

Context and uncertainty in decisions from experience

Author:

Holwerda, Joel

Publication Date:

2023

DOI:

<https://doi.org/10.26190/unsworks/24747>

License:

<https://creativecommons.org/licenses/by/4.0/>

Link to license to see what you are allowed to do with this resource.

Downloaded from <http://hdl.handle.net/1959.4/101040> in <https://unsworks.unsw.edu.au> on 2024-05-05

Context and uncertainty in decisions from experience

Joel Holwerda

A thesis in fulfilment of the requirements for the degree of
Doctor of Philosophy



School of Psychology
Faculty of Science
The University of New South Wales

December 2022

ORIGINALITY STATEMENT

☒ I hereby declare that this submission is my own work and to the best of my knowledge it contains no materials previously published or written by another person, or substantial proportions of material which have been accepted for the award of any other degree or diploma at UNSW or any other educational institution, except where due acknowledgement is made in the thesis. Any contribution made to the research by others, with whom I have worked at UNSW or elsewhere, is explicitly acknowledged in the thesis. I also declare that the intellectual content of this thesis is the product of my own work, except to the extent that assistance from others in the project's design and conception or in style, presentation and linguistic expression is acknowledged.

COPYRIGHT STATEMENT

☒ I hereby grant the University of New South Wales or its agents a non-exclusive licence to archive and to make available (including to members of the public) my thesis or dissertation in whole or part in the University libraries in all forms of media, now or here after known. I acknowledge that I retain all intellectual property rights which subsist in my thesis or dissertation, such as copyright and patent rights, subject to applicable law. I also retain the right to use all or part of my thesis or dissertation in future works (such as articles or books).

For any substantial portions of copyright material used in this thesis, written permission for use has been obtained, or the copyright material is removed from the final public version of the thesis.

AUTHENTICITY STATEMENT

☒ I certify that the Library deposit digital copy is a direct equivalent of the final officially approved version of my thesis.

UNSW is supportive of candidates publishing their research results during their candidature as detailed in the UNSW Thesis Examination Procedure.

Publications can be used in the candidate's thesis in lieu of a Chapter provided:

- The candidate contributed **greater than 50%** of the content in the publication and are the "primary author", i.e. they were responsible primarily for the planning, execution and preparation of the work for publication.
- The candidate has obtained approval to include the publication in their thesis in lieu of a Chapter from their Supervisor and Postgraduate Coordinator.
- The publication is not subject to any obligations or contractual agreements with a third party that would constrain its inclusion in the thesis.

☒ The candidate has declared that **some of the work described in their thesis has been published and has been documented in the relevant Chapters with acknowledgement.**

A short statement on where this work appears in the thesis and how this work is acknowledged within chapter/s:

Shorter versions of Chapter 3 and Chapter 6 were published as conference papers in the Proceedings of the Annual Meeting of the Cognitive Science Society. The contribution of the other author (my primary supervisor) was emphasised in the Acknowledgements section and in a Publications section at the beginning of my thesis.

Candidate's Declaration



I declare that I have complied with the Thesis Examination Procedure.

Abstract

From the moment we wake up each morning, we are faced with countless choices. Should we press snooze on our alarm? Have toast or cereal for breakfast? Bring an umbrella? Agree to work on that new project? Go to the gym or eat a whole pizza while watching Netflix? The challenge when studying decision-making is to collapse these diverse scenarios into feasible experimental methods. The standard theoretical approach is to represent options using outcomes and probabilities and this has provided a rationale for studying decisions using gambling tasks. These tasks typically involve repeated choices between a single pair of options and outcomes that are determined probabilistically. Thus, the two sections in this thesis ask a simple question: are we missing something by using pairs of options that are divorced from the context in which we make choices outside the psychology laboratory?

The first section focuses on the impact of extreme outcomes within a decision context. Chapter 2 addresses whether there is a rational explanation for why these outcomes appear in decisions from experience and numerous other cognitive domains. Chapters 3-5 describe six experiments that distinguish between plausible theories based on whether they measure extremity as categorical, ordinal, or continuous; whether extremity refers to the centre, the edges, or neighbouring outcomes; whether outcomes are represented as types or tokens; and whether extreme outcomes are defined using temporal or distributional characteristics. In the second section, we shift our focus to how people perceive uncertainty. We examine a distinction between uncertainty that is attributed to inadequate knowledge and uncertainty that is attributed to an inherently random process. Chapter 6 describes

three experiments that examine whether allowing participants to map their uncertainty onto observable variability leads them to perceive it as potentially resolvable rather than purely stochastic. We then examine how this influences whether they seek additional information. In summary, the experiments described in these two sections demonstrate the importance of context and uncertainty in understanding how we make decisions.

Acknowledgements

Context and uncertainty have played a major role in this thesis, both in its content and in its creation. The world is a completely different place than it was at the beginning of my candidature and these uncertain times have made me realise that completing this thesis was only possible thanks to the contributions of countless other people. For me to adequately capture the role of this incredible context, this section would need to be longer than the thesis. Therefore, I'll keep this short and thank everyone properly over a long-awaited celebratory drink.

To my supervisor, Ben Newell, I am immensely grateful for your guidance throughout this entire process. In addition to the expertise and wisdom that were instrumental in shaping this thesis, you were an excellent role model, you challenged and inspired me, you were extremely generous, and you were endlessly patient. I simply could not have asked for a better supervisor.

To the members of the UNSW cognition lab, you provided me with a supportive and stimulating environment. From the conversations about philosophy of science to the Bayesian statistics reading group, you made me into the researcher that I am today. And to the other PhD students, your friendship made this whole thing worth it.

Finally, to my wonderful partner, Rebecca, I am so grateful for your presence in my life. You have always been supportive and understanding, you bring out the best in me, and you give me something to look forward to at the end of the day. Thank you from the bottom of my heart.

Publications

Chapter 3

Holwerda, J., & Newell, B. R. (2020). What is an extreme outcome in risky choice? *Proceedings of the Annual Meeting of the Cognitive Science Society* (Vol. 42, No. 42).

Chapter 6

Holwerda, J., & Newell, B. R. (2021). Epistemic and aleatory uncertainty in decisions from experience. *Proceedings of the Annual Meeting of the Cognitive Science Society* (Vol. 43, No. 43).

Contents

Contents	vii
List of Figures	xiii
List of Tables	xv
1 General introduction	1
1.1 Context and extreme outcomes	10
1.1.1 Mechanical explanations	13
1.1.2 Rational explanations	15
1.1.3 Integrative explanations	18
1.1.4 Summary of Chapter 2: Rational explanations	22
1.1.5 Summary of Chapters 3-5: Mechanical explanations	23
1.2 Aleatory and epistemic uncertainty	26
1.2.1 Summary of Chapter 6: Variability and uncertainty	29
1.3 Thesis overview	31
I Context	33
2 Rational explanations	34

2.1	Utility-weighted sampling	36
2.1.1	Rational explanation	36
2.1.2	Mechanical explanation	40
2.2	Reconsidering utility-weighted sampling	42
2.2.1	Attribute 1: Counterintuitive predictions	45
2.2.2	Attribute 2: Flexible parameters	47
2.2.3	Attribute 3: Limited domain of rationality	51
2.3	An alternative rational explanation	62
2.3.1	Ordinal criterion	65
2.3.2	Continuous criterion	69
2.4	Chapter discussion	73
2.4.1	Multiple options	74
2.4.2	Risky options	77
2.4.3	Conclusion	79
3	Levels of measurement	81
3.1	Categorical extreme outcomes	83
3.2	Ordinal extreme outcomes	85
3.3	Continuous extreme outcomes	87
3.3.1	Centre of the distribution:	88
3.3.2	Edges of the distribution:	89
3.3.3	Neighbouring outcomes:	90
3.4	Overview of experiments	92
3.4.1	Terminology	94
3.4.2	Manipulation 1: Value	95
3.4.3	Manipulation 2: Variance	96

3.4.4 Manipulation 3: Skewness	96
3.5 General method	97
3.5.1 Participants	97
3.5.2 Design and procedure	98
3.5.3 Analysis	101
3.6 Experiment 1: Value and skewness (non-extreme shared options)	103
3.6.1 Results and discussion	104
3.7 Experiment 2: Value and variance	109
3.7.1 Results and discussion	110
3.8 Experiment 3: Value and skewness (extreme shared options)	114
3.8.1 Results and discussion	117
3.9 Chapter discussion	121
3.9.1 Value manipulation	121
3.9.2 Variance manipulation	124
3.9.3 Skewness manipulation	125
3.9.4 Conclusion	127
4 Types and tokens	128
4.1 General method	133
4.1.1 Participants	133
4.1.2 Design and procedure	134
4.1.3 Analysis	136
4.2 Experiment 4: Skewed distribution (token-based)	137
4.2.1 Results and discussion	139
4.3 Experiment 5: Skewed distribution (type-based)	143
4.3.1 Results and discussion	146

4.4 Chapter discussion	148
4.4.1 Conclusion	152
5 Temporal and distributional	153
5.1 Experiment 6: Temporal extremity	155
5.1.1 Method	158
5.1.2 Analysis	158
5.1.3 Results and discussion	159
II Uncertainty	164
6 Epistemic and aleatory uncertainty	165
6.1 Uncertainty as a psychological concept	167
6.2 What is the nature of this duality?	168
6.2.1 Exploration	170
6.2.2 Competition	171
6.3 An example with two dice	173
6.4 Decisions from experience	176
6.5 Experiment 7: Partial-feedback	178
6.5.1 Method	180
6.5.2 Results and discussion	184
6.6 Experiment 8: Information or reward	189
6.6.1 Method	190
6.6.2 Results and discussion	192
6.7 Experiment 9: Free sampling	197
6.7.1 Method	197

6.7.2 Results and discussion	199
6.8 Chapter discussion	203
6.8.1 Epistemic and aleatory uncertainty	203
6.8.2 Outcome and observable variability	205
6.8.3 Measuring uncertainty using the EARS	208
6.8.4 Information seeking	210
6.8.5 Competition and risk	211
6.8.6 Conclusion	212
7 General discussion	214
7.1 Context and extreme outcomes	216
7.1.1 Categorical extreme outcomes	216
7.1.2 Continuous extreme outcomes	218
7.1.3 Ordinal extreme outcomes	220
7.1.4 Temporal extreme outcomes	221
7.1.5 Evidence and rebuttals	222
7.1.6 Future directions	226
7.2 Representation	228
7.3 Idealisation	232
7.4 Theory and experiment	236
7.5 Summary and conclusion	242
References	243
A Additional details for utility-weighted sampling simulations	287
B Details on the epistemic and aleatory rating scale (EARS)	289

C	Factor analysis for the epistemic and aleatory rating scale (EARS)	293
D	Details of the dynamic Ellsberg urn task	297
E	Variance reduction using outcome sequences	300

List of Figures

2.1	Schematic depiction of the experiment conducted by Ludvig et al. (2014)	43
2.2	Estimates and choices for the utility-weighted sampling simulations of Ludvig et al. (2014)	49
2.3	Bias, variance, and error for the utility-weighted sampling simulations of Ludvig et al. (2014)	53
2.4	Error for the utility-weighted sampling simulations using various distributions	58
2.5	Heatmap of error for the utility-weighted sampling simulations using various distributions	60
3.1	Screenshot of a decision trial between two options	93
3.2	Terminology use to describe each manipulation	94
3.3	Diagram of the designs of each experiment in Chapter 3	100
3.4	Choices and responses to the memory questions for Experiment 1	106
3.5	Choices and responses to the memory questions for Experiment 2	111
3.6	Choices and responses to the memory questions for Experiment 3	118
4.1	Diagram of the designs of each experiment in Chapter 4	133
4.2	Choices and responses to the memory questions for Experiment 4	141
4.3	Choices and responses to the memory questions for Experiment 5	147
5.1	Choices and responses to the memory questions for Experiment 6	160

6.1	Examples of the images used to represent options in Chapter 6	181
6.2	Responses to the EARS questionnaire and choices for Experiment 7	186
6.3	Responses to the EARS questionnaire and choices for Experiment 8	193
6.4	Responses to the EARS questionnaire, sampling, and choices for Experiment 9	201
6.5	The relationship between observable and outcome variability	207

List of Tables

3.1	Summary of the predictions associated with each class of theories	97
3.2	Summary of demographics in Chapter 3	98
4.1	The number of points associated with each option in Experiment 4	139
4.2	The number of points associated with each option in Experiment 5	146
5.1	Example of sequences presented in a random order or alternating between gains and losses	157
5.2	The number of points associated with each option in Experiment 6	157
B.1	Reliability estimates for each subscale	292
C.1	Goodness of fit indicators for confirmatory factor analysis	294
C.2	Confirmatory factor analysis loadings for Experiment 7	295
C.3	Confirmatory factor analysis loadings for Experiment 8	296

Chapter 1

General introduction

The most recent US presidential election was shrouded in speculation about an event that would usually be considered unlikely. What would happen if one of the major candidates dropped from the race during the campaign? As one newspaper headline contemplated, “Trump drops out. Biden gets sick. Pence is fired. What if 2020 gets really crazy?” (J. F. Harris & Lippman, 2020). Despite its salience during this particular election, this question was addressed nearly seventy years prior when Arrow (1951) asked how an analogous situation should influence constituents’ preferences regarding the other candidates. His proposed answer was that the winner should be determined by “blotting out the dead candidate’s name, and considering only the orderings of the remaining names” (p. 26).

In essence, why should your preference for one candidate over another change based on the availability of a third? If you were never going to choose that candidate either way or it has become unavailable, there is simply no reason why it should influence the option that you eventually select. To illustrate this further, consider a scenario in which a customer in a restaurant orders pizza from a menu that includes risotto as the daily special. Upon collecting the menu the waiter realises that it contains the special from the previous day and informs the customer that the special for today is, in fact, the gnocchi. Upon receiving this information, it would be strange if the customer were to respond with,

“In that case, I’ll take the lasagne”.

This process of blotting out options that are inferior or unavailable has been proposed as one of the foundational axioms of decision theory (Luce, 1959). It governs how people *should* make decisions and is an appealing picture for the scientist attempting to make sense of decision-making. It seemingly allows them to carve nature at its joints—into pairs of options—examine each preference in isolation, and construct an explanation by treating isolated preferences like Lego bricks. Whilst this is enticing, there is considerable evidence that this approach fails to adequately capture the decisions that people actually make.

Already by the early 1970s, Luce had conceded that “it is clear from many experiments that the conditions under which the choice axiom hold are surely delicate” (Luce, 1977, p. 215). Countless experiments have examined these conditions throughout the course of the subsequent decades and the unambiguous conclusion is that the way people evaluate options is influenced by the environments in which they are encountered (Huber et al., 1982; Parducci, 2011; Ronayne & Brown, 2017; Simonson, 1989; Spektor et al., 2018; Stewart et al., 2003; Tversky, 1972). Therefore, the axiom that options are evaluated as hermetically sealed units appears to describe a world that is materially different from the one that cognitive scientists are attempting to explain (Cartwright, 1999; Giere, 1999).

To make things worse, this affliction is not unique to the role of context and extends more generally to theories of decision-making under uncertainty. Simon (1982) observed that the expected utility model “is a beautiful object deserving a prominent place in Plato’s heaven of ideas”, but continued that “vast difficulties make it impossible to employ it in any literal way in making actual human decisions” (p. 13). Although these models are valuable in the small worlds that conform to their assumptions, this is an exception rather than the rule (Ellsberg, 1961; Keynes, 1921; Knight, 1921; Savage, 1954). Many decisions involve long run frequencies that are not stationary, possible outcomes that are not defined, situations that are difficult to categorise, small inaccuracies that lead to large errors, mechanisms that are not understood, and other assumptions that are violated in

practice (A. Gelman, 2018).

