

Heterogeneous Costs and the Decision to Learn^{*}

David Walker-Jones[†]

July 31, 2026

Abstract

Analysing demand curves provides some of the most foundational tools in economics. If information is difficult to obtain, however, such analysis requires understanding how the information obtained changes with price. Subjects are shown to be surprisingly flexible when choosing what information to acquire, achieving accuracies that differ across states of the world, and spending longer learning when the realized state is unexpected based on their prior. A novel data enrichment shows that many experiment subjects violate even the most flexible model of costly learning if costs are assumed to be constant, providing an important caveat for how learning is modelled. If heterogeneous costs are permitted, the resultant model correctly predicts how the probability of subjects learning changes with price.

Keywords: Rational Inattention, Costly Learning, Heterogeneous Information Costs

^{*}I would like to thank Andrew Caplin, Mark Dean, Rahul Deb, Tommaso Denti, Henrique de Oliveira, Yoram Halevy, Johannes Hoelzemann, Daniel Martin, and Colin Stewart, for their helpful advice and support. Funding was provided by the University of Toronto. The experiment was approved by the U of T Social Sciences, Humanities & Education REB (protocol number 39850). This paper was previously called “Optimal Choice Mistakes.” For the purpose of open access, the author has applied a Creative Commons attribution license (CC BY) to any Author Accepted Manuscript version arising from this submission.

[†]University of Surrey, d.walker-jones@surrey.ac.uk, ORCID: 0000-0001-6229-9958.

1 Introduction

In our modern world decision makers (DMs) typically have access to more information than it is practical or efficient for them to process. The resultant scarcity of attention motivates the study of rational inattention (Sims, 2003), a model in which DMs weigh the benefits of better choices against the opportunity costs of acquiring more information. Understanding how the information a DM obtains changes with the price of a good is thus an important pre-requisite for conducting the welfare and counterfactual analysis that are the common goal of economic modelling.

Downward sloping demand, for instance, is a core assumption in economics and the foundation of standard tools in welfare and counterfactual analysis. Costly learning, however, can actually generate locally increasing demand for some versions of a good.¹ If a DM, that wants to purchase a good iff it is of high quality, does more learning when the price of the good increases, then they might be more aware of the good's quality and there could be a higher probability of them purchasing the high quality good after the price increase. To a researcher, this could appear to be locally upward sloping demand for high quality goods.

This paper introduces an experiment to explore how the information acquisition of subjects changes with parameters like price. The experimental design is novel as it allows me to observe if a subject acquires information in a choice problem, or simply makes a decision based on their prior. This data enrichment shows that most subjects fluctuate back and forth between learning and not learning in a way that violates even the most flexible model of costly learning if it is assumed that the cost function for information is constant across choice problems in which subjects have the same prior belief. Fatigue is unsurprisingly shown to be an issue, but even allowing for fatigue cannot rationalize the data of many subjects as choice problems that are novel, in the sense that their parameters differ from previous choice problems, seem to increase subjects' willingness to invest effort into learning. This demonstrates the significance

¹See Example 1 in Section 2.2.

of heterogeneity when studying costly learning.

The experiment is arguably a setting where a simple “ q accuracy” model seems especially appropriate. The predictions for demand from a representative agent version of such a model are characterized, and it is shown that such a model can rationalize upward sloping demand for one of the types of good. Aggregate demand in the experiment, however, is downward sloping for both types of good in large part because of the apparent heterogeneity of costs, and aggregate demand violates the representative agent characterization of the “ q accuracy” model. Subjects often make decisions based on their prior and, reassuringly, changes in the probability of them acquiring information are correctly predicted by a model that allows for heterogeneous costs, but not by any representative agent model considered.

I also do see evidence that subjects are surprisingly flexible when choosing what information to acquire. When subjects do choose to acquire information they are, in contrast with the assumptions of the “ q accuracy” model, able to achieve accuracies that differ across the states of the world, and spend longer learning when the realized state is unexpected based on their prior. It seems they know they could be making mistakes and invest more effort into learning when their initial learning does not lead to the conclusion they anticipated, a degree of flexibility that is assumed by standard examples of Posterior Separable cost functions (Caplin & Dean, 2013) like Shannon Entropy (Shannon, 1948). The Posterior Separable model does, in fact, do well when predicting aggregate changes in demand in the data.

2 Model

Consider a DM choosing between two options, option X and option Y. Option Y is referred to as the **uncertain option** because there is some uncertainty about the payoff the DM gets from selecting it: The DM’s payoff from option Y, $u(\omega) \in \mathbb{R}$, is determined by the **state** $\omega \in \Omega = \{\underline{\omega}, \bar{\omega}\}$, with $u(\underline{\omega}) < u(\bar{\omega})$. I focus on a binary

state space because this is what is relevant for the experiment and eases exposition. The DM knows the distribution $\mu \in \Delta(\Omega)$ over states, which is referred to as their **prior belief**.² Option X is referred to as the **safe option** because the DM knows they receive a payoff of $p \in \mathbb{R}$ from selecting it, and p is referred to as the **price** of selecting option Y because the DM must give up the opportunity to have a payoff of p to select option Y.

Given μ , the **behavior** at a finite set of prices $\mathcal{P} \subset \mathbb{R}$ (a finite set of values for option X), is a mapping $s : \mathcal{P} \rightarrow S \equiv [0, 1] \times [0, 1]$. Namely, at each price $p \in \mathcal{P}$ the researcher observes the **information outcome** $s(p) \equiv (\underline{s}(p), \bar{s}(p)) \equiv (\Pr(Y|\underline{\omega}, p), \Pr(Y|\bar{\omega}, p))$, that describes the probability of the DM selecting option Y at each ω , denoted $\Pr(Y|\omega, p) \equiv 1 - \Pr(X|\omega, p)$. The DM is said to **learn** at a price p if $\Pr(Y|\underline{\omega}, p) \neq \Pr(Y|\bar{\omega}, p)$ and $\mu(\bar{\omega}) \in (0, 1)$. So, the DM is learning whenever their probability of selecting option Y changes with the realization of ω , because this is indicative of the DM at least partially differentiating between $\underline{\omega}$ and $\bar{\omega}$.

2.1 Costly Learning

Given a belief μ , a cost function $C_\mu : S \rightarrow \mathbb{R}_+$ describes the cost of the learning required for achieving each information outcome. The difficulty with demand analysis in settings with costly learning is the information acquired might change with price, so the realized cost of the learning $C_\mu(s(p))$ might change when p changes, but it is assumed that the function C_μ does not change if p changes. The cost function has a subscript of μ because the costs of different information outcomes are allowed change with the prior.

As is discussed by Walker-Jones (2026a), it is natural to think that outcomes that can be achieved without learning should not have any learning costs associated with them, and randomizing over different information outcomes should not cause additional learning costs, which leads to the following two assumption. **Assumption**

²The belief μ assigns a strictly positive probability to each state unless stated otherwise.

1: $\forall x \in [0, 1]$, $C_\mu(x, x) = 0$, and **Assumption 2:** C_μ is a weakly convex function.³

Assumption 1 and Assumption 2 represent the minimal structure that is assumed of all cost functions for information in this paper. Given the belief of the DM, μ , the **cost function** for information outcomes is defined as a mapping from S onto the positive reals, $C_\mu : S \rightarrow \mathbb{R}_+$, which satisfies Assumption 1 and Assumption 2.

Definition: Given μ and a set of prices $\mathcal{P}$, the behavior is **rationalized by a costly learning model** if there are payoffs for option Y, $u(\underline{\omega})$ and $u(\bar{\omega})$ with $u(\underline{\omega}) < u(\bar{\omega})$, and a cost function for information outcomes C_μ such that, $\forall p \in \mathcal{P}$,

$$s(p) = (\underline{s}(p), \bar{s}(p)) \in \arg \max_{(\underline{s}, \bar{s}) \in S} \left(p + \underline{s}\mu(\underline{\omega})(u(\underline{\omega}) - p) + \bar{s}\mu(\bar{\omega})(u(\bar{\omega}) - p) - C_\mu(\underline{s}, \bar{s}) \right).$$

Given p and μ , the expected **gain in payoff** the DM receives when they choose an information outcome $s = (\underline{s}, \bar{s})$ instead of picking option X with certainty is defined by:

$$\underline{s}\mu(\underline{\omega})(u(\underline{\omega}) - p) + \bar{s}\mu(\bar{\omega})(u(\bar{\omega}) - p) - C_\mu(\underline{s}, \bar{s}).$$

Maximizing the gain in payoff eases exposition and is equivalent to maximizing the expected payoff as the difference between them is a simple re-normalization of option payoffs. When p increases the gain in payoff from selecting option Y decreases. As a result, when p increases, the gain in payoff from an information outcome $s = (\underline{s}, \bar{s}) \in S$ decreases in proportion to the **unconditional probability** of selecting option Y that the information outcome results in, which is $\underline{s}\mu(\underline{\omega}) + \bar{s}\mu(\bar{\omega})$. So, the gain in payoff from information outcomes that create lower unconditional probabilities of selecting option Y decrease less quickly when p increases. An obvious requirement for behavior being rationalized by a costly learning model is thus that increases in the price cannot result in a higher average chance of option Y being selected. This notion is formalized

³Given Assumption 1, Lemma 1 from the work of Cheng and Kim (2025) implies that Assumption 2 imposes that more information in a Blackwell (1951, 1953) sense is weakly more costly.

by the following lemma, and is also part of Theorem 1 from the work of Walker-Jones (2026a).

Lemma 0. Option Y is weakly less likely to be selected when its price increases:

$\Pr(Y|p) \equiv \mu(\bar{\omega})\Pr(Y|\bar{\omega}, p) + \mu(\underline{\omega})\Pr(Y|\underline{\omega}, p)$ is weakly decreasing in p .

Proof. Suppose the result does not hold: so $p_1 < p_2$ and $\Pr(Y|p_1) < \Pr(Y|p_2)$. But, this leads to a contradiction since if $s(p_1)$ is optimal at p_1 then it provides a weakly higher payoff at p_1 than $s(p_2)$ does, and, when price increases to p_2 the decrease in payoff from $s(p_1)$ is strictly smaller than the decrease in payoff from $s(p_2)$: $(p_2 - p_1)\Pr(Y|p_1) < (p_2 - p_1)\Pr(Y|p_2)$. ■

If I further assume that the researcher knows the **expected payoff** from option Y, which is $\mathbb{E}[u(\omega)|\mu] \equiv \mu(\underline{\omega})u(\underline{\omega}) + \mu(\bar{\omega})u(\bar{\omega})$, then predictions can be made for if a single price change impacts whether or not the DM will learn, as is formalized in Corollary 1. Many results in the rational inattention literature, for instance in the work of Caplin, Dean, and Leahy (2022), assume that the utility function of the DM is known, and thus assuming that the expected payoff of Y is known is a relatively weak assumption, and in the experiment presented in this paper the expected payoff from option Y in terms of probability points is known.

