43	0.41	0.29	0.29
	0.27	0.25	0.44	0.55	0.22	0.30	0.53	0.56	0.20	0.26
	0.22	2.79	1.05	1.75	6.01	5.64	13.19	10.70	4.20	1.22
	(1.57)	(1.62)	(2.19)	(2.19)	(1.69)	(1.64)	(1.11)	(1.37)	(0.51)	(0.78)
AMT $S^{1.0-l}$ noinfo	1.00	1.00	4.00	3.50	2.00	3.00	1.00	0.00	0.00	2.00
	0.32	0.26	0.36	0.41	0.28	0.29	0.42	0.42	0.30	0.28
	0.30	0.15	0.41	0.44	0.20	0.29	0.46	0.54	0.21	0.26
	1.64	2.17	2.29	1.80	5.80	4.97	12.50	10.24	4.24	1.70
AMT $S^{1.0-l}$ noinfo	(1.53)	(1.63)	(2.18)	(2.19)	(1.70)	(1.67)	(1.09)	(1.37)	(0.53)	(0.80)
	0.00	0.00	5.00	2.50	2.00	3.00	1.00	0.00	0.00	2.00
	0.31	0.28	0.34	0.41	0.28	0.29	0.43	0.41	0.29	0.30
	0.29	0.24	0.34	0.45	0.26	0.30	0.49	0.54	0.21	0.29
AMT $S^{1.0-r}$	3.87	0.58	2.51	2.25	3.78	3.13	17.19	14.73	6.37	3.45
	(1.62)	(1.66)	(2.22)	(2.23)	(1.72)	(1.68)	(1.18)	(1.45)	(0.61)	(0.85)
	3.00	1.00	2.00	1.50	1.00	2.00	3.00	2.00	1.00	1.00
	0.35	0.31	0.33	0.42	0.30	0.27	0.45	0.44	0.34	0.29
AMT $S^{1.0-r}$ noinfo	0.34	0.29	0.33	0.51	0.25	0.25	0.54	0.57	0.27	0.23
	2.89	1.39	2.46	2.94	4.25	3.56	17.36	16.00	5.85	1.64
	(1.60)	(1.69)	(2.23)	(2.23)	(1.70)	(1.67)	(1.20)	(1.45)	(0.59)	(0.79)
	2.00	2.00	4.00	0.50	1.00	2.00	3.00	3.00	1.00	1.00
PRO $S^{1.0-l}$	0.34	0.30	0.35	0.40	0.29	0.27	0.45	0.46	0.32	0.29
	0.33	0.25	0.39	0.45	0.24	0.23	0.51	0.64	0.24	0.25
	4.92	7.16	10.47	9.81	10.63	9.98	2.43	2.21	0.66	2.74
	(1.48)	(1.55)	(2.08)	(2.09)	(1.62)	(1.58)	(0.79)	(1.19)	(0.34)	(0.65)
PRO $S^{1.0-l}$ noinfo	0.00	1.00	7.00	4.50	3.00	3.00	1.00	2.00	0.00	2.00
	0.26	0.27	0.35	0.41	0.26	0.29	0.33	0.39	0.20	0.28
	0.17	0.32	0.38	0.49	0.19	0.30	0.33	0.52	0.10	0.34
	5.02	6.48	10.63	8.35	9.68	9.98	4.60	3.26	1.03	2.59
PRO $S^{1.0-l}$ noinfo	(1.48)	(1.56)	(2.07)	(2.11)	(1.63)	(1.57)	(0.88)	(1.22)	(0.33)	(0.65)
	0.00	0.00	7.00	3.50	2.00	3.00	1.00	2.00	0.00	2.00
	0.28	0.26	0.33	0.39	0.26	0.24	0.34	0.37	0.22	0.28
	0.34	0.22	0.37	0.42	0.18	0.12	0.34	0.46	0.12	0.38
PRO $S^{1.0-r}$	5.96	8.51	13.17	11.71	10.75	10.28	0.99	1.41	0.46	3.63
	(1.47)	(1.54)	(2.04)	(2.07)	(1.62)	(1.58)	(0.77)	(1.11)	(0.33)	(0.63)
	0.00	1.00	8.00	5.50	3.00	4.00	2.00	3.00	1.00	3.00
	0.29	0.24	0.38	0.44	0.26	0.28	0.32	0.38	0.22	0.30
PRO $S^{1.0-r}$ noinfo	0.33	0.15	0.53	0.56	0.16	0.23	0.34	0.52	0.14	0.41
	1.65	4.46	4.73	3.30	7.58	7.21	10.74	8.15	3.07	0.30
	(1.52)	(1.58)	(2.13)	(2.15)	(1.66)	(1.60)	(1.00)	(1.27)	(0.42)	(0.71)
	2.00	1.00	5.00	1.50	1.00	2.00	2.00	1.00	0.00	1.00
PRO $S^{1.0-r}$ noinfo	0.32	0.28	0.32	0.38	0.30	0.28	0.38	0.39	0.29	0.28
	0.30	0.24	0.34	0.44	0.34	0.31	0.36	0.44	0.22	0.26

Notes: Differences between predicted and realized distribution over the number of searches when the search model is sequential search. In each cell, the first number is the prediction error in means; the second number (in brackets) is the standard error of the point estimate of the prediction error in means; the third number is the prediction error in medians; the fourth (fifth) number indicates the Hellinger distance (Kullback-Leibler divergence) between the predicted and realized distribution.

Table B4: Observed vs. Predicted Purchase Prices (Sequential Search)

$G^{[2]}$ $\hat{G}^{[2]}$	AMT $S 1.0$	AMT $S 1.0$ noinfo	AMT $S 1.0 - l$	AMT $S 1.0 - l$ noinfo	AMT $S 1.0 - r$	AMT $S 1.0 - r$ noinfo	PRO $S 1.0 - l$	PRO $S 1.0 - l$ noinfo	PRO $S 1.0 - r$	PRO $S 1.0 - r$ noinfo
AMT $S 1.0$	0.00	0.00	0.32	0.14	0.16	0.18	0.08	0.05	0.06	0.11
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.08)	(0.08)	(0.06)	(0.06)
	0.10	0.06	0.29	0.05	0.10	0.06	0.13	0.17	0.13	0.02
	0.30	0.29	0.38	0.43	0.28	0.25	0.51	0.46	0.30	0.27
	0.22	0.25	0.41	0.52	0.18	0.15	0.78	0.68	0.28	0.21
AMT $S 1.0$ noinfo	0.04	0.02	0.21	0.14	0.09	0.12	0.09	0.07	0.12	0.05
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.08)	(0.08)	(0.06)	(0.05)
	0.11	0.09	0.13	0.08	0.05	0.03	0.15	0.17	0.18	0.04
	0.30	0.30	0.38	0.45	0.25	0.24	0.50	0.46	0.33	0.26
	0.19	0.26	0.42	0.62	0.18	0.13	0.76	0.62	0.37	0.18
AMT $S 1.0 - l$	0.23	0.24	0.00	0.08	0.01	0.04	0.35	0.25	0.20	0.01
	(0.06)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.08)	(0.08)	(0.06)	(0.05)
	0.19	0.16	0.08	0.15	0.01	0.01	0.37	0.40	0.22	0.08
	0.27	0.26	0.36	0.41	0.24	0.23	0.53	0.45	0.34	0.25
	0.16	0.14	0.43	0.50	0.15	0.14	0.84	0.60	0.40	0.21
AMT $S 1.0 - l$ noinfo	0.14	0.19	0.03	0.06	0.05	0.07	0.36	0.26	0.17	0.04
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.08)	(0.08)	(0.06)	(0.05)
	0.14	0.15	0.02	0.10	0.01	0.00	0.48	0.28	0.20	0.08
	0.30	0.27	0.35	0.42	0.22	0.22	0.53	0.46	0.32	0.25
	0.27	0.20	0.31	0.53	0.10	0.10	0.81	0.60	0.30	0.19
AMT $S 1.0 - r$	0.33	0.31	0.15	0.18	0.06	0.07	0.49	0.42	0.27	0.11
	(0.06)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.23	0.20	0.34	0.41	0.01	0.03	0.71	0.67	0.24	0.11
	0.26	0.29	0.34	0.38	0.24	0.20	0.54	0.48	0.36	0.27
	0.11	0.25	0.35	0.42	0.16	0.11	0.83	0.66	0.48	0.25
AMT $S 1.0 - r$ noinfo	0.27	0.30	0.12	0.25	0.11	0.07	0.50	0.46	0.27	0.10
	(0.06)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.21	0.21	0.25	0.47	0.02	0.04	0.70	0.72	0.24	0.11
	0.27	0.28	0.34	0.39	0.25	0.20	0.54	0.50	0.35	0.25
	0.15	0.21	0.30	0.44	0.17	0.11	0.84	0.74	0.46	0.22
PRO $S 1.0 - l$	0.09	0.12	0.27	0.15	0.08	0.06	0.10	0.01	0.12	0.03
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.06	0.04	0.26	0.12	0.13	0.09	0.10	0.08	0.12	0.00
	0.31	0.29	0.38	0.42	0.31	0.25	0.46	0.43	0.28	0.21
	0.22	0.24	0.42	0.53	0.30	0.17	0.60	0.53	0.23	0.12
PRO $S 1.0 - l$ noinfo	0.12	0.15	0.22	0.05	0.03	0.05	0.15	0.12	0.20	0.02
	(0.06)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.06	0.06	0.25	0.06	0.08	0.07	0.14	0.17	0.14	0.01
	0.30	0.29	0.39	0.41	0.29	0.25	0.46	0.42	0.28	0.23
	0.22	0.23	0.48	0.50	0.25	0.19	0.57	0.49	0.24	0.15
PRO $S 1.0 - r$	0.02	0.02	0.42	0.31	0.10	0.13	0.08	0.12	0.08	0.12
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.04	0.06	0.38	0.23	0.13	0.12	0.01	0.04	0.08	0.04
	0.32	0.31	0.42	0.44	0.31	0.30	0.46	0.43	0.29	0.25
	0.24	0.27	0.61	0.58	0.28	0.31	0.62	0.55	0.24	0.18
PRO $S 1.0 - r$ noinfo	0.34	0.37	0.10	0.26	0.14	0.12	0.48	0.45	0.34	0.17
	(0.06)	(0.06)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.19	0.16	0.04	0.33	0.01	0.03	0.48	0.47	0.22	0.10
	0.29	0.29	0.33	0.38	0.23	0.20	0.51	0.46	0.33	0.25
	0.20	0.19	0.31	0.38	0.14	0.11	0.73	0.62	0.42	0.19

Notes: Differences between predicted and realized distribution over purchase prices when the search model is sequential search. In each cell, the first number is the prediction error in means; the second number (in brackets) is the standard error of the point estimate of the prediction error in means; the third number is the prediction error in medians; the fourth (fifth) number indicates the Hellinger distance (Kullback-Leibler divergence) between the predicted and realized distribution.

