



Belief adjustment: a double hurdle model and experimental evidence

Timo Henckel^{1,5} · Gordon D. Menzies^{2,5} · Peter G. Moffatt³ · Daniel J. Zizzo^{4,5}

Received: 4 May 2019 / Revised: 19 November 2020 / Accepted: 20 January 2021 /
Published online: 28 April 2021
© The Author(s) 2021

Abstract

We present an experiment where subjects sequentially receive signals about the true state of the world and need to form beliefs about which one is true, with payoffs related to reported beliefs. We attempt to control for risk aversion using the Offerman et al. (Rev Econ Stud 76(4):1461–1489, 2009) technique. Against the baseline of Bayesian updating, we test for belief adjustment underreaction and overreaction and model the decision making process of the agent as a double hurdle model where agents with inferential expectations first decide whether to adjust their beliefs and then, if so, decide by how much. We also test the effects of increased inattention and complexity on belief updating. We find evidence for periods of belief inertia interspersed with belief adjustment. This is due to a combination of random belief adjustment; state-dependent belief adjustment, with many subjects requiring considerable evidence to change their beliefs; and quasi-Bayesian belief adjustment, with aggregate insufficient belief adjustment when a belief change does occur. Inattention, like complexity, makes subjects less likely to adjust their stated beliefs, while inattention additionally discourages full adjustment.

Keywords Belief revision · Expectations · Overreaction · Underreaction · Inattention · Complexity

JEL Classification C34 · C91 · D03 · D84 · E03

✉ Daniel J. Zizzo
d.zizzo@uq.edu.au

Peter G. Moffatt
p.moffatt@uea.ac.uk

- ¹ Research School of Economics, Australian National University, Canberra, Australia
- ² University of Technology Sydney, Sydney, Australia
- ³ CBESS and School of Economics, University of East Anglia, Norwich, UK
- ⁴ School of Economics, University of Queensland, St Lucia, Australia
- ⁵ Centre for Applied Macroeconomic Analysis, Canberra, Australia

1 Introduction

Agents form and update their beliefs when they receive new information. In the presence of rational expectations, new information leads to belief updating every period according to Bayes rule. In reality many agents do not behave according to basic statistics, and task complexity and inattention may contribute to deviations from Bayesian predictions. Such violations of rational expectations have been studied in static settings, where all the information is presented at once to subjects who discount priors (Kahneman and Tversky 1973; Tversky and Kahneman 1982; El-Gamal and Grether 1995). In this paper we study a dynamic setting in which new information arrives sequentially and consider the *frequency* as well as the *extent* of belief adjustment, referring to *sticky* belief adjustment when it is insufficient in either domain.

Relative to previous urn experiments where agents need to state guesses, such as Khaw et al. (2017), our study is innovative in a number of dimensions. First, it is simple, in that our experiment does not have an evolving state of nature. We are interested instead in the basic question of how beliefs are updated dynamically, and this basic question can be answered in a way that is easiest for experimental participants and most interpretable for researchers by having a simple dynamic environment with new information flowing in. Second, we attempt to control for risk aversion and therefore are able to measure beliefs more accurately than in previous research. We try to do so by using the Offerman et al.'s (2009) technique. Third, we do not use cumulative earnings, which may lead to uncontrolled factors such as income effects or portfolio diversification.

Fourth, we develop a double hurdle econometric model to combine in a single framework different types of belief adjustment we may observe in the laboratory: time-dependent (random) belief adjustment and state-dependent (Bayesian, Quasi-Bayesian) belief adjustment.¹ Each type has been utilized in macro and microeconomic modelling, though the tendency has been to focus on only one type.

Within macroeconomic research, sticky belief adjustment can be seen as a possible microfoundation of sticky price adjustment, for example as a result of inattention and observation costs (Alvarez et al. 2016), information costs (Abel et al. 2013), cognitive costs (Magnani et al. 2016) and the consultation of experts by inattentive agents (Carroll 2003), or some combination of these factors (Cohen et al. 2019). State-dependence in beliefs implies a dependence of belief adjustment on the economic state, which in turn may depend on new information flowing in. Time-dependence in beliefs is often viewed stochastically [as for example in Caballero (1989)] and therefore yields random belief adjustment, following some underlying data generating process. One useful way of conceptualizing state- and

¹ None of these points should be interpreted as criticisms of regime switching studies such as Khaw et al. (2017). The focus of these studies is on getting a better understanding of how agents handle regime switching, and so an evolving economy or a more specialised modelling approach are appropriate for what they are trying to achieve, and potentially confounding factors (such as risk aversion) are less of an issue than they are when trying to get a basic understanding of belief updating. The latter is our different and more fundamental focus.

time-dependent beliefs is the inferential expectations (IE) model of Menzies and Zizzo (2009): that is, agents hold a belief until enough evidence has accumulated for a statistical test of a given test size α to become significant, at which point beliefs switch. Furthermore, if agents' stance towards evidence is modelled by a probabilistic draw on their α , then the probability that the test size is unity (which results in an update for any evidence, or none²) is the probability of a random belief adjustment.

Within microeconomic research, Quasi-Bayesian (QB) belief adjustment has been the preferred route to think about bounded-rational belief adjustment. Rabin (2013) distinguishes between warped Bayesian models which encapsulate a false model of how signals are generated, for example ignoring the law of large numbers (Benjamin et al. 2015); and information-misreading Bayesian models that misinterpret signals as supporting agents' hypotheses, thus giving rise to confirmation bias (Rabin and Schrag 1999), and therefore lead to underweighting of information (for early evidence, see Phillips and Edwards 1966). In static problems where priors and 'new information' were given, Kahneman and Tversky (1973) and Tversky and Kahneman (1982) made the contrasting finding of base rate neglect, with more weight being put on the new information; reviews of the literature on base rate neglect can be found in Koehler (1996), Barbey and Sloman (2007) and Benjamin (2019). One simple way of modelling QB adjustment, which we follow, is that the agent adjusts beliefs every period in response to new information, but this adjustment is either too big or too small (Massey and George 2005; Ambuehl and Li 2014). That is, if the posterior probability is the prior multiplied by (*likelihood*) ^{β} , Quasi-Bayesian models are marked by departures of β from unity.³

In this paper we use a double hurdle model to consider different perspectives about belief adjustment emphasized within macro- and microeconomics. Full rationality requires clearing two hurdles in a very specific way: fully rational agents must adjust every period as they clear hurdle 1, and they must use Bayes rule with $\beta = 1$ as they clear hurdle 2.

Table 1 describes the first hurdle using different values of the IE test size (α) and the columns describe the second hurdle using different values of the QB parameter β . Fully rational agents are fully attentive ($\alpha = 1$) and Bayesian ($\beta = 1$). The double hurdle model is formalized in Sect. 4.2 and generates a distribution for α and β , parameterizing both the frequency and extent of adjustment.

Fifth, we provide the first study that looks at how increased task complexity or scope for inattention affects belief updating. Task complexity and inattention are two factors that have been independently identified as playing a potentially substantial role in bounded-rational decision making.

² A standard (Neyman–Pearson) hypothesis test minimizes the probability of falsely believing a null (a type II error) subject to an upper bound on the probability of falsely rejecting that same null (a type I error). If α is unity, the constraint does not bind, so the best way to minimize the probability of falsely believing a null is to automatically (for any evidence against it, or none) reject the null. Under IE, rejecting a null is the same as 'updating' because a new value has to be chosen.

³ In Bayes rule, $P(A|B) = [P(B|A)/P(B)]P(A)$, the ratio $[P(B|A)/P(B)]$ is sometimes called the likelihood.

Among others, Simon (1979), Gigerenzer and Gaissmaier (2011) and Caplin et al. (2011) have identified complexity of decision settings as a key reason for ‘satisficing’ and heuristic-based decision making. This is neurobiologically plausible (Bossaerts and Murawski 2017) and leads to different ways information is processed (Payne 1979) and lotteries selected in binary choices (Wilcox 1993). Examples of practical applications where complexity can be important is consumer exploitation by firms to achieve greater profits (Carlin 2011; Huck et al. 2011; Sitzia and Zizzo 2011; Sitzia et al. 2015); decisions to engage in vertical integration or outsourcing (Tadelis 2002); climate change inaction as linked to the complexity of the relevant task environment (Slawinski et al. 2017); defaults becoming more attractive as an omission bias (e.g., Baron and Ritov 2004).

Inattention has been independently identified as a key source of bounded-rational decision making (e.g., Alvarez et al. 2016; Magnani et al. 2016; Carroll 2003), in ways that may but do not necessarily reflect rational inattention trade-offs. (See Caplin et al. 2020, for a discussion of this point.) It has a wide ranging and growing set of applications. Examples of applications in macroeconomics include the New Keynesian Philips Curve (Mankiw and Reis 2002), business cycle dynamics (Mackowiak and Wiederholt 2015) and the failure of uncovered interest rate parity (Bacchetta and van Wincoop 2010). Examples of applications in microeconomics include strategic product pricing (Martin 2017), corporate strategy (Dessein et al. 2016) and portfolio selection (Huang and Liu 2007).

Surprisingly, given the importance that both complexity and inattention have been stated to have in a wide range of settings with risk and imperfect information, we are not aware of papers that have looked at the effect of either on belief updating. A key contribution of this paper is to address this gap.

Regarding task complexity, we expect it to potentially reduce the frequency of stated belief changes in the first hurdle (α in Table 1) as well as the extent of the stated belief change when this takes place in the second hurdle (β in Table 1). This is because complexity makes subjects less likely to wish to make an ‘active’ choice and therefore more likely to stick to the default (see Gerasimou 2018); and because, if they do change their stated beliefs, as they perceive the task as more uncertain, they are likely to be more conservative in the degree to which they do so (see Brainard 1967).

Regarding inattention, if it matters in the way that the literature has suggested, then a simple experimental manipulation increasing the likelihood of inattention will lead experimental subjects to be less likely to update their beliefs regarding the variable to which they are not paying attention. Importantly, while this would not be surprising if the alternative distracting task were incentivized, in our experiment (as in Sitzia et al. (2015)) it is not. There is therefore unlikely to be any preference-based reason why a rational agent should ignore the guessing task on the basis of which payments are wholly made, and deviations from Bayes can be more precisely identified as being due to cognitive costs in information processing. Inattention should reduce the likelihood of agents switching their belief, and therefore enter the first hurdle of the model.

In brief, our results are as follows. Subjects change their beliefs about half the time, which is consistent with random belief adjustment, but they also consider the

Table 1 Updating taxonomy

Hurdle 1: Decision to Update		Hurdle 2: Extent of Update			
		Interpretation of Quasi-Bayesian β in $Posterior = (LikelihoodRatio)^\beta Prior$			
		Irrational $\beta < 0$	Underuse Information $0 \leq \beta < 1$	Bayes $\beta = 1$	Overuse Information $\beta > 1$
Inferential Expectations α is attention towards evidence	$\alpha = 0$ (Inattentive)				
	$0 < \alpha < 1$ (Partly attentive)				
	$\alpha = 1$ (Fully attentive)			Rational Expectations	

amount of evidence available, which is consistent with state-dependent belief adjustment. When subjects do change beliefs, they do so by around 80 per cent of the full Bayesian update, which is consistent with our version of Quasi-Bayesian belief adjustment. There is substantial heterogeneity in our results and the frequency and extent of belief adjustment are negatively correlated: agents who update with low frequency do so by more than 80 percent of the full Bayesian update. Furthermore, we find evidence that inattention reduces the propensity to update, as predicted, as well as the extent of update. Complexity is less important, as it only affects the propensity to update and does so by less than inattention. We do not find that task confusion explains belief stickiness to an important degree, nor is there any financial incentive to explain why beliefs are stickier if we add an alternative distracting task. Rather, inattention and cognitive costs are likely to explain the infrequent belief adjustment, to different degrees, by half of our subjects. Only a small fraction of agents have rational expectations where that is understood as full Bayesian updating each period.

Our paper is structured as follows: in Sect. 2 we construct a balls-and-urn experiment with treatments for complexity and inattention. Section 3 describes the balls-and-urn environment, and predicts behavior, under different expectational assumptions. Section 4 analyzes the experimental results using: nonparametric (model free) statistics for the raw data; the double hurdle econometric model for the risk-adjusted data; and a subject-specific density of test sizes derived from the double hurdle model. Section 5 draws together the main results, and concludes.

2 Experimental design and treatments

Our experiment was fully computerized in JavaScript and run with undergraduate and postgraduate students in the experimental laboratory of the University of East Anglia with $n = 245$ subjects in 16 sessions conducted between July and December 2013.⁴ Everyone in each session participated in the same treatment, and sessions were conducted in mixed order. ORSEE was used as experimental recruitment software for the assignment of subjects to sessions. Subjects were separated by partitions. The experiment was divided in two parts, labelled the *risk attitude part* (Stage 1) and the *main part* (Stage 2).

⁴ Dates, times and treatments of experimental sessions are listed in online appendix 2.

Experimental instructions were provided at the beginning of each part for the tasks in that part. Online appendices 3 and 4 contain a copy of the instructions; a file with more details on the software and with computer screens is also provided as supplementary material. A questionnaire was administered to ensure understanding after each batch of instructions. If a subject got an answer wrong, a brief and simple explanation was provided explaining the correct answer (see the computer screens file for the text in each case) and, if anything was still unclear, subjects were given the opportunity to obtain further clarification from an experimenter.

2.1 Main part of the experiment

After playing the risk attitude part described in more detail below, in the main part of the experiment subjects played 7 stages, each with 8 rounds, thus generating $T = 56$ observations. At the beginning of each stage the computer randomly chose one of two urns (Urn 1 or Urn 2), with Urn 1 being selected at a known probability of 0.6. Each urn represents a different state of the world. While this prior probability was known and it was known that the urn would remain the same throughout the stage, the chosen urn was not known to subjects. It was known that Urn 1 had seven white balls and three orange balls, and Urn 2 had three white balls and seven orange balls. At the beginning of each of the 8 rounds (round = t), there was a draw from the chosen urn (with replacement) and subjects were told the color of the drawn ball. These were therefore signals that could be used by subjects to update their beliefs.⁵ It was made clear to the subjects that the probability an urn was chosen in each of the seven stages was entirely independent of the choices of urns in previous stages. A visual representation of the urns was provided on the computer screens to facilitate understanding (see the computer screens file).