Corollary 1. Given a prior belief μ and an expected payoff from option Y, suppose the behavior is rationalized by a costly learning model with payoffs for option Y that are consistent with the given expected payoff and that there are two prices $p_1, p_2 \in \mathcal{P}$ with $p_1 < p_2$.

Then: **(i)** if $p_1 < p_2 \leq \mathbb{E}[u(\omega)|\mu]$, then the DM learns at p_2 if they learn at p_1 : $\Pr(Y|\underline{\omega}, p_1) \neq \Pr(Y|\bar{\omega}, p_1) \Rightarrow \Pr(Y|\underline{\omega}, p_2) \neq \Pr(Y|\bar{\omega}, p_2)$,
and **(ii)** if $\mathbb{E}[u(\omega)|\mu] \leq p_1 < p_2$, then the DM learns at p_1 if they learn at p_2 : $\Pr(Y|\underline{\omega}, p_2) \neq \Pr(Y|\bar{\omega}, p_2) \Rightarrow \Pr(Y|\underline{\omega}, p_1) \neq \Pr(Y|\bar{\omega}, p_1)$.

Proof. Suppose the DM learns at p_1 . If $p_1 < p_2 \leq \mathbb{E}[u(\omega)|\mu]$, then if the DM does no learning at p_2 then selecting option Y is an optimal choice, but this contradicts

Lemma 0, so the DM must learn at p_2 .

Suppose instead the DM learns at p_2 . If $\mathbb{E}[u(\omega)|\mu] \leq p_1 < p_2$, then if the DM does no learning at p_1 then selecting option X is an optimal choice, but this contradicts Lemma 0, so the DM must learn at p_1 . ■

If the researcher can observe if the DM is acquiring information at different prices, which is a binary variable and is observable in the experiment, Corollary 1 provide ways of testing a very general model of costly learning without needing to estimate choice probabilities, a variable that is continuous and inherently difficult to estimate. Further, even if the researcher does not know the expected payoff from option Y, it is not hard to show that learning at two prices implies learning at any price between them.

2.2 The q Accuracy Model

If more structure is desired then a natural starting point is to consider a model in which learning results in a symmetric accuracy, a probability q of selecting the best option that is the same in each state of the world. In environments where the DM does not have much flexibility deciding what information to acquire it may be compelling to make such an assumption about the DM when they learn. Remember though, the DM is not required to learn at a price since if, for instance, $p > u(\bar{\omega})$ then the DM has no incentive to learn, and Assumption 1 ensures the DM has the option to not learn, incur no cost of learning, and make a decision based on their prior. An immediate consequence of such a model, however, is that if the DM chooses different “accuracies” at different prices, then there is a state of the world in which their demand for option Y is increasing in price, as is shown in the following example.

Example 1. Consider a DM with information outcomes such that at each of two prices p_1 and p_2 with $p_1 < p_2$ there is $q(p) \in (\frac{1}{2}, 1]$ such that $s(p) = (1 - q(p), q(p))$. If $q(p_1) < q(p_2)$ then the demand for option Y goes up with p if the state of the world

is $\bar{\omega}$, while if $q(p_1) > q(p_2)$ the demand for option Y goes up with p if the state of the world is $\underline{\omega}$.

Definition: Given a prior belief μ , the cost function C_μ is a q **accuracy cost function** for information outcomes if $\forall s \in S$ with $\underline{s} \leq \bar{s}$, $\exists q \in (\frac{1}{2}, 1]$, $\exists x \in [0, 1]$, and $\exists \alpha \in [0, 1]$ such that: $\underline{s} = \alpha(1 - q) + (1 - \alpha)x$, $\bar{s} = \alpha q + (1 - \alpha)x$, and $C_\mu(s) = \alpha C_\mu(1 - q, q)$.

In other words, C_μ is a q accuracy cost function if the cost of each information outcome is the mixture of the costs of an information outcome where the DM does not learn (so $\underline{s} = \bar{s}$) and an information outcome with symmetric accuracies ($1 - \underline{s} = \bar{s} > \frac{1}{2}$).

Definition: Given a prior belief μ and a set of prices $\mathcal{P}$, the behavior is **rationalized by a q accuracy model** if it is rationalized by a costly learning model with a q accuracy cost function for information outcomes C_μ .

If the DM learns according to a q accuracy model and there is a price where the DM learns but $\Pr(Y|\bar{\omega}, p) \neq \Pr(X|\underline{\omega}, p)$, one can infer an optimal information outcome for the DM where $1 - \underline{s} = \bar{s}$ (the candidate points are on the yellow lines in Figure 3). This is shown formally by Lemma 1 in Appendix 1. If the behavior of the DM is rationalized by a q accuracy model, then at each $p \in \mathcal{P}$ where the DM learns define their **implicit learning accuracy**, denoted $q(p)$, to be the q that is optimal according to Lemma 1. This implicit learning accuracy helps us describe the behavior that characterizes q accuracy models in the following proposition.

Proposition 1. Given a prior belief μ and a finite and non-empty set of prices $\mathcal{P}$, the behavior is rationalized by a q accuracy model if and only if it is rationalized by a costly learning model and the following three properties are satisfied: **(i)** There is at most one price at which the DM randomizes over learning and selecting option Y without learning. There is at most one p such that $\Pr(Y|\bar{\omega}, p) > \Pr(X|\underline{\omega}, p) > 0$.

(ii) There is at most one price at which the DM randomizes over learning and selecting option X without learning: There is at most one p such that $\Pr(X|\underline{\omega}, p) > \Pr(Y|\underline{\omega}, p) > 0$.

(iii) How the DM's implicit accuracy changes is determined by the prior: If $\mu(\bar{\omega}) > \mu(\underline{\omega})$, then $q(p)$ is weakly decreasing in p wherever it is defined, while if $\mu(\bar{\omega}) < \mu(\underline{\omega})$, then $q(p)$ is weakly increasing in p wherever it is defined.

Proof. See Appendix 1.

For a graphical depiction of the structure imposed by Proposition 1, see Figure 3 and the discussion in Section 3.1.

2.3 Heterogeneous Learning Costs

Datasets and demand schedules used for analysis are not typically generated by a single individual. Typical datasets aggregate behavior over decision makers with different abilities, expertise, and familiarities with the choice environment, and thus potentially heterogeneous costs for information outcomes. Even when a dataset is generated by a single individual at different points in time, that individual might face different costs for acquiring the same pieces of information at different times. The individual may fatigue and face higher learning costs in later decisions, or gain experience and face lower costs in subsequent decisions. There could also be variation in the choice environment that is not observable to the researcher that changes the costs of information such as variation in the how many members of a store's staff are available for answering questions at a given point in time. It is thus natural to consider a model of costly learning that features a distribution over cost functions for information and study the impact that this has on the predictions that can be made by the resultant model.

Definition: Given μ and a set of prices $\mathcal{P}$, the behavior is **rationalized by a costly learning model with heterogeneous costs** if there are $T \in \mathbb{N}$ types of DM each

of which has probability of occurring $\pi_t > 0$, and behavior s_t , which is $\Pr_t(Y|\omega, p)$ for each state and price, such that each type's behavior is rationalized by a costly learning model with payoffs for option Y of $u(\underline{\omega})$ and $u(\bar{\omega})$ with $u(\underline{\omega}) < u(\bar{\omega})$ and belief μ , the probabilities of the different types sum to one, and given any pair of price p and state ω , the mean behavior is the behavior observed by the researcher:

$$\forall p \in \mathcal{P}, \forall \omega \in \Omega : \sum_{t=1}^T \Pr_t(Y|\omega, p) \pi_t = \Pr(Y|\omega, p).$$

Such a model is more difficult to make predictions with, but leads to the natural generalization of Corollary 1 below. Given a belief μ and a costly learning model with heterogeneous costs that rationalizes the data, use the indicator function χ to define the probability of a given DM i learning at a price p :

$$\Pr(\underline{s}_i(p) \neq \bar{s}_i(p)) \equiv \sum_t \pi_t \chi(\underline{s}_t(p) \neq \bar{s}_t(p)).$$

Corollary 2. Given a belief μ and expected payoff from option Y, suppose the behavior is rationalized by a costly learning model with heterogeneous costs and payoffs for option Y that are consistent with the given expected payoff, and that there are two prices $p_1, p_2 \in \mathcal{P}$ with $p_1 < p_2$.

Then **(i)** if the prices are below the expected payoff from option Y, i.e. $p_1 < p_2 \leq \mathbb{E}[u(\omega)|\mu]$, then learning is more likely at p_2 than at p_1 : $\Pr(\underline{s}_i(p_1) \neq \bar{s}_i(p_1)) \leq \Pr(\underline{s}_i(p_2) \neq \bar{s}_i(p_2))$,

and **(ii)** if the prices are above the expected payoff from option Y, i.e. $\mathbb{E}[u(\omega)|\mu] \leq p_1 < p_2$, then learning is more likely at p_1 than at p_2 : $\Pr(\underline{s}_i(p_2) \neq \bar{s}_i(p_2)) \leq \Pr(\underline{s}_i(p_1) \neq \bar{s}_i(p_1))$.

Proof. Trivial given Corollary 1. ■

3 The Experiment

In the experiment, as is explained below, whether or not a subject chooses to learn in a decision problem is observed, which is not typical in the literature on rational inattention. This data enrichment is useful as the choice to acquire information results in a binary outcome, and using it for analysis with Corollary 1 does not require the estimation of a probability.

In the experiment, which was completed by 243 undergraduate students from the University of Toronto, each subject faced 92 rounds of **decision problems** in which they chose between a safe option X and an uncertain option Y, and, if they desired, they could choose to “get more information” (GMI) about the payoff from option Y before selecting an option.

In each of these decision problems the payoff from option X is always a $p \in \{0.25, 0.5\}$, which is clearly displayed to the subject. In the first 40 and last 40 decision problems $p = 0.25$, and in the middle 12 decision problems (decision problems 41 through 52) $p = 0.5$.⁴ The payoff of option Y is a $u(\omega) \in \{0, 1\}$, which is not displayed.