Table B5: Observed vs. Predicted Number of Searches (Non-Sequential Search)

$G^{[1]}$ $\hat{G}^{[1]}$	AMT $S1.0$	AMT $S1.0$ noinfo	AMT $S1.0 - l$	AMT $S1.0 - l$ noinfo	AMT $S1.0 - r$	AMT $S1.0 - r$ noinfo	PRO $S1.0 - l$	PRO $S1.0 - l$ noinfo	PRO $S1.0 - r$	PRO $S1.0 - r$ noinfo
AMT $S1.0$	0.72	1.74	4.28	4.06	6.67	6.45	10.68	8.36	4.43	0.87
	(1.56)	(1.64)	(2.17)	(2.18)	(1.68)	(1.62)	(1.11)	(1.36)	(0.53)	(0.79)
	1.00	0.00	7.00	4.50	2.00	3.00	1.00	2.00	0.00	2.00
	0.36	0.32	0.39	0.48	0.30	0.33	0.44	0.43	0.32	0.32
AMT $S1.0$ noinfo	0.41	0.31	0.47	0.61	0.29	0.36	0.59	0.62	0.26	0.36
	1.77	0.15	0.72	1.32	4.75	4.61	12.27	11.21	6.28	2.56
	(1.59)	(1.68)	(2.25)	(2.21)	(1.71)	(1.66)	(1.12)	(1.41)	(0.63)	(0.83)
	2.00	1.00	6.00	4.50	1.00	2.00	0.00	1.00	1.00	1.00
AMT $S1.0 - l$	0.36	0.30	0.40	0.47	0.29	0.30	0.43	0.44	0.35	0.30
	0.38	0.29	0.52	0.63	0.26	0.29	0.54	0.62	0.30	0.32
	4.20	2.68	4.19	4.27	3.16	3.71	16.56	17.12	6.98	3.97
	(1.61)	(1.68)	(2.24)	(2.25)	(1.71)	(1.65)	(1.16)	(1.48)	(0.58)	(0.82)
AMT $S1.0 - l$ noinfo	4.00	4.00	1.00	1.50	1.00	1.00	4.00	3.00	2.00	0.00
	0.40	0.35	0.37	0.47	0.31	0.29	0.45	0.45	0.38	0.30
	0.49	0.40	0.44	0.66	0.27	0.29	0.52	0.59	0.38	0.25
	3.03	0.84	2.24	3.52	3.88	4.48	15.89	14.40	7.29	2.92
AMT $S1.0 - l$ noinfo	(1.58)	(1.66)	(2.22)	(2.25)	(1.70)	(1.64)	(1.16)	(1.44)	(0.61)	(0.80)
	4.00	2.00	3.00	0.50	2.00	3.00	2.00	2.00	2.00	0.00
	0.41	0.31	0.39	0.45	0.31	0.31	0.45	0.44	0.38	0.28
	0.55	0.30	0.47	0.59	0.31	0.28	0.54	0.58	0.38	0.23
AMT $S1.0 - r$	4.32	2.56	7.29	8.16	3.26	2.92	20.82	18.63	9.32	4.60
	(1.59)	(1.67)	(2.26)	(2.28)	(1.71)	(1.66)	(1.25)	(1.48)	(0.66)	(0.83)
	5.00	4.00	1.00	2.50	0.00	1.00	5.00	5.00	3.00	0.00
	0.42	0.34	0.39	0.46	0.30	0.31	0.47	0.47	0.43	0.32
AMT $S1.0 - r$ noinfo	0.54	0.38	0.48	0.57	0.28	0.33	0.57	0.63	0.50	0.26
	5.98	4.08	6.92	8.44	1.06	1.65	21.21	19.44	8.39	6.59
	(1.62)	(1.69)	(2.26)	(2.29)	(1.73)	(1.68)	(1.25)	(1.49)	(0.63)	(0.88)
	5.00	5.00	0.00	3.50	1.00	0.00	6.00	5.50	2.00	1.00
PRO $S1.0 - l$	0.41	0.38	0.37	0.46	0.34	0.31	0.50	0.47	0.41	0.36
	0.53	0.45	0.43	0.66	0.34	0.31	0.65	0.65	0.44	0.35
	4.93	6.95	10.96	11.02	9.86	9.59	2.99	1.06	0.94	2.46
	(1.46)	(1.54)	(2.05)	(2.05)	(1.62)	(1.57)	(0.77)	(1.12)	(0.29)	(0.63)
PRO $S1.0 - l$ noinfo	2.00	1.00	6.00	3.50	1.00	2.00	0.00	1.00	1.00	1.00
	0.34	0.29	0.37	0.44	0.29	0.28	0.30	0.32	0.24	0.25
	0.47	0.20	0.53	0.62	0.16	0.14	0.25	0.28	0.19	0.25
	4.41	6.48	9.74	9.73	9.28	9.57	4.85	1.71	1.27	2.14
PRO $S1.0 - l$ noinfo	(1.46)	(1.54)	(2.05)	(2.06)	(1.63)	(1.57)	(0.81)	(1.14)	(0.31)	(0.63)
	2.00	1.00	5.00	3.50	1.00	2.00	0.00	1.00	1.00	1.00
	0.36	0.28	0.36	0.42	0.29	0.29	0.32	0.33	0.26	0.26
	0.49	0.22	0.47	0.49	0.15	0.17	0.27	0.34	0.20	0.30
PRO $S1.0 - r$	6.56	8.16	14.12	13.84	10.95	10.72	0.27	2.13	0.10	3.69
	(1.45)	(1.54)	(2.02)	(2.03)	(1.62)	(1.57)	(0.69)	(1.08)	(0.29)	(0.62)
	1.00	0.00	8.00	5.50	2.00	3.00	2.00	3.00	0.00	2.00
	0.29	0.31	0.40	0.43	0.28	0.31	0.29	0.35	0.19	0.29
PRO $S1.0 - r$ noinfo	0.20	0.37	0.64	0.40	0.11	0.22	0.30	0.42	0.12	0.39
	2.00	4.41	5.48	6.15	8.10	8.14	7.56	6.27	2.79	0.79
	(1.48)	(1.55)	(2.10)	(2.09)	(1.63)	(1.58)	(0.85)	(1.21)	(0.33)	(0.65)
	4.00	2.00	3.00	0.50	0.00	1.00	3.00	2.00	2.00	0.00
PRO $S1.0 - r$ noinfo	0.38	0.35	0.37	0.42	0.32	0.27	0.36	0.37	0.34	0.28
	0.45	0.48	0.48	0.55	0.30	0.13	0.32	0.37	0.36	0.29

Notes: Differences between predicted and realized distribution over the number of searches when the search model is non-sequential search. In each cell, the first number is the prediction error in means; the second number (in brackets) is the standard error of the point estimate of the prediction error in means; the third number is the prediction error in medians; the fourth (fifth) number indicates the Hellinger distance (Kullback-Leibler divergence) between the predicted and realized distribution.

Table B6: Observed vs. Predicted Purchase Prices (Non-Sequential Search)

$G^{[2]}$ $\hat{G}^{[2]}$	AMT $S1.0$	AMT $S1.0$ noinfo	AMT $S1.0 - l$	AMT $S1.0 - l$ noinfo	AMT $S1.0 - r$	AMT $S1.0 - r$ noinfo	PRO $S1.0 - l$	PRO $S1.0 - l$ noinfo	PRO $S1.0 - r$	PRO $S1.0 - r$ noinfo
AMT $S1.0$	0.31	0.29	0.54	0.43	0.29	0.36	0.23	0.25	0.12	0.33
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.05)	(0.08)	(0.08)	(0.06)	(0.06)
	0.14	0.17	0.56	0.42	0.18	0.19	0.37	0.25	0.02	0.11
	0.35	0.36	0.43	0.45	0.34	0.32	0.51	0.46	0.34	0.34
	0.36	0.41	0.61	0.62	0.31	0.28	0.81	0.66	0.38	0.32
AMT $S1.0$ noinfo	0.14	0.14	0.46	0.39	0.19	0.24	0.10	0.13	0.04	0.21
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.08)	(0.08)	(0.06)	(0.06)
	0.00	0.09	0.48	0.36	0.11	0.11	0.08	0.02	0.10	0.02
	0.35	0.32	0.42	0.42	0.29	0.30	0.50	0.46	0.34	0.30
	0.36	0.29	0.55	0.55	0.23	0.23	0.73	0.62	0.39	0.25
AMT $S1.0 - l$	0.02	0.08	0.18	0.06	0.06	0.11	0.15	0.09	0.08	0.04
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.08)	(0.08)	(0.06)	(0.05)
	0.11	0.13	0.15	0.01	0.06	0.04	0.17	0.17	0.15	0.07
	0.30	0.30	0.37	0.39	0.26	0.24	0.51	0.43	0.36	0.28
	0.21	0.28	0.39	0.48	0.17	0.17	0.80	0.54	0.44	0.21
AMT $S1.0 - l$ noinfo	0.00	0.01	0.29	0.08	0.09	0.17	0.05	0.01	0.09	0.06
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.08)	(0.08)	(0.06)	(0.05)
	0.08	0.08	0.26	0.01	0.06	0.08	0.12	0.12	0.16	0.03
	0.32	0.31	0.39	0.40	0.30	0.25	0.51	0.45	0.34	0.27
	0.29	0.26	0.51	0.47	0.23	0.15	0.75	0.60	0.38	0.22
AMT $S1.0 - r$	0.07	0.11	0.10	0.03	0.08	0.09	0.24	0.16	0.17	0.03
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.08)	(0.08)	(0.06)	(0.05)
	0.15	0.13	0.09	0.02	0.05	0.04	0.24	0.24	0.21	0.05
	0.32	0.30	0.38	0.40	0.27	0.25	0.52	0.48	0.36	0.29
	0.31	0.26	0.50	0.47	0.18	0.20	0.79	0.66	0.43	0.25
AMT $S1.0 - r$ noinfo	0.08	0.11	0.13	0.01	0.01	0.08	0.19	0.21	0.12	0.00
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.08)	(0.08)	(0.06)	(0.05)
	0.15	0.16	0.11	0.03	0.02	0.02	0.24	0.25	0.18	0.08
	0.31	0.33	0.37	0.40	0.26	0.23	0.52	0.44	0.34	0.30
	0.27	0.36	0.44	0.50	0.15	0.13	0.81	0.57	0.37	0.26
PRO $S1.0 - l$	0.15	0.02	0.47	0.35	0.11	0.16	0.10	0.16	0.02	0.12
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.14	0.10	0.43	0.33	0.17	0.13	0.16	0.14	0.05	0.09
	0.34	0.33	0.42	0.41	0.33	0.28	0.49	0.44	0.31	0.31
	0.34	0.35	0.56	0.50	0.30	0.24	0.70	0.53	0.29	0.26
PRO $S1.0 - l$ noinfo	0.04	0.02	0.43	0.29	0.07	0.13	0.05	0.13	0.06	0.05
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.06	0.07	0.44	0.27	0.12	0.14	0.05	0.12	0.07	0.05
	0.32	0.34	0.41	0.42	0.32	0.27	0.47	0.43	0.31	0.27
	0.27	0.37	0.54	0.56	0.28	0.22	0.60	0.49	0.27	0.18
PRO $S1.0 - r$	0.31	0.24	0.67	0.60	0.28	0.34	0.39	0.41	0.11	0.25
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.06)
	0.28	0.26	0.77	0.69	0.28	0.27	0.53	0.50	0.05	0.15
	0.38	0.37	0.44	0.47	0.36	0.36	0.49	0.44	0.33	0.32
	0.40	0.41	0.63	0.72	0.39	0.41	0.67	0.60	0.31	0.25
PRO $S1.0 - r$ noinfo	0.13	0.17	0.19	0.12	0.01	0.04	0.12	0.06	0.18	0.04
	(0.06)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.06	0.05	0.23	0.10	0.07	0.07	0.10	0.10	0.14	0.02
	0.31	0.31	0.35	0.40	0.28	0.24	0.46	0.44	0.33	0.25
	0.28	0.28	0.39	0.48	0.24	0.15	0.57	0.51	0.37	0.19

Notes: Differences between predicted and realized distribution over purchase prices when the search model is non-sequential search. In each cell, the first number is the prediction error in means; the second number (in brackets) is the standard error of the point estimate of the prediction error in means; the third number is the prediction error in medians; the fourth (fifth) number indicates the Hellinger distance (Kullback-Leibler divergence) between the predicted and realized distribution.