Similar criticisms of decision theory have circulated since at least the early 1700s when Nicholas Bernoulli raised what has become known as the St. Petersburg problem. These arguments are nothing new and neither are their rebuttals. For example, Friedman (1953) argued that “the relevant question to ask about the ‘assumptions’ of a theory is not whether they are descriptively ‘realistic’, for they never are, but whether they are sufficiently good approximations for the purpose in hand” (p. 14). In other words, simply pointing out that *Homo economicus* deviates from *Homo sapiens* is almost as misguided as debunking the frictionless planes of Galileo (McMullin, 1985). These models were born refuted because they were conceived as idealisations and emphasising their falsehood would be merely tilting at windmills (Lakatos, 1978; Potochnik, 2017; Wimsatt & Wimsatt, 2007).

It would be similarly misguided to suggest that these criticisms have been overlooked or ignored. Indeed, the heuristics and biases research programme that has dominated the field over the last half-century focuses heavily on deviations from probability theory. We all know that humans are not intuitive statisticians. Instead, our motivation for discussing these assumptions was eloquently summarised by Lopes (1983):

Everyone knows about risk from experience, and most people would agree that it has to do with uncertainty and with the possibility of loss. But the best way to discover how psychologists define risk is to study their experiments on risk. One feature stands out clearly: The simple, static lottery or gamble is as indispensable to research on risk as is the fruit fly to genetics. The reason is obvious; lotteries, like fruit flies, provide a simplified laboratory model of the real world, one that displays its essential characteristics while allowing for the manipulation and control of important experimental variables (p. 137).

Our point, following Lopes, is that cognitive scientists *know* that these models are descriptively inadequate but that their methods betray their roots in early 20th century

economic theory. Our point is also that they use these methods for a reason. Experiments that study decisions from experience almost exclusively present participants with a single pair of options and assume that they interpret uncertainty as probability (for a review, see Wulff et al., 2018). This approach to studying decision-making, epitomised in the classic bandit task, is well-suited to the good experimentalist.¹ It allows them to manipulate the variables in their theories and has led to the discovery of numerous robust empirical regularities, such as the under-weighting of rare events (Hertwig et al., 2004).

We are drawn to tasks that allow us to precisely control these variables but we also need to reflect on what might be lurking in the shadows. We know that context influences our decisions and we know that people deviate from the laws of probability. These variables have been investigated in tasks where participants were given explicit descriptions about outcomes and probabilities but are usually neglected when this information is acquired through experience. One pragmatic reason is that learning through experience requires multiple encounters with each option and this quickly becomes infeasible when there are too many. Moreover, researchers are often interested in the asymptotic preferences that people acquire with sufficient knowledge but these preferences can take hundreds of trials to stabilise (Luce & Suppes, 1965).

As a consequence, there is no standard laboratory model for studying the influence of context in decisions from experience and this makes the assumption that options can be studied in isolation all the more enticing. There is similarly no good model for examining different interpretations of uncertainty and we usually resort to methods that manipulate uncertainty using static probabilities. This substitution does not arise because cognitive

¹The bandit task is a common experimental method that has been used since the 1950s to examine how people make decisions from experience (Goodnow, 1955). Its name refers to the one-armed bandit, which was a classic gambling machine that had a single handle (arm). Each option in the standard bandit task is analogous to a *one-armed bandit machine*. The participant can click on the option (pull the handle) and experience feedback on the number of points earned (see whether they have won). The one-armed bandit used static probabilities and this has carried over to its laboratory counterpart. The bandit task almost exclusively employs a single pair of options but some notable exceptions include the Iowa gambling task (Bechara et al., 1994) and experiments examining exploration (Daw et al., 2006; Schulz et al., 2018; Wu et al., 2018).

scientists are oblivious to the differences between bandit tasks and real world decisions. You would never accuse a geneticist of being unable to distinguish between their *Drosophila melanogaster* subjects and their *Homo sapiens* colleagues. Although there are countless differences, enough of their genome has been conserved that many human genes have matching *Drosophila* sequences.

The reason that model organisms can differ extensively from their targets is that they are similar in the ways that matter for the underlying theory. The nematode worm, *C. elegans*, is used to study the human brain because it is similar enough based on theories of neural development and the laboratory rat is used to study human psychology based on theories of reinforcement learning. What are the analogous theories that underlie the bandit task in decisions from experience? From a historical perspective, early experiments based on slot-machine designs were used to study reinforcement learning and theories of expected utility (for reviews, see Bush & Mosteller, 1955; Luce & Suppes, 1956). Although these experience-based tasks fell out of fashion around the 1970s, they were reintroduced into the experimental milieu by Barron and Erev (2003), Hertwig et al. (2004), and others to assess the predictions of prospect theory, which had been studied almost exclusively using described options.

Without even considering contemporary developments, we have already identified three underlying theories: reinforcement learning, expected utility theory, and prospect theory. Nonetheless, although these theories are usually described as fundamentally dissimilar, they share a common notion of uncertainty as probability. On one hand, “The objective chances of success were 75 percent to 25 percent on the respective sides” and on the other there is “25% chance to win \$150 and 75% chance to win \$50”. Without knowing that the former is a description of a reinforcement schedule for rats completing a T-maze (Brunswik, 1939, p. 177) and the latter was a gamble offered to graduate students at Stanford (Tversky & Kahneman, 1992, p. 305), these could easily be mistaken for two descriptions of the same experiment.

This common notion of probability provides a reason for using bandit tasks as a

model for decision-making. It explains why an undergraduate student selecting coloured squares on a computer screen might tell us something about how people choose between stocks and bonds, decide whether to purchase insurance, and select which career path they want to pursue. These scenarios can be described as a series of outcomes and probabilities that can be mapped onto those used in decisions from experience tasks. There are compelling pragmatic reasons for using these simplified models and their use is warranted when the model is similar enough in ways that matter to the underlying theory.

One consequence of this account is that our theories and experiments exist in a reciprocal relationship. The first component of this relationship is that our experiments influence our theories. We use them to make observations, test hypotheses, encounter anomalies, and update our beliefs. Some version of these activities forms the basis of every approach to conducting scientific experiments but there is also a second material influence of experiments on our theories. As we mentioned above, there are pragmatic reasons for selecting tasks that can be completed by undergraduate psychology students. Theories that are not well-suited to these conditions are unlikely to receive much attention regardless of their other virtues.

The second component is that our theories influence our experiments. Theories matter because we are active participants in the growth of knowledge. We conduct experiments “not in the capacity of a pupil who lets the teacher tell him whatever the teacher wants, but in the capacity of an appointed judge who compels the witnesses to answer the questions that he puts to them” (Kant, 1996, p. bxiii). In this way, our theories dictate the questions that we pursue and the methods that we use. Although questions and methods are rarely determined by the same theory in a single experiment, they are situated within an ecology of theories, experiments, and observations that influence each other (Duhem, 1991; Hacking, 1992).

A specific theory can provide the rationale for using an experimental task. An observation can provide evidence against a theory. The feasibility of an experimental task can lead the researcher to pursue one question rather than another. Within this

ecology, they can even evolve self-sustaining or self-refuting systems where one component influences another, which influences others until the series catches its own tail, influencing the longevity of the original component (Kauffman, 1993; Maturana & Varela, 1991). For example, an experiment that immediately undermines the rationale for its own methods might leave the researcher wondering how they missed such an obvious shortcoming. This researcher would almost certainly abandon the experiment and might even decide to pursue other questions with established experimental paradigms.

How might this apply to decisions from experience? We have emphasised that scientists often know at least some of the limitations of their theories and methods. As is often the case, these limitations were discussed explicitly when decision-making research was still in its infancy. Edwards (1956) employed a slot-machine task with static probabilities but emphasised how people make choices based on hypotheses about sequences of rewards. Goodnow (1955) presented participants with choices framed as gambles or problem-solving to investigate the effects of whether uncertainty is perceived as purely stochastic or potentially resolvable. Perhaps to a greater extent than anyone else, she grasped the implications of how people interpret uncertainty and her experiments influenced the ones presented in the second part of this thesis.

Nonetheless, these concerns gradually faded into the background as researchers focused on solving the puzzles offered by the emerging mathematical theories of decision-making. How compatible are people's decisions with the principle of stochastic transitivity? What are the properties of the function that maps utility onto the value of outcomes? How do people allocate their choices when they encounter options with different probabilities of success? These questions were usually concerned with asymptotic behaviour and bandit tasks offer a simple method to manipulate important variables. Other tasks such as the problem-solving scenarios of Goodnow (1955) manipulate variables that are orthogonal to these questions and seemingly incorporate unnecessary complications.²

²Although our main focus is on problems-solving and gambling tasks, a similar argument could be made regarding experience- and description-based choices. Prospect theory provides no rationale for examining experience and this might explain why the differences between these tasks were not uncovered until prospect theory was old enough to enter a

These theories reinforce the bandit task as the standard laboratory model and it has been used to examine everything from loss aversion (Yechiam & Hochman, 2013) to differences associated with age (Frey et al., 2015). Although the aim of these experiments was seldom to examine the theories that support their methods, this in no way suggests that the bandit task was unable to reciprocate the favour. As we mentioned above, it quickly becomes infeasible to use these methods to examine contexts that comprise multiple options. Consequently, most experience-based tasks use a single pair of options and are incapable of demonstrating that context matters. They ensure that evidence never conflicts with the enticing assumption that preferences can be studied in isolation.

Furthermore, bandit tasks usually involve options that are distinguished using static features—colour or position—and their outcomes are usually associated with static probabilities. The researcher who programmed the experiment knows that the choice is analogous to tossing a coin where it is impossible to improve over time. These scenarios, however, are rare outside the casino or the psychology lab and the participant—who has been asked to make 500 choices between two coloured squares—might interpret their uncertainty as being soluble. They might devise spurious hypotheses regarding outcome sequences or the flickering of a light bulb but unless these variables are included in the data, these strategies will be indistinguishable from noise.

Thus, we have identified two components that might exist in a symbiotic relationship. The theory gives rise to questions that are well-suited to the experiment and the experiment conceals aspects of the theory that would otherwise conflict with observation. In contrast with the experiment that undermines its own theoretical rationale, they would comprise a self-sustaining system that ensures its own stability. Insofar as we have given an accurate description of this relationship, we have reason to question whether our theories capture something about the outside world or whether they prevail as a system that has become impervious to refutation. Have these theories and experiments taken on a life of their own rather than serving the aims of the researcher?

nightclub.

Before reaching this conclusion we should consider the evidence for each component of this relationship. We asserted that theory influences experiment because the good experimentalist selects tasks that manipulate the variables in their theories and avoids tasks that include extraneous variables. We made the further assertion that theories that interpret uncertainty as probability are well-suited to methods such as the bandit task. We suspect that these assertions are relatively uncontroversial.

The same cannot be said for the influence of experiment on theory. We asserted that the influence of context on decisions from experience might be underestimated because the standard task consists of only two options. We suggested that a similar issue might affect how people interpret uncertainty but what evidence could we possibly have for this? Both of these assertions are about gaps in our knowledge and even though the standard task is unable to examine these variables, this would be inconsequential unless we are actually underestimating their influence. The problem is that we cannot know this until we have an accurate estimate.

Nonetheless, we are clearly missing *something* because the behavioural methods we use to elicit risk preferences fail to converge not only with self-report measures but they even contradict each other (L. R. Anderson & Mellor, 2009; Frey et al., 2017; Holzmeister & Stefan, 2021; Pedroni et al., 2017). Our account of the feedback loop between theory and experiment might provide an explanation for why these issues are so enduring. Furthermore, the few experiments that have examined multiple options in decisions from experience suggest that the impact of context might diverge from decisions from description (Ert & Lejarraga, 2018; Hadar et al., 2018; Ludvig et al., 2014; Spektor et al., 2018). Not only might we need to consider what we are missing in decisions from experience tasks, we might be unable to simply generalise from other tasks where these attributes have been studied more extensively.

Therefore, our aim throughout the remainder of this thesis will be to examine two specific attributes that cannot be adequately examined using the standard bandit task with two options. In the first section, we will examine the impact of introducing mul-

multiple options into the context so that different options are associated with the best and worst experienced outcomes. In the second section, we will examine how people interpret uncertainty when there is observable variability in the experienced options. Although these sections ostensibly examine unrelated attributes, they both contribute to our understanding of the feedback loop between theory and experiment that we identified in this section.

1.1 Context and extreme outcomes

The precariousness of generalising from bandit tasks with a single pair of options was demonstrated in a series of experiments conducted by Ludvig, Madan, and colleagues (Ludvig et al., 2018; Ludvig et al., 2014; Ludvig & Spetch, 2011; Madan et al., 2014, 2017). Their participants encountered numerous choices between pairs of options and were able to learn about them by observing their outcomes. So far, this description is consistent with the standard bandit task but instead of using a single pair of options they interspersed a second pair so that the context consisted of four options. This minor modification had the major consequence that the best and worst outcomes in the context were no longer associated with the same option.

As a concrete example, the first experiment by Ludvig et al. (2014) involved a safe option that always resulted in 20 points and a risky option that resulted in either 0 or 40 points with equal probability. The second pair also included safe and risky options but the outcomes were mirrored across the zero point. Namely, the safe option always resulted in -20 points and the risky option resulted in -40 or 0 points. As we mentioned above, using either one of these pairs in isolation would mean that the best outcome (e.g., 40 points) and the worst outcome (e.g., 0 points) would be associated with the same option. The second pair ensured that the low-value risky option resulted in the worst outcome (-40 points) and the high-value risky option resulted in the best outcome (40 points).

These choices always involved a single pair of options but they were interspersed so

that participants encountered every possible combination. The most interesting comparisons, however, were between the safe and risky options within each pair. These options had the same expected value so it would be reasonable to assume that participants would be relatively indifferent between them. Alternatively, prospect theory suggests that they would be risk-seeking with losses and risk-averse with gains (Kahneman & Tversky, 1979). Participants choices were not consistent with either of these patterns. Instead, they selected the risky option substantially more when deciding between the high-value pair than the low-value pair.

Ludvig et al. (2014) conducted a series of followup experiments to investigate this curious pattern. They eliminated the possibility that it resulted from the zero outcome that was shared between the risky options, the absolute magnitude of the outcomes, or whether they were positive or negative. The only alternative that remained was that the difference between the pairs of options was their value *relative* to the other options in the context. Importantly, they demonstrated that an option is evaluated differently depending on the other *experienced* options but they did not observe any of these phenomena when the same options were *described* to participants.

Madan et al. (2014) suggested that this difference arises because people must remember and aggregate multiple outcomes when learning from experience but not when the outcomes and probabilities are described. They examined this possibility by presenting participants with the experience-based task described above and then asked them to estimate both the frequency of each outcome and which outcome came to mind first (Ludvig et al., 2018; Madan et al., 2014, 2017). They observed that the best and worst outcomes were over-represented on both of these memory measures. This observation is consistent with their explanation because a bias towards these extreme outcomes would make the low-value risky option appear worse and the high-value risky option appear better.

Ludvig et al. (2014) labelled this pattern *the extreme-outcome effect* and devised an *extreme-outcome rule* that states that our decisions are disproportionately influenced by the best and worst outcomes. This interpretation entails that these outcomes differ

categorically from intermediate outcomes but there are several other plausible accounts. In addition to *categorical extremity*, extreme outcomes could be defined based on their *continuous* or *ordinal* distance from the centre or edges of the experienced distribution or even their *temporal* relationship with other outcomes. The scenarios used by Ludvig et al. (2014) always consisted of options that were symmetrical around the average outcome and this renders each of the alternatives indistinguishable.

In the first section of this thesis, we will describe six experiments that aim to differentiate between these explanations. Our motivation for conducting this investigation consists of two components: The first of these relates specifically to decisions from experience and the potential importance of context that we sketched above. The second relates to the influence of extreme items more broadly throughout cognition. The notion that extreme items are disproportionately influential appears in many explanations across numerous cognitive domains (e.g., Fredrickson, 2000; Kunar et al., 2017; Neath, 1993). This raises the question of whether the observations of Ludvig et al. (2014) are specific to decisions from experience or whether they reflect a broader phenomenon.