The payoffs of the options are measured in terms of percentage points. Subjects pick options to try to gradually increase their chance of winning a monetary prize at the end of the experiment.⁵ For example, if a subject always selects option Y, and in 50 decision rounds it has a payoff of 1, and the rest of the time it has a payoff of 0, then they have a 50% chance of winning the prize when the experiment ends. The prize the subject can win in the draw is either \$20, \$25, or \$30, depending on the

⁴As is typical in experiments on rational inattention, it is assumed that subjects know the value of the safe option. To aid the appropriateness of this assumption, there are only two points within the 92 decision problems at which the value of the safe option changes. Before each block of rounds with the same payoff for option X subjects are made aware of the payoff and the fact that it would stay the same for a certain number of rounds (see Figures 18 and 19 in the Online Appendix). Further, as can be seen in Figure 1, the payoff from option X is displayed in two places in each round.

⁵In this experiment subjects are paid with percentage points in a decision problem so that within subject analysis can be conducted in an incentive compatible way. See the Online Appendix for a discussion of incentive compatibility.

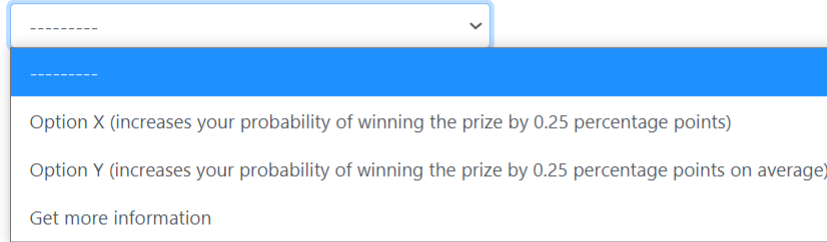
subject's treatment group.

Figure 1: Option X, Option Y, or Get more information (GMI):
Please respond to the following investment decision:

Round number: 9 of 100

Option X increases your probability of winning the \$30 prize by 0.25 percentage points

Please choose an option:



Option X (increases your probability of winning the prize by 0.25 percentage points)

Option Y (increases your probability of winning the prize by 0.25 percentage points on average)

Get more information

In each decision problem the subject is told the expected payoff from selecting option Y if no costly information is acquired (see Figure 1), which induces their prior belief. To help make changes in the prior more salient,⁶ in each decision problem

⁶It is standard in the experimental literature on rational inattention to assume that the subject perfectly learns the prior distribution. This assumption is, of course, at odds with the motivation behind the rational inattention literature, and so this experiment tries to make changes in the prior belief especially noticeable. I use big colorful dots to convey the prior so that the prior is as salient of a feature of the choice environment as possible. Further, as Figure 1 shows, in each round the numeric information about the prior contained in the colorful big dot, and its implication for the expected payoff of option Y, is repeated next to the choice options in the drop-down menu that subjects use to select alternatives, and only two relatively simple priors are used throughout the experiment. In addition, there are several quiz questions that all subjects face about the big dots (see, in particular, Figures 15 and 16, but also Figures 8 and 11, in the Online Appendix for examples), and thus subjects that did well on the quiz seem to be able to convert the big dots into a prior belief in the intended way. The robustness of data results to quiz performance is addressed in Table 3. Subjects are also reminded in every round, above the button that takes them to the next screen, about the information contained in the big dots (see Figure 17 in the Online Appendix).

the subject is presented with a big dot that is either **red** or **green**, as is depicted in Figure 1. If the big dot is **red** there is a $\frac{1}{4}$ chance that option Y is the better option ($\mu(\bar{\omega}) = \frac{1}{4}$), and if the big dot is **green** there is a $\frac{3}{4}$ chance that option Y is the better option ($\mu(\bar{\omega}) = \frac{3}{4}$). In each decision problem there is a $\frac{3}{4}$ chance that the big dot is **red**, and a $\frac{1}{4}$ chance that the big dot is **green**.

If a subject chooses to “get more information,” (GMI) 100 small dots appear, each of which is either **red** or **green**, as is depicted in Figure 2. There are always 49 of one color of small dot, and 51 of the other color.⁷ If 51 of the small dots are **red**, then the payoff from option Y is 0, and if 51 of the small dots are **green**, then the payoff from option Y is 1. Figure 1 and Figure 2 display an example of a decision problem before and after the subject has chosen to GMI. This learning environment seems like it may lead to symmetric q accuracies, as is assumed in Section 2.2, if a subject chooses to GMI because it seems that a subject may simply count the dots with a certain amount of accuracy, and the realized number of small dots of each color are always symmetrically located on either side of 50. It does not seem like there is much potential for a subject to flexibly adjust their accuracies across states of the world. In total, however, subjects only chose to GMI in 9,233 of 22,356 decision problems, with substantial variation in how many times different individuals chose to GMI.

In the context of the experiment a subject is said to **learn** in a decision problem if they choose to “get more information” (GMI), which is a prerequisite for their probability of selecting option Y differing across states. This is done because, while in the theory it is assumed that choice probabilities are perfectly observed, in any real dataset (such as the one generated in this paper’s experiment) only an estimate of choice probabilities can be obtained, and frequently the estimate relies on relatively few observations. Choice probabilities not being statistically significantly different is thus not a satisfactory indicator of whether or not learning occurred. Reaction times can also be used as a proxy for the decision to learn, but this seems to be a noisier

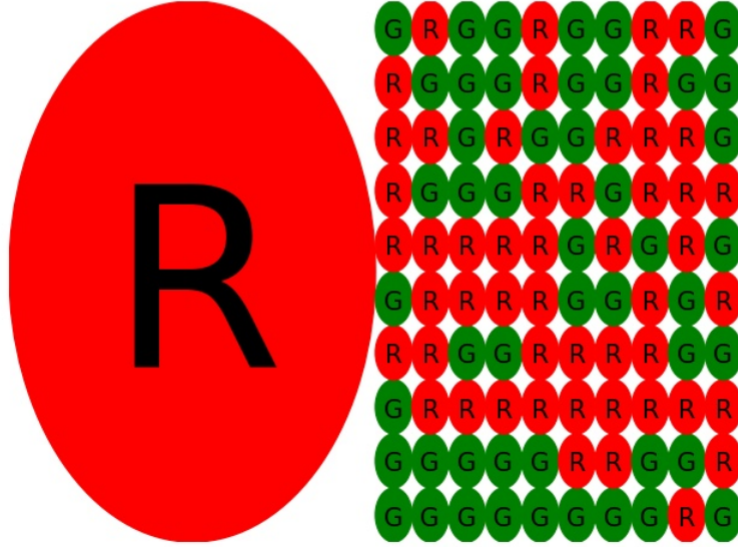
⁷Subjects are provided with this information in their training so that their uncertainty about the returns to effort are minimized.

Figure 2: After choosing “Get more information”:
Please respond to the following investment decision:

Round number: 9 of 100

Option X increases your probability of winning the \$30 prize by 0.25 percentage points

Please choose an option:



indicator of the decision to learn as there is nothing stopping subjects from taking short “breaks” during an experiment to check the time, stretch, rub their eyes, or look at their cellphones. I do, however, use reaction times to supplement the analysis in Section 3.2, and consider the possibility that subjects sometimes made mistakes when selecting if they wanted to GMI.

Further, even it is assumed that there is a fixed cost of choosing to GMI that must be paid before the convex cost function for information is applied to determine costs of information outcomes, it is easy to show that this does not change the implications of the costly learning models for how the choice to acquire information should depend on price (see the work of Walker-Jones (2026a)).

Subjects participated on-line due to COVID-19.⁸ In addition to the decision

⁸The experiment was programmed using oTree (Chen, Schonger, & Wickens, 2016).

problems discussed above, each subject was trained, completed a quiz, and did eight rounds of practice problems to familiarize themselves with the interface before the 92 rounds of decision problems. For a more detailed description of the experiment, see the Online Appendix.

3.1 Summary Statistics and Aggregate Demand

Table 1: Probability of Choosing “get more information” (GMI)

	Big Dot Red ($\mu(\bar{\omega}) = \frac{1}{4}$)	Big Dot Green ($\mu(\bar{\omega}) = \frac{3}{4}$)
$p = 0.25$	43.8%	36.5%
$p = 0.50$	36.0%	39.6%

Fixing either belief, the change in probability of choosing to GMI contradicts the representative agent model result of Corollary 1, but is predicted by Corollary 2, the heterogeneous cost version of the result.

Subjects often do, but also often do not, learn (choose to GMI) at each pair of p and μ that I observe in the experiment, as is shown by Table 1. If the information cost function of a subject is assumed to be constant then the fact that they sometimes do and sometimes do not learn at a given price indicates that they are indifferent between learning and not learning at that price. Thus, if a costly learning model rationalizes their observed behavior, it also rationalizes them never learning at that price, and then Corollary 1 implies they cannot ever be learning at another price that is further from the expected value of the uncertain option. Data that is aggregated over subjects, and features combinations of learning and not learning at prices that are both further from and closer to the expected value of the uncertain option thus rejects the predictions of Corollary 1 (see the next subsection for a more detailed discussion). If the cost function for information outcomes is assumed to be constant across rounds with the same prior belief the behavior thus rejects the most flexible model of costly learning this paper studies that only imposes Assumption 1 and Assumption 2. This

Table 2: Demand Schedules

Aggregate Behavior: Big Dot Red ($\mu(\bar{\omega}) = \frac{1}{4}$)	
$\Pr(Y \underline{\omega}, p = 0.25) = 0.182$	$\Pr(Y \underline{\omega}, p = 0.5) = 0.098$
$\Pr(Y \bar{\omega}, p = 0.25) = 0.513$	$\Pr(Y \bar{\omega}, p = 0.5) = 0.391$
Aggregate Behavior: Big Dot Green ($\mu(\bar{\omega}) = \frac{3}{4}$)	
$\Pr(Y \underline{\omega}, p = 0.25) = 0.645$	$\Pr(Y \underline{\omega}, p = 0.5) = 0.511$
$\Pr(Y \bar{\omega}, p = 0.25) = 0.943$	$\Pr(Y \bar{\omega}, p = 0.5) = 0.893$

is not surprising since it is natural to think that subjects might have changes in their cost function for information over the course of the experiment due to factors such as fatigue. If a distribution of cost functions is considered then Assumption 1 and Assumption 2 imply that the probability of a subject choosing to GMI should weakly increase as price moves towards the expected value of option Y (see Corollary 2), and this predicted pattern is present in the data.⁹

If the state and price dependent choice probabilities for option Y depicted in Table 2 are considered, fixing either prior belief, changes in aggregate demand are not consistent with the q accuracy model from Section 2.2 because in Figure 3, at each belief, both information outcomes are on the same side of the yellow line that indicates the symmetric accuracies, which contradicts the prediction of Proposition 1.¹⁰ Much of this violation is due to the fact that subjects fluctuate back and forth between learning and not learning, but there is also evidence that subjects that do choose to learn are actually achieving “accuracies” that differ across the states of the world (see Table 5). Subjects seem to be surprisingly flexible when choosing what

⁹The increase in the probability of GMI when $\mu(\bar{\omega}) = \frac{1}{4}$ and p increases from $p = 0.25$ to $p = 0.50$, which is predicted by Corollary 2, is statistically significant at the one percent level according to a two-sided Fisher Exact test). The difference between Corollary 2 and Corollary 1 is that the former allows for a non-constant cost function for information outcomes while the latter does not.