Once they saw the draw for the round, subjects were asked to make a probability guess between 0 and 100%, on how likely it was that the chosen urn was Urn 1. The corresponding variable for analysis is their probability guess expressed as a proportion, denoted g . Once a round was completed, the following round started with a new ball draw, up to the end of the 8th round.

Payment for the main part of the experiment was based on the guess made in a randomly chosen stage and round picked at the end of the experiment. A standard quadratic scoring rule (e.g. Davis and Holt 1993) was used in relation to this round to penalize incorrect answers. The payoff for each subject was equal to 18 GBP minus $18 \text{ GBP} \times (\text{guess} - \text{correct probability})^2$. Therefore, for the randomly chosen stage and round, subjects could earn between 0 and 18 GBP depending on the accuracy of their guesses. It was clarified to subjects that “if the chosen urn was Urn 1, then the correct probability of the chosen urn being Urn 1 is 100%; if the chosen urn was Urn 2, then the correct probability of the chosen Urn being Urn 1 is 0%”. While instructions were generally provided on the computer screen, a table with payoffs for each level of accuracy of the guesses was provided in print, to facilitate

⁵ All signals are thus informative. It would be an interesting extension to include non-informative signals in future research.

understanding (see online appendix 3). Table 8 in “[Appendix 5: Robustness and understanding](#)” includes a regression model with maths ability as an explanatory variable, which we measure in the C treatment. This variable is insignificant even in the treatment where potentially it should have mattered the most, which undermines its relevance. Furthermore, subjects could make use of a calculator. Specifically, a ‘calculate consequences’ button gave subjects ready information about the payoffs arising from their guess, depending on which urn was drawn.

2.2 Risk attitude part of the experiment

The risk attitude part was similar to the main part but simpler and therefore genuinely useful as practice. It was modelled after Offerman et al. (2009) to enable us to infer people’s risk attitude, as detailed in Sect. 3.

It consisted of 10 stages with one round each. In each stage a new urn was drawn (with probabilities 0.05, 0.1, 0.15, 0.2, 0.25, 0.75, 0.8, 0.85, 0.9, 0.95).⁶ Subjects were told the prior probability of Urn 1 being chosen but did not receive any further information. In particular, no balls were drawn. The guessing task (single round) was to nominate a probability that Urn 1 was chosen, where subjects could rely on the information available to them and the payment mechanism to identify which probability guess would maximize their expected utility. For subjects who are not risk neutral, the task is non-trivial as they should take account of the payoff structure rather than repeat the announced probabilities. Payment for the risk attitude of the experiment was based on the guess made in a randomly chosen round picked at the end of the experiment. A quadratic scoring rule was applied as in the main part, but this time this was equal to 3 GBP minus $3 \text{ GBP} \times (\text{guess} - \text{correct probability})^2$.⁷ Again, it was clarified to subjects that the correct probability was 0 or 1 depending on the urn being chosen, and again a table with payoffs for each level of accuracy of the guesses was provided in print (see online appendix 3) and a calculator was also available.

⁶ One might wonder why we did not use 0.6 as a possible value. The purpose of the risk attitude part is to estimate θ by the auxiliary regression in “[Appendix 3: Method for estimating CRRA risk parameter](#)”, and so it is desirable to have values close to 0 and 1, since the standard error in a regression is inversely related to the standard deviation of the independent variable. Adding extra values in the middle of the range of the independent variable may be useful for some purposes, but it would have a low leverage and therefore not help much in this context. This is not to deny that more data is always better than less, but 0.6 was not a natural choice for extra data. Furthermore, at the point in time when the subjects do the risk attitude part, the number 0.6 had no special significance for them; and we wished to avoid the danger of subjects potentially artificially anchoring in the main part to what they had done under 0.6 in the risk attitude part.

⁷ This ensured similar marginal incentives for each round in the risk attitude part (3 GBP prize picked up from 1 out of 10 rounds) and the main part (18 GBP prize picked up from 1 out of 56 rounds).

2.3 Experimental treatments

There were three treatments. The risk attitude parts were identical across all treatments, and the main part of the *Baseline* treatment (B) was as described.

In the main part (only) of the *Complexity* treatment (C), the information on the ball drawn from the chosen urn at the beginning of each round was presented as a statement about whether the sum of three numbers (of three digits each) is true or false. If true (e.g., $731 + 443 + 927 = 2101$), this meant that a white ball was drawn. If false (e.g., $731 + 443 + 927 = 2121$), this meant that an orange ball was drawn.

In the main part (only) of the *Inattention* treatment (I), subjects were given a non-incentivized alternative counting task which they could do instead of working on the probability. The counting task was a standard one from the real effort experimental literature (see Abeler et al. 2011, for an example) and consisted in counting the number of 1s in matrices of 0s and 1s. Subjects were told that they could do this exercise for as little or as long as they liked within 60 s for each round, and that we were not asking them in any way to engage in this exercise at all unless they wanted to.⁸

As in Caplin et al. (2020) and the key treatments in Sitzia et al. (2015), we see as important not to incentivize the alternative task. If the alternative task were incentivized—financially or in terms of doing something fun such as browsing the internet—it would be rational for an agent to split his or her time allocation between tasks, which would trivially imply worse decision making in the guessing task. This would make it difficult to precisely identify what is due to inattention as a psychological mechanism, and could be construed to simply reflect the fact that agents are bad at multitasking (e.g., Buser and Peter 2012). In our setting, instead, agents should just focus on a single task and not be distracted by the alternative task.⁹ Undoubtedly, future research could change the incentives associated to the alternative task.

3 Theoretical model

3.1 Model variables and risk attitude correction

In this section, we build a model of subject play in our experiment where the key driver is the way subjects form beliefs. Table 2 lays out the main variables from the experiment and the interrelationships between them, when the event is described in terms of the chosen urn (row 1) and when it is described in terms of the probability of a white ball being drawn (row 2). The two descriptions are equivalent since

⁸ They were also told that, if they did not make a guess in the guessing task within 60 s, they would automatically keep the guess from the previous round and move to the next round (or to the next stage). The length of 60 s was chosen based on piloting, in such a way that this would not be a binding constraint if subjects focused on the guessing task.

⁹ In their experiment on the selection of energy tariffs, Sitzia et al. (2015) did not find different results when the alternative task was the ability to browse the internet instead of a counting task like the one in this paper.

Table 2 Model variables

Event	Estimator of event probability	Estimator symbol	$P(\text{Event})$ Subject guesses	Transformed guess	Strength of evidence z_t against earlier choice at m
Urn 1 drawn	Bayes rule	P_t	g_t : optimal guess (observed) g_t^* : true guess (inferred)	$r_t^* = \Phi^{-1}(g_t^*)$	$z_t = \frac{2(\Phi^{-1}(P_t) - \Phi^{-1}(P_m))}{\sqrt{t}}$
White drawn	Prop. white balls out of t	P_t^w	Not guessed	Not guessed	$z_t \approx \frac{P_t^w - P_m^w}{\sqrt{0.5^2/t}}$

the subject's subjective guess of the probability that Urn 1 was chosen generates an implied subjective probability that a white ball is drawn.¹⁰ In our modelling of this experimental environment, we sometimes use the former probability—that Urn 1 was chosen—and it will be useful to transform this probability guess using the inverse cumulative Normal distribution (so that the support has the same dimensionality as a classic z -statistic). Alternatively, we sometimes describe agents' guesses in terms of the probability that a white ball is drawn because the sample proportion from repeated Bernoulli trials has a tractable sampling distribution for hypothesis testing.

Time is measured by t , the draw (round) number for the ball draws in each stage. We define the value of t for which subjects last moved their guess (viz. updated their beliefs) to be m (for 'last move'). Thus, for any sequence of ball draws at time t , the time that has elapsed since the last change in the guess is always $t - m$.

Along the top row the theoretical estimator for the probability that Urn 1 was drawn is provided by Bayes rule, which we denote by P_t after t ball draws. Many subjects do not use Bayes rule when they are guessing the probability that Urn 1 is chosen, though some guesses are closer to it than others.

As derived in Offerman et al. (2009), the elicited guess g_t in the fourth column is the result of maximizing expected utility based on a Constant Relative Risk Aversion (CRRA) utility function, $U\{\text{Payoff}\}$:

$$U\{\text{Payoff}\} = \frac{\text{Payoff}^{1-\theta} - 1}{1-\theta},$$

and a true guess g_t^* in $E[U\{\text{Payoff}\}] = g_t^* U\{1 - (1 - g_t)^2\} + (1 - g_t^*) U\{1 - g_t^2\}$ where the payoffs for Urn 1 and Urn 2 are proportional to $1 - (1 - g_t)^2$ and $1 - g_t^2$, according to the quadratic scoring rule, as explained in Sect. 2.¹¹ Expected utility is assumed to be maximized with respect to g_t and yields the following relationship between g_t^* and g_t :

$$\ln\left(\frac{g_t^*(1 - g_t)}{g_t(1 - g_t^*)}\right) = \theta \ln\left(\frac{g_t(2 - g_t)}{(1 + g_t)(1 - g_t)}\right). \quad (1)$$

In the risk attitude part, the prior probabilities given to the subjects (by way of reminder, 0.05, 0.1, 0.15, 0.2, 0.25, 0.75, 0.8, 0.85, 0.9, 0.95 for 10 separate stages/rounds) are in fact the correct probabilities P_t . We see no reason not to credit subjects with realizing this, and they possess no other information anyway, so we define their true guess to be $g_t^* = P_t$. Offerman et al. (2009) then interpret the deviations of g_t from g_t^* as being due to the subjects' risk preferences, and so do we. We use the ten datapoints (g_t^*, g_t) for each subject to estimate θ in a version of (1) appended with

¹⁰ For example, at the start of the experiment, before any ball is drawn, subjects know that the chance that Urn 1 was drawn is 0.6. It therefore follows that the chance of a white ball being drawn for the very first time is $0.7 \times 0.6 + 0.3 \times (1 - 0.6) = 0.54$.

¹¹ It is straightforward to derive the equivalent of (1) for CARA utility, and later in "Appendix 5: Robustness and understanding" the CARA g_t^* s are used in a robustness check.

a regression error.¹² Armed with a subject-specific value of θ from the risk attitude part, all the observable g_t values in the main experiment can be transformed to a set of inferred g_t^* . This transformation is accomplished by exponentiating both sides of (1), and solving for g_t^* . By taking the inverse cumulative Normal function, Φ^{-1} , of g_t^* we move it outside the $[0, 1]$ interval and give it the same dimensionality as a z-test statistic, namely $(-\infty, \infty)$. The variable in the penultimate column, $r_t^* = \Phi^{-1}(g_t^*)$, thus becomes the basis for our econometric analysis.¹³

We will explain the last column of Table 2 and why it is useful later.

3.2 Expectation processes

Using the notation of Table 2, we define three processes of expectation formation that will be relevant for our double hurdle model in Sect. 4.

3.2.1 Rational expectations

The rational expectations solution predicts straightforward Bayesian updating. The (conditional) probability the subject is being asked to guess is the rational expectation (RE), which is given by P_t . Calling $P_{initial}$ the initial prior probability and noting that the number of white balls is tP_t^w we can write down P_t in a number of ways:

$$P_t = \left(\frac{(0.7)^{tP_t^w} (0.3)^{t-tP_t^w}}{(0.7)^{tP_t^w} (0.3)^{t-tP_t^w} P_{initial} + (0.3)^{tP_t^w} (0.7)^{t-tP_t^w} (1 - P_{initial})} \right) P_{initial} \quad (2)$$

$$= \frac{1}{1 + \frac{(0.3)^{tP_t^w} (0.7)^{t-tP_t^w} (1 - P_{initial})}{(0.7)^{tP_t^w} (0.3)^{t-tP_t^w} P_{initial}}}.$$

The second line is a useful simplification (which we use in “[Appendix 1: Closeness of two strength-of-evidence measures](#)”) whereas the bracketed fraction in the first line is the probability of obtaining the tP_t^w white balls when Urn 1 is drawn versus the total probability of obtaining this number of white balls. We called this the likelihood ratio in Table 1.

3.2.2 Quasi-Bayesian updating

In our version of Quasi-Bayesian updating (QB), agents use Bayesian updating as each new draw is received, but they incorrectly weight the likelihood ratio:

¹² See “[Appendix 3: Method for estimating CRRA risk parameter](#)” for details. The estimated mean θ across subjects is 0.2. There is considerable heterogeneity across subjects, so in our econometric modeling we check for robustness by excluding subjects with extreme values ($|\theta| > 1.5$).

¹³ The non-parametric analysis of Sect. 4.1 uses unadjusted data, while the econometric analysis of Sect. 4.2 uses the risk-adjusted data.

$$P_t^{QB} = \left(\frac{(0.7)^{tP_t^w} (0.3)^{t-tP_t^w}}{(0.7)^{tP_t^w} (0.3)^{t-tP_t^w} P_{initial} + (0.3)^{tP_t^w} (0.7)^{t-tP_t^w} (1 - P_{initial})} \right)^\beta P_{initial} \quad (3)$$

The parameter β may be thought of as the QB parameter: if $\beta = 1$, agents are straightforward Bayesians; if $\beta > 1$ they overuse information and under-weight priors; if $0 \leq \beta < 1$ they underuse information and over-weight priors and if $\beta < 0$ they respond the wrong way to information—raising the conditional probability when they should be lowering it, and vice versa.