The categorical definition endorsed by the extreme-outcome rule was heavily influenced by the *peak-end rule* that explains how people integrate instantaneous experiences into a holistic evaluation of an event (Fredrickson, 2000; Fredrickson & Kahneman, 1993). Instead of amounting to a summation of all the moment-by-moment experiences, these evaluations appear to incorporate only a small number of salient moments. For example, the simple average of the eponymous peak and end is a much better predictor of global evaluations than algorithms that incorporate the duration of experiences (Ariely & Carmon, 2000; Kahneman et al., 1993; Redelmeier & Kahneman, 1996). This rule either categorically includes or excludes moments from these evaluations and advocates an equally categorical definition of the peak as the single most extreme moment.

The similarities between the peak-end and extreme-outcome rules are inescapable. Both theories describe how people aggregate their experiences and their shared definition of extremity establishes the latter as a natural extension of the former that considers risky

decisions. Nonetheless, this categorical conception disqualifies them as explanations for a number of other phenomena that are explained using other conceptions of extremity. For example, either ordinal or continuous extreme outcomes are identified with greater sensitivity (Berliner et al., 1977; D. L. Weber et al., 1977), capture more attention (Kunar et al., 2017; Pleskac et al., 2019; Zeigenfuse et al., 2014), and are retrieved more easily than mid-sequence items in free recall, forwards and backwards serial recall, and recognition tasks (Capitani et al., 1992; Li & Lewandowsky, 1995; Murdock, 1962; Neath, 1993; Wright et al., 1985).

The observations of Ludvig and colleagues and—we will argue—those underlying the peak-end rule can be explained using an ordinal or continuous definition but these other phenomena cannot be explained using a categorical one. Therefore, what is at stake is not merely our understanding of decisions from experience but also the plausibility of a unified explanation for the influence of extreme items across cognition. Our experiments will raise serious issues for the categorical conception but to understand the structure of the first section, we must begin by introducing a distinction between two explanatory modes: *mechanical* explanations that describe *how* a mechanism gives rise to the influence of extreme outcomes and *rational* explanations that describe *why* these outcomes are influential.

1.1.1 Mechanical explanations

The mechanical explanatory mode is used in the vast majority of explanations for the influence of extreme outcomes and arguably the majority of explanations throughout cognitive science. These explanations describe component entities (or parts) whose actions and organisation gives rise to the phenomena of interest (Bechtel & Abrahamsen, 2005; Craver, 2006; Glennan, 2017; Machamer et al., 2000). The distinctiveness account proposed by Brown and colleagues (2007) offers a concrete example of how these elements of the mechanical explanatory mode can be used to explain the influence of extreme outcomes. Their theory implies that *items* stored in memory (entities) *interfere* with the

retrieval of other memory traces (actions) based on the *distance* between these items along a psychological dimension (organisation). This mechanism explains the *influence of extreme outcomes* (phenomenon) because these outcomes have fewer close neighbours than intermediate outcomes. In other words, extreme outcomes are located in a sparser region of psychological space where there is less retrieval interference and this enhances memory performance.

The mechanical explanatory mode has the potential to provide a unifying explanation for a seemingly diverse range of phenomena using a common mechanism. The challenge with implementing this strategy is that we must differentiate between common and separate mechanisms. How is this accomplished? There is both an analytical and an empirical component to this process. The first component is analytical because the reach of a well-specified mechanistic explanation can be derived from the attributes of the mechanism itself.³ For example, a narrow explanation that opium causes sleepiness by virtue of its dormitive qualities has little applicability beyond this specific phenomenon (Molière & Laun, 1673). In contrast, many of the fundamental theories in physics have almost universal applicability (Deutsch, 2011). This component of the reach of an explanation depends on the attributes of the mechanism rather than the intentions of the researcher, and therefore, it often extends well beyond the phenomenon that it was designed to explain.

Returning to our example of the distinctiveness account, although it was primarily designed to explain memory-based phenomena, the mechanism could easily be used to explain downstream effects on decision-making. Even with phenomena that are less reliant on memory, this account might appeal to broader mechanisms that underlie processes such as discrimination between items based on psychological distance. Indeed, Brown and colleagues (2007) emphasise that their model imports mechanisms that were developed to

³When the explanation is not well-specified, this analysis results in ambiguity. For example, de Beauvoir (1949) employed this criticism against psychoanalysis: “If one criticises the doctrine to the letter, the psychoanalyst maintains that its spirit has been misunderstood; if one approves of the spirit, he immediately wants to limit you to the letter.” A well-specified explanation is one in which the spirit and the letter are inseparable.

explain phenomena related to identification (Murdock, 1960) and categorisation (Nosofsky, 1986). Such a common mechanism that encompasses discrimination-based activities might imply that extreme outcomes will be influential “whenever a set of items can be sensibly ordered along a particular dimension” (M. R. Kelley et al., 2015, p. 1715).

The empirical component examines whether the mechanism is able to parsimoniously explain the phenomena of interest without resorting to domain specific assumptions. For example, in the distinctiveness account, interference primarily occurs between outcomes that are close neighbours. This allows the model to explain numerous well-established phenomena in the memory literature, notably that memory recall is easier for items with fewer neighbours even when they are in the centre of the distribution (Hunt, 1995; Neath et al., 2006; Wallace, 1965). This attribute might, however, restrict the reach of the mechanism because there is some evidence that this effect is absent in decision-making (Ludvig et al., 2018). The presence of this effect in memory and its absence in choice might imply that a separate mechanism underlies the role of extreme outcomes in each domain.⁴ Therefore, beginning in the third chapter, we will conduct a comprehensive empirical examination of numerous candidate mechanisms for the influence of extreme outcomes.

1.1.2 Rational explanations

The mechanical explanatory mode addresses the question of *how* extreme outcomes influence cognition and this is often what is sought in an explanation. But focusing solely on mechanisms can lead us to “act as if God created the mind more or less arbitrarily, out of bits and pieces of cognitive mechanisms, and [that] our induction task is to identify an arbitrary configuration of mechanisms” (J. R. Anderson, 1990, p. 26). The question

⁴Admittedly, the presence of this effect may depend on task demands (requiring temporal or positional encoding) and is not always present in memory (e.g., forward serial recall) (Morin et al., 2010). For a second example, see Brown and colleagues (2008) for a discussion regarding empirical evidence for a common mechanism underlying serial and free recall.

remains of *why* these seemingly arbitrary mechanisms exist and the rational explanatory mode provides an answer by suggesting that they are in some sense optimal or rational.⁵ In evolutionary biology, optimality models are used to explain the prevalence of traits within a population (Parker & Smith, 1990). In economics, the behaviour of agents is explained with reference to expected utility theory and stable equilibrium points (Nash, 1950; von Neumann et al., 1944). Rational analysis provides a framework for the ACT-R cognitive architecture (J. R. Anderson et al., 2004), ideal observer models are used to explain visual perception (Geisler, 2011), optimal foraging theory explains the movement of animals (Charnov & Orians, 1973), and ecological rationality explains the benefits associated with using simple heuristics (Gigerenzer & Todd, 1999).

Fredrickson (2000) explained the motivation underlying the rational explanatory mode regarding the role of the most extreme moments (peaks) and the final moments (ends) in affective experiences. “Why? Of all the moments people could select to represent past affective experiences, why do they choose peaks and ends? Is it perceptual? Are peaks and ends simply more salient than other moments? This is unlikely to be the whole story” (p. 589). In response to this question, Fredrickson proposed a rational explanation that extreme outcomes contain self-relevant information about the capacity required to endure an event or achieve an outcome. In other words, as long as someone is capable of handling the worst-case scenario, they are unlikely to be overwhelmed by less extreme outcomes. Although future events sometimes exceed the experienced maxima, the most extreme outcome in each interval is commonly used to estimate the probability of extreme natural disasters, insurance losses, and stock market risks (Beirlant et al., 2004). It is plausible that encoding the most extreme moment of each event enables a similar technique for avoiding outcomes that would surpass capacity thresholds.

The rational explanatory mode has the potential to offer a unifying explanation

⁵There is a close relationship between these explanatory modes and Marr’s (1982) levels of analysis. Rational explanations and Marr’s computational level both address *what* a mechanism does and *why* this is appropriate whereas mechanical explanations and Marr’s algorithmic level both address *how* this is implemented (for further discussion of this relationship, see Bechtel & Shagrir, 2015; Kaplan, 2017; Milkowski, 2013; Zednik, 2017).

based on the attributes of a common reason. To accomplish this, we must distinguish between common and separate reasons, and similarly to our analysis of common mechanisms, this process combines analytical and empirical components. In Fredrickson’s capacity threshold account, described above, encoding the peaks enables people to avoid situations that would otherwise harm them by exceeding their capacity. This readily explains why unpleasant experiences are encoded based on the worst moment (e.g., Fredrickson & Kahneman, [1993]) but how might this account explain why positive experiences are encoded based on the best moment? Fredrickson suggests that positive outcomes also strain our personal capacity, and to the extent that we are willing to accept this attribute of their explanation, its reach might extend to positive outcomes. It is less plausible, however, that there are similar consequences associated with neutral perceptual stimuli—a tone with the lowest frequency or a circle with a largest diameter—and this limits the applicability of their explanation.

The empirical component examines the extent to which a given rational explanation offers an adequate description of the goals and capacities of the agent and the environment in which they function. An example in which this empirical component became exceedingly salient was the stock market crash of 2008 (Gigerenzer & Gaissmaier, [2011]). The axioms of decision theory provide the normative force for the majority of economic models but this does not guarantee that the small worlds depicted in them accurately correspond to the real world in which people live (Cartwright, [1999]; Giere, [1999]; Savage, [1954]). Observing that their models had failed to capture the complexity of the market, Lehman Brothers’ head of quantitative research lamented that “Events that models only predicted would happen once in 10,000 years happened every day for three days” (Whitehouse, [2007]). This demonstrated without a doubt that even models with rigorous mathematical foundations can be misleading unless their translation to the world is carefully examined.

1.1.3 Integrative explanations

The two explanatory modes offer two distinct strategies for understanding the pervasiveness of extreme outcomes in cognitive explanations: we can establish a common mechanism or a common reason. These strategies are not mutually exclusive.⁶ To demonstrate this, let us return once more to our previous example of the distinctiveness account. Although this theory derives much of its explanatory force from its mechanism, Brown and colleagues (2007) also assert that their mechanism might “arise for the same reasons” across multiple domains (p. 542). One of these common reasons is undoubtedly the intimate connection between their model and Shepard’s universal law of generalisation (Chater & Brown, 2008). Specifically, their similarity function is analogous to the one that Shepard (1987) derived from universal principles of probabilistic geometry and natural kinds. This provides evidence for a mechanism that incorporates this similarity function due to selective pressure towards an “increasingly close approximation in sentient organisms wherever they evolve” (Shepard, 1987, p. 1323).

Within explanations that integrate these strategies, the answers to “Why?” that are given by rational explanations constrain the plausible answers to “How?” that are given by mechanical explanations. These constraints are useful because the empirical component of mechanical explanations is typically under-determined. Nonetheless, one need only survey the enormous catalogue of irrational behaviour that has accumulated in the heuristics-and-biases research programme to conclude that responses to *normative* “Why should...?” questions must be distinguished from responses to *descriptive* “Why is...?” or “Why

⁶They are also not necessarily exhaustive. Following Cummings (2000), functional explanations are often presented as the canonical mode of explanation in cognitive science. Although the relationship with other explanatory modes is beyond the scope of this chapter, it should be noted that our definition of “mechanism” does not imply that psychological explanations are mechanical to the extent that they correspond to the brain (see Barrett, 2014; Kaplan & Craver, 2011; Piccinini & Craver, 2011; Shapiro, 2017; Weiskopf, 2011). Therefore, at least for our purposes, functional explanations can be interpreted as a special case of the mechanical explanatory mode that minimises the emphasis on *robust entities* whilst retaining attributes such as decomposition (Glennan, 1996; Piccinini & Craver, 2011).

will...?” questions (Maccrimmon, [1968]; Tversky & Kahneman, [1974]). Furthermore, from our modern scientific perspective, expecting the world to conform to our rational models seems to imply a problematic backwards causality in which the rationality of the final mechanism is the reason that it possesses certain attributes. These peculiarities of descriptive rational explanations has led many researchers to approach them with outright scepticism (see the numerous commentaries and special issues that discuss rationality: J. R. Anderson, [1991]; L. J. Cohen, [1981]; Dennett, [1983]; Elqayam & Evans, [2011]; Felin et al., [2017]; Jones & Love, [2011]; Kyburg, [1983]; Lieder & Griffiths, [2019]; Oaksford & Chater, [2009]; Schoemaker, [1991]; Stanovich & West, [2003]; Todd & Gigerenzer, [2000]).

Thus to avoid being dismissed as unscientific, we must develop a naturalistic account that tolerates deviations from rationality and is consistent with our understanding of causality. Towards this aim, consider the simple scenario of a marble released on the rim of a spherical bowl. From this minimal description, most people would agree on three things: 1) at least in principle, we could calculate the causal trajectory of the marble using the established laws of physics, 2) even without performing this calculation, we can be fairly certain that the marble will come to rest at the bottom of the bowl, and 3) we can explain this prediction without referring to dubious notions such as final causes or divine intervention. Instead, the shape of the bowl constrains the set of possible causal trajectories and this allows us to disregard most of the complexity while remaining confident that the actual trajectory will lead to the predicted equilibrium (Rice, [2015]; Strevens, [2003]). Rational explanations provide evidence regarding mechanisms in the same way. They constrain the set of causal trajectories such that there are often only two plausible outcomes: “organisms either know the elements of logic or become posthumous” (Fodor, [1981], p. 121).⁷

⁷One foreseeable counterargument is that our rational landscape analogy is *prima facie* implausible because rational explanations can be employed without interpreting them as constraints on a causal trajectory. For example, J. R. Anderson ([1990]) argues that rationality should be seen as a *scientific hypothesis* that cannot be proven or disproven using “a priori considerations” but instead must be evaluated on how well it does “leading to successful theory” (p. 29). This approach is not incompatible with our account. Instead, the scientific hypothesis approach allows the researcher to recognise that many cognitive mechanisms are approximately rational and the rational landscape analogy explains why

Admittedly, the easily recognisable constraints imposed by the spherical bowl contrast with the intricacy of those that arise from the ability of an agent to pursue their goals or the blind watchmaker of natural selection. These scenarios are analogous to a rugged landscape in which the marble might end up at the bottom of the deepest valley or a shallow crater in the peak of the highest mountain. Descriptive rational models serve as topographic maps of this landscape but when there is more than one equilibrium, it becomes impossible to predict the outcome without information regarding the actual trajectory (Gould & Lewontin, [1979](#); Kauffman & Levin, [1987](#); Marcus, [2008](#); Rellihan, [2012](#)). In other words, relying exclusively on rational models to identify mechanisms encounters the same problem of underdetermination that arises when using other empirical methods.

Have we merely ended up back where we started? Not necessarily. Although the mechanical and rational explanatory modes are typically underdetermined on their own, integrating them allows evidence from each one to constrain the other—like multiple cross-word clues—to narrow down plausible mechanisms. On one hand, we can conduct experiments that provide evidence regarding “the access to information and the computational capacities that are actually possessed by organisms” (Simon, [1955](#), p. 99). This information can be used to constrain the plausible area of the rational landscape in which the mechanism might reside. On the other hand, we can employ these boundedly rational models to further narrow down to the mechanisms that conform to constraints imposed by rationality. When employed in succession, these integrative constraints form a virtuous spiral that allows the initially inadequate explanations to effectively pull each other up by their bootstraps.

Whereas the rational landscape analogy does not provide specific instructions for using these constraints to identify mechanical explanations, Lieder and Griffiths ([2019](#)) have advanced *resource-rational analysis* as a concrete methodological framework. The first stage in their approach is to identify a problem that organisms might encounter and a

this regularity occurs. The former is analogous to estimating outcome probabilities by repeatedly tossing a coin whereas the latter considers the relationship between the geometry of the coin and these probabilities.

class of plausible algorithms that can solve it. These algorithms are subsequently narrowed down to the single algorithm with the highest performance given resource limitations. Although this strict optimality assumption is untenable as a substantive claim about the mind (e.g., Gould & Lewontin, 1979), it might serve as a “methodological device to efficiently search through the endless space of possible mechanisms” (Lieder & Griffiths, 2019, p. 45).⁸ It offers somewhere to begin searching but the process does not end there. The predictions of the optimal algorithm are evaluated against empirical evidence. Then the procedure is repeated with an increasingly accurate understanding of the constraints and mechanisms until an adequate explanation is discovered.