¹⁰Four two-sided Fisher Exact tests reject that each information outcome is on the yellow line at the one percent level.

information to acquire, and in part this may be explained by subjects spending more time learning when the state of the world is not what they expected based on their prior (see Table 6). Aggregate demand does change in line with the predictions of the Posterior Separable model of costly learning studied by Walker-Jones (2026a), as can be seen in Figure 3.

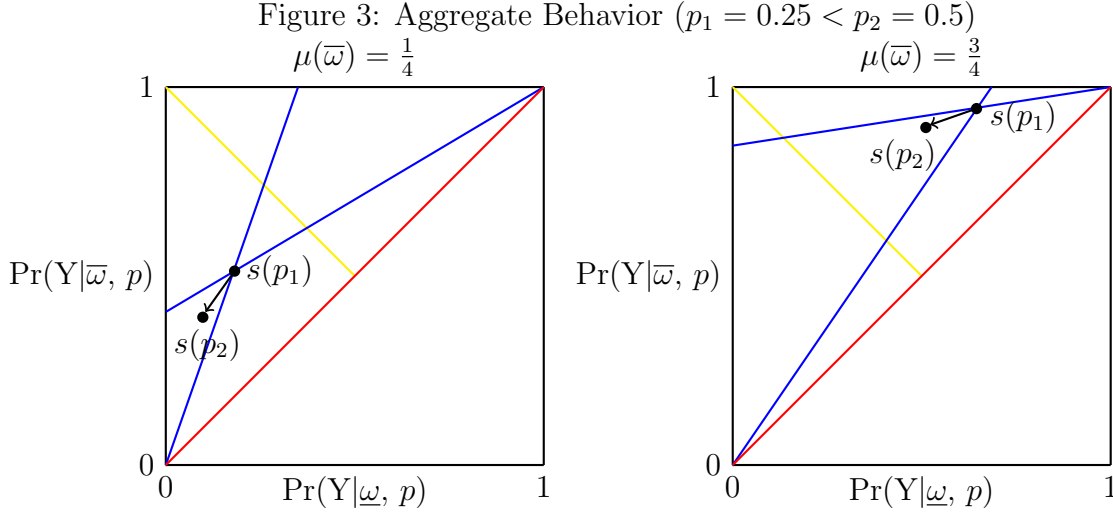


Figure 3 shows the aggregated behavior of subjects in the experiment with rounds with $\mu(\bar{\omega}) = \frac{1}{4}$ on the left side and rounds with $\mu(\bar{\omega}) = \frac{3}{4}$ on the right side. The changes in aggregate demand are not consistent with the q accuracy model from Section 2.2 because, at each belief, the changes contradict the prediction of Proposition 1. The changes in aggregate observed behavior are completely in line with the predictions of the PS model (Theorem 2) from Walker-Jones (2026a), which requires that $s(p_2)$ remains weakly above the blue line from $(0, 0)$ through $s(p_1)$ and $s(p)$ must remain weakly below the blue line from $(1, 1)$ through $s(p_1)$.

3.2 Within Subject Heterogeneity: Fatigue and Novelty

Corollary 1 implies that, if the cost function for information outcomes is constant across rounds with the same prior (as is assumed throughout this paragraph), if $\mu(\bar{\omega}) = \frac{1}{4}$ a subject has more incentive to learn if the price is $p = 0.25$ compared to $p = 0.5$. Corollary 1 thus implies that, in rounds with $\mu(\bar{\omega}) = \frac{1}{4}$, if a subject sometimes chooses GMI in rounds with $p = 0.5$ then they should choose GMI in all rounds with $p = 0.25$. While if $\mu(\bar{\omega}) = \frac{3}{4}$, Corollary 1 says a subject has more

incentive to learn if the price is $p = 0.5$. Corollary 1 thus implies that, in rounds with $\mu(\bar{\omega}) = \frac{3}{4}$, if a subject sometimes chooses GMI in rounds with $p = 0.25$ then they should choose GMI in all rounds with $p = 0.5$. Table 3 shows that many subjects thus violate the predictions of the most general model of costly learning that only imposes Assumption 1 and Assumption 2, and that this conclusion is robust to concerns about lack of understanding as it is true even for subjects that did quite well on the quiz (rows 2, 3, and 4, of Table 3).

To help ensure that the violations of Corollary 1 are not being generated by subjects either accidentally selecting or not selecting to GMI, Table 3 also conducts some analysis that allows for some “mistakes” and additionally controls for quiz performance (rows 5 through 8). Pervasive violations of Corollary 1 are still found when some choices are explained away as “mistakes.” This is not surprising because the experiment is designed so that these types of mistakes are unlikely. Subjects pick if they do or do not want to GMI from a drop-down menu (see Figure 1), which then displays their choice, and they must then subsequently hit a “Next” button to finalize their decision that is quite some distance from the drop-down menu, below both the big dot for the round and a paragraph that reminds them of the meaning of the big dot (see Figure 17 in the Online Appendix).

One natural explanation for why the predictions of Corollary 1, and the very general model of costly learning it studies, are violated by so many subjects is that subjects may fatigue over the course of the experiment and as a result choosing to GMI may be less appealing as the experiment progresses. This can be modelled by allowing C_μ to increase as a subject progresses through the decision problems. Let C_μ^r denote the cost function of a subject when their belief is μ and they are facing their r^{th} decision problem. A subject’s cost function for information outcomes is said to be **consistent with fatigue** if, given belief μ , and integer $t > 0$, $\forall \underline{s} \in S$ with $\underline{s} \leq \bar{s} : C_\mu^r(\underline{s}, \bar{s}) \leq C_\mu^{r+t}(\underline{s}, \bar{s})$. As is also demonstrated in Table 3, allowing for fatigue substantially reduces the percent of subjects that violate the very flexible

Table 3: Percent of subjects that violate Corollary 1
(violate the costly learning model that only imposes Assumptions 1 and 2)

Row: 1	If analysis ignores fatigue and “mistakes”	56%
2	And is restricted to subjects that got 5/10 or more on the quiz	55% (122/223)
3	And is restricted to subjects that got 7/10 or more on the quiz	54% (94/173)
4	And is restricted to subjects that got 9/10 or more on the quiz	49% (49/99)
5	If analysis allows for some “mistakes”	45%
6	And is restricted to subjects that got 5/10 or more on the quiz	46% (102/223)
7	And is restricted to subjects that got 7/10 or more on the quiz	46% (80/173)
8	And is restricted to subjects that got 9/10 or more on the quiz	46% (46/99)
9	If analysis allows for fatigue	37%
10	And is restricted to subjects that got 5/10 or more on the quiz	36% (81/223)
11	And is restricted to subjects that got 7/10 or more on the quiz	35% (60/173)
12	And is restricted to subjects that got 9/10 or more on the quiz	31% (31/99)

Row 1 indicates the percent of subjects that, when $\mu(\bar{\omega}) = \frac{1}{4}$, sometimes choose to acquire information (choose to GMI) if $p = 0.5$ but do not always choose GMI if $p = 0.25$, or, when $\mu(\bar{\omega}) = \frac{3}{4}$, sometimes choose GMI if $p = 0.25$ but do not always choose GMI if $p = 0.5$. Rows 2, 3, and 4, conduct the same exercise as row 1 but restrict attention to subjects that got quiz scores of at least 5, 7, and 9, out of 10 respectively. Row 5 conducts the same exercise as row 1 but a subject is only considered to have not GMI at a pairing of μ and p if there are at least two instances of them not choosing GMI at said pairing, and they are only considered to have GMI at a pairing of μ and p if there is at least one round in which they chose GMI at said pairing and then took at least 10 seconds to make a decision. Rows 6, 7, and 8, conduct the same exercise as row 5 but restrict attention to subjects that got quiz scores of at least 5, 7, and 9, out of 10 respectively. Row 9 indicates the percent of subjects that, in rounds with $\mu(\bar{\omega}) = \frac{1}{4}$, chose not to GMI for $p = 0.25$ and then subsequently (in a later round) chose to GMI for $p = 0.5$, or, in rounds with $\mu(\bar{\omega}) = \frac{3}{4}$, chose not to GMI for $p = 0.5$ and then subsequently chose to GMI for $p = 0.25$. Rows 10, 11, and 12, conduct the same exercise as row 9 but restrict attention to subjects that got quiz scores of at least 5, 7, and 9, out of 10 respectively.

model of costly learning introduced in Section 2.1, but many subjects still violate the predictions of Corollary 1, and this result holds even in the subsets of subjects that did best on the quiz.

It seems that fatigue is not sufficient for rationalizing the behavior of subjects, at least in part, because subjects are more likely to learn in rounds that are novel in the sense that they feature a price and or prior belief that differs from previous rounds. Table 4 reports the results of 3 different Logit regressions on the probability of subjects choosing to GMI in rounds. Notice that the sign of the significant coefficients on the difference between the price and the expected value of option Y ($|p - \mu(\bar{\omega})|$)

are consistent with a costly learning model with heterogeneous costs (see Corollary 2).

Table 4: Logit regressions on probability of GMI

Constant	−1.645 (0.050)	−1.676 (0.052)	−1.747 (0.054)
Round number	−0.026 (0.001)	−0.025 (0.001)	−0.025 (0.001)
$ p - \mu(\bar{\omega}) $	−2.957 (0.258)	−2.964 (0.258)	−3.033 (0.258)
Number of other rounds the subject chose GMI in	0.073 (0.001)	0.073 (0.001)	0.073 (0.001)
Dummy for $\mu(\bar{\omega}) = \frac{3}{4}$	0.737 (0.124)	0.698 (0.126)	0.649 (0.126)
Decision problem different than 1 round before	0.153 (0.047)	0.148 (0.047)	0.129 (0.047)
Decision problem different than 2 rounds before		0.096 (0.046)	0.072 (0.047)
Decision problem different than 3 rounds before			0.224 (0.046)
Pseudo R-squared	0.4988	0.4989	0.4997
Number of observations	22356	22356	22356

A decision problem is said to be different than the decision problem a certain number of rounds before if the price and or the belief differ across the rounds.