B.3 Search without Priors Model

B.3.1 Search Cost Estimation

The estimation of search costs under the search without priors model is similar to that under sequential search. In both cases the search strategy is a reservation price policy. Suppose the DM searches n shops and her reservation value at the n 'th shop equals r_n^{swp} . We assume $\theta = \frac{a}{b}$. For given parameter values ρ, κ, θ , her search cost parameter k from equation (23) equals

$$k^{\text{swp}}(r_n^{\text{swp}}, \rho, \kappa, n, \theta) = \frac{1 + \kappa}{n^{1+\kappa} - (n-1)^{1+\kappa}} \left(\frac{1}{\Delta_F^\rho} \frac{1}{n} \frac{r_n^{\text{swp}} - \theta b}{2} + w \right). \quad (63)$$

Note that this equation is independent of the true price distribution, i.e., it also holds for the left- and right-skewed distributions as we have them in our experiment. The rest of the empirical model proceeds in the same manner as under sequential search.

B.3.2 Computations for the Search without Priors Model

Derivation of the distribution over the number of searches for given search cost parameter k , any price distribution. We derive the distribution over the number of searches for any distribution F from equality (63). For convenience, we write r_n instead of r_n^{nop} . From equation (63) we then get the reservation price at the n 'th search

$$r_n = \theta b + 2\Delta_F^\rho n k. \quad (64)$$

We obtain an upper bound $\bar{k}_1$ for the search cost parameter when we set the reservation price at the first search equal to the highest possible price, $r_1 = b$:

$$\bar{k}_1 = \frac{b - \theta b}{2\Delta_F^\rho}. \quad (65)$$

It is optimal for the DM not to conduct any search if $k \geq \bar{k}_1$ and to conduct at least one search if $k < \bar{k}_1$. Denote by $\bar{n}_k$ the maximal number of searches the DM may conduct. This is the highest integer value n so that $\theta b + 2\Delta_F^\rho n k < b$. If $\bar{n}_k = 1$, the DM searches exactly one time. Suppose that $\bar{n}_k > 1$. In this case, the probability that the DM conducts exactly one search equals $\Pr(1; k) = F(r_2)$. The probability that she conducts $n \in \{2, \dots, \bar{n}_k - 1\}$ searches is

$$\Pr(n; k) = (1 - F(r_{n+1}))^{n-1} F(r_{n+1}) + \sum_{x=1}^{n-1} \binom{n-1}{x} (F(r_{n+1}) - F(r_n))^x (1 - F(r_{n+1}))^{n-1-x}, \quad (66)$$

and the probability that she conducts $\bar{n}_k$ searches equals

$$\Pr(\bar{n}_k; k) = (1 - F(r_{\bar{n}_k}))^{\bar{n}_k - 1}. \quad (67)$$

Derivation of the distribution over purchase prices for given search cost parameter k , any price distribution. We derive the distribution over purchase prices. The DM does not conduct any search if $k \geq \bar{k}_1$ as defined in (65). The expected purchase price then equals b . Define

$$\bar{k}_n = \frac{b - \theta b}{2n\Delta_F^\rho}. \quad (68)$$

We obtain $\bar{k}_n$ by setting $r_n = b$ in equation (64). If $\bar{k}_1 > k \geq \bar{k}_2$, the DM searches exactly once and we have $F(p; k) = F(p)$ and $f(p; k) = f(p)$. If $\bar{k}_n > k \geq \bar{k}_{n+1}$, the maximal number of searches is $n \equiv \bar{n}_k$ and we get the following distribution over purchase prices:

$$F(p; k) = \begin{cases} F_1(p; k) & \text{if } p \in [a, r_2) \\ F_m(p; k) & \text{if } p \in [r_m, r_{m+1}) \text{ and } m \in \{2, \dots, \bar{n}_k - 1\} \\ F_{\bar{n}_k}(p; k) & \text{if } p \in [r_{\bar{n}_k}, b] \end{cases}, \quad (69)$$

where

$$\begin{aligned} F_1(p; k) = & 1 - \sum_{n=1}^{\bar{n}_k - 1} \left[(1 - F(r_{n+1}))^{n-1} (F(r_{n+1}) - F(p)) \right. \\ & + \sum_{x=1}^{n-1} \binom{n-1}{x} (F(r_{n+1}) - F(r_n))^x (1 - F(r_{n+1}))^{n-1-x} (1 - F(p)) \Big] \\ & - (1 - F(r_{\bar{n}_k}))^{\bar{n}_k - 1} (1 - F(p)), \end{aligned} \quad (70)$$

and

$$\begin{aligned} F_m(p; k) = & 1 - \left[(1 - F(r_{m+1}))^{m-1} (F(r_{m+1}) - F(p)) \right. \\ & + \sum_{x=1}^{m-1} \binom{m-1}{x} (F(r_{m+1}) - F(p))^x (1 - F(r_{m+1}))^{m-1-x} (1 - F(p)) \Big] \\ & - \sum_{n=m+1}^{\bar{n}_k - 1} \left[(1 - F(r_{n+1}))^{n-1} (F(r_{n+1}) - F(p)) \right. \\ & + \sum_{x=1}^{n-1} \binom{n-1}{x} (F(r_{n+1}) - F(r_n))^x (1 - F(r_{n+1}))^{n-1-x} (1 - F(p)) \Big] \\ & - (1 - F(r_{\bar{n}_k}))^{\bar{n}_k - 1} (1 - F(p)), \end{aligned} \quad (71)$$

and

$$F_{\bar{n}_k}(p; k) = 1 - (1 - F(p))^{\bar{n}_k}. \quad (72)$$

From this, we obtain the density over purchase prices for $\bar{k}_n > k \geq \bar{k}_{n+1}$ and $n \equiv \bar{n}_k$:

$$f(p; k) = \begin{cases} f_1(p; k) & \text{if } p \in [a, r_2) \\ f_m(p; k) & \text{if } p \in [r_m, r_{m+1}) \text{ and } m \in \{2, \dots, \bar{n}_k - 1\} \\ f_{\bar{n}_k}(p; k) & \text{if } p \in [r_{\bar{n}_k}, b] \\ 0 & \text{else} \end{cases}, \quad (73)$$

where

$$f_1(p; k) = f(p) \left[\sum_{n=1}^{\bar{n}_k-1} (1 - F(r_{n+1}))^{n-1} + \sum_{x=1}^{n-1} \binom{n-1}{x} (F(r_{n+1}) - F(r_n))^x (1 - F(r_{n+1}))^{n-1-x} \right] + (1 - F(r_{\bar{n}_k}))^{\bar{n}_k-1}, \quad (74)$$

and

$$f_m(p; k) = f(p) \left[(1 - F(r_{m+1}))^{m-1} + \sum_{x=1}^{m-1} \binom{m-1}{x} \{x(F(r_{m+1}) - F(p))^{x-1} (1 - F(r_{m+1}))^{m-1-x} (1 - F(p)) + (F(r_{m+1}) - F(p))^x (1 - F(r_{m+1}))^{m-1-x}\} + \sum_{n=m+1}^{\bar{n}_k-1} \{(1 - F(r_{n+1}))^{n-1} + \sum_{x=1}^{n-1} \binom{n-1}{x} (F(r_{n+1}) - F(r_n))^x (1 - F(r_{n+1}))^{n-1-x}\} + (1 - F(r_{\bar{n}_k}))^{\bar{n}_k-1} \right], \quad (75)$$

and

$$f_{\bar{n}_k}(p; k) = f(p) \bar{n}_k (1 - F(p))^{\bar{n}_k-1}. \quad (76)$$

For simulating purchases through inverse transform sampling, we need to obtain $p = F^{-1}(u; k)$ in (69) for $u \in [0, 1]$, which yields

$$p = \begin{cases} F_1^{-1}(u; k) & \text{if } u \in [F(a), F(r_2)) \\ F_m^{-1}(u; k) & \text{if } u \in [F(r_m), F(r_{m+1})) \text{ and } m \in \{2, \dots, \bar{n}_k - 1\} \\ F_{\bar{n}_k}^{-1}(u; k) & \text{if } u \in [F(r_{\bar{n}_k}), F(b)] \end{cases} \quad (77)$$

which can be implemented sequentially by the fact that $F(p; k) = u$, so the initial case distinc-

tion in p -space can be used. The quantiles in (77) are obtained through piecewise inversion of the expressions in (70) to (72) in order to obtain $F(p) = u'$, and then to compute $p = F^{-1}(u')$.

B.3.3 Search Cost Estimation for the Search without Priors Model

We consider the same specifications as for the sequential search model and obtain estimated values of $\rho = 1.00$ and $\kappa = 0.00$. Table B7 provides an overview of the search cost estimates. To simplify the comparison, we also show the search cost estimates from the classic models as well as direct search costs in this table. Table B8 shows the detailed results from the search cost estimation with the search without priors model.

The estimated search costs for the November 2023 AMT sample vary between 0.15 USD and 0.18 USD, and for the June 2024 AMT sample they vary between 0.09 USD and 0.16 USD. In both cases, the differences between the respective information and no-information treatments are not statistically significant (p -values > 0.194). In the June 2024 PRO sample, search costs vary between 0.05 USD and 0.15 USD. The difference between the treatments PRO $S\ 1.0 - l$ and PRO $S\ 1.0 - l.noinfo$ is not significant (p -value = 0.211), while the difference between PRO $S\ 1.0 - r$ and PRO $S\ 1.0 - r.noinfo$ is statistically significant (p -value < 0.001). Therefore, the search without priors model does not equalize the search cost estimates between these treatments.