Agents' attitude towards the extent of belief change in the light of evidence can be summarized by the distribution $f(\beta)$ across subjects. If $f(\beta)$ has most probability mass between 0 and 1, most agents only partially adjust, and subjects converge to full adjustment at $\beta = 1$ to the extent that the probability mass in $f(\beta)$ converges towards unity.

3.2.3 Inferential expectations

Cohen et al. (2019) show that models with cost-based state-dependent sticky belief adjustment are equivalent to an inferential expectations (IE) model, where agents' hypothesis testing generates infrequent belief adjustment (Menzies and Zizzo 2009).¹⁴ We show this in our specific context in “Appendix 2: Relationship between inferential expectations and switching cost models”. We therefore are in a position to model the degree of state-dependent belief stickiness by the test size, where a low test size implies a higher cost of changing beliefs and therefore relatively infrequent adjustment.

We assume subjects hold a belief until enough evidence has accumulated to pass a threshold of statistical significance, at which point beliefs are updated. Agents form a belief and do not depart from that belief until the weight of evidence against the belief is sufficiently strong. Under IE, each agent is assumed to start with a belief about the probability of U (that is, $P_0 = 0.6$) and its implied probability of a white ball ($P_0^w = 0.54$), and conducts a hypothesis test that the latter is true after drawing a test size from his or her own distribution of α , namely $f_i(\alpha)$. Agents are assumed to draw this every round during the experiment.

In the first row of the final column of Table 2, we provide a measure z_t of the strength of evidence against the probability guess at the time of the last change. Agents change their guesses from time to time, and z_t tells us if the value of P at the last change, denoted P_m , seems mistaken in the light of subsequent evidence.

Importantly, as shown in “Appendix 1: Closeness of two strength-of-evidence measures”, the top-row z_t is approximately equal to the standard test statistic for a

¹⁴ Cohen et al.'s (2019) explanation relies on a cost of adjustment. This must be interpreted as a cognitive cost, since there is no financial penalty for adjustment in our experiment. A cognitive cost might also be described as ‘laziness’. However, we should note that cognitive effort in attention may well take place, rather than the cognitive cost being simply about lack of effort as the notion of ‘laziness’ and its associated implicit moral censure seems to imply. Either way, any explanation of inertia is provisional in our context given the difficulty of inferring mental states from inaction.

proportion, shown in the second row of the final column of Table 2, using the maximal value of the variance of the sampling distribution (namely $(\frac{1}{2})^2$):

$$z_t \approx \frac{P_t^w - P_m^w}{(0.5^2/t)^{1/2}}. \quad (4)$$

Thus, the p value for IE can be derived from (4) as the test statistic. We assume for simplicity that z_t is distributed as a standard Normal.

Later in the paper we derive the full distribution of α so as to let the data adjudicate each agent's attitude towards evidence. If $f_i(\alpha)$ has most probability mass near zero, agent i exhibits sticky belief adjustment. Probability mass in $f_i(\alpha)$ near unity implies a willingness to update for any evidence. In the limit, as α approaches unity, agents will update regardless of evidence (or, more precisely, even for zero evidence against the null). This is equivalent to (stochastic) time dependent updating. That is, if the probability mass at unity in $f_i(\alpha)$ is, say, 0.3, it implies that there is a thirty per cent chance that agent i will update regardless of what the evidence says. More formally, the decision rule in a hypothesis test is to reject H_0 , the status quo, if the p value $\leq \alpha$. A value for α of unity implies the status quo will be rejected, which is the same as updating in this context, for any p value whatsoever.

3.2.4 Relationship between expectations benchmarks

When agent i rejects H_0 within the IE framework we assume she updates her probability guess using Quasi-Bayesian updating.

Since each agent has a full distribution of α , namely $f_i(\alpha)$, we need a representative α_i to summarize the extent of sticky belief adjustment for agent i and to relate to her β_i . There are a number of possibilities, but a natural choice which permits analytic solutions is the median α_i from their $f_i(\alpha)$. For the purposes of our empirical analysis a fully rational (Bayesian) agent is one who has (median) $\alpha_i = \beta_i = 1$, whereas any other sort of agent does not have RE.

We now parameterize all three expectation processes in a double hurdle model. We find evidence for all of them in our data, and importantly we find that the IE representation of $f_i(\alpha)$ has non-zero measure at unity. As discussed above, this is the fraction of agents who undertake random belief adjustment.

4 Data analysis and model estimation

4.1 Nonparametric analysis

The baseline and complex treatments each had 82 subjects, and the inattention treatment had 81 subjects. In this sub-section, we motivate our model with nonparametric analysis.

First, we consider the number of times our subjects executed a no-change, meaning a guessed probability equal to that of the previous period. This is interesting

because, given the nature of the information and comparatively small number of draws, incidences of no-change are not predicted by either Bayesian or Quasi-Bayesian updating, and so, if such observations are widespread in the data, this is the first piece of nonparametric evidence that these standard models are incomplete.

The maximum of the number of no-changes for each subject is 49: seven opportunities for no change out of eight draws, times the seven stages. The distributions over subjects separately by treatment are shown in Fig. 1. The baseline distribution shows a concentration at low values; for both Complex and Inattention, there appears to be a shift in the distribution towards higher values, as one might expect. The means for each treatment are represented by the vertical lines on the right hand side. The vertical lines on the left hand side show the mean number of no-changes that would result from subjects rounding the Bayesian probabilities to two decimal places (or, equivalently, rounding percentages). Clearly, rounding cannot account for the prevalence of no-changes found in the empirical distribution.

The mean is higher under C (22.79) than under B (18.74) (Mann–Whitney test gives $p = 0.007$); and higher under I (26.97) than under B ($p < 0.001$).¹⁵ This is expected: complexity and inattention are both expected to increase the tendency to leave guesses unchanged. When C and I are compared, the p value is 0.06, indicating mild evidence of a difference between the two treatments.

In Fig. 1 it is clear from the nonparametric evidence of widespread incidence of no-changes that any successful model of our data will have to deal with the phenomenon of whether to adjust, before considering how much to adjust.

Second, when agents do change, it is of interest why they do. In Fig. 2 we plot the binary indicator for updating against the strength of evidence against the maintained beliefs $|z_{it}|$ (the absolute change in the Bayesian posterior since the last time the subject updated; see top of row Table 2). A Lowess smoother is superimposed, and this can be interpreted as the predicted probability of an update for a given value of $|z_{it}|$. This provides good nonparametric evidence that higher values of $|z_{it}|$ make change more likely, but, across subjects, agents who are more reluctant to change will exhibit both a relatively low probability of update and a higher value of $|z_{it}|$. Hence there is an econometric concern that the relationship may be affected by endogeneity bias.¹⁶

Third, Fig. 3 shows the extent of any updating on receipt of a white ball and an orange ball, both as a raw change and as a proportion of the absolute change dictated by Bayes rule. The upper panels indicate that updates are often in step sizes of 0.1 or 0.05 and the 0.6 prior for P_i was not so asymmetric as to generate artefacts.

In the lower panels, the updates relative to the Bayesian benchmark cluster between zero and one (for a white ball) and minus one and zero (for an orange

¹⁵ All p values in the paper are two tailed. All bivariate tests use subject level means so that the independent observations avoid the problem of dependence of within-subject choices.

¹⁶ On another econometric matter, the RHS plots by treatment provide suggestive evidence that treatment dummies shifting the probability of adjustment are warranted in the econometric modelling of Sect. 4.2.

ball).¹⁷ This shows that when agents adjust, they tend to do so in a reasonable direction for risk averse agents, raising their probability for Urn 1 when a white ball is drawn and lowering it when an orange ball is drawn.¹⁸ Furthermore, the clustering indicates that any Quasi-Bayesian representation of their adjustment will require a β parameter less than unity, reflecting insufficient belief adjustment.

These nonparametric statistics are consistent with more than one theoretical approach, but it is not clear that just one approach will explain all the features of the data. With that in mind, we now turn to a model which allows the different approaches to co-exist.

4.2 A double hurdle model of belief adjustment

In this section, we develop a parametric double hurdle model which simultaneously considers the decision to update beliefs and the extent to which beliefs are changed when updates occur. The purpose of the model is to act as a testing tool for state-dependent belief adjustment, namely Bayesian belief adjustment and Quasi-Bayesian belief adjustment in the simple version previously defined, as well as (stochastic) time-dependent belief adjustment.

Our econometric task is to model the transformed implied belief $r_t^* = \Phi^{-1}(g_t^*(\theta_i))$, which in turn requires an estimate for risk aversion. We estimate this at the individual level using the technique by Offerman et al. (2009). “Appendix 3: Method for estimating CRRA risk parameter” contains the subject-level details surrounding the estimation of θ_i . On average, agents are risk averse with a mean θ of 0.2.¹⁹

We will refer to r_{it}^* , subject i 's belief in period t , as shorthand for ‘transformed implied belief’. We will treat r_{it}^* as the focus of the analysis, because r_{it}^* has the same dimensionality as z_{it} , the test statistic defined in (4). That is, both have support $(-\infty, \infty)$. Sometimes r_{it}^* changes between $t-1$ and t ; other times, it remains the same. Let Δr_{it}^* be the change in belief of subject i between $t-1$ and t . That is, $\Delta r_{it}^* = r_{it}^* - r_{it-1}^*$.

In the following estimation we exploit the near equivalence between (4) and the scaled difference since the last update $2(\Phi^{-1}(P_t) - \Phi^{-1}(P_m))/\sqrt{t}$ (from Table 2). In round 1, P_m equals the prior 0.6 and the movement of the guess for a given

¹⁷ Plots by treatment for Fig. 3 are provided in “Appendix 5: Robustness and understanding”.

¹⁸ The true guess g^* should always rise when the draw is white, and vice versa when the draw is orange. Online appendix 1 has an extensive discussion of why a minority of contrarian adjustments (falls on white and increases on orange) is observed in Fig. 3. In brief, round 1 data suggests an error rate of around 5% and this is supported by the results of the double hurdle model presented later. That said, it can be optimal for g to move in the opposite direction to g^* . The variance of the scaled payoff linearized around g^* , $1 - (X - g)^2 \approx 1 - (g^* - g)^2 - 2(g^* - g)(X - g^*)$, is $4(g^* - g)^2 V(X)$, where $X \sim \text{Bernoulli}(g^*)$. Thus, risk averse agents will always want to move their g in line with g^* to minimize variance, but agents who love risk enough might possibly move g in the opposite direction to g^* .

¹⁹ This is in the neighborhood of the 0.3–0.5 range elicited by Holt and Laury (2002, p. 1649) using their Multiple Price List (MPL) task. In the results to follow, the heterogeneity of θ , along with a roughly even split between risk loving and risk averse agents, corroborates the usefulness of the robustness analysis in “Appendix 5: Robustness and understanding”.

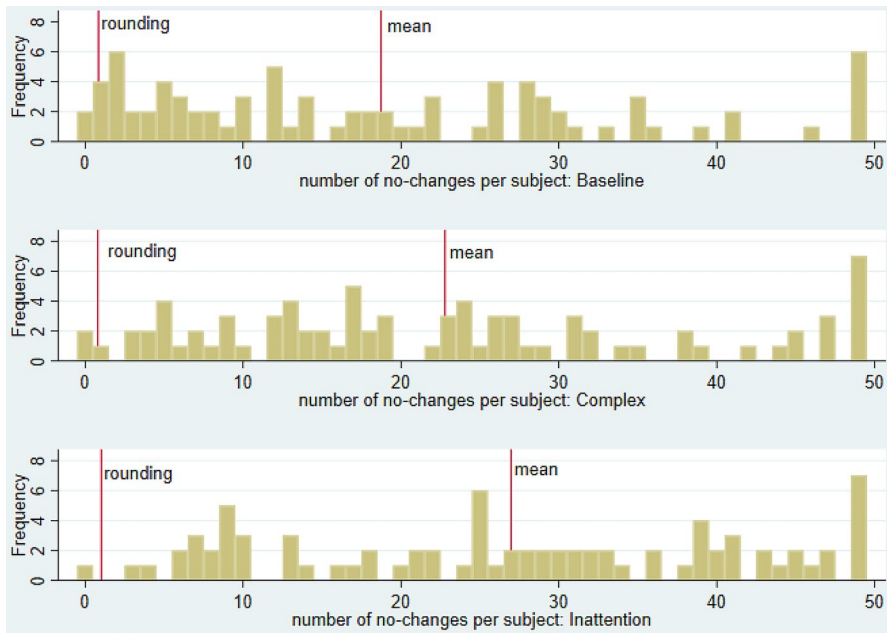


Fig. 1 Distributions of number of no-changes over subjects separately by treatment. In each panel, the vertical line on the right represents the mean number of no-changes per subject, while the vertical line on the left would be the outcome of Bayesian behaviour if subjects rounded their answers

subject is $\Delta r_{it}^* = r_1^* - \Phi^{-1}(0.6)$. That is, both the objective measure of the information change and the subjective guess of the agent are assumed to anchor onto the prior probability that Urn 1 is chosen, 0.6, in the first period.