If the model fitted in the second last column in Table 4 is used to predict the probability of a subject that has gone through half of the decision problems and chose to GMI in 26 of their other decision problems (26 is the median number of times to GMI) choosing GMI in a round with $p = \mu(\bar{\omega}) = 0.25$, then the decision problem being the same as the previous two rounds (same price and prior as the previous two rounds) results in a 23.6% chance of them choosing GMI whereas the decision problem being different from the previous two rounds (different price and or prior from the

previous two rounds) results in a 28.2% chance of them choosing GMI. If, instead, the model fitted in the last column in Table 4 is used to predict the probability of a subject that has gone through half of the decision problems and chose to GMI in 26 of their other decision problems choosing GMI in a round with $p = \mu(\bar{\omega}) = 0.25$, then the decision problem being the same as the previous three rounds results in a 22.5% chance of them choosing GMI whereas the decision problem being different from the previous three rounds results in a 30.8% chance of them choosing GMI. Proportionally, both of these changes are large. This is surprising given how little variation there is between different decision problems in the experiment. Further, the coefficients on the round number in Table 4 is one indication of the fatigue that subjects experience during the experiment.

The impacts of fatigue and the “novelty” of a decision problem that are demonstrated in Table 4 suggest that even if behavior is aggregated across the decisions of a single individual and not across multiple individuals, allowing for the cost function for information to vary across decision problems, as is done in Section 2.3, seems like a natural modelling decision.

Table 5: Non-symmetric Accuracies ($\mu(\bar{\omega}) = \frac{3}{4}$, $p = 0.25$)

	$\Pr(Y \bar{\omega}, p):$		$\Pr(X \underline{\omega}, p):$
Decision Problems with “get more information”	0.964	>**	0.905
Counting Time ≥ 15 seconds	0.977	>**	0.944
Note: ** indicates rejection at one percent level (two-sided Fisher Exact test)			

Table 5 shows that even in rounds where subjects chose GMI, and rounds where subjects chose GMI and then spent at least 15 seconds making a decision, there is evidence that accuracies are not constant across the states of the world, in contrast to what is assumed by the q accuracy model from Section 2.2. Choosing a threshold of 15 seconds is arbitrary, but similar thresholds also produce statistically significant differences in accuracies.

3.3 Flexibility of Learning

In contrast with the assumptions of the q accuracy model, in some types of decision problems there is evidence that subjects had probabilities of selecting the best available option that differed across states of the world, and this is true even when only rounds with learning in them are considered. One example of this type of asymmetry is provided in Table 5. Interestingly, even if heterogeneous costs are introduced into the q accuracy model, these accuracies that differ across the states of the world when subjects are choosing to learn still reject a model with a distribution over q accuracy cost functions for information.

One explanation for probabilities of selecting the best available option that differ across states of the world could be subjects having a correct count with a fixed accuracy, but changing the number of times they count depending on if the result of their count was expected or unexpected based on their prior. There also seems to be some evidence of this, as shown by Table 6.

Table 6: OLS Regression on Counting Time After GMI ($\mu(\bar{\omega}) = \frac{1}{4}$, $p = 0.25$)

	Counting Time ≤ 90 seconds	All Decision Problems
Constant	33.466** (0.502)	37.564** (0.682)
Round Number	-0.099** (0.007)	-0.128** (0.009)
Value of option $Y = 1$ ($u(\omega) = 1$, $\omega = \bar{\omega}$)	1.726** (0.470)	1.91** (0.641)
Adjusted R-squared	0.037	0.031
Number of observations	6169	6383

Note: ** indicates significance at one percent level. Table 6 shows that, after choosing GMI, learning times varied based on the value of option Y . Subjects spent more time learning when the state was unexpected based on the prior. This is true when only decision problems where less than 90 seconds was spent after choosing GMI and before selecting an option (to rule out that the effect is being generated by outliers) or when all decision problems with GMI selected are considered.

4 Literature Review

A growing literature demonstrates the importance of inattention in standard economic fields (Maćkowiak, Matějka, & Wiederholt, 2023). These fields include finance (Huberman, 2001), labor search (Acharya & Wee, 2020), trade (Dasgupta & Mondria, 2018), and voting behavior (Shue & Luttmer, 2009).

The work of Chetty, Looney, and Kroft (2009) is particularly relevant as they show that including sales tax in prices reduces demand, highlighting the impact of inattention even in settings where it seems information should be costless. Taubinsky and Rees-Jones (2018) further show that accounting for the heterogeneity of under-reaction to tax is crucial for accurately estimating the resultant efficiency loss.

There are also a number of other papers that experimentally test the implications of models of costly learning. The experiment in this paper is inspired by the work of Dean and Neligh (2023), but in their paper they do not have a change in parameters that is equivalent to the change in price that is the primary focus of this paper, and do not observe the outcome of the subjects' decision to "get more information." Dewan and Neligh (2020) study a different set of costly learning tasks experimentally, but again they do not observe the decision to "get more information," and do not have a change in parameters that is equivalent to a change in price. When these papers change option values they primarily do so in a multiplicative fashion, i.e. they do something analogous to multiplying p and $u(\omega)$ by a constant so the incentive to learn is unambiguously higher or unambiguously lower. Understanding how inattention changes when price changes is important because changes in price are so commonly observed in the real world, whereas changes in belief or a multiplicative change in payoffs are more difficult to observe. Ambuehl (2017) and Ambuehl, Ockenfels, and Stewart (2022) experimentally and theoretically explore environments with changes in parameters that are equivalent to changes in price, but they do not observe the decision to "get more information."

5 Conclusion

This paper uses a novel data enrichment to study how the information obtained by subjects about a good changes with the price of the good. Experiment subjects are more likely to invest effort into learning about the value of options if simple choice parameters, like price, differ from previous choice problems. This increase in effort in “unfamiliar” choice problems, and fatigue, mean that the behavior of many subjects violate even the most flexible model of costly learning if the cost for information is assumed to be constant across choice problems with the same prior beliefs. This observation motivates the introduction of heterogeneous decision makers when modelling the environment to better fit the data. Such a model of learning with heterogeneous costs is used to predict that when price is changed to be closer to the expected value of the good subjects should be more likely to acquire information about the good. This prediction is found to be accurate in the data.

The experiment is arguably a setting where a simple “ q accuracy” model seems especially appropriate. If subjects are assumed to learn in each decision problem and have the same cost function for information, then such a model would predict locally upward sloping demand for one of the versions of the good. Aggregate demand in the experiment ends up being downward sloping for both types of good in large part because of the apparent heterogeneity of costs. This point drives home the importance of considering the heterogeneity of costs when modelling costly learning environments.

There is also evidence that subjects are surprisingly flexible when choosing what information to acquire. When subjects do choose to acquire information they are, in contrast with the assumptions of the “ q accuracy” model, able to achieve accuracies that differ across the states of the world, and spend longer learning when the realized state is unexpected based on their prior. It seems subjects know they could be making mistakes and invest more effort into learning when their initial learning does not lead to the conclusion they anticipated, a degree of flexibility that is assumed by

standard examples of Posterior Separable cost functions (Caplin & Dean, 2013) like Shannon Entropy (Shannon, 1948). The Posterior Separable model is shown to do well predicting changes in aggregate demand in the data.

Appendix 1: Proof of Main Theorems

Lemma 1. Given price p , prior μ , and a q accuracy cost function for information outcomes C_μ , if $\underline{s}(p) < \bar{s}(p)$ and $s(p)$ maximizes the expected gain in payoff, I can find $q(p) \in (\frac{1}{2}, 1]$ such that $s = (1 - q(p), q(p))$ maximizes the expected gain in payoff:

(i) If $\Pr(Y|\bar{\omega}, p) = \Pr(X|\underline{\omega}, p) = q$, then $s = (1 - q, q)$ maximizes the expected gain in payoff.

(ii) If $\Pr(Y|\bar{\omega}, p) > \Pr(X|\underline{\omega}, p)$, then on the line through $s = (1, 1)$ and $s(p)$ the point where this line hits the line defined by $s = (1 - q, q)$ maximizes the expected gain in payoff: there are unique $q \in (\frac{1}{2}, 1]$ and $\alpha \in [0, 1]$ such that $\alpha(1 - q) + (1 - \alpha)1 = \Pr(Y|\underline{\omega}, p)$ and $\alpha q + (1 - \alpha)1 = \Pr(Y|\bar{\omega}, p)$, and $s = (1 - q, q)$ maximizes the expected gain in payoff.

(iii) If $\Pr(Y|\bar{\omega}, p) < \Pr(X|\underline{\omega}, p)$, then on the line through $s = (0, 0)$ and $s(p)$ the point where this line hits the line defined by $s = (1 - q, q)$ maximizes the expected gain in payoff: there are unique $q \in (\frac{1}{2}, 1]$ and $\alpha \in [0, 1]$ such that $\alpha(1 - q) + (1 - \alpha)0 = \Pr(Y|\underline{\omega}, p)$ and $\alpha q + (1 - \alpha)0 = \Pr(Y|\bar{\omega}, p)$, and $s = (1 - q, q)$ maximizes the expected gain in payoff.

Proof The first case is trivial. For the second case, consider the q defined by the equation in the statement of the second case, and denote it $q(p)$. For each $x \in [\frac{1}{2}, 1]$ such that $\exists q \in (\frac{1}{2}, 1]$ and $\alpha \in [0, 1]$ such that $\alpha(1 - q) + (1 - \alpha)x = \Pr(Y|\underline{\omega}, p)$ and $\alpha q + (1 - \alpha)x = \Pr(Y|\bar{\omega}, p)$, denote the the unique q that solves the equation by $\tilde{q}(x)$. Given the definition of q accuracy cost functions, there must be an $\tilde{x} \in [\frac{1}{2}, 1]$ such that $(\tilde{x}, \tilde{x})$ maximizes the expected gain in payoff and the associated $\tilde{q}(\tilde{x})$ is such that $(1 - \tilde{q}(\tilde{x}), \tilde{q}(\tilde{x}))$ also maximizes the expected gain in payoff. Note: $\tilde{q}(\tilde{x}) \geq q(p)$ (if

I move away from $(x, x) = (1, 1)$ the rotation gives us higher $\tilde{q}(x)$). Either $\tilde{x} = 1$ so $\tilde{q}(\tilde{x}) = q(p)$, or $\tilde{x} < 1$ and since $(\tilde{x}, \tilde{x})$ maximizes the expected gain in payoff, $(\frac{1}{2}, \frac{1}{2})$ must also maximize the expected gain in payoff, and as optimal sets need to be convex, $q(p)$ is then optimal.