Resource-rational analysis has demonstrated considerable promise and has been used to explain observations across numerous cognitive domains (for a review, see Lieder & Griffiths, 2019). Although we might attribute this success to the integration between mechanical and rational explanations, their approach is also notable for the inspiration that it draws from computer science. Many of the potential cognitive mechanisms suggested using the framework were originally designed to increase the efficiency of estimation algorithms. This translation between machines and humans relies on the observation that we encounter similar computational challenges and must solve them efficiently because we have finite resources. As such, algorithms that have been selected by human engineers might point us towards ones that were favoured by natural selection. This approach is the foundation on which Lieder et al. (2018) constructed their rational utility-weighted sampling model that we will examine in the second chapter.

⁸This methodological approach can be traced back to Kant’s Critique of the Power of Judgement. He emphasises the difficulty in explaining biological organisms “in accordance with merely mechanical principles of nature” (p. 270), and therefore, acknowledges that teleological notions are necessary as “a heuristic principle for researching the particular laws of nature” (p. 280).

1.1.4 Summary of Chapter 2: Rational explanations

Lieder et al. (2018) suggest that the disproportionate influence of extreme outcomes stems from an optimal Monte Carlo algorithm that uses fewer samples to accurately estimate the value of each option. Their algorithm is based on a statistical method that reduces estimation error by prioritising outcomes that have the greatest impact on the estimate (Kroese et al., 2013; Tokdar & Kass, 2010). The accuracy of these estimates depends on how well important outcomes are identified. Lieder and colleagues demonstrated that sampling outcomes based on their distance from the expected value of the option being estimated minimises the variance. The catch, however, is that this requires knowledge of the expected value of the option, which is the parameter being estimated.

Therefore, they suggest that this optimal algorithm can be approximated by replacing the expected value of the *option* with the average outcome in the *context*. This substitution underlies the potential of their utility-weighted sampling model to increase the accuracy of estimates and explains the influence of extreme outcomes. The resulting algorithm is biased towards outcomes based on their continuous extremity and is able to capture choices and memory in decisions from experience (Ludvig et al., 2014; Madan et al., 2014), the overestimation of extreme events (Lichtenstein et al., 1978), the temporal dynamics of the Technion choice prediction tournament (Erev et al., 2010), and the pattern of preferences described in prospect theory (Allais, 1953; Lichtenstein & Slovic, 1971; Tversky & Kahneman, 1992).

Despite these promising credentials, we will discuss three attributes of their model that undermine both the mechanical and rational components of their explanation. The first attribute demonstrates that the sampling mechanism produces highly counterintuitive predictions that are unlikely to capture actual behaviour. The second attribute uses the flexibility of its free parameters to avoid these implausible predictions but undermines the empirical evidence for the model. The third attribute demonstrates that the scenarios in which utility-weighted sampling improves the accuracy of estimates is limited to options that are close to the centre of the distribution. The combination of these attributes suggest

that utility-weighted sampling does not offer an adequate explanation for the influence of extreme outcomes.

In response to this conclusion, we will describe an alternate rational model based on the consequences associated with forgetting extreme outcomes. Compared with those near the centre of the experienced distribution, forgetting an extreme outcome has a greater impact on the probability of selecting the correct option and the expected value of the choice. We will begin this investigation with a formal analysis of the model using binary choices between safe options and then generalise our conclusions to choices involving multiple risky options. In contrast with utility-weighted sampling, our model suggests that there are *always* benefits associated with remembering extreme outcomes and this might contribute to an explanation for why these outcomes are influential. Finally, we will discuss the implications of this model for the definition of extremity and its relationship with other potential rational explanations.

1.1.5 Summary of Chapters 3-5: Mechanical explanations

The subsequent chapters in the section on context examine potential mechanical explanations for the influence of extreme outcomes. Numerous explanations have been suggested including goal-directed attention (Tsetsos et al., 2012), environmental characteristics (Lichtenstein et al., 1978), task relevance (Vanunu et al., 2020), outcome salience (Madan et al., 2014), distinctiveness (Murdock, 1960; Neath et al., 2006), edge-based encoding (Berliner et al., 1977), bias-variance optimisation (Lieder et al., 2018), selective attention (Luce et al., 1976), personal meaning (Fredrickson, 2000), and implicit reference points (Holyoak, 1978). It would not be feasible to test each of these theories separately so we will commence our investigation by developing a framework that categorises them along numerous dimensions.

Chapter 3 examines the dimension that differentiates between theories based on their *levels of measurement*. As we emphasised above, the peak-end and extreme-outcome rules employ a *categorical* definition in which the best and worst outcomes are considered

extreme and the intermediate outcomes are considered non-extreme (Fredrickson, [2000]; Ludvig et al., [2014]). There are also numerous explanations that define extreme outcomes based on their *ordinal* position within the distribution of experienced outcomes (Kunar et al., [2017]; Pleskac et al., [2019]; Tsetsos et al., [2012]; Vanunu et al., [2020]; Zeigenfuse et al., [2014]). This conceptualisation is based on the notion that many of our interactions with the world are based on rank: attending to one item entails neglecting another and selecting an option entails rejecting others. There are considerable differences between the members of this class but their ordinal definition gives rise to some shared empirical predictions.

The third level of measurement defines extreme outcomes based on a measure of continuous distance but the members of this class differ along a second dimension depending on their answer to the question: continuous distance from what? Some theories, such as utility-weighted sampling, define extreme outcomes based on their distance from the *centre* of the distribution (Lieder et al., [2018]). Others suggest that items are encoded with reference to the *edges* of the distribution and that extreme outcomes are remembered more easily because of their proximity to these anchors (Berliner et al., [1977]; Braida et al., [1984]; Farrell & Lelièvre, [2009]; Henson, [1998]; Jou, [2010]; Marley & Cook, [1984]). Finally, some theories suggest that items in memory interfere with the retrieval of similar items, and therefore, items that have fewer close *neighbours* are more likely to be remembered (M. C. Anderson & Neely, [1996]; Schmidt, [1991]). On average, outcomes located near the edges of the distribution have fewer close neighbours than outcomes near the centre and this would explain their influence in memory (Berliner et al., [1977]; Bower, [1971]; Eriksen & Hake, [1957]; Lacouture & Marley, [2004]; Luce et al., [1982]; Murdock, [1960]; Neath et al., [2006]; D. L. Weber et al., [1977]).

Chapter 4 examines a distinction between extreme outcomes based on *types* and *tokens*. This distinction has been employed primarily in linguistics (Richards, [1987]; Templin, [1957]) and philosophy (Fodor, [1974]; Putnam, [1975]; Quine, [1987]) but can be illustrated using the following passage from *The Bells* by Edgar Allan Poe:

To the swinging and the ringing
Of the bells, bells, bells,
Of the bells, bells, bells, bells,
Bells, bells, bells—
To the rhyming and the chiming of the bells!

How many words are in this section of the poem? One reasonable response would be that there are 29 words in the passage but another would be to notice that some are repeated and that there are nine unique words: “to the swinging rhyming chiming and ringing of bells”. There is a sense in which both of these responses are correct: the first is the number of word tokens whereas the second is the number of word types.

When someone experiences numerous instances of the same type of outcome, the type-token distinction can influence how we define extreme outcomes. As a concrete example, imagine that you were offered a role with a salary of \$75000 in a company that has a pay transparency policy. Of the ten existing employees, seven of your potential colleagues receive \$50000, one receives \$95000, and one receives \$100000. The only salary level lower than yours is \$50000 and there are two higher salary levels, so based on types, you might conclude that your offer is on the lower end of the pay scale. Conversely, your potential salary would be better than seven of your colleagues and only worse than two, so based on tokens, you might conclude that the offer is on the higher end of the scale. This demonstrates that whether outcomes are considered extreme can depend on how they are represented.

Finally, Chapter 5 examines whether extreme outcomes are conceptualised based on their location within the experienced *distribution* or their *temporal* relationships. Every single one of the distinctions examined in Chapters 3 and 4 are based on distributional characteristics. Whether an outcome is the best or worst, how many neighbours it has, and its rank within the distribution do not depend on whether the other outcomes occurred immediately beforehand or preceded them by several minutes. Nonetheless, extreme outcomes can also be conceptualised by whether they constitute a *temporal peak* relative to

the outcomes immediately before and after them.

One reason that events are remembered based on moments such as the peak and end might be that many experiences are not made up of discrete outcomes (Langer et al., 2005). For example, painful medical procedures usually involve some moments that are better and others that are worse but the overall experience is continuous rather than discrete. Integrating all the moments that comprise this experience poses a computational challenge that could be overcome by summarising events using a small number of salient aspects (Ariely & Carmon, 2000). The temporal peaks—the moments where the experience stops getting worse and starts getting better—might offer one such particularly salient aspect of continuous experiences.

1.2 Aleatory and epistemic uncertainty

The concept of probability has become inescapable when discussing how we interpret uncertainty. This has not always been the case. Its dominance can be crudely described as acquired through two revolutions that were separated by three hundred years. In fact, although there were several *qualitative* notions of evidence and authority, attempts to *quantify* uncertainty were essentially non-existent before the first probabilistic revolution in the middle of the seventeenth century (Hacking, 1975; Hacking et al., 1990). Beginning with Pascal, Huygens, Leibniz, and the other early mathematical probabilists, the concept immediately became indispensable for a broad range of questions from games of chance to jurisprudence (Daston, 1995).

What are the fair odds for a gamble? Is it rational to accept a hypothesis? Should we believe in God? Probability soon acquired a normative character as underlying rational beliefs in response to evidence. It was described as “good sense reduced to calculus” (Laplace, 1814) and as “the very guide of life” (Butler, 1736). But in addition to this *normative* status, probability was seen as equally *descriptive* of how people make decisions (Daston, 1995). It was naturally interpreted in light of the early associationist idea that

degrees of belief are attributed “in proportion as we have found it to be more or less frequent” (Hume, 1748, p. 41).

A similar pattern played out in the computational and inferential revolution of the 1950s. Following the newly axiomatised systems of von Neumann et al. (1944) and Savage (1954), the notion of a utility-maximising Homo economicus was reminiscent of the normative status of probability at the height of the European Enlightenment. There was also a radical shift in the way that we described the mind. As emphasised by Gigerenzer et al. (1989), “Once psychologists came to view statistics as an indispensable method, it was not long before they began to conceive of the mind itself as an intuitive statistician” (p. 203). In other words, the computational metaphor and the development and application of statistical procedures meant that uncertainty within our psychological theories was represented almost exclusively as probability.

This presents us with a curious puzzle. How can a concept that was conceived a few centuries ago explain something that is fundamental to the human condition? Can probability explain how people dealt with uncertainty before the 1600s? What about people in 2022 that have never studied mathematics? We at least need to take seriously the possibility that the experience of uncertainty might not be captured using probability. Uncertainty was a psychological reality long before probability was a set of axioms and we might even need to consider the possibility that uncertainty comprises numerous qualitatively distinct concepts.

Hacking (1975) emphasised one such distinction between what he labelled *epistemic uncertainty* (derived from the Greek word for “knowledge”) and *aleatory uncertainty* (derived from the Latin word for “dice”). The difference between these two notions of uncertainty can be grasped in the following questions: “Who is the prime minister of the UK?” and “What will be the outcome when I toss this coin?” Given the current political landscape, both of these questions would likely entail some degree of uncertainty. You might even exclaim that there is a 50/50 chance that your answer to the first question is correct but would this render your uncertainty regarding these questions equivalent?

Although the probabilities in both cases are identical, your uncertainty regarding the first question would be perceived as arising from your own *lack of knowledge*. You could resolve this uncertainty by reading today’s newspaper or doing a quick Google search. The question is about a specific instance for which the truth is knowable, in principle. In contrast, you would likely perceive the second question as reflecting an *inherently stochastic* process. In some sense, this answer is also knowable in advance, but unless you have some very high-tech equipment, the only way to learn the outcome would be to simply toss the coin.

These two faces of uncertainty can be recognised throughout the history of probability and different versions of the distinction have been proposed by philosophers, scientists, and statisticians (Carnap, [1945]; Cournot, [1843]; Fox & Ülkümen, [2011]; Kahneman & Tversky, [1982]; Poisson, [1837]; Russell, [1948]; Savage, [1954]). There are almost as many names for knowledge-based and stochastic forms of uncertainty as there are people who have written about them. In addition to epistemic and aleatory uncertainty, they have been referred to as probability and chance, subjective and objective possibility, probability₁ and probability₂, credibility and probability, personalistic and objectivistic probability, and internal and external uncertainty. Hacking ([1975]) suggested that the distinction has a life of its own so that “the same idea crops up everywhere, on the pens of people who have never heard of each other” (p. 16).

He also noted, however, that most people who use probability are oblivious to the distinction and this raises the question of whether it only surfaces for a select few academics who have thought about uncertainty for too long. Nonetheless, statements expressing epistemic and aleatory uncertainty diverge across a wide range of attributes in natural language (Juanchich et al., [2017]; Løhre & Teigen, [2016]; Olson & Budescu, [1997]; Teigen, [1988]; Ülkümen et al., [2016]). People respond differently to tasks involving inadequate knowledge about previous events and those involving a stochastic process that will occur in the future (Beck et al., [2011]; Chua Chow & Sarin, [2002]; A. J. Harris et al., [2011]; Heath & Tversky, [1991]; Robinson et al., [2009]; Robinson et al., [2006]). Investors are more willing to pay a financial advisor when they interpret their uncertainty regarding the stock

market as reflecting their own ignorance (Walters et al., 2022). And finally, distinct fMRI activation patterns have been observed in tasks that involve inadequate knowledge and those that involve stochastically determined outcomes (Volz et al., 2004, 2005).

1.2.1 Summary of Chapter 6: Variability and uncertainty

The unitary concept of probability appears to neglect the expansive and consequential duality of uncertainty. The distinction between epistemic and aleatory uncertainty has emerged numerous times and influences the behaviour of people that have never even heard of the distinction. We *know* that the concept of probability might not capture the nature of uncertainty. Nonetheless, most decisions from experience tasks involve static probabilities and options are represented using coloured squares that are identical each time they are encountered. These slot-machine-style tasks are well-suited to studying aleatory uncertainty but they usually neglect its epistemic counterpart—we might be missing half the picture.

To illustrate this possible issue, consider the example with which Hertwig et al. (2004) began their classic article on decisions from experience: Patients and their doctors often make decisions using information that is similar in its content but is acquired through different sources. The patient will often google the procedure and read a written *description* of the probabilities associated with each outcome. The doctor has access to this information but they also have extensive personal *experience*, gathered across many patients, and use this information to estimate the probabilities associated with each outcome. Thus, the doctor in this scenario offers an archetypical example of a decision from experience and we might learn something about uncertainty by examining its relationship with their experimental task.

Similarly to the options encountered by participants, when a patient comes in with a sore throat, the doctor can either prescribe them antibiotics or bed rest, plenty of fluids, and a good series on Netflix. They encounter hundreds of patients with the same condition and receive feedback regarding the effectiveness of each treatment. The patient prescribed

antibiotics is more likely to come back complaining of digestive issues but is less likely to return with a persistent strep infection. Eventually, the doctor will learn the probabilities associated with each of these outcomes and use this to inform their choices.

This description is more or less analogous to the participant who is repeatedly presented with two unlabelled buttons and receives feedback when one is selected. But in contrast with the experimental task, the doctor also encounters observable variability *within* each class of options. When they prescribe someone antibiotics, they are not merely choosing between a generic Option A (antibiotics) and Option B (bed rest) based on the probability of success. They are choosing between options for an individual patient who varies along myriad potentially relevant dimensions. They have access to medical records with the patient’s age, chronic illnesses, family history, whether they smoke or drink, and prior treatments for the present condition.