4.2.1 First hurdle

The probability that a belief is updated (in either direction) in period t is given by:

$$P(\Delta r_{it}^* \neq 0) = \Phi[\delta_i + x_i' \Psi_1 + \gamma |z_{it}|], \quad (5)$$

where $\Phi[\cdot]$ is the standard Normal cdf and δ_i represents subject i 's idiosyncratic propensity to update beliefs, and therefore models random probabilistic belief adjustment (time-dependent belief adjustment). The probability of an update is assumed to depend (positively) on the absolute value of z_{it} , the test statistic. The vector x_i contains treatment and gender dummy variables together with an age variable and a score on two questions from the comprehension questionnaire administered after the

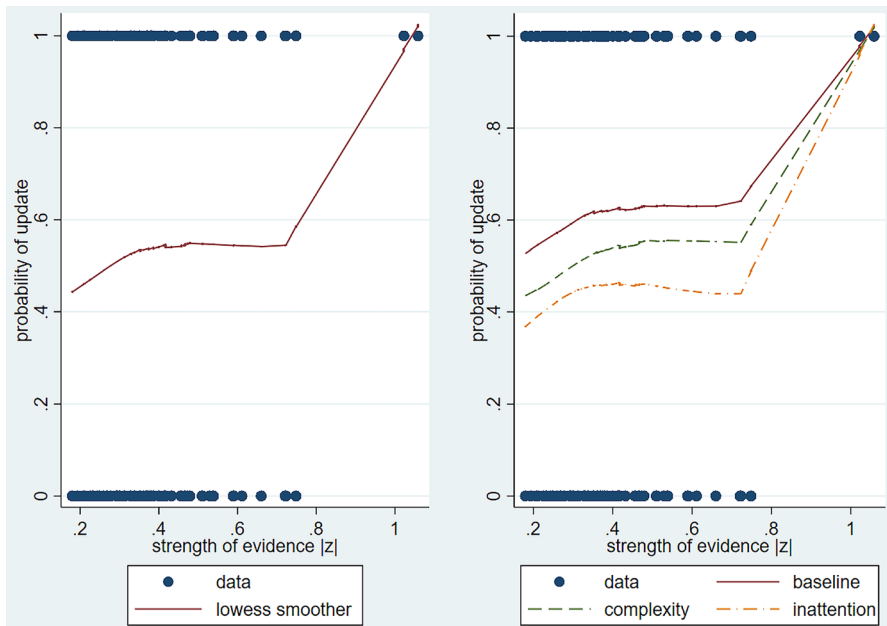


Fig. 2 Predicted probability of updating against strength of evidence (the latter measured as the absolute value of z , defined in Table 2). The dots represent individual decisions to update (1 = update; 0 = no update). The lines are Lowess smoothers, obtained using a tricube weighting function and bandwidth 0.8 (both STATA defaults). Left panel: smoother obtained for full sample. Right panel: smoother obtained separately by treatment

main part of the experiment was explained,²⁰ all of which are time invariant and can be expected to affect the propensity to update.

One econometric issue flagged in the last sub-section is the endogeneity of the variable $|z_{it}|$: subjects who are averse to updating tend to generate large values of $|z_{it}|$ while subjects who update regularly do not allow it to grow beyond small values. This could create a downward bias in the estimate of the parameter γ in the first hurdle. To deal with this concern we use an instrumental variables (IV) estimator which uses the variable $\widehat{|z_{it}|}$ in place of $|z_{it}|$, where $\widehat{|z_{it}|}$ comprises the fitted values from a regression of $|z_{it}|$ on a set of suitable instruments.²¹

4.2.2 Second hurdle

Conditional on subject i choosing to update beliefs in draw t , the next question relates to how much they do so. This is given by:

²⁰ These are questions 1 and 2 in the main part questionnaire as provided in online appendix 4. A third question was used on subjects for all treatments but a software coding error prohibited its use for analysis.

²¹ See “Appendix 4: IV Estimator”.

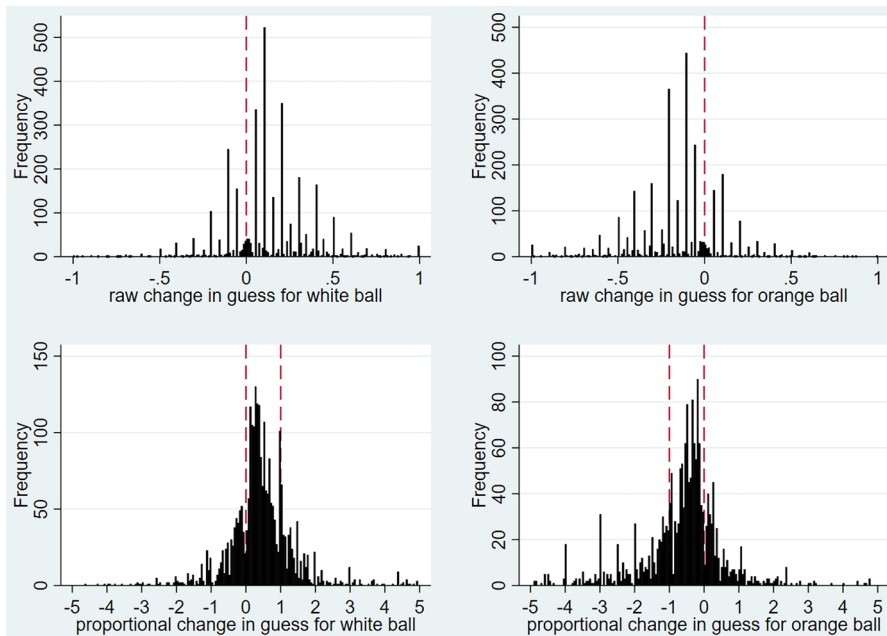


Fig. 3 Change in guess: raw, and as a proportion of a Bayesian benchmark. The top panels show the raw size of the updates on receiving a ball of each color. The bottom panels show the proportional size of the updates on receiving a ball of each color: this is defined as the actual update on receiving a ball as a proportion of the absolute correct Bayesian update (assuming the subject starts from the Bayesian prediction). A vertical line is drawn in correspondence to 0 (no update) and, for the bottom panels, in correspondence to the Bayesian proportional change in guess (so 1 on receiving a white ball and -1 on receiving an orange ball). Data from all treatments are used

$$\Delta r_{it}^* = (\beta_i + x_i' \Psi_2) \frac{\sqrt{t} z_{it}}{2} + \varepsilon_{it}, \quad \varepsilon_{it} \sim N(0, \sigma^2). \quad (6)$$

As a reminder, the Quasi-Bayesian belief adjustment parameter β_i represents subject i 's idiosyncratic responsiveness to the accumulation of new information: if $\beta_i = 1$, subject i responds fully; if $\beta_i = 0$, subject i does not respond at all. Remember that β_i is not constrained to $[0, 1]$. In particular, a value of β_i greater than one would indicate the plausible phenomenon of overreaction. Again, treatment variables are included: the elements of the vector Ψ_2 tell us how responsiveness differs by treatment.

Considering the complete model, there are two idiosyncratic parameters, δ_i and β_i . These are assumed to be distributed over the population of subjects as follows:

$$\begin{pmatrix} \delta_i \\ \beta_i \end{pmatrix} \sim N \left[\begin{pmatrix} \mu_1 \\ \mu_2 \end{pmatrix}, \begin{pmatrix} \eta_1^2 & \rho \eta_1 \eta_2 \\ \rho \eta_1 \eta_2 & \eta_2^2 \end{pmatrix} \right]. \quad (7)$$

In total, there are seventeen parameters to estimate: $\mu_1, \eta_1, \mu_2, \eta_2, \rho, \gamma, \sigma$, four treatment effects (two in each hurdle); two gender effects (one in each hurdle); two scores from the comprehension questionnaire (one in each hurdle); and two age effects (one

in each hurdle). Estimation is performed using the method of maximum simulated likelihood (MSL), with a set of Halton draws representing each of the two idiosyncratic parameters appearing in (7). Following estimation of the model, Bayes rule is used to obtain posterior estimates (denoted $\hat{\delta}_i$ and $\hat{\beta}_i$) of the idiosyncratic parameters for each subject.²²

The results are presented in Table 3 for four different models. The last column shows the preferred model. Model 1 estimates the QB benchmark, in which it is assumed that the first hurdle is crossed for every observation—that is, updates always occur. Zero updates are treated as zero realizations of the update variable in the second hurdle, and their likelihood contribution is a density instead of a probability. Because of this difference in the way the likelihood function is computed, the log-likelihoods and AICs cannot be used to compare the performance of QB to that of the other models.

Model 2 estimates the IE benchmark, in which the update parameter (β_i) is fixed at 1 for all subjects. Consequently the extra residual variation in updates is reflected in the higher estimate of σ . The parameters in the first hurdle are free.

Model 3 combines IE and QB, but constrains the correlation (ρ) between δ and β to be zero. Model 4 is the same model with ρ unconstrained.

The overall performance of a model is judged initially using the AIC; the preferred model being the one with the lowest AIC. Using this criterion, the best model is the most general model 4 (model 1 not being subject to the AIC criterion): IE-QB with ρ unrestricted, whose results are presented in the final column of Table 3.

To confirm the superiority of the general model over the restricted models, we conduct Wald tests of the restrictions implied by the three less general models. We see that, in all three cases, the implied restrictions are rejected, implying that the general model is superior. Note in particular that this establishes the superiority of the general model 4 (IE-QB with ρ unrestricted) over the QB model 1 (a comparison that was not possible on the basis of AIC).

Further confirmation is furnished by measuring the sample predictive accuracy on a subsample of data (the ‘cross validation’ approach). ROC (Receiver Operating Characteristic) is the model’s out of sample predictive accuracy for hurdle 1 (the frequency of updates) and R^2 (out of sample) is the model’s predictive accuracy for hurdle 2 (the extent of updates).²³ Model 4 is no worse than model 3 on an ROC

²² For detailed examples of the use of MSL applied to similar models, including the process of extracting the posterior estimates, see Moffatt (2015). For a more general discussion of this estimation approach, see Train (2009).

²³ ROC is a methodology that is applied in binary data settings (e.g. our first hurdle). Clearly, when the outcome is binary, the prediction is in the form of a “predicted probability of a 1”, so there is no such thing as a “correct prediction”. A starting point is to define a prediction to be correct if the predicted probability of the observed outcome is greater than 0.5. However, the threshold need not be 0.5. ROC forms a test statistic by finding the number of correct predictions at all possible thresholds. The outcome in the second hurdle is continuous, so standard measures of predictive performance (e.g. predictive R-squared) are applicable. Both measures of predictive performance are out-of-sample measures. For the purpose of obtaining them, a 50% sample was used for estimation, and the remaining 50% of the observations were predicted. The number we report for ROC is the area under the ROC curve. This curve compares the true positive rate of prediction with the false positive rate of prediction for the universe of

criterion, but predicts the extent of adjustment better. Model 1 is best of all at predicting the extent of adjustment, but it fails to predict ‘no-change behavior’, by construction. Thus, on both in and out of sample criteria, model 4 is best overall.²⁴

We interpret the results from model 4 as follows. Consider the first hurdle (propensity to update). The intercept parameter in the first hurdle (μ_1) tells us that a typical subject has a predicted probability of $\Phi(0.061) = 0.524$ of updating in any task, in the absence of any evidence (i.e. when $|z_{it}| = 0$). We note that this estimate is not significantly different from zero, which would imply a 50% probability of updating. The Inattention treatment effect is significant and negative, suggesting that the probability of update is lower when subjects are not paying attention. So is the Complexity treatment but not by as much. The effect of the questionnaire score is negative and significant, though it is not large. The negative coefficient is consistent with Cohen et al.’s (2019) model if subjects have cognitive costs. This gives us our first result:

Result 1 *There is evidence of time-dependent (random) belief adjustment. Subjects update their beliefs idiosyncratically around half the time.*

The large estimate of η_1 tells us that there is considerable heterogeneity in the propensity to update (see Fig. 4), something we will explore further in Sect. 4.3. The parameter γ is estimated to be significantly positive, and this tells us, as expected, that the more cumulative evidence there is, in either direction, the greater the probability of an update:

Result 2 *There is evidence of state-dependent belief adjustment. Subjects are more likely to adjust if there is more evidence to suggest that an update is appropriate (thus making it costlier not to update).*

In the second hurdle, the intercept (μ_2) is estimated to be 0.819 in our preferred model 4: when a typical (baseline) subject does update, she updates by a proportion 0.819 of the difference from the Bayes probability. The large estimate of η_2 tells us that there is considerable heterogeneity in this proportion also (see Fig. 4). Interestingly, on the basis of the posterior estimates from model 4, only 13 out of 245 subjects appear to have $\beta < 0$, which indicates noise or confused subjects who adjusted in the wrong

Footnote 23 (continued)

possible thresholds. The R^2 (out of sample) is computed by comparing predictions from the second hurdle with actual decisions contingent on an update occurring.

²⁴ In “Appendix 5: Robustness and understanding” we provide a number of robustness checks on model 4, and check for comprehension by the subjects more generally. We supplement this section by (1) running model 4 using CARA preferences to correct for risk aversion; (2) estimating model 4 on a subset of (near) risk neutral subjects; (3) separately running model 4 using the subset of data where subjects correctly answered both comprehension questions; (4) exploring the Inattention and Complexity treatments on their own to check for coherence between the comprehension questionnaires and the model results; (5) providing the marks for the comprehension questionnaire; and (6) providing data on the risk attitude part which helps explain the distribution of θ .