For the third case, consider the q defined by the equation in the statement of the third case, and denote it $q(p)$. For each $x \in [0, \frac{1}{2}]$ such that $\exists q \in (\frac{1}{2}, 1]$ and $\alpha \in [0, 1]$ such that $\alpha(1 - q) + (1 - \alpha)x = \Pr(Y|\underline{\omega}, p)$ and $\alpha q + (1 - \alpha)x = \Pr(Y|\bar{\omega}, p)$, denote the the unique q that solves the equation by $\tilde{q}(x)$. Given the definition of q accuracy cost functions, there must be an $\tilde{x} \in [0, \frac{1}{2}]$ such that $(\tilde{x}, \tilde{x})$ maximizes the expected gain in payoff and the associated $\tilde{q}(\tilde{x})$ is such that $(1 - q(\tilde{x}), q(\tilde{x}))$ also maximizes the expected gain in payoff. Note: $\tilde{q}(\tilde{x}) \geq q(p)$ (if I move away from $(x, x) = (0, 0)$ the rotation gives us higher $\tilde{q}(x)$). Either $\tilde{x} = 0$ so $\tilde{q}(\tilde{x}) = q(p)$, or $\tilde{x} > 0$ and since $(\tilde{x}, \tilde{x})$ maximizes the expected gain in payoff, $(\frac{1}{2}, \frac{1}{2})$ must also maximize the expected gain in payoff, and as optimal sets need to be convex, $q(p)$ is then optimal. ■

Proof of Proposition 1. (i) is necessary because if there are two prices $p_1 < p_2$ such that $\Pr(Y|\bar{\omega}, p) > \Pr(X|\underline{\omega}, p) > 0$, then at each price $(1, 1)$ is optimal (because of the construction of q accuracy cost functions), and the DM could have selected it at p_2 , but then Theorem 1 tells us they were not selecting an optimal strategy at p_1 . (ii) is necessary because if there are two prices $p_1 < p_2$ such that $\Pr(X|\bar{\omega}, p) > \Pr(Y|\underline{\omega}, p) > 0$, then at each price $(0, 0)$ is optimal (because of the construction of q accuracy cost functions), and the DM could have selected it at p_1 , but then Theorem 1 tells us they were not selecting an optimal strategy at p_2 . (iii) is necessary because if not, I get a violation of Theorem 1 as the respective $q(p)$'s were both optimal, the DM could have selected information outcomes for the two prices such that at each $s(p) = (1 - q(p), q(p))$, and then the unconditional chance of selecting Y given either $q(p)$ is $\mu(\bar{\omega})q(p) + \mu(\underline{\omega})(1 - q(p))$, which means the unconditional chance of selecting Y increased when price increased, contradicting Theorem 1.

Order the prices in $\mathcal{P} = \{p_1, p_2, \dots, p_n\}$, so $p_1 < p_2 < \dots < p_n$. Further,

assume that there is a price in $\mathcal{P}$ where the DM learns, because if not I can easily rationalize the behavior by defining a C_μ for $\underline{s} < \bar{s}$ using a steep enough plane through the diagonal of S (points (x, x) with $x \in [0, 1]$). Theorem 1 thus tells us that there is not a price $p_i \in \mathcal{P}$ with $s(p_i) = (x, x)$ and $x \in (0, 1)$. It is further without loss to assume that $s(p_1) = (1, 1)$ and $s(p_n) = (0, 0)$. It is further without loss to assume that if the DM learns at a price p_i , then $s(p_i)$ is such that $1 - \underline{s}(p_i) = \bar{s}(p_i)$. If not, and $1 - \underline{s}(p_i) > \bar{s}(p_i)$, then change $s(p_i) = (0, 0)$ and enrich the behavior so there is a $p_j \in \mathcal{P}$ so $p_j \in (p_{i-1}, p_i)$ and $s(p_j) = (1 - q(p_i), q(p_i))$ where $q(p_i)$ is selected as in Lemma 1. If not, and $1 - \underline{s}(p_i) < \bar{s}(p_i)$, then change $s(p_i) = (1 - q(p_i), q(p_i))$ where $q(p_i)$ is selected as in Lemma 1. Then, I can just use the same recursive definition for the costs of information outcomes as in the proof of Theorem 1 to rationalize the behavior, and the resultant C_μ is a q accuracy cost function for information outcomes (see the proof of Theorem 1). ■

References

- Acharya, S., & Wee, S. L. (2020). Rational inattention in hiring decisions. *American Economic Journal: Macroeconomics*, 12(1), 1–40.
- Allais, M. (1953). Le comportement de l’homme rationnel devant le risque: critique des postulats et axiomes de l’école américaine. *Econometrica: Journal of the Econometric Society*, 503–546.
- Ambuehl, S. (2017). An offer you can’t refuse? incentives change how we inform ourselves and what we believe. *CESifo Working Papers*(6296).
- Ambuehl, S., Ockenfels, A., & Stewart, C. (2022). Who opts in? composition effects and disappointment from participation payments. *The Review of Economics and Statistics*, 1–45.
- Azrieli, Y., Chambers, C. P., & Healy, P. J. (2018). Incentives in experiments: A theoretical analysis. *Journal of Political Economy*, 126(4), 1472–1503.

- Azrieli, Y., Chambers, C. P., & Healy, P. J. (2020). Incentives in experiments with objective lotteries. *Experimental Economics*, 23(1), 1–29.
- Blackwell, D. (1951). Comparison of experiments. In *Proceedings of the second berkeley symposium on mathematical statistics and probability* (Vol. 2, pp. 93–103).
- Blackwell, D. (1953). Equivalent comparisons of experiments. *The annals of mathematical statistics*, 265–272.
- Caplin, A., & Dean, M. (2013). *Behavioral implications of rational inattention with shannon entropy* (Tech. Rep.). National Bureau of Economic Research.
- Caplin, A., Dean, M., & Leahy, J. (2022). Rationally inattentive behavior: Characterizing and generalizing shannon entropy. *Journal of Political Economy*, 130(6), 1676–1715.
- Chen, D. L., Schonger, M., & Wickens, C. (2016). otree—an open-source platform for laboratory, online, and field experiments. *Journal of Behavioral and Experimental Finance*, 9, 88–97.
- Cheng, X., & Kim, Y. (2025). On the monotonicity of information costs.
- Chetty, R., Looney, A., & Kroft, K. (2009). Salience and taxation: Theory and evidence. *American economic review*, 99(4), 1145–77.
- Dasgupta, K., & Mondria, J. (2018). Inattentive importers. *Journal of International Economics*, 112, 150–165.
- Dean, M., & Neligh, N. (2023). Experimental tests of rational inattention. *Journal of Political Economy*, 131(12), 3415–3461.
- Dewan, A., & Neligh, N. (2020). Estimating information cost functions in models of rational inattention. *Journal of Economic Theory*, 187(105011).
- Greiner, B. (2015). Subject pool recruitment procedures: organizing experiments with orsee. *Journal of the Economic Science Association*, 1(1), 114–125.
- Huberman, G. (2001). Familiarity breeds investment. *The Review of Financial Studies*, 14(3), 659–680.

- Maćkowiak, B., Matějka, F., & Wiederholt, M. (2023). Rational inattention: A review. *Journal of Economic Literature*, 61(1), 226–273.
- Shannon, C. E. (1948). A mathematical theory of communication. *The Bell System Technical Journal*, 27(3), 379–423.
- Shue, K., & Luttmer, E. F. (2009). Who misvotes? the effect of differential cognition costs on election outcomes. *American Economic Journal: Economic Policy*, 1(1), 229–57.
- Sims, C. A. (2003). Implications of rational inattention. *Journal of monetary Economics*, 50(3), 665–690.
- Taubinsky, D., & Rees-Jones, A. (2018). Attention variation and welfare: theory and evidence from a tax salience experiment. *The Review of Economic Studies*, 85(4), 2462–2496.
- Walker-Jones, D. (2026a, July). *Heterogeneous costs and demand*.
- Walker-Jones, D. (2026b, July). *Heterogeneous costs and the decision to learn*.

Online Appendix: More Experiment Details

There is some repetition of Section 3 from Walker-Jones (2026b) to make things more clear, but the explanation of the experiment in this online appendix is not complete, and is meant to complement the description in Section 3 of Walker-Jones (2026b).

Participants were recruited for the experiment using ORSEE (Greiner, 2015) from the Toronto Experimental Economics Laboratory recruitment pool. Subjects signed up ahead of time for a particular day, either the 21st, 22nd, 23rd, 28th, 29th, or 30th of September 2020.¹¹ Subjects were told ahead of time they would be sent a link at 8 AM EDT, and would have until 8 PM EDT to finish the experiment.¹² In total 270 undergraduate students from the University of Toronto attempted to partake in the experiment, and 243 of them completed the experiment.¹³ I refer to the 243 students that completed the experiment as the ‘subjects.’

Each subject began by consenting to participate in the experiment. They then saw a welcome screen that explained that they were about to be trained and quizzed. They were told: “Please read the following instructions carefully. Understanding what is going on will help you earn more money. You will first be trained, and then you will be quizzed, to make sure you understand how everything works. There will be 10 questions in the quiz. You must answer each question correctly before you can move on to the next question. When you complete the quiz you will earn the \$5

¹¹I had 24 subjects on the 21st, 39 on the 22nd, 48 on the 23rd, 43 on the 28th, 46 on the 29th, and 70 on the 30th.

¹²On the 22nd the session was extended until 9 PM EDT because students experienced crashes. Then on September 28th the Economic Department’s servers were unexpectedly down until almost 10 AM EDT, so students were given until 10 PM EDT to complete the experiment.

¹³There were 315 undergraduate students that signed up to receive a link for the experiment but 45 of them did not use the link. This proportion is in line with the “turnout” rates of other experiments run with the recruitment pool. A handful of students experienced one of two errors on September 22nd. One seemed to be caused by a server problem, while the other was caused by a missing image file. Four of the students that experienced errors did not want to re-start. Two students started the experiment with less than 20 minutes left before their session finished and could not finish as a result. This left 264 students, but some of them chose to drop out, as is explained later, so 243 ended up completing the experiment.

training payment. In addition, each quiz question you answer correctly on the first try earns you another \$0.5 training bonus payment, so you can earn up to \$10 just by doing well during the training.” I wanted to incentivize students to pay attention during the training so that they understood the environment and their choices would be indicative of their preferences. All payments were in Canadian dollars.