Table B7: Estimated Mean Search Costs and Direct Search Costs

Sample Treatment	Estimated Search Costs (sequential)	Estimated Search Costs (non-sequential)	Estimated Search Costs (s. w/o priors)	Direct Search Costs
AMT $S\ 1.0$	0.185 (0.310)	0.219 (0.334)	0.179 (0.303)	0.228 (0.298)
AMT $S\ 1.0 - noinfo$	0.169 (0.298)	0.192 (0.316)	0.153 (0.281)	0.207 (0.312)
AMT $S\ 1.0 - l$	0.130 (0.298)	0.141 (0.314)	0.141 (0.266)	0.191 (0.273)
AMT $S\ 1.0 - l.noinfo$	0.136 (0.306)	0.153 (0.327)	0.160 (0.283)	0.187 (0.223)
AMT $S\ 1.0 - r$	0.100 (0.258)	0.129 (0.299)	0.093 (0.215)	0.197 (0.292)
AMT $S\ 1.0 - r.noinfo$	0.101 (0.259)	0.123 (0.292)	0.096 (0.219)	0.163 (0.198)
PRO $S\ 1.0 - l$	0.149 (0.285)	0.138 (0.257)	0.147 (0.232)	0.148 (0.126)
PRO $S\ 1.0 - l.noinfo$	0.127 (0.258)	0.120 (0.233)	0.120 (0.205)	0.154 (0.198)
PRO $S\ 1.0 - r$	0.184 (0.322)	0.202 (0.324)	0.136 (0.221)	0.144 (0.129)
PRO $S\ 1.0 - r.noinfo$	0.067 (0.172)	0.081 (0.180)	0.048 (0.112)	0.123 (0.128)

Notes: Standard errors in parentheses. The search cost estimates for each model originate from the regressions in Table 3, Table 4, and Table B8, respectively. The last column shows the average direct search costs.

Table B8: Search Cost Estimates – Search without Priors

	(1)	(2)	(3)	(4)	(5)
<i>Parameter Estimates</i>					
$S\ 1.0$	-2.33*** (0.21)		-2.62*** (0.34)		
$S\ 5.0$	-0.55* (0.25)		-2.89*** (0.34)		
$S\ 1.0 - \text{noinfo}$	-2.69*** (0.21)		-3.14*** (0.32)		
$S\ 1.0 - l$				-3.35*** (0.26)	-2.85*** (0.18)
$S\ 1.0 - r$				-4.37*** (0.24)	-2.97*** (0.16)
$S\ 1.0 - l.\text{noinfo}$				-3.01*** (0.27)	-3.15*** (0.16)
$S\ 1.0 - r.\text{noinfo}$				-4.29*** (0.23)	-4.30*** (0.14)
β		-2.17*** (0.13)			
σ	2.71*** (0.11)	2.35*** (0.10)	3.14*** (0.17)	2.87*** (0.11)	1.68*** (0.06)
κ		-0.00 (0.03)			
ρ		1.00*** (0.00)			
Data	AMT	AMT	AMT	AMT	PRO
Treatments	$S\ 5.0, S\ 1.0,$ $S\ 1.0 - \text{noinfo}$	$S\ 5.0, S\ 1.0,$ $S\ 1.0 - pr4,$ $S\ 1.0 - pr8,$ $S\ 1.0 - pr12$	$S\ 5.0, S\ 1.0,$ $S\ 1.0 - \text{noinfo}$	$S\ 1.0 - l, S\ 1.0 - r,$ $S\ 1.0 - l.\text{noinfo},$ $S\ 1.0 - r.\text{noinfo}$	$S\ 1.0 - l, S\ 1.0 - r,$ $S\ 1.0 - l.\text{noinfo},$ $S\ 1.0 - r.\text{noinfo}$
ρ est.	fixed 0	est.	fixed at est.	fixed at est.	fixed at est.
Observations	613	1015	613	802	562
Log.Lik.	-1565.32	-2934.17	-1536.51	-2198.06	-1367.91
AIC	3138.64	5876.34	3081.01	4406.12	2745.82
BIC	3156.32	5896.03	3098.69	4429.55	2767.48
<i>Implied Search Costs</i>					
SC $S\ 5.0$	2.174 (4.001)		0.165 (0.292)		
SC $S\ 1.0$	0.944 (2.566)		0.179 (0.303)		
SC $S\ 1.0 - \text{noinfo}$	0.772 (2.288)		0.153 (0.281)		
SC $S\ 1.0 - l$				0.141 (0.266)	0.147 (0.232)
SC $S\ 1.0 - l.\text{noinfo}$				0.160 (0.283)	0.120 (0.205)
SC $S\ 1.0 - r$				0.093 (0.215)	0.136 (0.221)
SC $S\ 1.0 - r.\text{noinfo}$				0.096 (0.219)	0.048 (0.112)

Notes: Ordered probit regressions with truncation for the search without priors model. In the top panel, standard errors of the estimated parameters are in parentheses. The bottom panel presents the search cost estimates for each treatment, derived from the parameter estimates shown in the top panel. The first value represents the mean of the implied search cost distribution, with the standard deviation of the distribution provided in parentheses. Significance at * $p < 0.1$, ** $p < 0.05$, and *** $p < 0.01$.

B.3.4 Evaluation of Search Cost Estimates for the Search without Priors Model

Table B9: Predicted versus Observed Distribution over Number of Searches

	Prediction Error		Difference predicted and realized distribution	
	Means (SE)	Medians	HD	KL
<i>sequential search (in-sample)</i>				
information treatments	1.84 (1.32)	1.40	0.30	0.27
no-information treatments	2.52 (1.48)	1.70	0.32	0.33
all treatments	2.18 (1.40)	1.55	0.31	0.30
<i>non-sequential search (in-sample)</i>				
information treatments	2.25 (1.31)	0.40	0.30	0.30
no-information treatments	1.56 (1.48)	0.50	0.33	0.36
all treatments	1.91 (1.40)	0.45	0.32	0.33
<i>search without priors (in-sample)</i>				
information treatments	2.09 (1.26)	0.60	0.28	0.25
no-information treatments	2.57 (1.44)	1.20	0.30	0.30
all treatments	2.33 (1.35)	0.90	0.29	0.27
<i>sequential search (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	2.83 (1.41)	1.85	0.32	0.33
any treatment → any treatment	3.43 (1.38)	1.87	0.32	0.33
<i>non-sequential search (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	2.64 (1.40)	1.35	0.35	0.38
any treatment → any treatment	3.67 (1.36)	1.48	0.34	0.36
<i>search without priors (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	2.95 (1.36)	1.40	0.31	0.31
any treatment → any treatment	4.30 (1.33)	1.87	0.31	0.33

Notes: Average differences between predicted and observed distribution over number of searches, as outlined in Subsection 4.2, for the classic search models as well as the search without priors model. The detailed values for each treatment combination for the latter model are presented in Table B11.

Table B10: Predicted versus Observed Distribution over Purchase Prices

	Prediction Error		Difference predicted and realized distribution	
	Means (SE)	Medians	HD	KL
<i>sequential search (in-sample)</i>				
information treatments	0.05 (0.06)	0.07	0.33	0.33
no-information treatments	0.09 (0.06)	0.10	0.32	0.32
all treatments	0.07 (0.06)	0.09	0.32	0.32
<i>non-sequential search (in-sample)</i>				
information treatments	0.16 (0.06)	0.11	0.36	0.39
no-information treatments	0.09 (0.06)	0.05	0.33	0.31
all treatments	0.13 (0.06)	0.08	0.34	0.35
<i>search without priors (in-sample)</i>				
information treatments	0.05 (0.06)	0.05	0.33	0.32
no-information treatments	0.07 (0.06)	0.05	0.31	0.31
all treatments	0.06 (0.06)	0.05	0.32	0.32
<i>sequential search (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	0.09 (0.06)	0.09	0.33	0.32
any treatment → any treatment	0.13 (0.06)	0.14	0.33	0.33
<i>non-sequential search (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	0.15 (0.06)	0.10	0.36	0.39
any treatment → any treatment	0.13 (0.06)	0.11	0.35	0.38
<i>search without priors (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	0.11 (0.06)	0.07	0.34	0.35
any treatment → any treatment	0.15 (0.06)	0.13	0.34	0.36

Notes: Average differences between predicted and observed distribution over purchase prices, as outlined in Subsection 4.2, for the classic search models as well as the search without priors model. The detailed values for each treatment combination for the latter model are presented in Table B12.

Table B11: Observed vs. Predicted Number of Searches (Search w/o Priors)