Table 3 Results of hurdle model with risk adjustment

	QB (1)	IE (2)	IE-QB ($\rho = 0$) (3)	IE-QB (4)
Propensity to update				
μ_1	$+\infty$	0.171 (0.216)	0.143 (0.118)	0.061 (0.109)
[Baseline time-dependent update probability $\Phi(\mu_1)$]	[1]	[0.57]	[0.56]	[0.52]
η_1	0	1.243** (0.063)	1.315** (0.059)	1.28** (0.057)
γ	0	0.603** (0.102)	0.598** (0.102)	0.597** (0.102)
Complex	0	-0.199 (0.181)	-0.273** (0.081)	-0.211** (0.075)
Inattention	0	-0.276 (0.174)	-0.398** (0.077)	-0.305** (0.076)
Male	0	-0.324** (0.137)	-0.419** (0.088)	-0.382** (0.079)
Comprehension	0	-0.126 (0.107)	-0.147** (0.064)	-0.133** (0.056)
Age-22	0	-0.037** (0.009)	-0.026** (0.004)	-0.025** (0.005)
Extent of update				
μ_2	0.522** (0.085)	1	0.609** (0.099)	0.819** (0.123)
η_2	0.498** (0.022)	0	0.807** (0.033)	0.876** (0.044)
Complex	-0.022 (0.078)	0	0.049 (0.085)	0.116 (0.093)
Inattention	-0.074 (0.075)	0	0.240** (0.076)	-0.254** (0.092)
Comprehension	-0.007 (0.046)	0	0.014 (0.054)	-0.023 (0.062)
Male	-0.027 (0.068)	0	0.271** (0.072)	0.203** (0.078)
Age-22	0.021** (0.005)	0	0.019** (0.005)	0.019** (0.007)
σ	0.516** (0.003)	0.731** (0.006)	0.660** (0.006)	0.660** (0.006)
ρ	N/A	N/A	0	-0.18* (0.081)
LogL	-8,790	-12,328	-11,924	-11,923
AIC (= $2k - 2\text{LogL}$)	N/A	24,690	23,882	23,880
Wald test (df, p value)	66,599 (9, 0.000)	430 (8, 0.000)	4.36 (1, 0.037)	N/A
ROC (out of sample)	Fail by definition	0.551	0.588	0.588
R^2 (out of sample)	0.090	0.082	0.081	0.083
n (subjects)	245	245	245	245
T (observations per subject)	56	56	56	56

LogL for QB cannot be compared with that of other columns

direction. Moreover, 87 out of 245 subjects (around one third) display overreaction to the evidence. We summarize this in the following result:

Result 3 *There is evidence of Quasi-Bayesian partial belief adjustment. On average, subjects who adjust do so by around 80%. There is evidence of prior information under-weighting: around one third of the subjects overreact to evidence once they decide to adjust.*

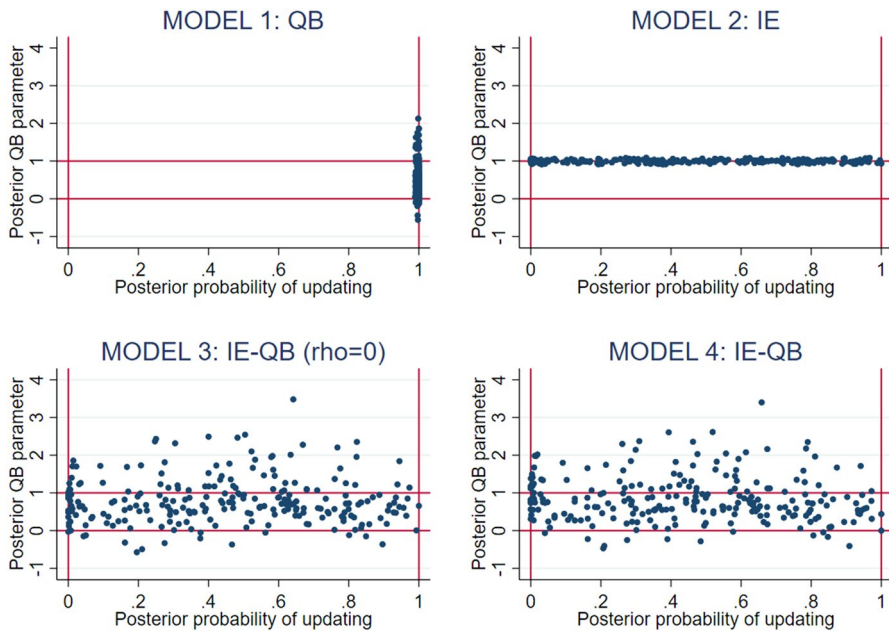


Fig. 4 Jittered scatter of posterior QB parameter against posterior probability of updating for the four models whose estimates are reported in Table 3. Scatters based on Models 1 (top left), 2 (top right), 3 (bottom left) and 4 (bottom right) in Table 3 respectively. Each dot corresponds to a subject

The estimate of ρ is negative, indicating that subjects who have a higher propensity to update, tend to update by a lower proportion of the difference from the Bayes probability. Inattention is important for both hurdles:

Result 4 *Inattention lowers the probability of updating from 50 to 40% and lowers the extent of update from 80 to 56% of the amount prescribed by Bayes rule. Complexity also lowers the probability of update in the first hurdle.*

4.3 The empirical distribution of α and β

To get a better sense of the population heterogeneity in belief adjustment, this subsection maps out the empirical distribution of the IE α_i and QB β_i parameters across subjects against each other. The estimated distribution $f(\beta)$ can be seen from the distribution of the posterior estimates $\hat{\beta}_i$ from Model 4, and this distribution is the marginal distribution of the extent of update on the vertical axis on the bottom-right panel of Fig. 4. We next use the first hurdle information to generate $f_i(\alpha)$, the empirical distribution of α_i .

As we flagged earlier, each agent has a full distribution of α and so we need a representative α_i to summarize the extent of sticky belief adjustment for agent i , to then relate to their β_i . As will be clear below, the choice that permits analytic solutions is the median α_i from $f_i(\alpha)$.

The econometric equation for the first hurdle is equivalent to the probability of rejecting the null under IE. We omit the dummy variables and begin by re-writing the first hurdle, namely (5):

$$\Pr(\text{reject } H_0)_{it} = \Phi(\delta_i + \gamma|z_{it}|), \quad (8)$$

where δ_i and γ are estimated parameters and $|z_{it}|$ is the test statistic based on the proportion of white balls:

$$z_{it} = \frac{P_{it}^w - P_{im}^w}{\sqrt{0.5^2/t}}. \quad (9)$$

For any $|z_{it}|$ it is possible to work out an implied p value and we do so by assuming that (9) is approximately distributed $N(0, 1)$. This in turn allows us to work out $f_i(\alpha)$ from the econometric equation for the first hurdle. When $|z_{it}| = 0$, the p value for a hypothesis test is unity, and so the equation says that a fraction of agents will reject H_0 if the p value is unity. Since the criterion for rejecting H_0 in a hypothesis test is always $\alpha \geq p$ -value, the observed behavior of rejecting H_0 when $|z_{it}| = 0$ implies that there must be a non-zero probability mass on $f_i(\alpha)$ at the value of α *exactly* equal to 1. The pdf of α_i will thus have a discrete ‘spike’ at unity and be continuous elsewhere. We know what that spike is from Eq. (8) with $|z_{it}| = 0$ substituted in, namely $\Phi(\delta_i)$.

The probability of rejecting H_0 depends on the probability that the test size is greater than the p value, but this is also equal to the econometric equation for the first hurdle.

$$\Pr(\text{reject } H_0)_{it} = \int_{p\text{-value}_{it}}^1 f_i(\alpha) d\alpha_i = 1 - F_i(p\text{-value}_{it}) = \Phi(\delta_i + \gamma|z_{it}|) \quad (10)$$

Upper case F in the last equality is the anti-derivative of the density. We define $F_i(1)$ to be unity since 1 is the upper end of the support of α but we also note that there is a discontinuity such that F jumps from $1 - \Phi(\delta_i)$ to 1 at $\alpha = 1$, as a consequence of the non-zero probability mass on $f_i(\alpha)$ at unity. To solve the equation we use an expression for the p value of $|z_{it}|$ on a two-sided Normal test.

$$p\text{-value}_{it} = 2(1 - \Phi(|z_{it}|)). \quad (11)$$

We use a ‘single parameter’ approximation to the cumulative Normal (see Bowling et al. 2009). For our purposes $\sqrt{3}$ is sufficient for the single parameter.

$$\Phi(|z_{it}|) \approx \frac{1}{1 + \exp(-\sqrt{3}|z_{it}|)}. \quad (12)$$

We can now write down $|z_{it}|$ as a function of the p value using (11) and (12):

$$|z_{it}| \approx \frac{1}{\sqrt{3}} \ln \left(\frac{2 - p\text{-value}_{it}}{p\text{-value}_{it}} \right). \quad (13)$$

Intuitively, a p value of zero implies an infinite $|z_{it}|$ and a p value of unity implies $|z_{it}|$ is zero, and (13) confirms this. We can now use the relationship between $F_i(p\text{-value}_{it})$ and our estimated first hurdle to generate $F_i(\alpha)$.

$$\begin{aligned} 1 - F_i(p\text{-value}_{it}) &= \Phi(\delta_i + \gamma|z_{it}|) \\ \therefore F_i(p\text{-value}_{it}) &= 1 - \Phi(\delta_i + \gamma|z_{it}|) \\ &\approx 1 - \Phi\left(\delta_i + \gamma \left\{ \frac{1}{\sqrt{3}} \ln \left(\frac{2 - p\text{-value}_{it}}{p\text{-value}_{it}} \right) \right\}\right). \end{aligned} \quad (14)$$

In the above expression the variable ' $p\text{-value}_{it}$ ' is just a place-holder and can be replaced by anything with the same support leaving the meaning of (14) unchanged. Thus, it can be replaced by α giving the cumulative density of α .

$$\begin{aligned} F_i(\alpha) &\approx 1 - \Phi\left(\delta_i + \gamma \left\{ \frac{1}{\sqrt{3}} \ln \left(\frac{2 - \alpha}{\alpha} \right) \right\}\right) \\ &\approx 1 - \frac{1}{1 + \left[\frac{\alpha}{2-\alpha}\right]^\gamma \exp(-\sqrt{3}\delta_i)}. \end{aligned} \quad (15)$$

Substitution of $\alpha = 1$ does not give unity, which is what we earlier assumed for the value of $F_i(1)$. However, it does give $1 - \Phi(\delta_i)$, which of course concurs with the econometric equation for the first hurdle when $|z_{it}| = 0$. This discontinuity in F_i is consistent with a discrete probability mass in $f_i(\alpha)$ at unity, as we noted earlier. It now just remains to differentiate F_i to obtain the continuous density $f_i(\alpha)$ for α strictly less than unity. The description of the function at the upper end of the support (unity) is completed with a discrete mass at unity of $\Phi(\delta_i)$.

$$\left. \begin{aligned} f_i(\alpha) &\approx \frac{2\gamma\alpha^{\gamma-1} \exp(-\sqrt{3}\delta_i)}{(2-\alpha)^{1+\gamma} \left\{ 1 + \left[\frac{\alpha}{2-\alpha}\right]^\gamma \exp(-\sqrt{3}\delta_i) \right\}^2}, & \alpha < 1 \\ \Pr_i(\alpha = 1) &\approx \frac{1}{1 + \exp(-\sqrt{3}\delta_i)}, & \alpha = 1 \end{aligned} \right\} \quad (16)$$

Figure 5 illustrates the distribution $f_i(\alpha)$ for $\delta_i = 0.1$ and $\gamma = 0.6$ together with the distributions one standard deviation either side of δ_i . The former is the mean of δ across subjects, from our estimation (from the last column of Table 3, rounded). On the right-most of the chart is the probability mass when $\alpha = 1$. As discussed earlier, this corresponds to the proportion of agents who update on vanishingly small evidence ($|z_{it}| = 0$). There is clearly a great deal of interesting heterogeneity. One distribution has a near-zero probability of a random update (10%) and when the agent uses information they are very conservative, with α close to zero. We might

call them ‘classical statisticians’ given the large probability mass around 1%, 5% and 10%. Another distribution has a virtually certain probability of a random update (90%) and we might call these agents ‘fully attentive’. The central estimate of δ describes an agent who updates roughly half the time, and otherwise has a more or less uniform distribution over α .

Since there are idiosyncratic values of δ_i there will be a separate distribution for every subject varying over δ_i . So we must use a summary statistic for $f_i(\alpha)$, and the one which comes to hand is the median α value, obtained by solving $F_i(\alpha) = 0.5$ in Eq. (15). In Fig. 6, we plot the collection of subject i ’s (median α , β) duples for model 4, our preferred equation. Table 4 lists the percentage of subjects in each (median α_i , β_i) 0.2 bracket.²⁵

Roughly half the subjects update regardless of evidence, so the median α ’s cluster at unity along the bottom axis with half of them (49%) in the range at or above 0.8. Just under one quarter (22%) of agents could be described as classical statisticians with median α ’s around the 1–10% level and a similar figure (28%) have ‘conservative belief adjustment’, with α -values no more than 0.20.

Regarding the size of updating, we already know from Result 3 that it is less than complete. In Table 4, 22% update no more than 40 per cent of what they should.

Result 5 *Estimated test sizes spread over the whole support [0, 1] but are clustered at zero and unity. The extent-of-update distribution has a large probability mass around 50% but an even larger mass for values over unity.*

In supplementary analysis (see online appendix 5), we find that infrequent updaters (α_{med} : 0–0.2) have larger mean square deviations (MSD) from Bayesian’s guesses than other subjects. Frequent updaters (α_{med} : 0.8–1) tend to have larger MSD as each stage progresses, which can be explained by comparative underadjustment or overadjustment of beliefs in the second hurdle of our model (see Table 4).²⁶

5 Discussion and conclusion

The double hurdle model we have developed in this paper allows us to integrate both time- and state-dependent belief adjustment in a unified econometric framework. Our experiment uses a quadratic scoring rule with monetary payoffs to incentivize subjects, and we operationalize Offerman et al. (2009) in order to attempt to control for risk aversion. Further research could use different belief elicitation methods to verify the robustness of our findings.

Our econometric model found evidence for considerable heterogeneity in both the propensity and extent of updating, with the majority of agents departing from the

²⁵ Rounding implies occasional discrepancies in the row and column totals shown.

²⁶ Within each stage, MSD increases as the stage progresses since the behavioral impact of the deviations from Bayesian behaviour becomes progressively more significant. We find instead no evidence that MSD changes across stages, i.e. with general experience.