This observable variability would elicit a completely different interpretation of uncertainty for the doctor and the participant choosing between unlabelled buttons. The uncertainty regarding an individual patient can often be resolved by mapping potential outcomes onto observable variables. The doctor is likely to interpret their uncertainty as reflecting insufficient knowledge (epistemic uncertainty) whereas the participant is more likely to interpret their uncertainty as inherently stochastic (aleatory uncertainty). Given that the doctor believes that their uncertainty is resolvable—at least in principle—they are also more likely to seek additional information by asking questions, doing a physical exam, or ordering tests (Walters et al., 2022).

Therefore, Chapter 6 examines whether introducing observable variability to a standard bandit task impacts how participants interpret uncertainty. As we mentioned above, our experiments were partly inspired by the gambling and problem-solving tasks devised by Goodnow (1955), but whereas her experiments focused on the propensity of participants to exhibit probability matching, our tasks examine information-seeking. The experiments in this chapter will examine whether participants that are able to map potential outcomes onto observable variability are more likely to experience epistemic uncertainty.

They will examine whether—analogue to the doctor ordering tests—participants that interpret their uncertainty as epistemic are more likely to seek additional information.

1.3 Thesis overview

The overarching aim of this thesis is to contribute to our understanding of how context and uncertainty influence decisions from experience. To comprehend the role of these attributes, we first needed to broaden our perspective to encompass the system in which they are situated. Although context and uncertainty could be examined separately, the relationship between them illuminates the symbiotic relationship between theories and experiments. Specifically, we argued that interpreting uncertainty as probability gives rise to questions that are conducive to bandit task experiments. These experiments reciprocate by obscuring aspects of these theories that would otherwise conflict with observation. Our exploration of context and uncertainty, therefore, aims to perturb this self-sustaining system to see whether the theories and experiments can stand alone.

Second, we needed to narrow our focus to a tractable question regarding each of these attributes. Our examination of context in the first section focuses on the disproportionate influence of extreme outcomes. Chapter 2 addresses *why* these outcomes are influential by reevaluating the utility-weighted sampling model and then developing an alternative with a broader domain of applicability. The subsequent chapters examine the empirical adequacy of numerous mechanical explanations by organising them according to three primary dimensions. Chapter 3 examines whether extreme-outcome phenomena reflect a categorical, ordinal, or continuous level of measurement. Chapter 4 examines whether extreme outcomes are represented as types or tokens. And finally, Chapter 5 examines whether these phenomena are effectuated by temporal or distributional characteristics.

In the second section, our examination of uncertainty focuses on epistemic and aleatory uncertainty. This distinction has been observed across numerous domains but the bandit tasks used to study decisions from experience were not designed to capture

CHAPTER 1. GENERAL INTRODUCTION

epistemic uncertainty. Therefore, Chapter 6 examines whether introducing observable variability to these tasks promotes an epistemic interpretation of uncertainty, and in turn, whether attributing uncertainty to inadequate knowledge elicits information-seeking behaviour. Finally, Chapter 7 summarises the main findings regarding context and uncertainty and integrates them into a unified system.

Part I

Context

Chapter 2

Rational explanations

Items that are located on the edges of a given distribution appear to exert a disproportionate influence across numerous cognitive domains. These extreme outcomes are discriminated with greater sensitivity (Berliner et al., [1977]; D. L. Weber et al., [1977]) and capture more attention (Kunar et al., [2017]; Pleskac et al., [2019]; Zeigenfuse et al., [2014]). The earliest and most recent items along the temporal dimension are retrieved more easily than mid-sequence items in free recall, forwards and backwards serial recall, and recognition tasks (Capitani et al., [1992]; Li & Lewandowsky, [1995]; Murdock, [1962]; Neath, [1993]; Wright et al., [1985]). The most intense moments of affective experiences are strong predictors of how people subsequently evaluate them (Fredrickson & Kahneman, [1993]; Kahneman et al., [1993]; Redelmeier & Kahneman, [1996]). People also tend to select options that indicate an excessive influence of the best and worst outcomes (Ludvig et al., [2018]; Ludvig et al., [2014]) and describe them as occurring with a higher frequency than they were experienced (Lichtenstein et al., [1978]; Madan et al., [2014]).

These phenomena have been observed in distributions ranging from basic perceptual features to higher-level semantic attributes (M. R. Kelley et al., [2015]; Neath et al., [2006]). It has been observed in numerous perceptual modalities, including sight (Tsetsos et al., [2012]), sound (G. D. Brown et al., [2002]; Schäfer et al., [2014]), temperature (Kahneman

et al., [1993], pressure (Ariely, [1998]), and smell (Scheibehenne & Coppin, [2020]). Their influence is largely invariant to the scale of the distribution and has been observed in short-term and long-term memory (Ludvig et al., [2014]; Nairne & Dutta, [1992]; Neath & Brown, [2006]). Finally, these phenomena have been observed across cultures (Wagner, [1975]), at multiple stages of development (Capitani et al., [1992]; Koppenol-Gonzalez et al., [2014]), and even amongst non-human animals (Kesner & Novak, [1982]; Sands & Wright, [1980]; Wright et al., [1985]).

How can we explain this seeming pervasiveness of extreme outcomes in cognition? Extreme outcomes are present within theories of perception, attention, affect, memory, and decision-making, but does this merely reflect a superficial resemblance or might it point towards a deeper regularity? Across these cognitive domains, a bewildering assortment of distinct and frequently incompatible explanations have been offered regarding extreme outcomes. For example, their influence has been attributed to environmental characteristics (Lichtenstein et al., [1978]), outcome salience (Madan et al., [2014]), task relevance (Vanunu et al., [2020]), distinctiveness (Murdock, [1960]; Neath et al., [2006]), edge-based encoding (Berliner et al., [1977]), bias-variance optimisation (Lieder et al., [2018]), selective attention (Luce et al., [1976]), personal meaning (Fredrickson, [2000]), goal-directed attention (Tsetsos et al., [2012]), and implicit reference points (Holyoak, [1978]).

Most of these theories were developed to explain specific phenomena within a single cognitive domain, and because of this, theories across domains are commonly perceived as separate or complementary rather than competing. Nonetheless, several of these theories can account for observations associated with multiple phenomena. This raises the possibility of an overarching explanation for the influence of extreme outcomes that parsimoniously encompasses some of the more constrained explanations. In the first section of this chapter, we will evaluate a particularly expansive explanation for the role of extreme outcomes that “provides a unifying mechanistic and teleological explanation for a wide range of seemingly disparate cognitive biases” (Lieder et al., [2018] p. 2).

The rational component of this explanation employs a variance-reduction algorithm

from computer science that improves the efficiency of estimation but that is biased towards extreme outcomes. Although this bias contributes to the overall error in the estimation, Lieder and colleagues argue that their algorithm is rational because it enables a greater reduction in the error caused by variance. They use this rational explanation to narrow down the set of plausible mechanisms to the one that approximates the optimal algorithm and that only uses information that would be available to the organism. This mechanism purports to capture the influence of extreme outcomes in numerous experiments across multiple cognitive domains.

Despite these promising credentials, the second section will examine attributes of the model that are troublesome for both its rational foundation and its underlying mechanism. Finally, in the third section, we will develop an alternative rational explanation based on the assertion that prioritising extreme outcomes increases 1) the probability of making the correct response and 2) the utility of the selected option. We will demonstrate this using a formal analysis of binary choices across numerous domains and then examine attributes that might influence the applicability of these idealised choices. This will allow us to evaluate the degree to which this rational model offers a viable alternative explanation that constrains the set of plausible mechanical explanations.

2.1 Utility-weighted sampling

2.1.1 Rational explanation

To understand the rational foundation of the utility-weighted sampling model proposed by Lieder et al. (2018), we must begin by reflecting on the way we learn from experience. The explicit outcomes and probabilities that appear in many theories of decision-making are seldom encountered in our daily lives. Instead, options are typically evaluated by aggregating experienced outcomes and this may be accomplished in numerous different ways. One of the most common explanations for this capacity is that experienced items are randomly sampled from memory (Denison et al., 2013; Fiedler, 2000; Hau et al., 2010;

Shadlen & Shohamy, [2016]; Stewart et al., [2006]). This method allows parameters such as the mean of the distribution to be estimated with arbitrarily high precision and does not rely on simplifying assumptions that are often necessary to obtain analytical solutions (Diaconis & Efron, [1983]; Sanborn et al., [2010]). The downside of the sampling method is that repeatedly retrieving items from memory is computationally expensive, which must be balanced against the performance obtained using larger samples (Bogacz et al., [2006]; Kareev, [1995]; Plonsky et al., [2015]; Vul et al., [2014]).

These computational costs increase the importance of sampling in the most efficient manner and algorithms developed in computer science might offer plausible ways in which this was addressed by natural selection. We have grown so accustomed to having vast amounts of computation at our fingertips that it is hard to imagine digital computers encountering similar resource limitations to human brains. Nonetheless, the foundational Monte Carlo sampling method was developed in the 1940s when the most advanced computers were slow and expensive punched-card machines (Ulam, [1976]). As a consequence, there is a long history of attempts to improve the efficiency of sampling algorithms. These include rejection sampling (von Neumann, [1951]), the Metropolis-Hastings algorithm (Hastings, [1970]; Metropolis et al., [1953]), Gibbs sampling (S. Geman & Geman, [1987]), Hamiltonian Monte Carlo (Duane et al., [1987]), hit-and-run sampling (Smith, [1984]), data augmentation (Tanner & Wong, [1987]), slice sampling (Neal, [2003]), and this is just the tip of the iceberg. A wide range of other variance reduction strategies have been developed using controlled, adaptive, conditional, stratified, multilevel, and sequential algorithms (Botev & Ridder, [2017]; Fishman, [2013]; Kroese et al., [2013]; Luengo et al., [2020]; Rubinstein & Kroese, [2016]).

Importance sampling is one of the most frequently used variance-reduction techniques. It deviates from the standard Monte Carlo method by prioritising outcomes based on their influence on the estimation error—their importance (Kroese et al., [2013]; Robert & Casella, [2004]; Tokdar & Kass, [2010]). This process involves two broad conceptual stages: 1) outcomes are sampled using a biased *importance-based* probability distribution and then 2) these samples are aggregated by weighting them according to the difference

between the importance distribution and the unbiased (representative) probability distribution. Importance sampling specifies a method for sampling and aggregation but allows for numerous possible importance distributions. Choosing the right distribution can result in estimation with zero variance whereas choosing the wrong one can result in estimation with infinite variance (Owen & Zhou, 2000). Needless to say, the former scenario is preferable to the latter but there is one considerable catch. The minimum-variance importance distribution samples outcomes based on their distance from the expected value of the option being evaluated but this expected value is the very parameter being estimated (Lieder et al., 2018).

In other words, it only makes sense to use importance sampling when there is uncertainty regarding the minimum-variance importance distribution. Nonetheless, even though this distribution is unavailable in practice, it might be possible to increase sampling efficiency by selecting one that is sufficiently similar to the minimum-variance distribution (Kroese et al., 2013; Rubino & Tuffin, 2009). This raises the inevitable question of how we determine sufficient similarity and Lieder et al. (2018) provide us with a concrete recommendation: we can replace the expected value in the minimum-variance distribution with the average outcome of previous decisions made in a similar context. Their utility-weighted sampling model—and thus their explanation for the influence of extreme outcomes—is simply an implementation of importance sampling based on this assumption. Concretely, the first stage consists of sampling s outcomes based on the distribution $\tilde{q}(o)$, which is defined as

$$\tilde{q}(o) \propto p(o) \cdot |\Delta u(o) - \bar{\Delta u}| \quad (2.1)$$

where the probability of sampling each outcome using the standard unbiased sampling method $p(o)$ is weighted by its distance from the average experienced outcome $|\Delta u(o) - \bar{\Delta u}|$.

The second stage in their implementation of importance sampling generates an estimate $\Delta \hat{U}_{\tilde{q},s}^{IS}$ of the expected value by weighting the samples so that outcomes that were

over-represented in the sampling stage are equivalently under-weighted in the aggregation stage. Specifically,

$$\Delta \hat{U}_{\hat{q},s}^{IS} = \frac{1}{\sum_{j=1}^s 1/|\Delta u(o_j) - \bar{\Delta}u|} \cdot \sum_{j=1}^s \frac{\Delta u(o_j)}{|\Delta u(o_j) - \bar{\Delta}u|} \quad (2.2)$$

Whilst the derivation of this equation is unnecessary for our purposes (see Lieder et al., 2018), the crucial attribute is that the utility of each sampled outcome $\Delta u(o_j)$ is divided by its distance from the average $|\Delta u(o_j) - \bar{\Delta}u|$. In the previous stage, the importance distribution was derived by multiplying the outcome probabilities with these same distances. Therefore, by performing the opposite operation in the aggregation stage, the effect of biased sampling is corrected in the subsequent stage. The effectiveness of this process, however, depends on the sample size parameter s such that estimates based on small samples remain biased towards extreme outcomes.

We now have the required information to grasp the utility-weighted sampling explanation for the influence of extreme outcomes. Lieder et al. (2018) employ an importance distribution that is biased towards extreme outcomes as a proxy for the minimum-variance distribution. With a large enough sample, this bias would disappear in the aggregation stage but estimates produced by the rational mechanism might remain biased due to the costs associated with sampling (Bogacz et al., 2006; Kareev, 1995; Plonsky et al., 2015; Vul et al., 2014). In other words, the disproportionate influence of extreme outcomes might reflect an optimal balance between estimation error and computational costs. The bias is only tolerated because the importance sampling algorithm permits an even greater reduction in the variance and this reduces the overall estimation error for a given number of samples.

Since we lack the computational omnipotence of Homo economicus or Laplace's demon, one of the fundamental challenges that we encounter is minimising estimation error using our limited cognitive resources. This gives the utility-weighted sampling explanation broad applicability across domains. It might explain why extreme outcomes are

over-represented in memory and influential in choice (Ludvig et al., 2014; Madan et al., 2014). It might explain why these phenomena are observed for items that provide minimal information about capacity thresholds (Fredrickson, 2000). It might even employ attention as a mechanism for increasing the probability that important outcomes are sampled from memory, and therefore, explain why extreme outcomes capture more attention (Kunar et al., 2017).

2.1.2 Mechanical explanation

Importance sampling forms the rational foundation of utility-weighted sampling but acquiring a rational explanation is not the only purpose of the resource-rational analysis framework. Lieder et al. (2018) used this analysis to suggest a mechanism that might give rise to the influence of extreme outcomes. As we discussed in the *Introduction* section of this thesis, resource-rational analysis progresses under the assumption that evolution, learning, and cognitive development shaped mechanisms that are approximately rational, and therefore, tentatively accepts the rational algorithm as the actual mechanism. This simplifies the transition from rational to mechanical explanation: Lieder and colleagues use importance sampling as a model of the mechanism underlying our capacity to evaluate options based on experience.

Although we have focused entirely on the estimation mechanism derived from importance sampling, Lieder et al. (2018) integrate this component within a broader mechanism that also includes a second component that encodes the utility associated with experienced outcomes. The mechanism that implements this encoding can be decomposed into a normalisation process that ensures that the range is always the same regardless of the scale of the distribution and a stochastic component that reflects uncertainty regarding the utility of the normalised outcomes. Formally, the utility $u(o)$ associated with an outcome is defined as

$$u(o) = \frac{o}{o_c^{max} - o_c^{min}} + \epsilon \quad (2.3)$$

where $o_c^{max} - o_c^{min}$ is the range of the experienced outcomes and ϵ represents normally distributed encoding noise that has a standard deviation estimated as a free parameter.