$G^{[1]}$ $\hat{G}^{[1]}$	AMT $S1.0$	AMT $S1.0$ noinfo	AMT $S1.0 - l$	AMT $S1.0 - l$ noinfo	AMT $S1.0 - r$	AMT $S1.0 - r$ noinfo	PRO $S1.0 - l$	PRO $S1.0 - l$ noinfo	PRO $S1.0 - r$	PRO $S1.0 - r$ noinfo
AMT $S1.0$	2.70	3.40	2.90	4.89	8.39	7.01	11.07	8.29	2.31	0.93
	(1.52)	(1.62)	(2.15)	(2.14)	(1.66)	(1.63)	(1.03)	(1.30)	(0.44)	(0.72)
	0.00	0.00	4.00	2.50	2.00	3.00	2.00	0.00	0.00	2.00
	0.32	0.27	0.36	0.40	0.31	0.31	0.40	0.41	0.29	0.30
AMT $S1.0$ noinfo	0.32	0.24	0.37	0.45	0.29	0.29	0.42	0.50	0.23	0.33
	0.84	2.23	1.62	0.84	6.80	6.81	13.69	11.24	2.96	0.48
	(1.55)	(1.64)	(2.16)	(2.19)	(1.68)	(1.63)	(1.08)	(1.37)	(0.47)	(0.76)
	1.00	0.00	4.00	1.50	2.00	3.00	1.00	0.00	0.00	2.00
AMT $S1.0 - l$	0.32	0.26	0.34	0.41	0.29	0.31	0.41	0.40	0.29	0.28
	0.31	0.20	0.33	0.48	0.26	0.30	0.42	0.48	0.22	0.26
	0.74	3.68	2.33	1.90	7.35	7.50	12.31	9.02	3.96	0.31
	(1.55)	(1.61)	(2.15)	(2.17)	(1.67)	(1.61)	(1.04)	(1.31)	(0.51)	(0.75)
AMT $S1.0 - l$ noinfo	1.00	0.00	3.00	0.50	2.00	3.00	3.00	1.00	0.00	2.00
	0.33	0.28	0.35	0.43	0.29	0.31	0.40	0.38	0.30	0.28
	0.33	0.24	0.39	0.57	0.29	0.33	0.41	0.39	0.23	0.24
	1.83	4.29	5.57	3.93	8.71	8.87	9.48	9.06	2.18	0.40
AMT $S1.0 - l$ noinfo	(1.53)	(1.60)	(2.12)	(2.15)	(1.65)	(1.59)	(0.98)	(1.33)	(0.41)	(0.72)
	1.00	0.00	5.00	2.00	2.00	3.00	1.00	1.00	1.00	2.00
	0.33	0.27	0.34	0.41	0.30	0.29	0.38	0.38	0.27	0.26
	0.35	0.21	0.35	0.47	0.31	0.24	0.36	0.43	0.18	0.23
AMT $S1.0 - r$	2.50	0.10	3.84	6.37	4.57	3.46	19.11	17.89	6.38	3.46
	(1.59)	(1.65)	(2.22)	(2.24)	(1.69)	(1.66)	(1.14)	(1.43)	(0.58)	(0.83)
	2.50	2.00	0.00	4.50	0.00	1.50	8.00	6.00	1.00	1.00
	0.36	0.31	0.34	0.42	0.29	0.29	0.48	0.46	0.35	0.31
AMT $S1.0 - r$ noinfo	0.37	0.29	0.35	0.55	0.24	0.23	0.57	0.60	0.30	0.27
	2.22	0.71	4.50	5.20	4.66	4.69	17.56	17.22	6.45	2.15
	(1.58)	(1.64)	(2.22)	(2.23)	(1.70)	(1.64)	(1.14)	(1.43)	(0.59)	(0.80)
	3.00	2.00	0.00	3.50	1.00	2.00	6.00	6.00	1.00	1.00
PRO $S1.0 - l$	0.36	0.29	0.34	0.42	0.29	0.28	0.45	0.45	0.35	0.28
	0.37	0.25	0.35	0.56	0.27	0.23	0.49	0.56	0.30	0.23
	6.80	9.26	13.21	12.47	11.54	11.22	0.48	1.07	0.66	4.02
	(1.46)	(1.53)	(2.02)	(2.03)	(1.62)	(1.57)	(0.68)	(1.07)	(0.28)	(0.62)
PRO $S1.0 - l$ noinfo	0.00	1.00	6.00	3.50	3.00	4.00	0.00	1.00	1.00	2.00
	0.22	0.27	0.36	0.39	0.28	0.28	0.25	0.30	0.17	0.31
	0.08	0.21	0.59	0.44	0.12	0.14	0.19	0.28	0.11	0.45
	6.46	8.48	11.21	11.16	10.91	11.00	2.29	0.80	0.12	3.55
PRO $S1.0 - l$ noinfo	(1.46)	(1.53)	(2.03)	(2.04)	(1.62)	(1.57)	(0.72)	(1.11)	(0.28)	(0.62)
	0.00	1.00	5.00	3.50	2.00	3.00	1.00	1.00	0.00	2.00
	0.27	0.22	0.36	0.39	0.25	0.28	0.27	0.32	0.17	0.26
	0.19	0.08	0.54	0.52	0.07	0.19	0.20	0.32	0.09	0.30
PRO $S1.0 - r$	6.51	8.62	12.52	11.69	11.08	11.15	3.13	0.01	0.35	4.16
	(1.46)	(1.53)	(2.02)	(2.03)	(1.62)	(1.57)	(0.75)	(1.10)	(0.28)	(0.62)
	0.00	1.00	6.00	3.50	2.00	3.00	1.00	1.00	0.00	2.00
	0.27	0.24	0.34	0.37	0.27	0.27	0.30	0.32	0.17	0.30
PRO $S1.0 - r$ noinfo	0.20	0.14	0.46	0.30	0.14	0.11	0.23	0.32	0.09	0.42
	3.28	4.78	4.17	3.02	8.82	8.75	9.80	8.55	2.19	1.19
	(1.48)	(1.56)	(2.08)	(2.09)	(1.63)	(1.58)	(0.84)	(1.20)	(0.35)	(0.66)
	2.00	2.00	0.00	2.50	1.00	2.00	5.50	4.00	1.00	1.00
PRO $S1.0 - r$ noinfo	0.34	0.32	0.38	0.43	0.29	0.25	0.41	0.42	0.29	0.25
	0.32	0.35	0.53	0.72	0.27	0.11	0.44	0.50	0.24	0.26

Notes: Differences between predicted and realized distribution over the number of searches when the search model is search without priors. In each cell, the first number is the prediction error in means; the second number (in brackets) is the standard error of the point estimate of the prediction error in means; the third number is the prediction error in medians; the fourth (fifth) number indicates the Hellinger distance (Kullback-Leibler divergence) between the predicted and realized distribution.

Table B12: Observed vs. Predicted Purchase Prices (Search w/o Priors)

$G^{[1]}$ $\hat{G}^{[1]}$	AMT $S1.0$	AMT $S1.0$ noinfo	AMT $S1.0 - l$	AMT $S1.0 - l$ noinfo	AMT $S1.0 - r$	AMT $S1.0 - r$ noinfo	PRO $S1.0 - l$	PRO $S1.0 - l$ noinfo	PRO $S1.0 - r$	PRO $S1.0 - r$ noinfo
AMT $S1.0$	0.12	0.10	0.20	0.12	0.32	0.32	0.20	0.08	0.11	0.29
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.05)	(0.07)	(0.08)	(0.06)	(0.06)
	0.05	0.03	0.18	0.06	0.13	0.13	0.20	0.18	0.09	0.04
	0.32	0.33	0.36	0.38	0.32	0.30	0.50	0.47	0.34	0.29
	0.30	0.33	0.42	0.45	0.28	0.26	0.68	0.64	0.43	0.24
AMT $S1.0$ noinfo	0.02	0.04	0.10	0.03	0.20	0.27	0.29	0.21	0.08	0.13
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.05)	(0.07)	(0.08)	(0.06)	(0.06)
	0.05	0.06	0.10	0.01	0.09	0.08	0.29	0.24	0.11	0.02
	0.31	0.29	0.35	0.40	0.29	0.29	0.52	0.46	0.32	0.28
	0.29	0.28	0.38	0.45	0.23	0.25	0.76	0.61	0.36	0.23
AMT $S1.0 - l$	0.02	0.04	0.12	0.01	0.16	0.25	0.27	0.16	0.01	0.10
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.05)	(0.07)	(0.08)	(0.06)	(0.06)
	0.08	0.07	0.14	0.02	0.08	0.08	0.27	0.20	0.14	0.03
	0.32	0.30	0.35	0.39	0.28	0.27	0.52	0.46	0.32	0.27
	0.32	0.28	0.38	0.48	0.21	0.18	0.78	0.61	0.31	0.23
AMT $S1.0 - l$ noinfo	0.03	0.01	0.24	0.09	0.23	0.27	0.13	0.12	0.07	0.19
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.05)	(0.08)	(0.08)	(0.06)	(0.06)
	0.06	0.00	0.22	0.04	0.12	0.11	0.14	0.18	0.06	0.01
	0.31	0.29	0.39	0.41	0.30	0.30	0.49	0.47	0.32	0.28
	0.29	0.23	0.46	0.52	0.23	0.25	0.67	0.67	0.33	0.23
AMT $S1.0 - r$	0.14	0.20	0.03	0.22	0.01	0.05	0.50	0.41	0.17	0.00
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.17	0.16	0.07	0.37	0.02	0.01	0.64	0.62	0.21	0.08
	0.30	0.28	0.35	0.38	0.24	0.22	0.55	0.47	0.34	0.26
	0.21	0.21	0.38	0.48	0.15	0.11	0.85	0.64	0.38	0.22
AMT $S1.0 - r$ noinfo	0.19	0.20	0.08	0.21	0.02	0.05	0.42	0.39	0.15	0.03
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.18	0.14	0.19	0.37	0.02	0.02	0.47	0.50	0.20	0.07
	0.31	0.29	0.32	0.37	0.26	0.21	0.51	0.47	0.34	0.24
	0.30	0.27	0.33	0.37	0.19	0.10	0.73	0.66	0.37	0.17
PRO $S1.0 - l$	0.12	0.11	0.36	0.22	0.28	0.24	0.00	0.06	0.01	0.20
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.15	0.18	0.39	0.22	0.26	0.19	0.03	0.00	0.01	0.13
	0.36	0.35	0.40	0.42	0.38	0.32	0.44	0.42	0.29	0.29
	0.34	0.36	0.51	0.56	0.45	0.34	0.53	0.47	0.25	0.22
PRO $S1.0 - l$ noinfo	0.05	0.02	0.23	0.14	0.14	0.18	0.10	0.06	0.06	0.09
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.09	0.07	0.27	0.13	0.19	0.15	0.09	0.09	0.05	0.07
	0.34	0.32	0.40	0.40	0.35	0.30	0.45	0.43	0.29	0.26
	0.36	0.31	0.46	0.49	0.40	0.29	0.51	0.51	0.27	0.18
PRO $S1.0 - r$	0.08	0.03	0.32	0.17	0.17	0.24	0.12	0.00	0.02	0.19
	(0.07)	(0.07)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.08)	(0.06)	(0.05)
	0.11	0.15	0.38	0.17	0.21	0.17	0.10	0.04	0.00	0.12
	0.37	0.33	0.38	0.37	0.35	0.31	0.46	0.43	0.29	0.29
	0.42	0.32	0.48	0.43	0.38	0.33	0.58	0.46	0.26	0.23
PRO $S1.0 - r$ noinfo	0.30	0.32	0.12	0.24	0.09	0.05	0.51	0.43	0.29	0.13
	(0.06)	(0.06)	(0.07)	(0.07)	(0.05)	(0.04)	(0.07)	(0.07)	(0.06)	(0.05)
	0.13	0.10	0.00	0.18	0.05	0.02	0.51	0.38	0.18	0.05
	0.28	0.29	0.35	0.39	0.28	0.20	0.51	0.49	0.32	0.23
	0.14	0.24	0.40	0.50	0.24	0.10	0.70	0.68	0.40	0.16

Notes: Differences between predicted and realized distribution over purchase prices when the search model is search without priors. In each cell, the first number is the prediction error in means; the second number (in brackets) is the standard error of the point estimate of the prediction error in means; the third number is the prediction error in medians; the fourth (fifth) number indicates the Hellinger distance (Kullback-Leibler divergence) between the predicted and realized distribution.