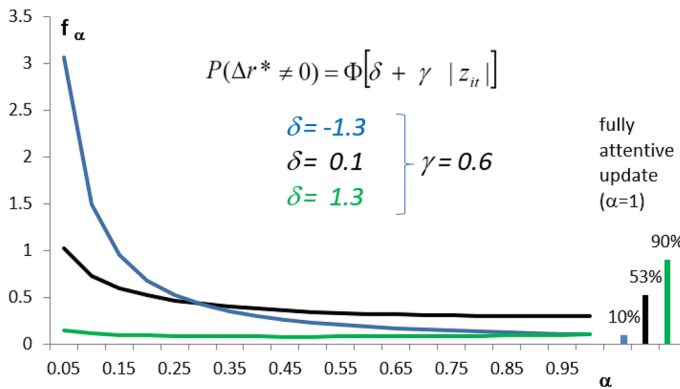


Fig. 5 Distribution of $f_i(\alpha)$ for $\delta_i = 0.1$ and $\gamma = 0.6$ together with the distributions one standard deviation either side of δ_i

rational expectations benchmark of $\alpha = \beta = 1$. Yet deviations from this benchmark are systematic, predictable and can be understood within our modelling framework. We observe random belief adjustment around half the time, which is consistent with stochastic time-dependent belief adjustment. Deviations from Bayesian updating are systematically in the direction of under-adjustment, with a mean of 80% of full adjustment. The likelihood of a belief change increases as the amount of evidence against the no-change status quo increases, which is consistent with state-dependent belief adjustment.

Our aggregate findings are broadly in favor of under- as opposed to over-adjustment to information, which is consistent with prior belief conservatism findings (such as Phillips and Edwards 1966) and the overall finding of under-inference in Benjamin's (2019) review. It is however in apparent contrast with the base rate neglect from static experimental settings such as Kahneman and Tversky (1973) and Tversky and Kahneman (1982). That said, we note based on Table 4 that, when subjects are at the second hurdle, while 52% of subjects have a β_i less than 1, 36% of them do have a β_i above 1, i.e. around a third over-weight rather than under-weight new information. One possible explanation for the greater proportion of under-weighting relative to some other research is that, where the prior is not actually perceived as a genuine and meaningful anchor by subjects—or at least it is perceived as a less reliable source of information than the new information [as in Goodie and Fantino (1999), or some of the settings by Massey and George (2005)]—, then it is more likely to be under-weighted. This is likely to be more the case in static experimental settings, or (following a conjecture by Benjamin (2019)) where priors are based on extreme probabilities.²⁷ Clearly, more research is needed.

Another interesting finding from our double hurdle model is that around half our subjects have α_i less than 0.8. Unlike our dynamic setup, the nature of tasks in static experimental settings as well as others where regime change is to be detected [as in

²⁷ See Benjamin (2019) for a recent discussion of some of the other factors that may be at work.

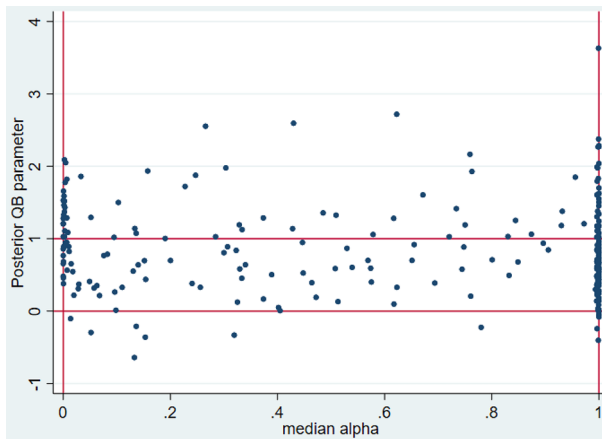


Fig. 6 Median β against median α for Model 4. Each point represents a subject

Table 4 Percentage of subjects in each (α_i, β_i) bracket in model 4

α_{med}	β							Total (%)
	below 0 (%)	0–0.2 (%)	0.2–0.4 (%)	0.4–0.6 (%)	0.6–0.8 (%)	0.8–1 (%)	Above 1 (%)	
0–0.2	2	0	4	3	3	4	12	28
0.2–0.4	0	1	1	1	1	1	3	8
0.4–0.6	0	2	0	2	1	1	2	7
0.6–0.8	0	0	1	0	1	1	3	7
0.8–1	2	3	5	11	7	5	16	49
Total	5	6	11	17	13	12	36	100

Massey and George (2005)], arguably nudges people to make an active choice and to pay the required cognitive and inattention costs.

We believe that a setting where there is no such nudge is a more accurate reflection of many real world decision settings, such as investment portfolio choice over time (Cohen et al. 2019). Cost-based state-dependent sticky belief adjustment can be micro-founded on adjustment costs, for example inattention costs (Alvarez et al. 2016), information costs (Abel et al. 2013), or cognitive costs (Magnani et al. 2016); these, in turn, can be modelled in a stylized way using a hypothesis testing framework, as shown by Cohen et al. (2019) and explained in “[Appendix 2: Relationship between inferential expectations and switching cost models](#)”. We avoid incentivizing the distractor task to help with the interpretability of the findings in a first experiment studying the effect of inattention on belief updating; changing this could be an interesting direction of future research.

We have parameterized the degree of state-dependent belief stickiness by the distribution of the test size. This is informative in two respects. First, the probability mass at $\alpha = 1$ measures the extent of time-dependent adjustment. Second,

where agents instead adopt state-dependent adjustment, the density over the support $\alpha = [0, 1]$, as shown in Fig. 5, informs us about the extent of belief conservatism. On this note, we estimate that roughly one quarter of agents are highly belief conservative with $\alpha \leq 0.2$.

Although the negative correlation between α and β in Fig. 6 is not significant, the relevant pair of parameters in the double hurdle model do have a significantly negative correlation. Thus we have established that agents who tend to update with a low frequency may update a little more when they do update. Subjects with a good understanding of the experiment may defer adjustment because of cognitive and inattention costs, but may be more rather than less Bayesian when an adjustment does take place.

We vary task complexity and likelihood of inattention in belief updating, as these are important dimensions that, we claim, affect how effectively agents process information in the real world. We conceive of inattention costs as the costs of departing from a default choice. This is a natural way of modelling many real world settings in which choices remain as default (e.g. portfolio choices) unless actively changed, and in this respect we follow Khaw et al. (2017) although they do not manipulate the likelihood of inattention (or task complexity). We do not find that task confusion explains belief stickiness to an important degree, nor is there any financial incentive to explain why beliefs are stickier if we add an alternative distracting task. Rather, inattention and cognitive costs are likely to explain infrequent adjustment, to different degrees, by half our subjects.

“Appendix 5: Robustness and understanding” provides descriptive statistics on the optional task, and specifically on how many counting tasks were answered by each subject correctly or incorrectly. It also contains a regression for the inattention treatment where ‘Correct counting tasks’ and ‘incorrect counting tasks’ (referring to the number of each per subject) both have a significant and negative effect on the probability of updating. This supports the conclusion that the effect of the optional task was to distract subjects, who therefore paid less attention than they should have. As noted at the end of Sect. 2, since the optional task was unincentivized, we have cleanly identified this as inattention as opposed to (for example) being bad at multitasking.

Table 8 in “Appendix 5: Robustness and understanding” includes a regression model with maths ability as an explanatory variable, which we measure in the C treatment. This variable is insignificant even in the treatment where potentially it should have mattered the most, which undermines its relevance. The answer of how increased complexity affects updating is instead provided by the preferred general models (3) and (4) in Table 3, as well as the model in Table 8 in “Appendix 5: Robustness and understanding” restricting the sample to subjects who answered correctly the understanding questions. Specifically, complexity makes subjects keener to stick to the default rather than making an active choice (see Gerasimou 2018). This leads to guesses that are on average more distant from the Bayesian predictions (see online appendix 5).

We conclude by populating the cells from Table 1, which provided a taxonomy for agents’ updating, with our empirical results, to give Table 5. Such an exercise is tentative, since ours is the first study to combine the frequency and extent of

Table 5 Updating taxonomy with results (rational expectations corresponds to $\alpha = \beta = 1$)

Hurdle 1: Decision to Update		Interpretation of Quasi-Bayesian β in <i>Posterior = (LikelihoodRatio)$^\beta$Prior</i>			
		Hurdle 2: Extent of Update			
		Irrational $\beta < 0$	Underuse of Information $0 \leq \beta < 0.9$	Bayes $\beta \approx 1^*$	Overuse of Information $\beta > 1.1$
Inferential Expectations α is attention towards evidence	$\alpha_{med} \approx 0 (< 0.01)$ (Inattentive)	0%	4%	3%	7%
	$0.01 < \alpha_{med} < 1$ (Partly attentive)	3%	22%	4%	12%
	$\alpha_{med} = 1$ (Fully attentive)	2%	28%	3%	11%

* The value of one is taken to be $0.9 \leq \beta \leq 1.1$.

adjustment in a unified econometric framework. We have also had to make minor adjustments to the table: with respect to the frequency of adjustment, agents are inattentive if the median α is less than the smallest classical test size (0.01) and, with respect to the extent of adjustment, β is continuous in our model so we must deem it to be unity when it is within a small range (0.2) of that value.

Based on the prevalence of agents, the Quasi-Bayesian modelling strategy of assuming period-by-period updating of information, but with less than full adjustment, commends itself by the behaviour of 28% of agents. The next most common behaviour is Quasi-Bayesian adjustment combined with Inferential Expectations (22%). An advantage of the latter, not shown in the table but shown in Fig. 5, is that an empirical distribution of inferential expectations α 's allows for both time- and state-dependent adjustment, where the probability mass on unity implies stochastic time-dependent adjustment. Among the least common behaviours is full rational expectations (3%), defined as the fully attentive use of Bayes rule ($\alpha = \beta = 1$).

It is not clear yet how generalizable these proportions are and future research might profitably probe them with our double hurdle model in different experimental environments. However, the non-dominance of any row or column within our taxonomy is evidence that both the extent and frequency of adjustment should be taken into account in economic modelling.

Appendix 1: Closeness of two strength-of-evidence measures

Our measure of evidence that the guess should change between times t and m is:

$$z_t = \frac{2(\Phi^{-1}(P_t) - \Phi^{-1}(P_m))}{\sqrt{t}},$$

where the cumulative Normal Φ , and its inverse, are approximated.

$$\Phi(P_t) \approx \frac{1}{1 + \exp(-\sqrt{3}P_t)} \quad \Leftrightarrow \quad \Phi^{-1}(P_t) \approx \frac{1}{\sqrt{3}} \ln \left(\frac{P_t}{1 - P_t} \right).$$

From (2) in Sect. 3,

$$\begin{aligned}
 P_t &= \frac{1}{1 + \frac{(0.3)^{P_t^w} (0.7)^{t-P_t^w} \frac{2}{3}}{(0.7)^{P_t^w} (0.3)^{t-P_t^w} \frac{2}{3}}} \\
 \therefore \frac{P_t}{1-P_t} &= \frac{3}{2} \left(\frac{7}{3} \right)^{2tP_t^w-t} \\
 z_t &= \frac{2(\Phi^{-1}(P_t) - \Phi^{-1}(P_m))}{\sqrt{t}} \\
 &\approx \frac{2}{\sqrt{3t}} \left[\ln \left(\frac{3}{2} \left[\frac{7}{3} \right]^{2tP_t^w-t} \right) - \ln \left(\frac{3}{2} \left[\frac{7}{3} \right]^{2mP_m^w-m} \right) \right] \\
 &= \frac{2 \ln \frac{7}{3}}{\sqrt{3t}} [(2tP_t^w - t) - (2mP_m^w - m)].
 \end{aligned}$$

If $t \approx m$, we have

$$z_t \approx \frac{P_t^w - P_m^w}{\sqrt{\frac{0.5^2}{t}}}.$$

Appendix 2: Relationship between inferential expectations and switching cost models

Cohen et al. (2019) prove that, for a sequential inference problem, such as our updating of the sample proportion of white draws P_t^w , the presence of a cost for switching the inferred value results in a classic confidence-interval band of inaction. This enables us to model sticky adjustment due to switching costs as a two-sided hypothesis test, since it is reasonable to posit a cognitive cost for our subjects changing their declared probability of Urn 1 (which is the same as changing their estimate of P_t^w). While the fully-fledged proof uses a large sample size and a Bayesian filtering framework, this sub-section illustrates the intuition of their argument in an environment like our experiment. We prove that the proportion of white balls is approximately a random walk, and then provide the key insight of Cohen et al. (2019) that a symmetric band of inaction corresponds to a confidence interval.

Our P_t^w is the average of individual Bernoulli draws x_t (with probability of white p) and we now show how the average is approximated by a random walk. We create a synthetic variable s_t with the same moments as a standard Normal random variable (zero mean and unit variance) such that the difference in means between P_t^w and P_{t-1}^w is approximately a scaled version of s_t . First, it is exactly true that:

$$P_t^w = P_{t-1}^w + \frac{x_t - P_{t-1}^w}{t} \quad (17)$$

For this to be a random walk we need to remove P_{t-1}^w from the last term in (17). We replace it by the true proportion p as a large t approximation. Note that since we are

describing the approximate evolution of P_t^w we do not need to concern ourselves with the fact that agents do not know p .

$$\frac{x_t - P_{t-1}^w}{t} \approx \frac{x_t - p}{t} = \frac{\sqrt{pq}}{t} s_t, \quad \text{where } s_t = \frac{x_t - p}{\sqrt{pq}} \quad (18)$$

Second, we now demonstrate that the final stochastic term in (18) has a similar mean and variance to the correct last term in (17).

The means are the same:

$$E\left(\frac{\sqrt{pq}}{t} s_t\right) = 0 \quad \text{and} \quad E\left(\frac{x_t - P_{t-1}^w}{t}\right) = 0.$$

And for large t :

$$\text{Var}\left(\frac{\sqrt{pq}}{t} s_t\right) = \frac{pq}{t^2}; \quad \text{Var}\left(\frac{x_t - P_{t-1}^w}{t}\right) = \frac{pq}{t^2} \left(1 + \frac{1}{t}\right) \approx \frac{pq}{t^2}.$$

Our random walk approximation is then:

$$P_t^w \approx P_{t-1}^w + \frac{\sqrt{pq}}{t} s_t. \quad (19)$$

If agents face a fixed adjustment cost to change their estimate of this proportion and (as in this experiment) a cost of getting their estimate of p wrong, we may posit a band of inaction around the current value of P_t^w .

Crucially, for a random walk the timeless (unconditional) distribution for a given band will be symmetric just like a confidence interval.