The training consisted of three pages, which are figures 4, 5, and 6, below. In the second training page subjects were required to select “get more information” (GMI) so that they knew what happened if they did. If they did not, then they got an error message asking them to do so. The 10 quiz questions can be seen at the end of this appendix in figures 7 through 16. If a subject got a quiz question wrong the order of the answers in the drop down menu randomly permuted so that subjects could not just try them in order. The quiz verified that subjects understood the environment, but also gave me another opportunity to train them. In the data I can see how many times each subject attempted each question. I also know how long subjects spent on each page. Among subjects the average time spent on the welcome page, training pages, and quiz pages was 17.9 minutes, and the average quiz score was 7.6/10. So, subjects in general put a lot of effort into understanding the environment, and were quite successful in doing so.

After the quiz subjects faced 100 rounds of ‘investment decisions’ (the last 92 of which are the decision problems I study in the body of the paper). In each round the subject selects one of two options, option X or option Y. In all 100 rounds the DM selects from their options in a drop down menu, and the order of the options is randomly generated in each round.

The subject earned probability points in each round. Subjects gradually increased their collection of probability points so as to have a higher chance of winning the draw at the end of the 100 rounds. If they won the draw they would receive a monetary prize of either \$20, \$25, \$30.¹⁴ The prize they would receive if they won

¹⁴This amount of variation in the prizes across treatments ended up being fairly inconsequential compared to the apparent variation in learning costs across subjects. The maximal mean and median

was determined by the subject’s treatment group, and was displayed to the subject throughout the training pages and rounds.¹⁵ In a round the smallest increase the subject can get in their probability of winning the prize is zero percentage points, and the largest increase is one percentage point. In each round option X is worth an amount of percentage points strictly between 0 and 1, and option Y is worth 0 or 1 percentage points.

I paid subjects in probability points so I could attempt within subject analysis in an incentive compatible way. The ideal when it comes to incentive compatibility is that in each of the 100 rounds the subject chooses what they would have chosen if they had only had to make a decision in that one round (Azrieli, Chambers, & Healy, 2018, 2020). In the setting of my experiment if I pay the subject money in each round then their marginal value for money could decrease as the rounds progress (wealth effects). Further, it would be difficult to separate this change in preferences from fatigue. The standard solution is to use a “random incentive system” (RIS), and pay the DM based on their decision in that round (Allais, 1953). In this paper’s experiment setting this strategy is not incentive compatible. The DM needs to consider the cost of their learning in the setting studied in this paper, and they cannot defer their learning until they know which round is selected. If I increase the number of rounds they see, and keep the monetary payments the same, their accuracy should go down. The benefit (in expectation) from making a good selection in a round decreases while their cost of learning stays the same. Paying subjects in probability points is mathematically equivalent to RIS if subjects reduce compound lotteries and know each round is selected with equal chance.

If subjects do not accurately internalize an objective lottery over rounds, then paying in probability points might be better for eliciting preferences from a pedagogical perspective. This strategy is certainly not infallible, however. The value of a one

number of choices to GMI both occurred in the treatment with the \$25 prize.

¹⁵Because the subjects participated on-line, they were transferred the money electronically within 24 hours of their session ending.

percent increase in the probability of winning the prize might depend on the subject's current probability of winning the prize. I think this is of particular concern in two cases, when the subject's probability of winning the prize is moving away from 0%, and when it is moving to 100%. I think a probability point should have essentially the same value in two different rounds if neither of these two cases are occurring. This means that if I can move a subjects chance of winning away from 0% and 100% I could reduce certainty effects, while certainty effects would be persistent with RIS.

The first 8 rounds of investment decisions familiarize subjects with the experiment interface and allow me to move a subject's chance of winning away from 0% and 100%. In the first 8 rounds subjects could not "get more information" and they had to make a decision based on the payoff from the safe option X and the big dot. If the subject chooses the safe option in any of the first 8 rounds then they know they cannot achieve a 100 percent chance of winning the prize, and certainty effects are mitigated at least partially. If the subject chooses X in any of the first 8 rounds their probability of winning the prize is strictly above zero, and I do not need to worry about them making a decision later purely because it means they guarantee a chance of winning the prize. Further, the first 8 rounds test the understanding of the DMs and if they are reducing compound lotteries in the rounds.

In the first 8 rounds each subject saw the same sequence of big dots (which form beliefs) and payoffs for the safe option X (prices). The big dots were green, green, red, red, red, red, red, and red, and the payoffs for option X were 0.8, 0.7, 0.2, 0.5, 0.4, 0.3, 0.24, and 0.26.

In each decision problem subjects had the option to stop doing decision problems. They were told that if they decided to 'exit' they would maintain their current probability of winning the prize, and immediately find out if they had won or not, forfeiting the chance to further increase their chance of winning. If a DM fatigues, and they no longer think it is worth their time to select the better option, but they cannot stop without losing any probability of winning they had accrued, then they

might rush through the experiment, selecting options not because they are indicative of the DM’s preferences, but because the DM is trying to finish as fast as they can. So, if a subject prefers to stop doing decision problems, I let them, so my data would be less noisy and more indicative of preferences.

I had 21 subjects decide not to finish the experiment, which left me with a sample of 243 subjects. Less than ten percent of students dropping out is not overly concerning since they were participating over the internet and could have had any number of distractions arise.

It was hoped that **red** and **green** would have intuitive value to the subjects. The downside with red and green is that about one in twenty people have red/green color blindness. To combat this issue, each dot displayed has a black letter in its middle, either R or G, so that subjects can still differentiate between red and green dots even if they are color blind. See Figure 2 for an example. I did not receive any complaints about the dots being hard to differentiate.

Before any subjects participated in the experiment ten different “treatments” of image files were generated ($100 \times 10 = 1000$ rounds of images, each consisting of a big dot image and an image of 100 small dots). This means that 1000 times a big dot color was drawn. The chance of the big dot being **green** was $\frac{1}{4}$ and the chance of the big dot being **red** was $\frac{3}{4}$. After the big dot color for the round was determined the composition of the small dots was drawn. In each instance there was a $\frac{3}{4}$ chance that 51 of the small dots would match the color of the big dot and a $\frac{1}{4}$ that 51 of the small dots would not match the color of the big dot. Either way, the order of the small dots was randomly generated given the drawn proportion. After this process was completed, the first eight rounds of images were used for all ten treatments, but other than that the images were not altered. This means, for instance, that if a subject was assigned to image “treatment” nine, their first eight rounds of images were the first eight rounds of images generated, and their last 92 rounds of images were the 809th through 900th rounds of images generated.

Figure 4:

Training (1 of 3):

The experiment consists of 100 rounds of investment decisions, which does not include the training round (Round 0). Each of the 100 rounds provides you with the opportunity to increase your probability of winning a prize of \$25 at the end of the experiment. This prize is in addition to whatever you earn during the training process.

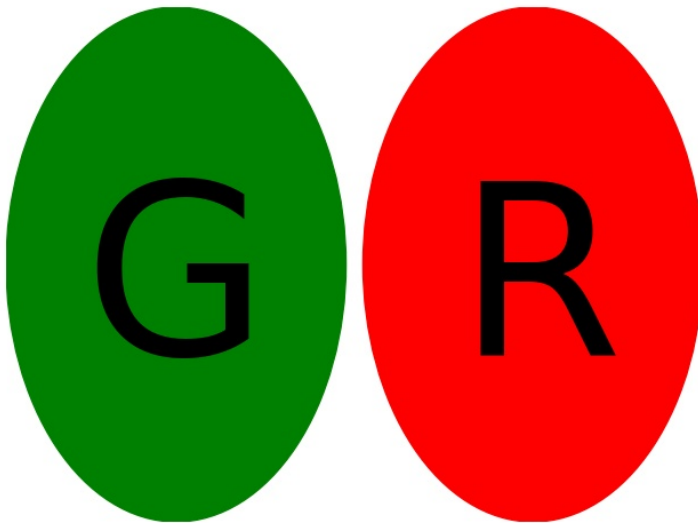
You will not be able to increase your probability of winning the prize by more than 1 percentage point in any one round. Ideally, you would thus like to increase your probability of winning the prize by 1 percentage point in each round, so that you win the prize with a 100% probability when the experiment is over, but it is extremely unlikely that this will be possible.

In each investment decision you get to choose between two investment options, option X and option Y.

Option X always increases your probability of winning the prize by a displayed positive amount, which is always less than 1 percentage point.

Option Y, in contrast, either does not increase your probability of winning the prize at all, or it increases your probability of winning the prize by 1 percentage point.

To help provide you with some information about the chance that option Y increases your probability of winning the prize, in each round there is one large dot that is either **red** with a black R in the middle of it (example below on right), or **green** with a black G in the middle of it (example below on left).



When the large dot is **red**, this is a bad sign for the value of option Y, and there is only a 1 out of 4 chance (25% chance) that option Y increases your probability of winning the prize by 1 percentage point, and a 3 out of 4 chance (75% chance) that option Y leaves your probability of winning the prize unchanged. This means that when the large dot is **red**, on average option Y increases your probability of winning the prize by 0.25 percentage points.

When the large dot is **green**, this is a good sign for the value of option Y, and there is a 3 out of 4 chance (75% chance) that option Y increases your probability of winning the prize by 1 percentage point, and only a 1 out of 4 chance (25% chance) that option Y leaves your probability of winning the prize unchanged. This means that when the large dot is **green**, on average option Y increases your probability of winning the prize by 0.75 percentage points.

In each round there is a 3 out of 4 chance (75% chance) the big dot is **red**, and a 1 out of 4 chance (25% chance) the big dot is **green**. In the first 8 rounds (rounds 1 through 8) the large dot is the only source of information available to you.

Next

Figure 5:

Training (2 of 3):

In the last 92 rounds (rounds 9 through 100), you can get more information if you choose to. When you choose to 'Get more information,' you are shown 100 small dots, each of which is **red** with a black R in the middle, or **green** with a black G in the middle. Together, the small dots tell you for sure whether or not option Y increases your probability of winning the \$25 prize by 1 percentage point in the round you are in.

In each round that you can 'Get more information,' there are either 49 small **red** dots and 51 small **green** dots, or 51 small **red** dots and 49 small **green** dots.