B.4 Restricted Number of Searches

Table B13: Average Search Outcomes – Restricted Number of Searches

Sample Treatment	Share Searchers	Mean and Median Rest. No. Searches	Diff. <i>p</i> -value
<i>Information and No-Information Treatments</i>			
AMT <i>S</i> 1.0	0.972	2.49 (1.45)	0.451
AMT <i>S</i> 1.0 – noinfo	0.957	2.60 (1.46)	
AMT <i>S</i> 1.0 – <i>l</i>	0.980	3.08 (1.34)	0.342
AMT <i>S</i> 1.0 – <i>l</i> .noinfo	1.000	2.95 (1.39)	
AMT <i>S</i> 1.0 – <i>r</i>	0.975	2.77 (1.43)	0.674
AMT <i>S</i> 1.0 – <i>r</i> .noinfo	0.995	2.83 (1.38)	
PRO <i>S</i> 1.0 – <i>l</i>	0.922	3.02 (1.39)	0.510
PRO <i>S</i> 1.0 – <i>l</i> .noinfo	0.918	3.13 (1.40)	
PRO <i>S</i> 1.0 – <i>r</i>	0.954	2.20 (1.37)	0.000
PRO <i>S</i> 1.0 – <i>r</i> .noinfo	0.963	2.86 (1.40)	
<i>Piece Rate, Scale, and Search Cost Treatments</i>			
AMT <i>S</i> 0.2 – <i>pr</i> 4	0.917	2.25 (1.47)	–
AMT <i>S</i> 0.2 – <i>pr</i> 8	0.970	2.50 (1.46)	
AMT <i>S</i> 0.2 – <i>pr</i> 12	0.974	2.38 (1.44)	
AMT <i>S</i> 5.0	0.974	2.50 (1.42)	
PRO <i>S</i> 1.0 – <i>sc</i> 4	0.968	2.93 (1.36)	0.001
PRO <i>S</i> 1.0 – <i>sc</i> 16	0.920	2.43 (1.48)	

Notes: Standard deviation in parentheses. The *p*-value in the last column originates from a two-sided t-test which compares the mean restricted number of searches between an information treatment and the no-information treatment with the same price distribution.

Table B14: Predicted versus Observed Distribution over Restricted Number of Searches

	Prediction Error		Difference predicted and realized distribution	
	Means	Medians	HD	KL
<i>sequential search (in-sample)</i>				
information treatments	0.27 (0.12)	0.60	0.14	0.08
no-information treatments	0.30 (0.12)	0.80	0.17	0.09
all treatments	0.29 (0.12)	0.70	0.16	0.09
<i>non-sequential search (in-sample)</i>				
information treatments	0.12 (0.12)	0.20	0.17	0.12
no-information treatments	0.16 (0.11)	0.20	0.20	0.15
all treatments	0.14 (0.12)	0.20	0.18	0.13
<i>search without priors (in-sample)</i>				
information treatments	0.11 (0.12)	0.00	0.13	0.06
no-information treatments	0.09 (0.12)	0.40	0.17	0.11
all treatments	0.10 (0.12)	0.20	0.15	0.08
<i>sequential search (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	0.37 (0.12)	0.80	0.17	0.12
any treatment → any treatment	0.34 (0.12)	0.92	0.17	0.11
<i>non-sequential search (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	0.30 (0.12)	0.60	0.20	0.16
any treatment → any treatment	0.26 (0.12)	0.58	0.19	0.14
<i>search without priors (out-of-sample)</i>				
inf./no-inf. → no-inf./inf.	0.26 (0.12)	0.55	0.17	0.12
any treatment → any treatment	0.35 (0.12)	0.71	0.18	0.13

Notes: Average differences between predicted and observed distribution over the restricted number of searches, as outlined in Subsection 4.2, for the classic search models as well as the search without priors model. The detailed values for each treatment combination are presented in Table B15 to Table B17.

Table B15: Observed vs. Predicted Restricted Number of Searches (Sequential Search)

$G^{[1A]}$ $\hat{G}^{[1A]}$	AMT $S1.0$	AMT $S1.0$ noinfo	AMT $S1.0 - l$	AMT $S1.0 - l$ noinfo	AMT $S1.0 - r$	AMT $S1.0 - r$ noinfo	PRO $S1.0 - l$	PRO $S1.0 - l$ noinfo	PRO $S1.0 - r$	PRO $S1.0 - r$ noinfo
AMT $S1.0$	0.14	0.26	0.74	0.49	0.62	0.64	0.66	0.71	0.03	0.67
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	0.00	1.00	2.00	1.00	2.00	2.00	1.00	1.00	0.00	2.00
	0.17	0.16	0.25	0.35	0.19	0.21	0.18	0.17	0.08	0.17
	0.11	0.09	0.20	0.29	0.14	0.15	0.12	0.11	0.02	0.12
AMT $S1.0$ noinfo	0.07	0.21	0.62	0.50	0.48	0.57	0.58	0.66	0.14	0.57
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	0.00	1.00	1.00	1.00	2.00	2.00	1.00	1.00	0.00	2.00
	0.16	0.13	0.23	0.36	0.17	0.19	0.17	0.16	0.07	0.15
	0.09	0.06	0.16	0.31	0.11	0.12	0.12	0.10	0.02	0.09
AMT $S1.0 - l$	0.22	0.09	0.26	0.24	0.30	0.39	0.21	0.40	0.21	0.42
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	1.00	1.00	0.00	0.00	2.00	2.00	0.00	0.00	0.00	2.00
	0.15	0.10	0.18	0.29	0.14	0.16	0.10	0.11	0.08	0.12
	0.09	0.04	0.11	0.20	0.07	0.08	0.04	0.05	0.02	0.06
AMT $S1.0 - l$ noinfo	0.08	0.03	0.32	0.25	0.37	0.45	0.23	0.41	0.20	0.38
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	1.00	0.00	0.00	0.00	2.00	2.00	0.00	0.00	0.00	2.00
	0.15	0.12	0.19	0.29	0.14	0.18	0.12	0.11	0.09	0.11
	0.08	0.06	0.12	0.20	0.07	0.10	0.06	0.05	0.03	0.05
AMT $S1.0 - r$	0.50	0.18	0.07	0.08	0.14	0.18	0.04	0.17	0.41	0.21
	(0.11)	(0.11)	(0.10)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	2.00	1.00	0.00	0.00	1.00	1.00	0.00	0.00	1.00	1.00
	0.17	0.11	0.13	0.29	0.13	0.12	0.08	0.09	0.12	0.07
	0.13	0.04	0.06	0.22	0.06	0.05	0.03	0.03	0.06	0.02
AMT $S1.0 - r$ noinfo	0.39	0.32	0.12	0.02	0.13	0.20	0.05	0.15	0.37	0.32
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	2.00	1.00	0.00	0.00	1.00	1.00	0.00	0.00	0.00	1.00
	0.15	0.13	0.14	0.27	0.11	0.11	0.11	0.09	0.11	0.11
	0.10	0.07	0.07	0.19	0.04	0.04	0.05	0.03	0.05	0.04
PRO $S1.0 - l$	0.06	0.20	0.66	0.52	0.73	0.69	0.59	0.73	0.14	0.77
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	0.00	1.00	1.00	1.00	3.00	2.00	1.00	1.00	0.00	2.00
	0.15	0.14	0.22	0.32	0.22	0.23	0.16	0.19	0.05	0.21
	0.08	0.07	0.16	0.25	0.20	0.21	0.10	0.13	0.01	0.19
PRO $S1.0 - l$ noinfo	0.03	0.04	0.57	0.28	0.54	0.61	0.48	0.61	0.04	0.67
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	0.00	0.00	1.00	0.00	2.00	2.00	1.00	1.00	0.00	2.00
	0.11	0.14	0.19	0.29	0.20	0.20	0.14	0.17	0.06	0.21
	0.05	0.07	0.13	0.22	0.15	0.16	0.08	0.11	0.01	0.17
PRO $S1.0 - r$	0.29	0.41	0.85	0.71	0.80	0.84	0.77	0.94	0.24	0.96
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	0.00	1.00	2.00	2.00	3.00	3.00	2.00	2.00	1.00	3.00
	0.19	0.15	0.26	0.35	0.24	0.26	0.19	0.24	0.08	0.26
	0.13	0.09	0.22	0.30	0.24	0.26	0.14	0.22	0.03	0.29
PRO $S1.0 - r$ noinfo	0.41	0.26	0.11	0.12	0.23	0.18	0.03	0.08	0.35	0.23
	(0.11)	(0.11)	(0.10)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	2.00	1.00	0.00	0.00	1.00	1.00	0.00	0.00	0.00	1.00
	0.17	0.12	0.11	0.23	0.16	0.11	0.07	0.12	0.12	0.13
	0.12	0.06	0.05	0.17	0.09	0.05	0.02	0.05	0.06	0.06

Notes: Differences between predicted and realized distribution over the restricted number of searches when the search model is sequential search. In each cell, the first number is the prediction error in means; the second number (in brackets) is the standard error of the point estimate of the prediction error in means; the third number is the prediction error in medians; the fourth (fifth) number indicates the Hellinger distance (Kullback-Leibler divergence) between the predicted and realized distribution.

Table B16: Observed vs. Predicted Restricted Number of Searches (Non-Sequential Search)