Theorem 1 of Cohen et al. (2019) proves that there is an optimal band width which balances the fixed cost of adjustment against the cost of the expected distance from the true value of p . Wide bands will incur few adjustment costs because the boundaries are rarely hit, but they lead to a large expected distance from p . Tight bands mean the adjustment cost is paid repeatedly as the boundaries are hit, but the expected distance from p becomes very small. Cohen et al. derive the optimal band width which minimizes the sum of these costs, and in their Theorem 2 establish the equivalence between this band and a confidence interval. This band of inaction increases with belief conservatism which, in the inferential expectations framework used here, corresponds to a smaller test size α .

Appendix 3: Method for estimating CRRA risk parameter

The log relationship between transformations of g_t and g_t^* in the text (Eq. 1) has an i.i.d. error added to it and is run with 10 observations as an OLS regression. Figure 7 shows the (g_t, g_t^*) duples from subjects. The variable t in (20) refers to the 10 rounds in the

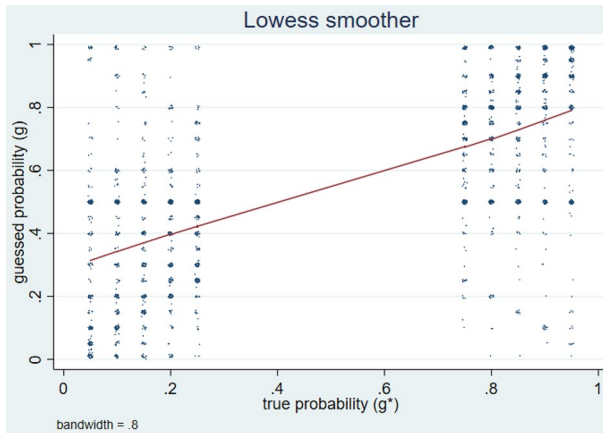


Fig. 7 Jittered scatterplot of guessed probability against true probability. Lowess smoother shown [not Eq. (20)]

risk attitude part, not to the rounds in the main part. The estimated parameter θ_i is subscripted for subjects, because (20) is run for each subject to provide her own θ .

$$\ln \left(\frac{g_t^* (1 - g_t)}{g_t (1 - g_t^*)} \right) = \theta_i \ln \left(\frac{g_t (2 - g_t)}{(1 + g_t)(1 - g_t)} \right) + \eta_t, \quad t = 1, 2, \dots, 10. \quad (20)$$

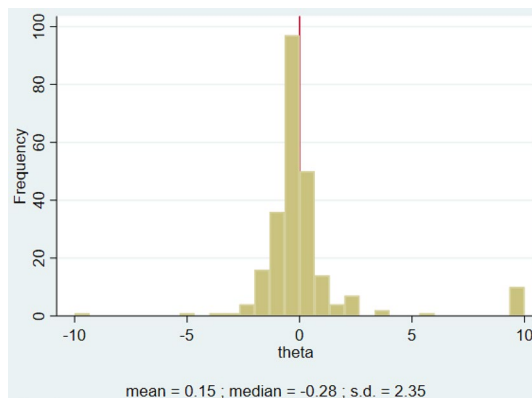
We note the following:

1. The regression has no intercept. If an intercept is included the θ_i estimates are inefficient.
2. For some subjects $g_t = 0.5$ in every period. In this case, the RHS variable is $\ln(1)$ in every period. This means that θ approaches $+\infty$. These estimates need to be re-coded to a high positive number, and we use $+10$.
3. If an agent declared $g_t = 0$ or 1 , (20) cannot be run, so these were replaced with 0.01 and 0.99 respectively.

The above procedure gives rise to a distribution of θ over the 245 subjects, which is shown in Fig. 8.

The empirical distribution is right-skewed because of subject play outlined in step 2 above. The choices in step 2 speak of high risk aversion, giving us probability mass on high values of θ . Choices of g near zero or unity imply θ approaches -1 . In “Appendix 5: Robustness and understanding” we exclude both extreme risk averse and risk loving tails in Fig. 8 by restricting the sample ($|\theta| < 1.5$), and the results are robust.

Fig. 8 Distribution of θ among subjects



Appendix 4: IV Estimator

As mentioned in Sect. 4.2, there is potentially an endogeneity problem with using the strength of evidence against the previously chosen value as an explanatory variable in the first hurdle, and in this appendix we guard against this. The problem is that the variable is endogenous, because subjects with a low propensity to update are clearly likely to generate large values of $|z_{it}|$ simply by virtue of rarely updating. Hence $|z_{it}|$ pooled across subjects may appear to have a negative effect on the propensity to update.

We proceed using an IV estimator. We first create a prediction of the absolute value of z_{it} , $|z_{it}|$, for use in the second stage of a two-stage least squares estimation. The two instruments for $|z_{it}|$ that we use to create this predictor are the round number (t), and the absolute value of the contribution to $|z_{it}|$ in the *current* round $|\Delta z_{it}|$. This is not to be confused with the difference built into z_{it} which spans the current period to period m , namely, $\Phi^{-1}(P_{it}) - \Phi^{-1}(P_{im})$. We note that, because $\Phi^{-1}(P_{im})$ is fixed, a difference operator will eliminate it, leaving Δz_{it} as the change in $\Phi^{-1}(P_{it})$ over the last period.

The stage 1 OLS regression therefore is:

$$|z_{it}| = \pi_0 + \pi_1 t + \pi_3 |\Delta z_{it}| + \varepsilon_{it}. \quad (21)$$

It is important that the dependent variable in (21) is $|z_{it}|$ and not z_{it} . If z_{it} were used as the dependent variable, and consequently the absolute value of the prediction, $\hat{z}_{it}$, were used in the second stage, we would have what Wooldridge (2010, pp. 267–268) refers to as a “forbidden regression”. Using $|z_{it}|$ as the dependent variable in (21) avoids this problem.

Appendix 5: Robustness and understanding

In this appendix we further explore the level of subjects' understanding and the robustness of the regression results. In the main text we established that, on average, subjects raised their guess on receipt of a white ball and dropped it on receipt of an orange ball, and that subjects were more likely to update when the evidence suggested they should. Both features of the data, confirmed in the preferred model, are evidence that subjects understood their environment.

Behavior also seemed sensible in the risk attitude part. In Table 6, we note the frequencies of the chosen probability g at the key choices: zero, 0.5, unity and the probability given to them (g^*), as well as the ranges in between. We separate the cases where the given probability (g^*) was higher, and lower, than 0.5.

Any extreme choices (e.g. 0% or 100%) can be attributed to risk loving preferences, but it is noteworthy that in both columns many subjects placed their declared probability somewhere between the true probabilities g^* and 0.5, making them observationally, and reasonably, risk neutral or risk averse.²⁸

The subject pool was split evenly by gender (49% male), and had an average age of 22 (ranging from 18 to 50). As outlined in the text, a questionnaire was administered after the experimental instructions were explained. Incorrect answers were corrected and explained. There were two questions asked across all treatments. Encouragingly, the distribution of correct answers had a mode of 2 out of 2. Table 7 tallies the number of correct and incorrect answers to each of the two questions. Furthermore, for the Complexity treatment only, we had an extra set of 11 questions related to mathematical abilities.²⁹ Every question was answered correctly by approximately 70% of the subjects (between 67.76 and 72.65%). The density of the number of questions answered correctly is shown in Fig. 9 and the mode for correct answers was 11.

Regression results for our alternative models are given in Table 8. We found reasonable results when we ran model 4 using CARA, and over different samples. The coefficients on $|z|$ in the first hurdle and the Quasi-Bayesian parameter are both positive and significant. The CARA utility model gives similar results to the main text, though CRRA is preferable on a predictive ability criterion. For CRRA (Table 3) $ROC = 0.5880$ and $R^2 = 0.083$ while for CARA (Table 8), $ROC = 0.5041$ and $R^2 = 0.058$. We then censored subjects who exhibited extreme risk behaviors out of the sample (leaving 204/245 subjects with $-1.5 < \theta < 1.5$) and separately ran the model using the subset of data where subjects correctly answered both

²⁸ A value close to 50%, or 50% with rounding, is compatible with strong risk aversion. This is not an unreasonable interpretation, noting for example that in finance data it is often the case that better calibration results are achieved assuming very high CRRA parameters, such as 30 (see Cecchetti and Mark 1990; Boguth and Kuehn 2013). Rounding is common in experiments and does not of course imply that one cannot use models to predict behavior while accepting the additional noise it entails. For example, it is standard to predict ultimatum game results using social preference models which often predict offering a little less than half of the pie, even though a typical modal offer is 50% of the pie.

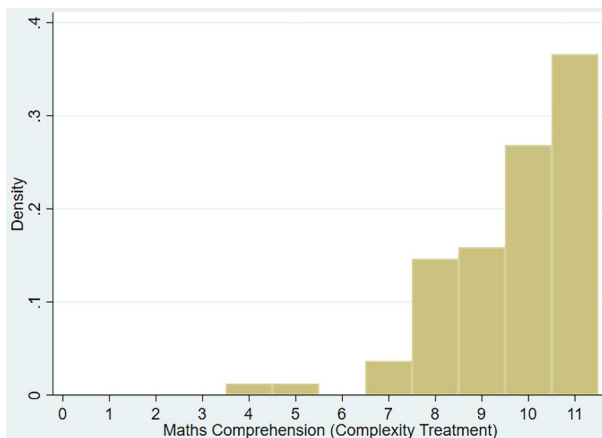
²⁹ Technically, there were three questions, but two of these had five subparts, making a total of 11 questions.

Table 6 Choices in risk attitude part

Low True Probability $g^* = 5, 10, 15, 20, 25\%$		High True Probability $g^* = 75, 80, 85, 90, 95\%$	
(risk neutral or risk averse in <i>italics</i>)			
g Choices	% making choice	g Choices	% making choice
0%	5	100%	16
in between	11	in between	14
g^*	<i>8</i>	g^*	<i>15</i>
in between	<i>35</i>	in between	<i>29</i>
<i>50%</i>	<i>24</i>	<i>50%</i>	<i>18</i>
in between	14	in between	8
100%	3	0%	0

Table 7 Correct answers: comprehension checks in Q1 and Q2 of the main part questionnaire

Q1 \ Q2	Incorrect	Correct	Total
Incorrect	20 (8%)	49 (20%)	69 (28%)
Correct	34 (14%)	142 (58%)	176 (72%)
Total	54 (22%)	191 (78%)	245 (100%)

**Fig. 9** Maths comprehension (complexity treatment): Proportions of subjects who answered x number of questions correctly

comprehension questions. Finally, the model was run using just the data from the Complexity and Inattention treatments.³⁰

The estimated models are broadly similar to the preferred model. Random updating is more common in CARA, but it is in the neighborhood of 50% for most models. Importantly, the coefficient on $|z|$ is positive and significant in all models.

In Fig. 10 below, we provide a version of Fig. 2 in a bin scatter format, which is able to represent the ‘average’ of the choices in each of a predetermined number of bins. (We chose 10.) The pattern seen here is the same as in Fig. 2, and similarly suggests the desirability of protecting the estimates from any endogeneity bias. Otherwise, it is given for completeness.

In Fig. 11 below, we show the extent of updating depicted in the main text in Fig. 3, by treatment. They are not particularly informative, except to confirm that the direction of change observed in the pooled data is true of individual treatments as well.

Finally, we have also checked how subjects played the counting task in the inattention treatment. Fully 63 out of 81 subjects (77.8%) did at least one counting task. These 63 subjects solved counting tasks on average 12.2% of the times. When they solved them, they did so correctly around nine out of ten times (88.6%).

³⁰ The extent of update is not significantly different from zero in the Complexity treatment. However, the positive and significant maths ability coefficient means that there is an indirect updating effect.

Table 8 Robustness checks

	CARA	CRRA ($\theta \approx 0$)	Both comprehension ques- tions right	Complexity only	Inattention only
Propensity to update		$(-1.5 < \theta < 1.5)$			
μ_1	0.720** (0.085) [0.76]	0.033 (0.107) [0.51]	-0.223** (0.97) [0.41]	-0.332 (0.292) [0.37]	-0.382* (0.160) [0.35]
[Baseline time-dependent update probability $\Phi(\mu_1)$]					
η_1	1.302** (0.035)	1.109** (0.053)	1.321** (0.68)	-1.336** (0.065)	1.376** (0.118)
γ	0.607** (0.101)	0.679** (0.109)	0.659** (0.134)	0.858** (0.178)	0.809** (0.186)
Complex	-0.116 (0.059)	-0.130 (0.116)	-0.268** (0.091)		
Inattention	-0.514** (0.057)	-0.216 (0.206)	-0.301** (0.082)		
Male	-0.063 (0.049)	-0.462** (0.133)	-0.417** (0.086)	0.391** (0.081)	-0.415** (0.129)
Comprehension	-0.304** (0.040)	0.048 (0.070)		-0.295** (0.083)	-0.109 (0.086)
Age-22	0.006 (0.004)	-0.035** (0.010)	-0.020** (0.008)	-0.024** (0.007)	-0.001 (0.006)
Maths ability				0.033 (0.028)	
Correct counting tasks					-0.312** (0.028)
Incorrect counting tasks					-0.564** (0.125)
Extent of update					
μ_2	0.617** (0.093)	0.707** (0.170)	0.749** (0.103)	-0.279 (0.448)	0.687** (0.198)
η_2	1.440** (0.046)	0.789** (0.041)	0.812** (0.043)	1.479** (0.146)	0.610** (0.082)
Complex	-0.656** (0.078)	-0.113 (0.107)	0.009 (0.106)		
Inattention	-0.151* (0.062)	-0.271** (0.102)	-0.339** (0.100)		
Comprehension	0.419** (0.054)	0.047 (0.076)		-0.296 (0.154)	-0.132 (0.105)
Male	0.042 (0.061)	0.418** (0.101)	0.314** (0.086)	0.101 (0.128)	0.078 (0.143)
Age-22	-0.030** (0.005)	0.025** (0.009)	0.023** (0.008)	0.033** (0.011)	0.039** (0.010)
Maths ability				0.176** (0.040)	
Correct counting tasks					-0.026 (0.064)
Incorrect counting tasks					0.051 (0.256)

Table 8 (continued)

	CARA	CRRA ($\theta \approx 0$)	Both comprehension ques- tions right	Complexity only	Inattention only
σ	0.532** (0.004)	0.674** (0.007)	0.663** (0.008)	0.730** (0.012)	0.635** (0.011)
ρ	- 0.557** (0.024)	- 0.164** (0.77)	- 0.115 (0.087)	- 0.526** (0.045)	0.220 (0.177)
LogL	- 10,805.391				
ROC (out of sample)	0.5041		0.538		
R^2 (out of sample)	0.058		0.108		
n (subjects)	245	204	142	82	81
T (observations per subject)	56	56	56	56	56

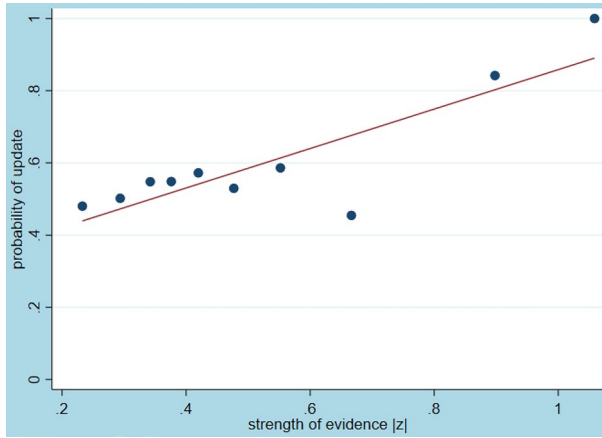


Fig. 10 Estimated probability of updating. (Fig. 2 as a bin scatter plot; number of bins = 10.)