The normalisation process encodes outcomes *relative* to the local context of experienced outcomes rather than encoding their *absolute* utility. This relative encoding scheme is based on psychophysical and neural evidence (Carandini & Heeger, 2012; Louie et al., 2013; Rangel & Clithero, 2012) and explains the approximate scale invariance observed in extreme outcome phenomena (Ludvig et al., 2014; Neath & Brown, 2006). Consistent with the resource-rational analysis framework, this empirical evidence is supplemented with a rational explanation. Namely, assuming that the precision of encoded information is constrained by the finite bandwidth of neural firing rates, efficient use of these computational resources will maximise the representational bandwidth employed in each comparison (Heeger, 1992; Summerfield & Tsetsos, 2015; Tsetsos et al., 2016). Based on a further assumption that outcomes encountered in a given context are usually correlated, normalising outcomes relative to the local context will increase the sensitivity of the mechanism to differences between outcomes.

Therefore, the full utility-weighted sampling model comprises an encoding mechanism derived from efficient normalisation and an estimation mechanism derived from optimal importance sampling. Lieder et al. (2018) examined the empirical component of this mechanical explanation using a diverse range of experimental evidence. Their model is able to capture experience-based choices and memory recall (Ludvig et al., 2014; Madan et al., 2014), the overestimation of extreme events, such as murder or dying in an accident (Lichtenstein et al., 1978), numerous description-based choice phenomena including the pattern of risk preferences described in prospect theory (Allais, 1953; Lichtenstein & Slovic, 1971; Tversky & Kahneman, 1992), and the temporal dynamics of participants' choices in the Technion choice prediction tournament (Erev et al., 2010). Although the resource-rational framework specifies a recursive process and Lieder et al. (2018) emphasise several deviations from the predictions of their model, the empirical evidence appears broadly consistent with their conclusion that “the neural mechanisms of decision-making share some of the abstract properties of utility-weighted sampling” (p. 24).

2.2 Reconsidering utility-weighted sampling

Our primary objective in the previous section was to offer an overview of utility-weighted sampling that was faithful to the original derivation by Lieder et al. (2018) while emphasising its potential as a unifying explanation. We have described its rational foundation in optimal importance sampling and explained why prioritising extreme outcomes might reflect a rational trade-off between bias and variance. We discussed evidence in favour of its mechanical explanation, which encompasses encoding and estimation. In short, if utility-weighted sampling lives up to its promise, it would explain both *why* and *how* extreme outcomes are influential throughout cognition, leaving little room or requirement for any other explanation.

As a consequence, this section is necessarily adversarial and we will demonstrate three attributes that severely limit its applicability. Our goal is not to provide a refutation of the utility-weighted sampling model, but instead, given its persuasiveness, to explain why it will be necessary to develop an alternative rational explanation in the subsequent section. To prevent this critique from drifting into abstraction, we will anchor it on the concrete example of the first experiment conducted by Ludvig et al. (2014), which has served as a paradigm for subsequent experiments (e.g., Ludvig et al., 2014; Madan et al., 2014, 2017).¹ Once we have described our alternative explanation—in the final section—we will attempt to provide a more even-handed appraisal that draws upon utility-weighted sampling and numerous other rational explanations.

In the experiment conducted by Ludvig et al. (2014), participants made numerous choices between pairs of options that were depicted as coloured doors. In contrast with decision tasks in which outcomes and probabilities are explicitly described, they were required to learn about options by selecting them and observing the outcomes. There

¹Assuming that outcomes are normalised during encoding so that range of utilities is identical across contexts, Experiment 1, 3 (extreme options), 4G, and 4L conducted by Ludvig et al. (2014) as well as both experiments conducted by Madan et al. (2014) are structurally equivalent. In other words, they consist of options with identical normalised utilities.

were four options: two *safe options* that always resulted in a single outcome and two *risky options* that resulted in a better or worse outcome with equal probability. Specifically, there was a safe option (the red door in Figure 2.1) that caused participants to *gain* 20 points and a risky option (the yellow door) that caused them to *gain* either 0 or 40 points. These outcomes were mirrored across the zero-point for the other options: there was a safe option (the green door) that caused participants to *lose* 20 points and a risky option (the blue door) that caused them to *lose* either 0 or 40 points.

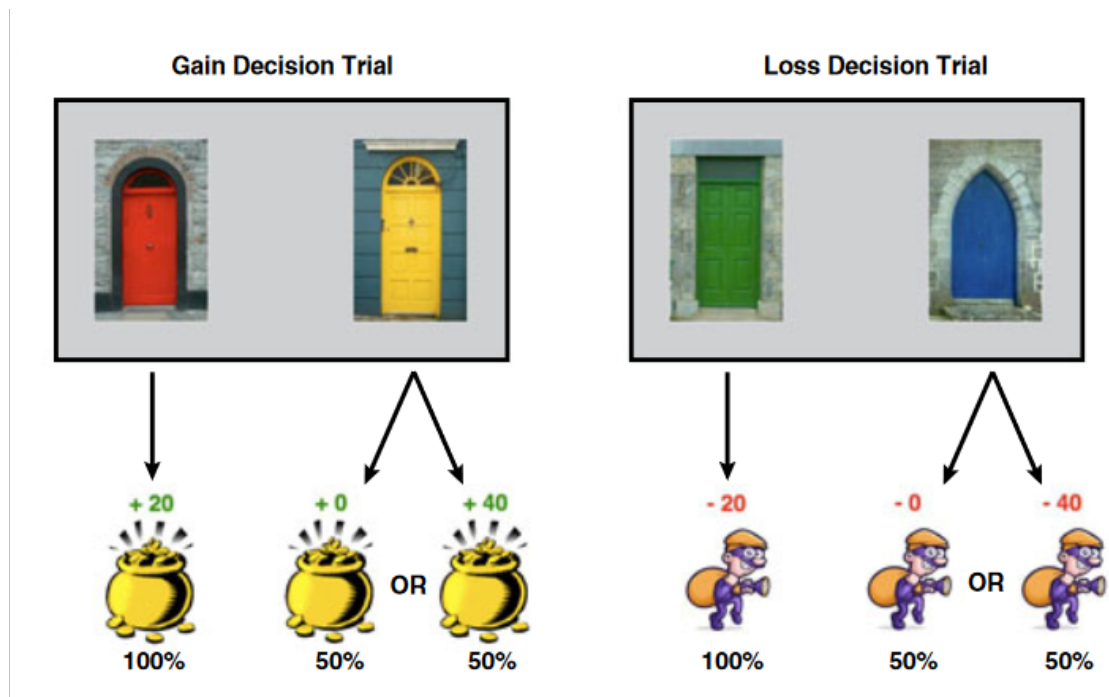


Figure 2.1: **Schematic depiction of the experiment conducted by Ludvig et al. (2014).** Participants repeatedly made choices between pairs of options represented using coloured doors. They received feedback about the number of points resulting from their choice.

Let us take a moment to consider the choices that participants might make when presented with these options. When they encounter a choice between an option that allows them to gain points and one that causes them to lose points, we would expect them to prefer the positive option—gaining points is generally preferable to losing them. What about when they encounter a choice between a safe and risky option when both are gains or both are losses? These pairs of options have the same expected value so it might

be reasonable to assume that participants would be relatively indifferent between them. Then again we might invoke the reflection effect described in prospect theory and suggest that participants would be risk-averse with potential gains and risk-seeking with potential losses (Kahneman & Tversky, 1979). Participants' choices were not consistent with either of these patterns. Instead, they were risk-seeking for the options that allowed them to gain points and risk-averse for the options that caused them to lose points.

Ludvig et al. (2014) explained this curious pattern by pointing out that the positive risky option was associated with the best outcome experienced in the experiment whereas the negative risky option was associated with the worst outcome. Assuming that extreme outcomes are disproportionately influential in memory, the best outcome would make the positive risky option seem better than the positive safe option and the worst outcome would make the negative risky option seem worse than the negative safe option. They examined numerous alternative explanations and demonstrated that the pattern of choices does not depend on whether there is a shared zero outcome (Experiment 2), the absolute magnitude of outcomes (Experiment 3), or whether there are only positive or negative outcomes (Experiment 4). Instead, participants' choices appear to be influenced by the extremity of outcomes relative to the experienced distribution.

Lieder et al. (2018) demonstrated that utility-weighted sampling can capture the qualitative patterns observed in the four experiments conducted by Ludvig et al. (2014) using a single set of parameters. They also demonstrated that its predictions were broadly consistent with participants' memory for experienced outcomes (Madan et al., 2014). These experiments constitute a large proportion of the evidence for the utility-weighted sampling mechanism and their design was highly similar to the first experiment that was described above. This will facilitate the transition from the following attributes in the context of this specific experiment to an evaluation of the model as a whole.

2.2.1 Attribute 1: Counterintuitive predictions

Given that our interest in utility-weighted sampling is as an explanation for the influence of extreme outcomes, it makes sense to begin our evaluation by identifying the component of the model that is responsible for this capacity. Although other components impact the magnitude of their influence, we can narrow down the source of this bias to the presence of $|\Delta u(o) - \bar{\Delta u}|$ in the importance distribution. This term corresponds to how extreme an outcome is within a given context and has a major impact on the probability of sampling each outcome. To state this more concretely, in the case where outcomes are experienced with equal frequency, the probability of sampling an outcome is simply how extreme it is relative to the average experienced outcome.

How might this component of the importance distribution influence predictions regarding the experiment by Ludvig et al. (2014)? Recall that the experienced outcomes were symmetrical around zero and that the risky options were both associated with a non-extreme outcome worth 0 points. If each outcome were experienced with the same frequency, the symmetry across positive and negative outcomes would ensure that the average outcome would be exactly 0 points. There would be no difference between the non-extreme outcome and the average, and therefore, the probability of retrieving the non-extreme outcome from memory would be exactly zero. The estimate for the risky options would merely be the value of the extreme outcome—either -40 or 40 points. Therefore, even though the corresponding safe option had an equivalent expected value, this model predicts that participants would *always* select the positive risky option and *always* avoid the negative risky option.

The above scenario involved options that have the same expected value and preferring one option over the other was relatively inconsequential. When this is not the case, utility-weighted sampling can lead to absurd predictions. Imagine a second scenario in which someone experienced both positive and negative outcomes so that the average is very close to \$0. In this context, they encounter a choice between two options that they have encountered many times before. One of the outcomes always gave them \$99 and the

other option gave them \$1 half the time and \$100 the other half of the time. This decision should be rather straightforward because the expected value of the first option is almost twice as large as the second option.

Nonetheless, although the probability of sampling the non-extreme outcome is now greater than zero, the \$100 outcome would be more likely to be sampled than the \$1 outcome. To make this statement more concrete we will briefly omit the encoding noise from the model and additionally assume that the risky outcomes were experienced the same number of times. In this simplified scenario, the \$100 outcome is 100 times more likely to be retrieved from memory on each sample than the \$1 outcome. Therefore, if the sample size parameter were equal to two samples, which was the best fitting parameter estimate for the choices in the experiment conducted by Ludvig et al. (2014), roughly 98% of estimates would conclude that the worse option is worth exactly \$100. If we increase this to four samples, which was the estimate for the memory responses in the experiment conducted by Madan et al. (2014), roughly 96% of estimates would reach this same erroneous conclusion and select the worse option.²

Without running this experiment, we suggest that few people would make this choice but this alone does not entail that we should discard the model. Most scientific explanations are literally false but are sufficiently accurate within a domain that is limited by numerous explicit and implicit conditions (Cartwright, 1983; Wimsatt & Wimsatt, 2007). The standard example concerns Newton’s laws of motion, which were superseded by relativity and yet remain broadly applicable unless objects are approaching the speed of light. In the case of utility-weighted sampling, the predictions are most implausible when outcomes are close to the experienced average and the model might perform reasonably for outcomes outside this class. This approach might salvage utility-weighted sampling but necessarily diminishes its reach and leaves it vulnerable to alternative explanations that

²Our analysis in this section focuses on small values of the sample size parameter that were used by Lieder et al. (2018) and that are consistent with other empirical evidence (e.g., Vul et al., 2014). This is not essential to our conclusions. It is not until the sample size reaches 70 samples that a narrow majority of estimates for the worse option are not just the value of the extreme (\$100) risky outcome.

have a broader scope.

2.2.2 Attribute 2: Flexible parameters

When examining the previous attribute, our focus was on the role of extreme outcomes in the importance distribution of utility-weighted sampling. We isolated this element using a simplified model in which each risky outcome was experienced the same number of times and outcomes were encoded noiselessly. This allowed us to calculate the exact probability of sampling each outcome but also potentially gives rise to a compelling objection to our critique: we cannot rule out the possibility that our simplifying assumptions amount to nothing more than a straw man who lacks an encoding mechanism. In other words, the stochastic elements of the complete model might generate predictions that are far more plausible than those described above.

Given that the analytical method employed above is not practicable when examining the complete stochastic model, we will outline an alternative simulation-based approach in this section. Firstly, each estimate was based on simulated experienced outcomes from the experiment conducted by Ludvig et al. (2014) instead of assuming that each outcome was experienced with equal frequency. Secondly, instead of calculating sampling probabilities using a simplified model, we simulated the behaviour of the complete utility-weighted sampling model using numerous values of the sample size and noise parameters. We simulated a total of 1.5 million utility-weighted sampling estimates and choices for both the positive (blue: 40 or 0 points) and negative (orange: -40 or 0 points) risky options.³

These simulations are summarised visually in Figure 2.2. If we begin by focusing on the top left-hand corner of Figure 2.2a we can observe similar predictions to those described in the previous section. These estimates are based on small samples and there was minimal encoding noise so the average estimate is very close to the more extreme outcomes (+/-40 points). As the number of samples increases, however, this bias gradually decreases so that

³Further details of these simulations are presented in Appendix A. The full simulation code is available at <https://github.com/joelholwerda>

the average estimate can be arbitrarily close to the expected value of the option (± 20 points). Likewise, as the standard deviation parameter of the encoding noise increases from 0.05 to 0.2 across the three vertical panels, the average estimate for a given sample size steadily moves closer to the expected value.

A similar pattern can be observed in the simulated choices displayed in Figure 2.2b. When both the sample size and standard deviation parameters are small, the positive risky option is almost always selected and the negative risky option is almost always rejected. When either of these parameters increases, preferences gradually become weaker and eventually approach indifference. This attribute appears to salvage the model from the most counterintuitive predictions described in the previous section. Moreover, at least one combination of these parameters should be able to generate behaviour that is compatible with the observations of Ludvig et al. (2014).

As such, instead of a single set of predictions, we can identify numerous patterns that we might attribute to utility-weighted sampling: the predictions can be highly counterintuitive when the parameters are small, relatively unbiased when the parameters are large, and conform to participants' behaviour somewhere in between. Which of these should we use to assess the empirical support for the mechanism? The answer to this question remains ambiguous because these parameters are entirely unconstrained by the rational explanations that underlie the importance sampling and efficient coding components of the model—they are independent empirical assumptions that require their own rational or mechanical justifications.