$G^{[1A]}$ $\hat{G}^{[1A]}$	AMT $S 1.0$	AMT $S 1.0$ noinfo	AMT $S 1.0 - l$	AMT $S 1.0 - l$ noinfo	AMT $S 1.0 - r$	AMT $S 1.0 - r$ noinfo	PRO $S 1.0 - l$	PRO $S 1.0 - l$ noinfo	PRO $S 1.0 - r$	PRO $S 1.0 - r$ noinfo
AMT $S 1.0$	0.02	0.05	0.66	0.58	0.44	0.42	0.70	0.76	0.19	0.61
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	1.00	0.00	1.00	1.00	2.00	2.00	1.00	1.00	0.00	2.00
	0.24	0.19	0.28	0.39	0.19	0.22	0.20	0.19	0.14	0.17
AMT $S 1.0$ noinfo	0.22	0.14	0.24	0.38	0.12	0.15	0.16	0.14	0.08	0.11
	0.30	0.08	0.56	0.50	0.20	0.25	0.43	0.55	0.34	0.30
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	2.00	1.00	1.00	1.00	1.00	1.00	0.00	0.00	1.00	1.00
AMT $S 1.0 - l$	0.20	0.17	0.26	0.37	0.14	0.17	0.18	0.15	0.16	0.11
	0.17	0.12	0.21	0.32	0.07	0.09	0.14	0.08	0.10	0.05
	0.58	0.49	0.04	0.07	0.18	0.05	0.07	0.14	0.64	0.07
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
AMT $S 1.0 - l$ noinfo	2.00	1.00	0.00	0.00	0.00	0.00	0.00	0.00	2.00	0.00
	0.25	0.22	0.20	0.32	0.17	0.16	0.12	0.09	0.21	0.10
	0.29	0.22	0.15	0.30	0.11	0.09	0.06	0.03	0.18	0.04
	0.57	0.39	0.20	0.05	0.10	0.08	0.19	0.32	0.65	0.04
AMT $S 1.0 - l$ noinfo	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	2.00	1.00	0.00	0.00	0.00	1.00	0.00	0.00	2.00	0.00
	0.26	0.19	0.23	0.32	0.15	0.18	0.12	0.12	0.20	0.08
	0.30	0.15	0.18	0.30	0.09	0.11	0.06	0.05	0.17	0.03
AMT $S 1.0 - r$	0.62	0.52	0.06	0.16	0.18	0.13	0.10	0.12	0.84	0.07
	(0.11)	(0.11)	(0.10)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	2.00	1.00	0.00	0.00	0.00	0.00	0.00	0.00	2.00	0.00
	0.26	0.21	0.20	0.32	0.17	0.17	0.13	0.11	0.25	0.09
AMT $S 1.0 - r$ noinfo	0.30	0.20	0.16	0.33	0.11	0.10	0.08	0.05	0.27	0.03
	0.67	0.60	0.05	0.15	0.27	0.19	0.01	0.02	0.75	0.18
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	2.00	1.00	0.00	0.00	0.00	0.00	0.00	0.00	2.00	0.00
PRO $S 1.0 - l$	0.25	0.23	0.19	0.33	0.17	0.18	0.14	0.08	0.23	0.09
	0.30	0.25	0.14	0.34	0.11	0.12	0.09	0.03	0.22	0.04
	0.36	0.31	0.33	0.16	0.18	0.24	0.29	0.34	0.42	0.27
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
PRO $S 1.0 - l$ noinfo	2.00	1.00	0.00	0.00	1.00	1.00	0.00	0.00	1.00	1.00
	0.22	0.19	0.18	0.32	0.25	0.22	0.10	0.12	0.15	0.17
	0.21	0.14	0.11	0.28	0.21	0.18	0.04	0.05	0.10	0.11
	0.53	0.38	0.23	0.12	0.01	0.16	0.15	0.30	0.46	0.16
PRO $S 1.0 - l$ noinfo	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	2.00	1.00	0.00	0.00	1.00	1.00	0.00	0.00	1.00	1.00
	0.23	0.18	0.17	0.31	0.19	0.23	0.06	0.12	0.16	0.17
	0.24	0.14	0.10	0.25	0.13	0.18	0.02	0.05	0.11	0.11
PRO $S 1.0 - r$	0.05	0.06	0.77	0.70	0.63	0.65	0.71	0.82	0.06	0.71
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	1.00	0.00	2.00	2.00	2.00	2.00	2.00	2.00	0.00	2.00
	0.23	0.22	0.28	0.38	0.28	0.28	0.18	0.21	0.13	0.26
PRO $S 1.0 - r$ noinfo	0.20	0.17	0.25	0.36	0.29	0.28	0.12	0.18	0.07	0.26
	0.77	0.69	0.09	0.22	0.21	0.20	0.14	0.01	0.82	0.18
	(0.11)	(0.11)	(0.10)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	2.00	1.00	0.00	0.00	0.00	0.00	0.00	0.00	2.00	0.00
PRO $S 1.0 - r$ noinfo	0.25	0.24	0.17	0.28	0.22	0.18	0.10	0.07	0.24	0.19
	0.30	0.25	0.11	0.27	0.19	0.14	0.04	0.02	0.26	0.14

Notes: Differences between predicted and realized distribution over the restricted number of searches when the search model is non-sequential search. In each cell, the first number is the prediction error in means; the second number (in brackets) is the standard error of the point estimate of the prediction error in means; the third number is the prediction error in medians; the fourth (fifth) number indicates the Hellinger distance (Kullback-Leibler divergence) between the predicted and realized distribution.

Table B17: Observed vs. Predicted Restricted Number of Searches (Search w/o Priors)

$G^{[1A]}$ $\hat{G}^{[1A]}$	AMT $S1.0$	AMT $S1.0$ noinfo	AMT $S1.0 - l$	AMT $S1.0 - l$ noinfo	AMT $S1.0 - r$	AMT $S1.0 - r$ noinfo	PRO $S1.0 - l$	PRO $S1.0 - l$ noinfo	PRO $S1.0 - r$	PRO $S1.0 - r$ noinfo
AMT $S1.0$	0.04	0.13	0.21	0.10	0.64	0.63	0.13	0.33	0.10	0.75
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	0.00	0.00	0.00	0.00	2.00	2.00	0.00	0.00	0.00	2.00
	0.18	0.15	0.19	0.28	0.22	0.24	0.09	0.10	0.12	0.20
	0.12	0.08	0.12	0.20	0.16	0.19	0.04	0.04	0.05	0.15
AMT $S1.0$ noinfo	0.09	0.05	0.10	0.01	0.44	0.48	0.01	0.16	0.00	0.47
	(0.11)	(0.11)	(0.10)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	1.00	0.00	0.00	0.00	2.00	2.00	0.00	0.00	0.00	2.00
	0.16	0.13	0.14	0.28	0.18	0.22	0.10	0.08	0.12	0.13
	0.10	0.06	0.07	0.22	0.11	0.14	0.04	0.02	0.06	0.07
AMT $S1.0 - l$	0.20	0.07	0.07	0.05	0.43	0.50	0.01	0.15	0.16	0.43
	(0.11)	(0.11)	(0.10)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	1.00	0.00	0.00	0.00	2.00	2.00	0.00	0.00	0.00	2.00
	0.17	0.14	0.15	0.29	0.17	0.22	0.06	0.10	0.13	0.13
	0.11	0.08	0.08	0.24	0.10	0.14	0.02	0.03	0.06	0.06
AMT $S1.0 - l$ noinfo	0.08	0.01	0.21	0.07	0.55	0.65	0.12	0.27	0.01	0.54
	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	1.00	0.00	0.00	0.00	2.00	2.00	0.00	0.00	0.00	2.00
	0.17	0.15	0.18	0.29	0.20	0.22	0.08	0.10	0.11	0.15
	0.11	0.08	0.12	0.24	0.14	0.16	0.03	0.04	0.05	0.08
AMT $S1.0 - r$	0.41	0.31	0.09	0.35	0.01	0.13	0.32	0.17	0.52	0.21
	(0.11)	(0.11)	(0.10)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	2.00	1.00	0.00	0.00	0.00	0.50	0.00	0.00	1.00	1.00
	0.19	0.14	0.14	0.27	0.12	0.14	0.14	0.08	0.16	0.08
	0.15	0.08	0.08	0.27	0.05	0.06	0.08	0.03	0.11	0.03
AMT $S1.0 - r$ noinfo	0.49	0.38	0.17	0.32	0.07	0.09	0.23	0.13	0.49	0.24
	(0.11)	(0.11)	(0.10)	(0.11)	(0.11)	(0.11)	(0.14)	(0.13)	(0.12)	(0.12)
	2.00	1.00	0.00	0.00	1.00	1.00	0.00	0.00	0.00	2.00
	0.18	0.15	0.14	0.27	0.13	0.16	0.13	0.08	0.17	0.10
	0.14	0.10	0.08	0.25	0.06	0.09	0.07	0.03	0.12	0.04
PRO $S1.0 - l$	0.17	0.40	0.26	0.16	0.95	0.94	0.21	0.38	0.32	0.92
	(0.11)	(0.11)	(0.10)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	0.00	1.00	0.00	0.00	3.00	3.00	0.00	0.00	1.00	2.00
	0.17	0.20	0.18	0.28	0.30	0.30	0.09	0.15	0.12	0.28
	0.10	0.16	0.11	0.21	0.37	0.36	0.03	0.08	0.06	0.33
PRO $S1.0 - l$ noinfo	0.02	0.12	0.10	0.00	0.68	0.79	0.03	0.16	0.11	0.71
	(0.11)	(0.11)	(0.10)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	0.00	1.00	0.00	0.00	2.00	2.00	0.00	0.00	0.00	2.00
	0.19	0.15	0.17	0.26	0.25	0.28	0.07	0.13	0.07	0.23
	0.13	0.08	0.10	0.20	0.25	0.31	0.02	0.06	0.02	0.22
PRO $S1.0 - r$	0.09	0.18	0.20	0.02	0.81	0.86	0.06	0.31	0.22	0.92
	(0.11)	(0.11)	(0.10)	(0.11)	(0.11)	(0.10)	(0.14)	(0.13)	(0.12)	(0.12)
	0.00	1.00	0.00	0.00	2.00	2.00	0.00	0.00	0.00	2.00
	0.18	0.17	0.16	0.28	0.28	0.28	0.09	0.13	0.09	0.27
	0.12	0.10	0.09	0.23	0.32	0.31	0.03	0.06	0.03	0.32
PRO $S1.0 - r$ noinfo	0.66	0.56	0.47	0.63	0.03	0.08	0.53	0.44	0.57	0.07
	(0.11)	(0.11)	(0.10)	(0.11)	(0.11)	(0.10)	(0.13)	(0.12)	(0.12)	(0.12)
	2.00	1.00	0.00	0.00	1.00	1.00	0.00	0.00	1.00	1.00
	0.23	0.21	0.19	0.28	0.19	0.15	0.18	0.17	0.18	0.13
	0.23	0.18	0.17	0.36	0.12	0.08	0.15	0.14	0.13	0.08

Notes: Differences between predicted and realized distribution over the restricted number of searches when the search model is search without priors. In each cell, the first number is the prediction error in means; the second number (in brackets) is the standard error of the point estimate of the prediction error in means; the third number is the prediction error in medians; the fourth (fifth) number indicates the Hellinger distance (Kullback-Leibler divergence) between the predicted and realized distribution.

B.5 Simulation Study

To complement the analysis of our experimental data we investigate the predictive performance of the classic search models in an environment in which the true data-generating process is fully known. To this end, we simulate synthetic search data under controlled conditions. Let α be the share of subjects who search non-sequentially and $1 - \alpha$ the share of subjects who search sequentially. We vary the fraction of subjects α from 0 to 100 percent in steps of 20 percentage points. This yields six specifications, specification 1 (with $\alpha = 0.00$) to specification 6 (with $\alpha = 1.00$). For each specification, we generate a dataset of search outcomes (sequence of observed prices, number of searches, and purchase prices). Then we apply the estimation and evaluation procedure from Section 4 to these datasets.

In all specifications, we consider the following search environment. The search cost parameter k is drawn for 10,000 simulated subjects from a lognormal distribution with (location, scale) parameters $(-2, 3)$, truncated at 3 as in our main analysis, implying an average search cost parameter of $\mathbb{E}(k) = 0.216$. Prices are drawn from the uniform distribution on the interval $[3.00, 6.00]$, which coincides with the *S* 1.0 treatment in our experiment. The resulting search cost distribution is close to the one estimated in our experimental sample, so that the simulation exercise speaks to an empirically relevant region of the parameter space.

Our objective is to evaluate how well the sequential and non-sequential search model perform when the assumed search protocol is (partly) incorrect. The sequential search model is correct only for the dataset from specification 1 and the non-sequential search model is correct only for the dataset from specification 6. We treat the simulated subjects' search outcomes exactly as we treat those of experimental subjects: We estimate search costs under both the sequential and the non-sequential search model using the ordered-probit regression framework of Section 4. This yields, for each specification and each model, a fully specified estimated distribution over the search cost parameter k . We then follow the evaluation procedure outlined in Section 4: We insert the estimated search cost distribution back into each model to simulate predicted distributions over the number of searches and purchase prices, and compare these predictions to the true (simulated) outcomes.