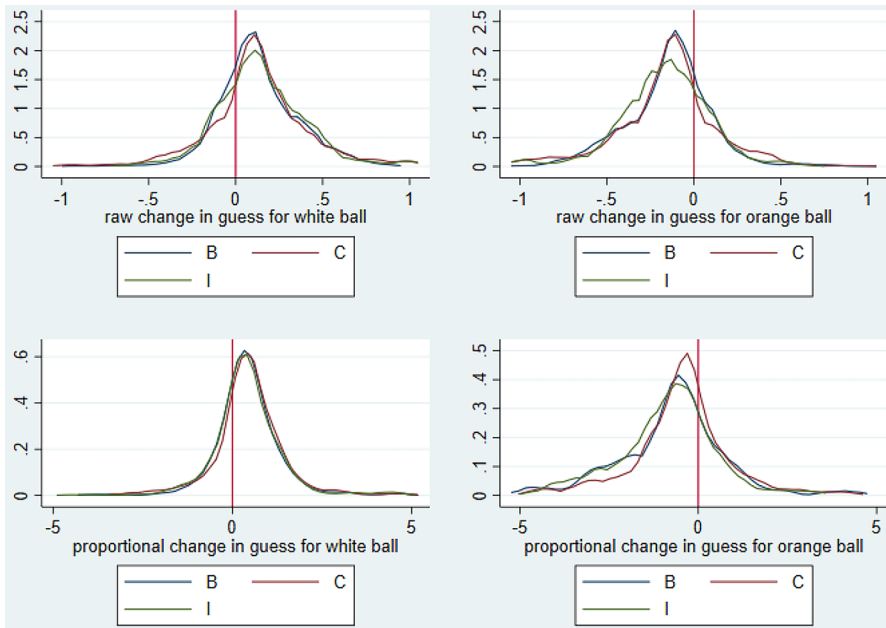


Fig. 11 Change in guess: raw, and as a proportion of a Bayesian benchmark; By treatment. This figure corresponds to Fig. 3 in the main text, using kernel densities to distinguish the treatments. As in Fig. 3, the top panels show the raw size of updates on receiving a ball of each color, and the bottom panels show the proportional size of the updates on receiving a ball of each colour

Supplementary Information The online version contains supplementary material available at <https://doi.org/10.1007/s10683-021-09701-2>.

Acknowledgements We acknowledge the support of the ESRC (Network for Integrated Behavioural Science, ES/K002201/1) and the ANU College of Business and Economics. We thank Rohan Alexander and Ailko van der Veen for research assistance, participants at the University of Technology Sydney seminar, the BENC/CEAR workshop on “Recent advances in modelling decision making under risk and uncertainty”, Durham University, the University of East Anglia, a Society for Experimental Finance workshop, University of Arizona, Tucson, an ESA conference in Sydney and a BEAM Workshop, Kyoto, as well as Han Bleichrodt, Chris Carroll, John Rust, two anonymous reviewers and the editors for useful feedback. The usual disclaimer applies. Ethical approval was obtained at the ANU and the University of East Anglia. The experimental instructions are in the online appendix.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

- Abel, A. B., Eberly, J. C., & Panageas, S. (2013). Optimal inattention to the stock market with information costs and transaction costs. *Econometrica*, 81(4), 1455–1481.
- Abeler, J., Falk, A., Goette, L., & Huffman, D. (2011). Reference points and effort provision. *American Economic Review*, 101(2), 470–492.
- Alvarez, F., Lippi, F., & Passadore, J. (2016). Are state and time dependent models really different? *NBER Macroeconomics Annual*, 31, 379–457.
- Ambuehl, S., & Li, S. (2014). *Belief updating and the demand for information*. SSRN Discussion Paper, June.
- Bacchetta, P., & van Wincoop, E. (2010). Infrequent portfolio decisions: A solution to the forward discount puzzle. *American Economic Review*, 100, 870–904.
- Barbey, A. K., & Sloman, S. A. (2007). Base-rate respect: From ecological rationality to dual processes. *Memory and Cognition*, 30(3), 241–254.
- Baron, J., & Ritov, I. (2004). Omission bias, individual differences, and normality. *Organizational Behavior and Human Decision Processes*, 94(2), 74–85.
- Benjamin, D. J. (2019). Errors in probabilistic reasoning and judgment biases. In B. Bernheim, S. D. Douglas, & D. Laibson (Eds.), *Handbook of behavioral economics* (Vol. 2, pp. 69–185). London: Elsevier.
- Benjamin, D. J., Rabin, M., & Raymond, C. (2015). A model of non-belief in the law of large numbers. *Journal of the European Economic Association*, 14(2), 515–544.
- Boguth, O., & Kuehn, L.-A. (2013). Consumption volatility risk. *Journal of Finance*, 68(6), 2589–2615.
- Bossaerts, P., & Murawski, C. (2017). Computational complexity and human decision-making. *Trends in Cognitive Sciences*, 21(12), 917–929.
- Bowling, S., Khasawneh, M., Kaewkuekool, S., & Cho, B. (2009). A logistic approximation to the cumulative normal distribution. *Journal of Industrial Engineering and Management*, 2(1), 114–127.
- Brainard, W. C. (1967). Uncertainty and the effectiveness of policy. *American Economic Review Papers and Proceedings*, 57(2), 411–425.
- Buser, T., & Peter, N. (2012). Multitasking. *Experimental Economics*, 15, 641–655.
- Caballero, R. J. (1989). *Time dependent rules, aggregate stickiness and information externalities*. Columbia University Department of Economics Discussion Paper 428.

- Calvo, G. A. (1983). Staggered prices in a utility-maximizing framework. *Journal of Monetary Economics*, 12(3), 383–398.
- Caplin, A., Csaba, D., Leahy, J., & Nov, O. (2020). Rational inattention, competitive supply, and psychometrics. *Quarterly Journal of Economics*, 135(3), 1681–1724.
- Caplin, A., Dean, M., & Martin, D. (2011). Search and satisficing. *American Economic Review*, 101(7), 2899–2922.
- Carlin, B. L. (2011). Strategic price complexity in retail financial markets. *Journal of Financial Economics*, 91, 278–287.
- Carroll, C. D. (2003). Macroeconomic expectations of households and professional forecasters. *Quarterly Journal of Economics*, 118, 269–298.
- Cecchetti, S. G., & Mark, N. C. (1990). Macroeconomic evaluating empirical tests of asset pricing models: Alternative interpretations. *American Economic Review Papers and Proceedings*, 80(2), 48–51.
- Cohen, S. N., Henckel, T., Menzies, G. D., Muhle-Karbe, J., & Zizzo, D. J. (2019). Switching cost models as hypothesis tests. *Economics Letters*, 175, 32–35.
- Davis, D. D., & Holt, C. A. (1993). *Experimental economics*. Princeton: Princeton University Press.
- Dessein, W., Galeotti, A., & Santos, T. (2016). Rational inattention and organizational focus. *American Economic Review*, 106(6), 1522–1536.
- El-Gamal, M. A., & Grether, D. M. (1995). Are people Bayesian? Uncovering behavioral strategies. *Journal of the American Statistical Association*, 90(432), 1137–145.
- Gerasimou, G. (2018). Indecisiveness, undesirability and overload revealed through rational choice deferral. *Economic Journal*, 128(614), 2450–2479.
- Gigerenzer, G., & Gaissmaier, W. (2011). Heuristic decision making. *Annual Review of Psychology*, 62, 451–482.
- Goodie, A. S., & Fantino, E. (1999). What does and does not alleviate base-rate neglect under direct experience. *Journal of Behavioral Decision Making*, 12(4), 307–335.
- Holt, C. A., & Laury, S. K. (2002). Risk aversion and incentive effects. *American Economic Review*, 92(5), 1644–1655.
- Huang, L., & Liu, H. (2007). Rational selection and portfolio selection. *Journal of Finance*, 62(4), 1999–2040.
- Huck, S., Zhou, J., & Duke, C. (2011). *Consumer behavioural biases in competition: A survey*. Office of Fair Trading.
- Kahneman, D., & Tversky, A. (1973). On the psychology of prediction. *Psychological Review*, 80(4), 237–251.
- Khaw, M. W., Stevens, L., & Woodford, M. (2017). Discrete adjustment to a changing environment: Experimental evidence. *Journal of Monetary Economics*, 91, 81–103.
- Koehler, J. J. (1996). The base rate fallacy reconsidered: Descriptive, normative and methodological challenges. *Behavioral and Brain Sciences*, 19(1), 1–53.
- Mackowiak, B., & Wiederholt, M. (2015). Business cycle dynamics under rational inattention. *Review of Economic Studies*, 82, 1502–1532.
- Magnani, J., Gorry, A., & Oprea, R. (2016). Time and state dependence in an SS decision experiment. *American Economic Journal: Macroeconomics*, 8(1), 285–310.
- Mankiw, N. G., & Reis, R. (2002). Sticky information versus sticky prices: A proposal to replace the New Keynesian Phillips curve. *Quarterly Journal of Economics*, 117, 1295–1328.
- Martin, D. (2017). Strategic pricing with rational inattention to quality. *Games and Economic Behavior*, 104, 131–145.
- Massey, C., & George, W. (2005). Detecting regime shifts: The causes of under- and overreaction. *Management Science*, 51(6), 932–947.
- Menzies, G. D., & Zizzo, D. J. (2012). Monetary policy and inferential expectations of exchange rates. *Journal of International Financial Markets, Institutions and Money*, 22(2), 359–380.
- Menzies, G. D., & Zizzo, D. J. (2009). Inferential expectations. *B.E. Journal of Macroeconomics*, 9(1) (Advances), Article 42.
- Moffatt, P. G. (2015). *Experimentics: Econometrics for experimental economics*. Melbourne: Macmillan International Higher Education.
- Offerman, T., Sonnemans, J., Van De Kuilen, G., & Wakker, P. P. (2009). A truth serum for non-Bayesians: Correcting proper scoring rules for risk attitudes. *Review of Economic Studies*, 76(4), 1461–1489.
- Payne, J. W. (1979). Task complexity and contingent processing in decision making: An information search and protocol analysis. *Organizational Behavior and Human Performance*, 16, 366–387.

- Phillips, L. D., & Edwards, W. (1966). Conservatism in a simple probability inference task. *Journal of Experimental Psychology*, 72(3), 346–354.
- Rabin, M. (2013). Incorporating limited rationality into economics. *Journal of Economic Literature*, 51(2), 528–43.
- Rabin, M., & Schrag, J. L. (1999). First impressions matter: A model of confirmatory bias. *Quarterly Journal of Economics*, 114(1), 37–82.
- Simon, H. A. (1979). Rational decision making in business organizations. *American Economic Review*, 69(4), 493–513.
- Sitzia, S., & Zizzo, D. J. (2011). Does product complexity matter for competition in experimental retail markets? *Theory and Decision*, 70(1), 65–82.
- Sitzia, S., Zheng, J., & Zizzo, D. J. (2015). Inattentive consumers in markets for services. *Theory and Decision*, 79(2), 307–332.
- Slawinski, N., Pinkse, J., Busch, T., & Banerjee, S. B. (2017). The role of short-termism and uncertainty avoidance in organizational inaction on climate change: A multi-level framework. *Business and Society*, 56(2), 253–282.
- Tadelis, S. (2002). Complexity, flexibility and the make-or-buy decision. *American Economic Review Papers and Proceedings*, 92(2), 433–437.
- Train, K. E. (2009). *Discrete choice methods with simulation*. Cambridge: Cambridge University Press.
- Tversky, A., & Kahneman, D. (1982). Judgments of and by representativeness. In D. Kahneman, P. Slovic, & A. Tversky (Eds.), *Judgment under uncertainty: Heuristics and biases*. Cambridge: Cambridge University Press.
- Wilcox, N. T. (1993). Lottery choice: Incentives, complexity and decision time. *Economic Journal*, 103(421), 1397–1417.
- Wooldridge, J. M. (2010). *Econometric analysis of cross section and panel data*. Cambridge: MIT Press.

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.