This attribute creates a problem for utility-weighted sampling. Insofar as the parameters are unconstrained—even though the model predicts preferences towards the positive risky option and against the negative risky option—it cannot determine the strength of this bias. This is problematic because it is not sufficient to demonstrate that the model can fit the data. As emphasised by Roberts and Pashler (2000), “Without knowing how much a theory constrains possible outcomes, you cannot know how impressed to be when observation and theory are consistent” (p. 359). In other words, the flexibility of utility-

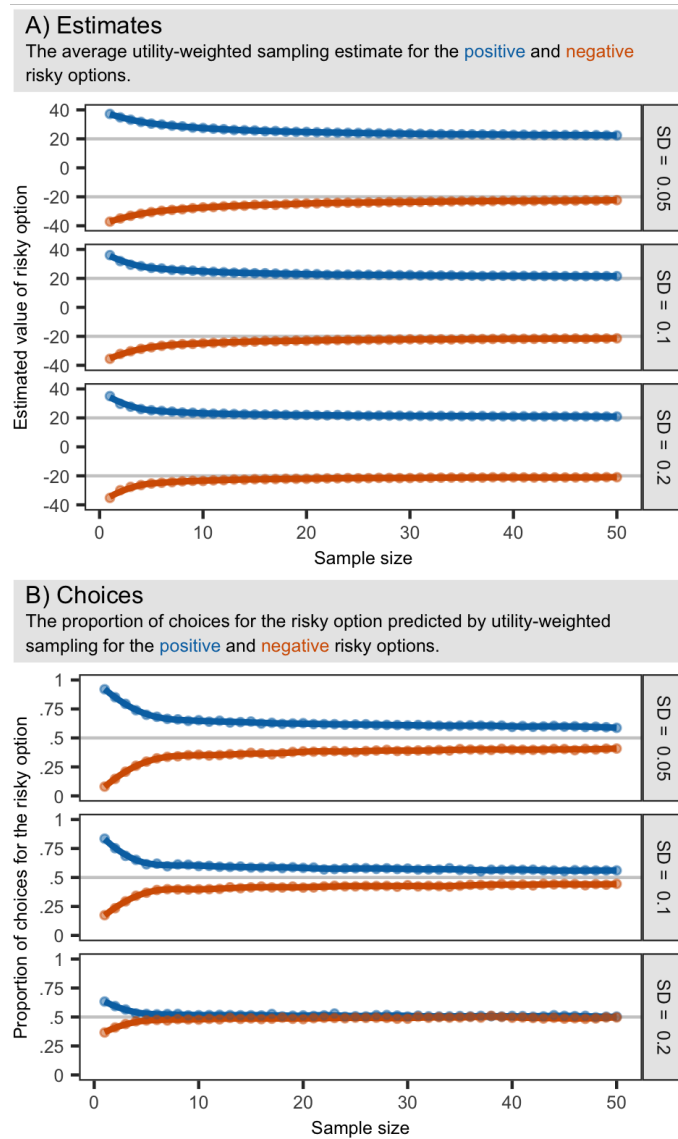


Figure 2.2: **Estimates and choices** for the utility-weighted sampling simulations of the first experiment by Ludvig et al. (2014). The positive risky option (blue) resulted in 0 or 40 points and the negative risky option resulted in -40 or 0 points. Each dot represents the mean of 10000 simulations.

weighted sampling means that we should not be surprised that the model captures the data and its ability to do so provides little empirical support (Mayo, [2018]; Roberts & Pashler, [2000]).

Having said that, if we were to end our critique here, we would rightly face the charge of attacking a straw man a second time. In addition to capturing the preferences observed in the first experiment by Ludvig et al. ([2014]), utility-weighted sampling also captures numerous qualitative patterns in the subsequent experiments. Firstly, the model correctly predicted that doubling the magnitude of each outcome would not change the strength of the bias. Secondly, it predicted that the bias observed in the first experiment would diminish when these outcomes were presented alongside others that were even more extreme. Thirdly, it predicted that a similar bias would be observed using only gains and only losses when risky options were associated with outcomes located near the lower or upper extremes of the distribution.⁴

Notably, these predictions are independent of the flexible parameters discussed above. They each provide support for the utility-weighted sampling mechanism but this support comes with an important caveat: they *do not* provide evidence for its importance sampling explanation for the influence of extreme outcomes. In the same way that these predictions are independent of the parameters, they are also largely unaffected by the characteristics of the sampling mechanism. Instead, they arise due to the efficient

⁴Lieder et al. ([2018]) also present evidence regarding decisions from description and the temporal dynamics of risk preferences. We have omitted this evidence because the models used to explain these phenomena differ markedly from the version described in this chapter. The model used for decisions from description employs an importance distribution that samples *pairs of outcomes* based on their differential utility rather than *individual outcomes* based on their extremity within the context. Our goal in this chapter is to explain the influence of extreme outcomes rather than providing a comprehensive refutation of utility-weighted sampling, and therefore, this version of the model is beyond the scope of our discussion. The temporal model uses the importance distribution described in this chapter and even relies on the same qualitative predictions for empirical support. The difference in this case is that the model uses five additional parameters to correctly predict that the bias towards extreme outcomes increases over time. Recalling the famous saying by von Neumann, we could use these parameters to fit an elephant, make him wiggle his trunk, and we would still have two parameters in reserve (Dyson, [2004]).

coding mechanism and this provides further evidence that some form of normalisation is carried out. This does not, however, provide a rational or mechanical explanation for extreme-outcome phenomena. Conversely, efficient coding is even cited as an explanation for situations where extreme outcomes are neglected (Payzan-LeNestour & Woodford, 2020; Summerfield & Tsetsos, 2015).

Where does this leave us? In the first attribute, we observed that the importance sampling mechanism makes counterintuitive predictions when examined in isolation. In the second, we recognised that the flexibility of the complete model undermines the evidence for this sampling mechanism. Therefore, although the amalgamation of these attributes does not necessarily falsify utility-weighted sampling we also have little reason to prefer it over the numerous alternative mechanical explanations (e.g., Madan et al., 2014; Neath et al., 2006). The flip side of this conclusion is that—especially in decisions from experience—there is little reason to prefer *any* explanation over the others. We will attempt to rectify this empirical issue in the subsequent chapters of this thesis, but for the time being, it suffices that the mechanism underlying the influence of extreme outcomes is still a matter of debate.

2.2.3 Attribute 3: Limited domain of rationality

The rational component of utility-weighted sampling is based on the assertion that over-representing extreme outcomes reduces the estimation variance and consequently leads to more efficient and accurate decisions. Lieder et al. (2018) established evidence for this claim using a mathematical proof that the estimation variance is minimised when outcomes are prioritised based on their distance from the expected value of the option. But despite the persuasiveness of these assertions, there are two potential issues that deserve further scrutiny: 1) minimising the variance component of the error neglects the contribution of the bias component and 2) the minimum-variance importance distribution identified in their mathematical analysis is not the same as the one used in utility-weighted sampling.

Whereas the first issue could be easily remedied by examining the two components

that comprise the overall estimation error, the resolution of the second issue is considerably less straightforward. It is an unavoidable consequence that arises because the minimum-variance distribution is derived from the expected value of the option being estimated—information that is unavailable whenever there is any reason to estimate this parameter. Lieder et al. (2018) were forced to propose an alternative and asserted that “the average utility of the outcomes of previous decisions made in a similar context could be used as a proxy for the expected utility gain” (p. 4). This is a plausible alternative but their mathematical analysis—despite its rigour—does not guarantee that this distribution is similar enough to improve the accuracy of the estimates.

Therefore, we compared the performance of the utility-weighted sampling estimates for the experiment by Ludvig et al. (2014) that were examined in the previous section (dark blue) with two alternative sampling algorithms, which are displayed visually in Figure 2.3. The first alternative is the standard Monte Carlo method (green). This algorithm is unbiased and serves as a baseline to evaluate the impact of introducing a bias towards extreme outcomes in utility-weighted sampling. The second is a version of utility-weighted sampling that uses the average sample in the aggregation phase rather than the bias correction algorithm (light blue). This simplified model allows us to evaluate whether correcting for the biased sampling algorithm improves the estimation accuracy, which will become important when we revisit utility-weighted sampling in the discussion section.

Looking first at the bias component displayed in Figure 2.3a, the utility-weighted sampling algorithm consistently produces estimates that are more biased than the standard Monte Carlo algorithm for every value of the sample size and noise parameters. This should not be a surprise. Utility-weighted sampling aims to reduce the variance by introducing a bias towards important outcomes. When the estimate is based on a single sample, utility-weighted sampling is equivalent to the version of the model that omits the bias correction algorithm. This equivalence occurs because reweighting a sample consisting of a single outcome is entirely ineffective. The divergent performance of these algorithms as the sample size increases, however, demonstrates that the bias steadily diminishes as the

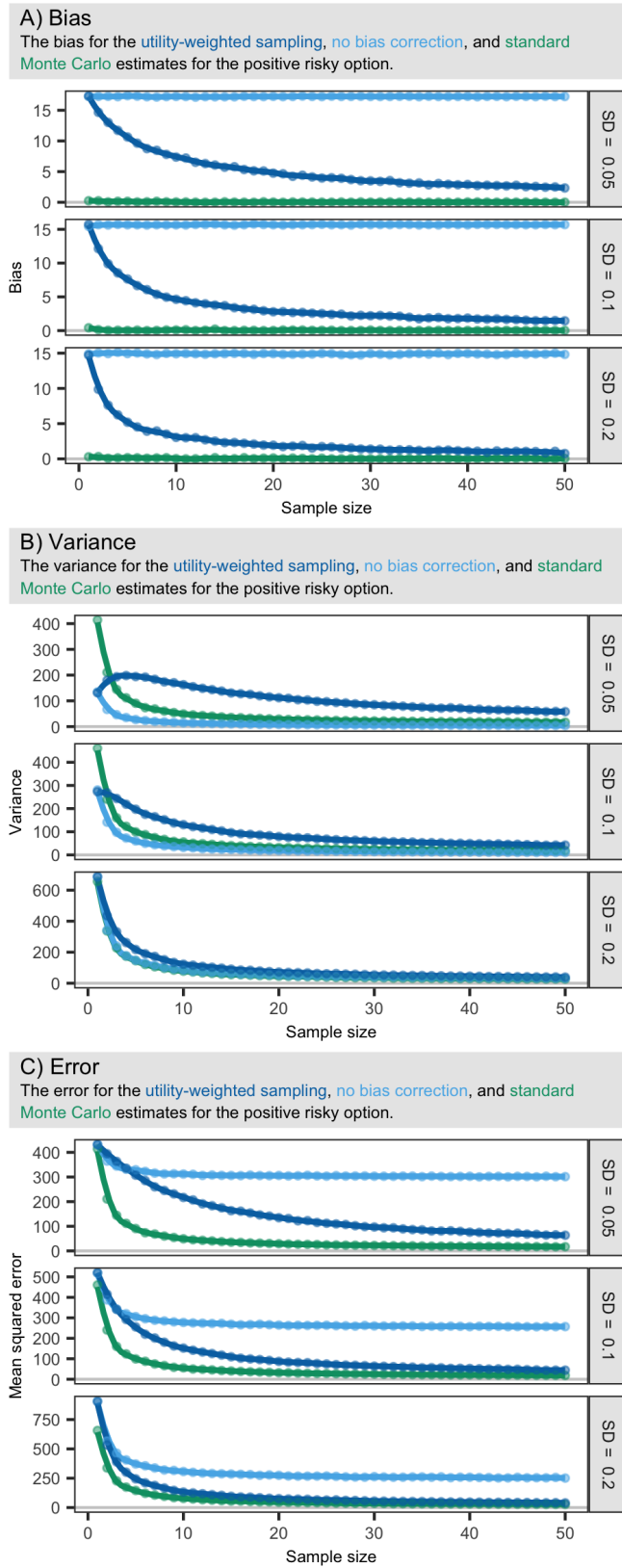


Figure 2.3: **Bias, variance, and error for the utility-weighted sampling simulations of the first experiment by Ludvig et al. (2014).** These estimates are compared with the standard Monte Carlo method and a version of utility-weighted sampling that does not re-weight outcomes to correct for biased sampling. Each dot represents the mean of 10000 simulations.

reweighting process becomes more effective.

The real question is whether utility-weighted sampling is similar enough to the minimum-variance importance distribution that the lower variance compensates for the biased sampling. Figure 2.3b indicates that this simply cannot be the case for most values of the sample size and noise parameters. Although the variance for utility-weighted sampling is lower than the standard Monte Carlo algorithm when the free parameters are small, the variance rapidly overtakes the standard algorithm as the sample size increases. Nonetheless, the best fitting parameter value for the sample size in the experiments by Ludvig et al. (2014) was only two samples and this might fall into the narrow range where utility-weighted sampling has lower variance.

For this reason, we must consider how the bias and variance are combined into the error displayed in Figure 2.3c. The interpretation of this figure is unambiguous. For every value of the sample size and noise parameters, the estimation error is higher for utility-weighted sampling than the standard Monte Carlo algorithm. Put simply, using a sampling algorithm that is biased towards extreme outcomes produces estimates that are worse than the unbiased algorithm. Furthermore, the estimation error decreases for utility-weighted sampling as the noise parameter increases because the influence of extreme outcomes—the main attribute that distinguishes the model from the standard Monte Carlo algorithm—is attenuated by the encoding noise.

It is exceedingly difficult to see how this could possibly be rational and this explains why the behaviour described in the first attribute was so counterintuitive. Nonetheless, we should resist forming strong conclusions regarding the rationality of the model based solely on these simulations. What we have demonstrated so far is that utility-weighted sampling produces inferior estimates for options that result with equal probability in either 0 or 20 points that were experienced in a context where the average was 0 points and the range was -40 to 40 points. This applies exclusively to binomially distributed options, to outcomes that are equiprobable, and arguably to the specific values used in the task—although the model is agnostic to the denomination so our analysis might also apply to

options that result in \$0 or \$20.

The forgoing simulations are analogous to noticing that penguins are short-sighted on land. Although this attribute might lead us to question whether their optical system is well-adapted, further analysis reveals that their myopia is merely a consequence of a mechanism that prioritises sharp focus underwater where they acquire their food (Neander, 1991). The seemingly irrational behaviour of utility-weighted sampling in the experimental task examined in this section might similarly be compensated by improving performance in scenarios that are broader or more consequential. This possibility is not unprecedented and numerous well-established biases have been reconceived as rational responses to environmental or cognitive constraints that differed from the assumptions of the researcher (Dawes & Mulford, 1996; Fawcett et al., 2014; Hertwig & Gigerenzer, 1999; Hertwig et al., 2005; Warren et al., 2018).

Therefore, given the narrow class of decisions that we have examined, we must resolve numerous impediments to the breadth of our critique or what appears to be irrationality might instead reflect an incomplete analysis. Firstly, even within the class of equiprobable binomially distributed outcomes, we have only examined one specific option whereas the number of possible options is infinite. This might initially appear to be insurmountable but the normalisation mechanism allows us to abstract much of this variability away from the predictions of the model. Specifically, the linear normalisation function that divides each outcome by the difference between the maximum and minimum outcomes allows us to perform three operations without altering the predictions of the model.

Firstly, we can change the value of the minimum and maximum outcomes as long as the difference between them remains the same. This allows us to change the *location* of the range so that, for example, a risky option that results in either \$1 or \$2 is equivalent when encountered in a context where the range is \$0 to \$10 and one where the range is -\$5 to \$5. Secondly, we can simultaneously multiply the maximum and minimum outcomes and divide the standard deviation of the encoding noise by the same value. This allows us to change the *scale* of the range so that the above decision would be identical in a context

where the range is \$0 to \$20 as long as the standard deviation of the noise parameter is halved.

Thirdly, we can multiply each outcome by a constant value so that the risky option that results in either \$1 or \$2 in a context where the mean is \$3 and the range is \$0 to \$10 is equivalent to an option that results in \$100 or \$200 in a context where the mean is \$300 and the range is \$0 to \$1000. In combination with the first and second operations that abstract the range of the context away—its location and scale—the third operation results in the following equivalence: increasing the scale of an option has the same impact on the utility-weighted sampling estimate as decreasing its distance from the average. For example, the distance between \$1 and the average in the scenario above is double the distance between \$2 and the average. We can modify this scenario so that the distance is four times greater by either changing the average to \$2.50 or scaling the outcomes of the option to -\$2 and \$2.

The implication of this equivalence is that varying the distance between an option and the average outcome gives rise to utility-weighted sampling estimates that are identical to those generated by rescaling the option or varying the minimum and maximum outcomes. This allowed us to evaluate the entire class of equiprobable binomial options by simulating an option that resulted in -1 or 1 and varying the distance between the expected value of this option and the average outcome. In principle, we would need to evaluate the entire range of average outcomes from negative to positive infinity but we can simplify this task by noting two characteristics: 1) This symmetrical option ensures that evaluating the positive averages will generalise to the negative averages and 2) Utility-weighted sampling is based on the ratio of distances from the average and converges towards the standard Monte Carlo algorithm when the option is sufficiently far from the average.