If there are 49 small **red** dots and 51 small **green** dots, then option Y increases your probability of winning the prize by 1 percentage point.

If there are 51 small **red** dots and 49 small **green** dots, then option Y does not increase your probability of winning the prize.

Under the dotted line below is an example of a round where you can choose to get more information. Notice that the round number is displayed, and the amount that option X increases your probability of winning the prize is displayed.

The large dot is **green**, so you know there is a 3 out of 4 chance (75% chance) option Y increases your probability of winning the prize by 1 percentage point, but you cannot know for sure unless you choose to 'Get more information.'

Please select 'Get more information' so you see what happens when you do.

.....

Round number: 0 of 100

Option X increases your probability of winning the \$25 prize by 0.5 percentage points

Please choose an option:

Remember:

If the large dot is **red**, then there is a 3 out of 4 chance (75% chance) option Y increase your probability of winning the prize by 0 percentage points, and a 1 out of 4 chance (25% chance) it increases your probability of winning the prize by 1 percentage point. This means that when the large dot is **red**, on average option Y increases your probability of winning the prize by 0.25 percentage points.

If the large dot is **green**, then there is a 3 out of 4 chance (75% chance) option Y increases your probability of winning the prize by 1 percentage point, and a 1 out of 4 chance (25% chance) it increases your probability of winning the prize by 0 percentage points. This means that when the large dot is **green**, on average option Y increases your probability of winning the prize by 0.75 percentage points.

If you choose to 'Get more information,' 100 small dots will appear that you can use to determine how much option Y increases your probability of winning the prize.

Next

Figure 6:

Training (3 of 3):

Under the dotted line below is an example of what you would see after requesting to 'Get more information' in a round. Selecting this option has made the 100 small dots for this round appear.

Try counting the number of small red dots and or the number of small green dots below. You should find there are 49 red dots and 51 green dots, which means option Y would increase your probability of winning the prize by 1 percentage point. So in this round, option Y increases your probability of winning the prize by more than option X. This training round (Round number 0) does not count towards your eventual probability of winning the prize.

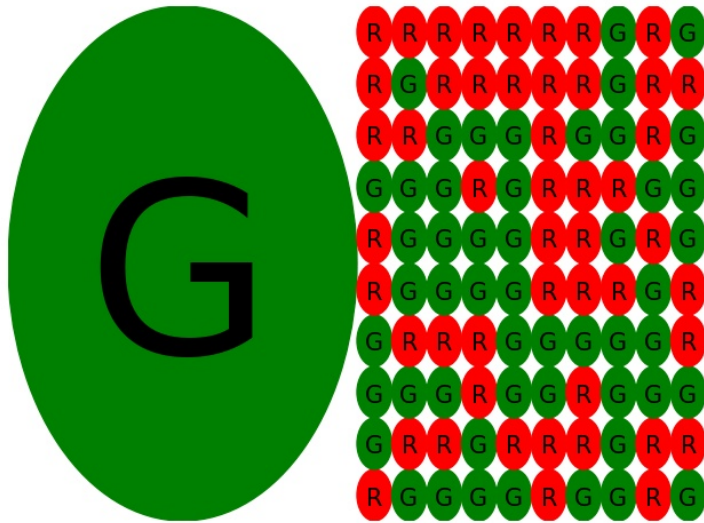
After you have finished the quiz, if you want to quit the experiment you will be given the option to do so at the bottom of the screen in each round. When you choose to exit, you still have the chance to win the \$25 prize, you do not lose any probability you have acquired of winning the prize, but you do lose the chance to further increase your probability of winning the prize.

.....

Round number: 0 of 100

Option X increases your probability of winning the \$25 prize by 0.5 percentage points

Please choose an option:

Remember:

In each round, option Y always increases your probability of winning the prize by either 0 or 1 percentage points.

In each round, there are always either 49 small red dots and 51 small green dots, or 51 small red dots and 49 small green dots.

If there are 49 small red dots and 51 small green dots, then option Y increases your probability of winning the prize by 1 percentage point.

If there are 51 small red dots and 49 small green dots, then option Y increases your probability of winning the prize by 0 percentage points.

Next

Figure 7:

Quiz time! (Question 1 of 10)

It is now time to do the quiz. Remember, to move on to the next question you must first answer the current question correctly. Completing the quiz earns you the \$5 training payment, and in addition to the training payment, you can earn another \$0.5 for each question you get correct on your first attempt to answer it.

Which of the following is correct?

Option X sometimes increases your probability of winning the prize by more than 1 percentage point

Sometimes selecting the wrong option can reduce your probability of winning the prize

In each round, option Y always increases your probability of winning the prize by 1 percentage point

In each round, option Y either increases your probability of winning the prize by 1 percentage point, or does not change it

Figure 8:

Quiz time! (Question 2 of 10)

If you see the big dot below in a round, which of the following is NOT correct?

If no other information is available then you must be in the first 8 rounds

If you are in the last 92 rounds you can get more information

If no other information is available then there is a 3 out of 4 chance that option Y increases your probability of winning the prize by 1 percentage point

If you are in the first 8 rounds you can get more information



Figure 9:
 Quiz time! (Question 3 of 10)

In the last 92 rounds, if you see the the big dot below, if you choose to 'Get more information,' how many small green dots with black Gs will be displayed?

All numbers between 0 and 100 occur with a strictly positive probability
 Either 49 or 51
 51
 100



Figure 10:
 Quiz time! (Question 4 of 10)

In the image below there are 51 small green dots. If you see such an image in a round, what happens when you select option Y?

Your probability of winning the prize stays the same
 Your probability of winning the prize goes up by 1 percentage point
 One of several things may happen with a strictly positive probability
 Your probability of winning the prize goes up by less than if you selected option X

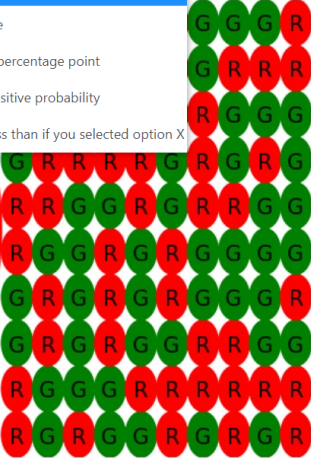



Figure 11:
Quiz time! (Question 5 of 10)

If the big dot below is visible in a round, which of the following is correct?

▼

If you are in the first 8 rounds you can get more information

If no other information is available then there is a 1 out of 4 chance that option Y increases your probability of winning the prize by 1 percentage point

If no other information is available then you must be in the last 92 rounds

If no other information is available then there is a 3 out of 4 chance that option Y increases your probability of winning the prize by 1 percentage point



Figure 12:
Quiz time! (Question 6 of 10)

In the image below there are 51 small red dots. If you see such an image in a round, what happens when you select option Y?

▼

Your probability of winning the prize goes up by 1 percentage point

Your probability of winning the prize stays the same

Your probability of winning the prize goes up by more than if you selected option X

It depends on the round number

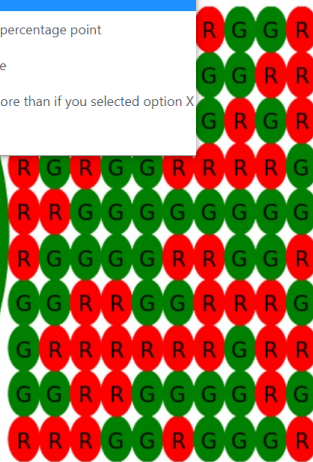


Figure 13:
Quiz time! (Question 7 of 10)

Count the number of small green dots in the image below. If you see such an image in a round, what happens when you select option Y?

It depends on the round number

Your probability of winning the prize goes up by less than if you selected option X

Your probability of winning the prize stays the same

Your probability of winning the prize goes up by more than if you selected option X

Figure 14:
Quiz time! (Question 8 of 10)

When you choose to 'Get more information' in a round, which of the following is correct?

Selecting option X in the same round causes your probability of winning the prize to stay the same

100 small dots appear that together tell if option Y increases your probability of winning the prize by 1 percentage point

Your probability of winning the prize goes up by 1 percentage point

Your probability of winning the prize may decrease

Figure 15:
Quiz time! (Question 9 of 10)

If the large dot below is the only information available, then which of the following is correct?

On average, option Y increases your probability of winning the prize by 0.25 percentage points

On average, option Y increases your probability of winning the prize by 0 percentage points

On average, option Y increases your probability of winning the prize by 1 percentage point

On average, option Y increases your probability of winning the prize by 0.75 percentage points



Figure 16:
Quiz time! (Question 10 of 10)

If the large dot below is the only information available, then which of the following is correct?

On average, option Y increases your probability of winning the prize by 0.25 percentage points

On average, option Y increases your probability of winning the prize by 0.75 percentage points

On average, option Y increases your probability of winning the prize by 0 percentage points

On average, option Y increases your probability of winning the prize by 1 percentage point



Figure 17:

Remember:

If the large dot is **red**, then there is a 3 out of 4 chance (75% chance) option Y increases your probability of winning the prize by 0 percentage points, and a 1 out of 4 chance (25% chance) it increases your probability of winning the prize by 1 percentage point. This means that when the large dot is **red**, on average option Y increases your probability of winning the prize by 0.25 percentage points.

If the large dot is **green**, then there is a 3 out of 4 chance (75% chance) option Y increases your probability of winning the prize by 1 percentage point, and a 1 out of 4 chance (25% chance) it increases your probability of winning the prize by 0 percentage points. This means that when the large dot is **green**, on average option Y increases your probability of winning the prize by 0.75 percentage points.

If you choose to 'Get more information' 100 small dots will appear that you can use to determine how much option Y increases your probability of winning the prize.

Next

Figure 18:

For your information: for the next 40 rounds option X will increase your probability of winning the prize by 0.25 percentage points

For your information: of the 92 rounds remaining, 12 of the rounds have an option X that increases your probability of winning the prize by 0.5 percentage points, and 80 of the rounds have an option X that increases your probability of winning the prize by 0.25 percentage points.

Remember: in all remaining rounds you can choose to 'Get more information' if you want to.

Next

Figure 19:

For your information: for the next 12 rounds option X will increase your probability of winning the prize by 0.5 percentage points

For your information: of the 52 rounds remaining, 12 of the rounds have an option X that increases your probability of winning the prize by 0.5 percentage points, and 40 of the rounds have an option X that increases your probability of winning the prize by 0.25 percentage points.

Remember: in all remaining rounds you can choose to 'Get more information' if you want to.

Next