We summarize our results in Table B18, Figure B1, and Figure B2 below. Table B18 shows the search cost estimates of both models for each specification. Figure B1 displays, in the left graph, the realized mean number of searches for each specification as well as the predicted mean according to the two models. The right graph shows the corresponding relative prediction errors. Figure B2 presents the same information for purchase prices. In the following, we first discuss the prediction performance of the classic search models. Then we examine why their prediction performance differs between number of searches and purchase prices.

Table B18: Simulated Data – Search Cost Estimates

Specification	(1)	(2)	(3)	(4)	(5)	(6)
Share α	0.00	0.20	0.40	0.60	0.80	1.00
<i>Sequential search</i>						
β	-2.08*** (0.06)	-1.85*** (0.08)	-1.64*** (0.10)	-1.42*** (0.11)	-1.24*** (0.13)	-1.07*** (0.14)
σ	2.83*** (0.04)	3.15*** (0.05)	3.39*** (0.06)	3.67*** (0.06)	3.87*** (0.07)	4.05*** (0.07)
Log.Lik.	-13493.84	-15240.55	-16840.09	-18493.78	-20031.13	-21563.69
SC	0.215 (0.324)	0.220 (0.332)	0.222 (0.336)	0.223 (0.339)	0.222 (0.341)	0.222 (0.342)
<i>Non-sequential search</i>						
β	-1.03*** (0.09)	-1.20*** (0.09)	-1.36*** (0.08)	-1.58*** (0.08)	-1.79*** (0.07)	-1.98*** (0.07)
σ	2.98*** (0.05)	3.00*** (0.05)	3.02*** (0.05)	3.01*** (0.05)	3.00*** (0.04)	2.98*** (0.04)
Log.Lik.	-26325.73	-27118.83	-27810.48	-28610.85	-29323.68	-29957.84
SC	0.276 (0.361)	0.264 (0.356)	0.253 (0.350)	0.240 (0.342)	0.228 (0.335)	0.217 (0.328)

Notes: Ordered probit regressions with truncation for six simulated specifications, using both sequential and non-sequential search. Standard errors of the estimated parameters are in parentheses. We also present the search cost estimates (SC) for each specification, derived from the parameter estimates shown above. The first value represents the mean of the implied search cost distribution, with the standard deviation of the distribution provided in parentheses. Significance at * $p < 0.1$, ** $p < 0.05$, and *** $p < 0.01$.

Sequential Search Model. We first consider the search cost estimates in the upper panel of Table B18. In specification 1 (where the true and assumed search protocol coincide), the estimated mean search cost is 0.215 and virtually identical to the true value of 0.216. As α increases and more subjects search non-sequentially, the estimated search cost adjust in order to rationalize the observed data under the maintained sequential search assumption. Since sequential search is the more efficient protocol, reconciling the less efficient outcomes generated by a population that partly searches non-sequentially requires attributing higher search costs to the subjects. Therefore, the search cost estimates rise and stabilize around 0.222.

The resulting predictions are shown by the red elements in Figure B1 and Figure B2. The predicted mean number of searches tracks the true value closely at the endpoints, but exhibits notable deviations for intermediate values of α , with a relative prediction error that peaks at around 13 percent. The predicted mean purchase price stays essentially flat near 3.90 across all specifications, reflecting the fact that, under sequential search, purchases occur at prices at or below the reservation price, largely independent of α . The relative prediction error therefore grows monotonically in α , but remains modest.

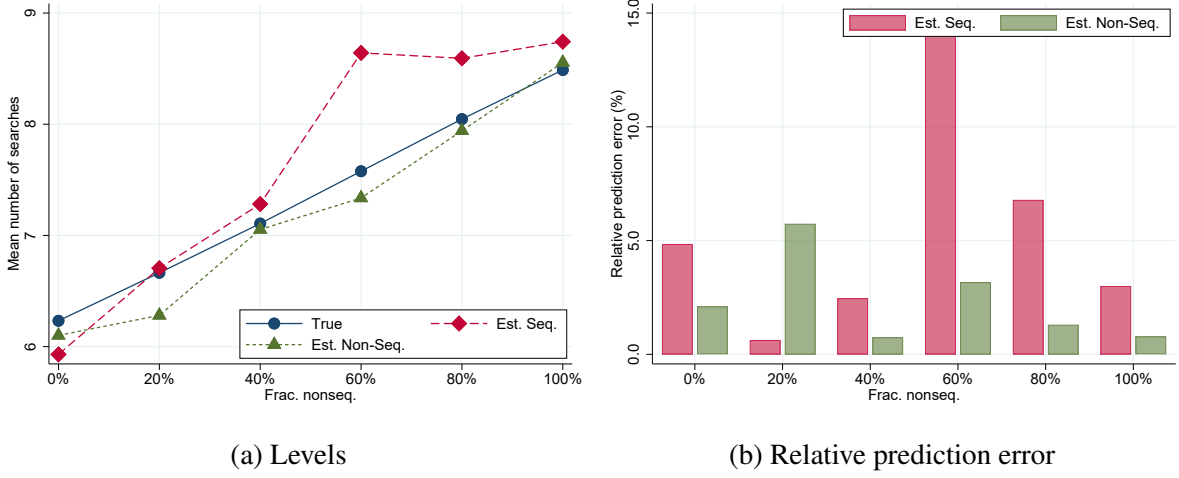


Figure B1: Simulated Data - Number of Searches

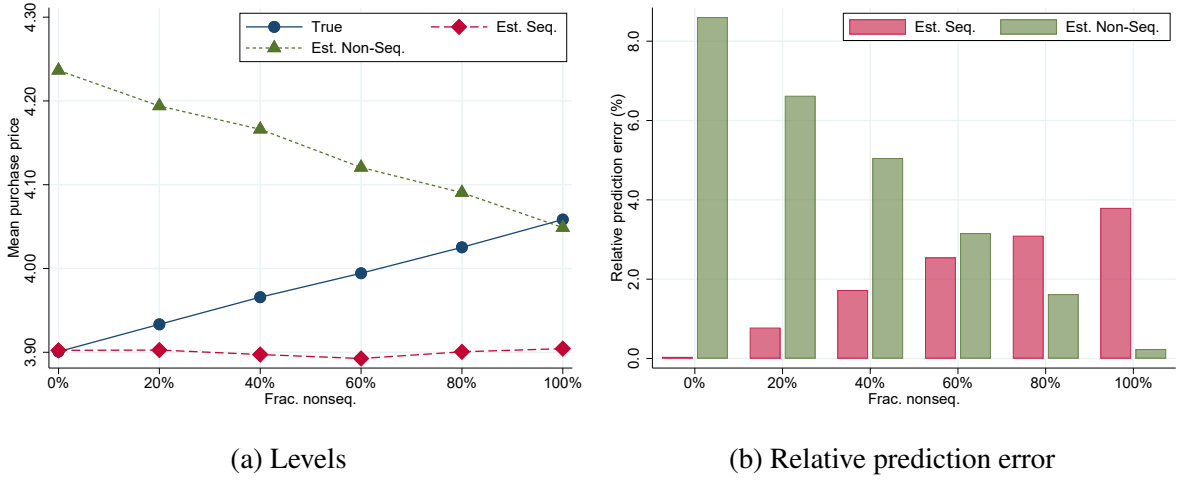


Figure B2: Simulated Data - Purchase Prices

Non-sequential Search Model. We first consider the search cost estimates from the non-sequential search model, shown in the lower panel of Table B18. In specification 6 (where the true and assumed search protocol are identical), the estimated mean search costs equal 0.217. Hence, the model recovers the true mean search costs accurately. As α decreases and a growing share of subjects searches sequentially, the estimated search costs adjust to rationalize outcomes generated by a more efficient protocol. Mirroring the logic above, the estimated search costs rise, from 0.217 in specification 6 to 0.276 in specification 1.

The predictions are displayed by the green elements in Figure B1 and Figure B2. The predicted mean number of searches tracks the true value closely across all specifications, with relative prediction errors below 7 percent throughout. The predicted mean purchase price, by contrast, systematically overshoots the true value for small values of α : It starts at roughly 4.25 in specification 1 and declines monotonically to 4.05 in specification 6. Hence, the relative

prediction error peaks at around 8 percent in specification 1 and vanishes as α approaches one.

Taken together, the simulation results show that both classic search models provide reasonable as-if descriptions of the data, even when their stated assumptions are violated. Across all six specifications, the estimated mean search costs remain in a narrow band around the true value of 0.216, and the relative prediction errors for both outcomes stay in single-digit to low double-digit ranges. Crucially, neither model's predictive performance collapses at the extremes of α , when the maintained protocol assumption is entirely counterfactual.

The two models display complementary rather than strictly ordered performance. The sequential search model delivers very tight predictions on mean purchase prices (the relative prediction errors stay below 4 percent throughout), but incurs larger prediction errors on the mean number of searches, peaking at roughly 13 percent at intermediate values of α . The non-sequential search model exhibits the mirror pattern: smaller prediction errors on the mean number of searches (below 7 percent throughout), but larger errors on mean purchase prices (up to roughly 8 percent). Neither model dominates the other overall; both remain close enough to the truth to be informative, consistent with the conclusion that sequential and non-sequential estimates perform similarly well, despite the stark differences in their behavioral assumptions.

Differences between Search Outcome Dimensions. Finally, we consider the differences in relative prediction errors between the two outcome dimensions, number of searches and purchase prices. Across both models and most specifications, the relative prediction errors for the mean number of searches exceed those for mean purchase prices. The simulations reveal three complementary reasons that rationalize this pattern.

First, outcomes vary substantially more for search intensity than for purchase prices: The true mean number of searches ranges from roughly 6.0 searches when all subjects search sequentially to around 8.5 searches when all subjects search non-sequentially, an increase of 41.7 percent. In contrast, the true mean purchase price only raises from 3.90 USD to 4.05 USD, an increase of 3.8 percent. There is simply more room for a misspecified model to generate errors in the former dimension than in the latter.

Second, the support of the price distribution mechanically limits the scope for prediction errors on purchase prices. With a uniform distribution on the interval [3.00, 6.00], the minimum of a sample of moderate size clusters near the lower bound, so that average purchase prices under either protocol occupy a narrow band. The number of searches, by contrast, can in principle range from 0 to 100, leaving both the truth and the estimated predictions considerably more room to diverge.

Third, and most importantly from an economic standpoint, the protocol assumption directly governs the stopping rule and hence the distribution over the number of searches. In contrast, purchase prices are only a derivative of the search process. For a given number of

searches and realized price draws, both protocols reduce to taking a low order statistic of a uniform sample, which renders the implied price distributions qualitatively similar. A misspecification of the search protocol therefore leaves its primary footprint on search intensity and only a muted one on purchase prices.

This observation carries a practical implication for applied work. Because purchase price moments are relatively insensitive to the maintained protocol, judging a search cost model by how well it reproduces prices alone can overstate its accuracy. Search intensity moments, being much more responsive to protocol misspecification, offer a sharper test and should carry greater weight when comparing competing search models.