We have made considerable progress but our aim is to evaluate utility-weighted sampling not merely its behaviour with equiprobable binomial options. Unfortunately, this raises a considerable challenge for our analysis. We cannot abstract away the complexity in the distribution of outcomes associated with the option—at least not using any method

we could conceive. Instead, we examined the generalisability of our analysis by generating estimates using some of the most frequent classes of distributions. In addition to binomial options, we simulated normal, uniform, exponential, right-skewed binomial, and right-skewed normal distributions. It is entirely possible that this leaves out a consequential option class that compensates for the examined distributions but this is far less plausible than when our critique was based on a single experimental task.⁵

The relationship between the average outcome and the estimation error for utility-weighted sampling is displayed in Figure 2.4. We normalised each distribution so the average outcome in the context (x-axis) is identical to the expected value of each option when the average equals zero. Three broad patterns appear to hold across each of these distributions: Firstly, the performance of utility-weighted sampling is relatively similar to the standard Monte Carlo algorithm when the average outcome equals zero. Secondly, when the distance between the average outcome and the expected value of the option is sufficiently large, its performance once again approaches the standard algorithm. Thirdly, the performance of utility-weighted sampling is considerably worse than the standard algorithm when the distance from the average is somewhere between these values.

Recall that we highlighted two potential issues for utility-weighted sampling at the beginning of this section: 1) that minimising the variance neglects the contribution of the bias to the overall error and 2) that in utility-weighted sampling, the expected value used to generate the minimum-variance distribution was replaced with the average outcome in the context. When the expected value of the option is equal to the average outcome in the context—in this case, when the average equals zero—utility-weighted sampling is

⁵We simulated each of these distributions using numerous values of the sample size and noise parameters. The predictions of the model change very gradually as the sample size exceeds 20 samples and continues with the trend depicted in Figure 2.4. As we observed with the experiment by Ludvig et al. (2014), increasing the encoding noise makes utility-weighted sampling approach the behaviour of the standard Monte Carlo algorithm. Unsurprisingly, this did not lead to any differences between the sampling algorithms that were absent using minimal encoding noise. Therefore, somewhat arbitrarily, the encoding noise for the simulations depicted in Figure 2.4 have a standard deviation of 0.05 on a scale where each option has a mean of 0 and a standard deviation of 1.

CHAPTER 2. RATIONAL EXPLANATIONS

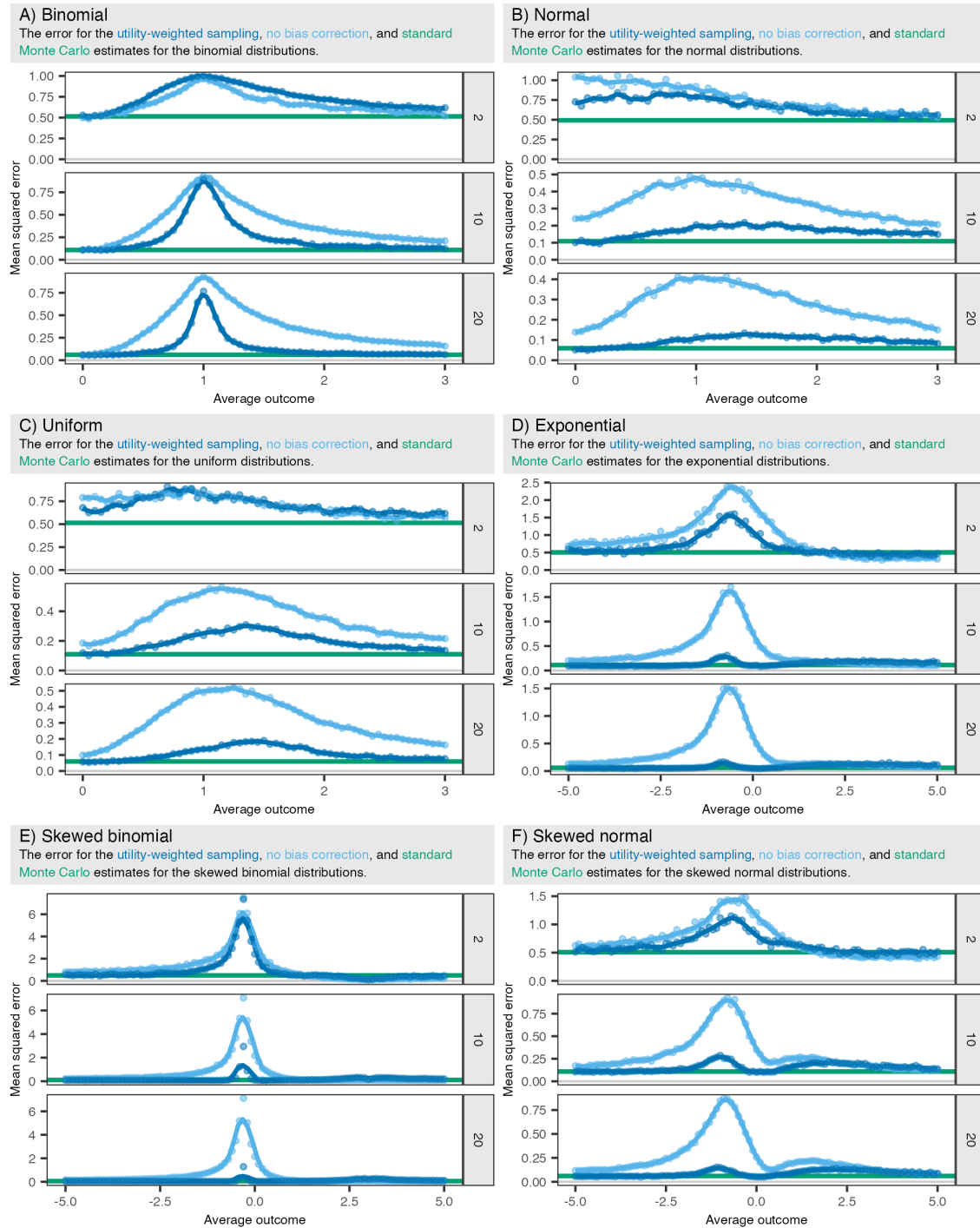


Figure 2.4: **Error for the utility-weighted sampling simulations using various distributions.** These estimates are compared with the standard Monte Carlo method and a version of utility-weighted sampling that does not re-weight outcomes to correct for biased sampling. Each dot represents the mean of 10000 simulations.

identical to the minimum-variance sampling algorithm. Therefore, insofar as the lower variance of this algorithm decreases the error more than any potential bias increases it, utility-weighted sampling should outperform the standard Monte Carlo algorithm when the average outcome equals zero.

Is this what we observed in our simulations? We mentioned above that utility-weighted sampling and the standard algorithm are similar when the average outcome is close to zero. It is not possible to discern the better algorithm only using Figure 2.4, however, because the y-axis must accommodate much larger differences in other regions. Therefore, we re-plotted them in Figure 2.5 to emphasise some of the subtler differences between the algorithms. This figure depicts a smaller region in which the expected value of the estimated option and the average outcome in the context are close together. The blue regions show where the error for utility-weighted sampling was lower than the standard algorithm and the red regions show where its error was higher.

In each panel of Figure 2.5, there is a narrow blue region around zero. For these values of the average outcome, utility-weighted sampling performs better than the standard algorithm when the sample size is large enough.⁶ Specifically, the width of this region is less than one standard deviation on the scale of each simulated distribution and the minimum sample size is roughly five samples. The magnitude of the difference in this region is many times smaller than the amount that utility-weighted sampling is worse in the adjacent regions. This creates two problems for utility-weighted sampling: 1) The sample size estimates in Lieder et al. (2018) were below the minimum sample size and 2) This algorithm will only be rational when the expected value of most options is very close to the average outcome.

This was the only region where utility-weighted sampling performed better than the standard algorithm when the distributions were symmetrical, but there were two additional regions for the skewed distributions. Utility-weighted sampling performs better than the standard algorithm when the sample size is *small* and the average is located on the *longer*

⁶The only exception was the symmetrical binomial distribution in Figure 2.5a. Utility-weighted sampling never outperformed the standard algorithm for this distribution.

CHAPTER 2. RATIONAL EXPLANATIONS

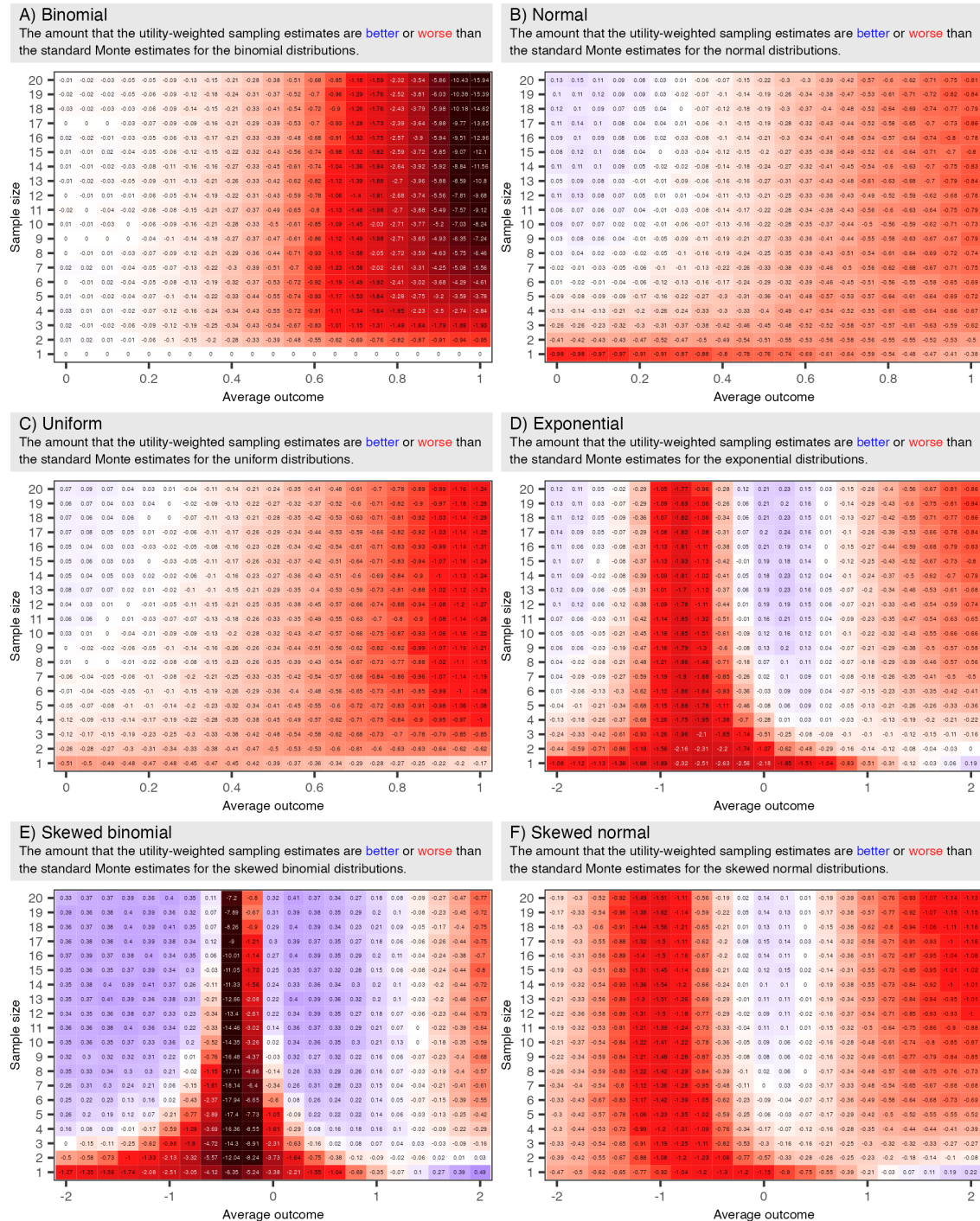


Figure 2.5: Heatmap of error for the utility-weighted sampling simulations using various distributions. These estimates are compared with the standard Monte Carlo method. Blue regions depict the parameter values where utility-weighted sampling outperformed the standard algorithm and red regions depict where its performance was worse. Each square represents the mean of 10000 simulations.

tail of the distribution. In our simulations, this occurred when the average outcome was positive. The second region is where the sample size is *large* and the average is located on the *shorter* tail. In our simulations, this occurred when the average outcome was negative. The relatively small magnitude of these differences means they are not easily identifiable in Figure 2.4 but they can be identified as blue regions in the lower-right and upper-left corners of Figure 2.5d to Figure 2.5f.

What is the cause of these patterns? Regarding the first pattern, when only a single sample is taken, the algorithm performs better when it exclusively samples outcomes that are closer to the mean of the distribution. Utility-weighted sampling increases the probability of sampling outcomes that are further from the average outcome in the context. Therefore, locating the average outcome on the longer tail of the distribution increases the probability of sampling an outcome from the shorter tail where outcomes are closer to the expected value. One problem that arises is that a context in which most options are consistent with this requirement would be necessarily bimodal because the average is a rare outcome for most options.

The second pattern is more interesting than the first pattern for three reasons: Firstly, this pattern corresponds to the main example that Lieder et al. (2018) used to motivate utility-weighted sampling. They present a hypothetical choice regarding whether someone should participate in a game of pistol roulette and argue that neglecting the less common but more lethal outcome might give rise to an unwise decision. Secondly, although importance sampling is used as a variance reduction method in multiple contexts, it is particularly useful for estimating distributions containing rare events (Luengo et al., 2020; Rubino & Tuffin, 2009). Thirdly, the previous pattern arises because the sample often contains a single outcome. In contrast, the larger sample size in the second pattern allows the reweighting function in utility-weighted sampling to correct for the biased importance distribution. In other words, this scenario is characteristic of importance sampling rather than merely biased sampling.

Despite these attributes, the second pattern faces two challenges when used as a

rational explanation for the influence of extreme outcomes: 1) Unless the average outcome in the context happens to be located on the shorter tail for most options, there will be large regions where utility-weighted sampling performs worse than the standard algorithm. In Figure 2.4, the difference between these algorithms was barely perceptible in the region where utility-weighted sampling outperforms the standard algorithm. How plausible is it that this small difference compensates for the large differences where utility-weighted sampling performed worse? 2) The second pattern only arises when the sample is much larger than the sample size parameter estimates reported by Lieder et al. (2018). Unless subsequent evidence discovers that this parameter is larger than these estimates suggest, the regions where utility-weighted sampling performs better than the standard algorithm might be limited to the region where the average outcome in the context is similar to the expected value of the option and the region outlined in the first pattern.

Therefore, we have finally reached the point where we can summarise our critique of the utility-weighted sampling model. The first attribute demonstrated that the sampling mechanism produces highly counterintuitive predictions that are unlikely to capture actual behaviour. The second attribute salvaged the model from its most implausible predictions but this was only possible because of the flexibility of the free parameters. This undermines the empirical evidence supporting the complete mechanism. Finally, the third attribute demonstrated that the domain in which utility-weighted sampling improves performance is enormously limited and this emphasises the potential benefits of developing an alternative rational explanation that has a broader scope.

2.3 An alternative rational explanation

Although our critique in the previous section was necessarily adversarial in response to the persuasiveness of utility-weighted sampling, our approach throughout the remainder of this chapter will be more integrative. It will become clear that we have not negated their explanation without remainder and the approach taken by Lieder et al. (2018) will be instructive as we attempt to develop an alternative rational explanation. To begin

with, recall that the most consequential element of utility-weighted sampling was its use of the average outcome in the context as a proxy for the expected value of the option. This substitution allowed the model to capture the influence of extreme outcomes but this capacity was purchased by sacrificing the minimum-variance importance distribution as its rational foundation.

This dilemma emphasises a question that needs to be answered by any plausible rational explanation for the role of extreme outcomes: why would it be beneficial for an estimate regarding one option to be influenced by the attributes of other options that comprise the context? Why does this make more sense than using the number of spectators at the most recent Wimbledon championship to estimate the height of the Eiffel Tower or using fluctuations in the price of crude oil to estimate the weight of the moon? The simplest explanation is that much of our world involves clustering and nested structure. As such, we can estimate someone’s salary based on their position within an organisation within a country or the cost of renting an apartment based on its location within a building within a neighbourhood within a city.

We can allocate some of the outcome variability to each level of this nested structure and it might be possible to exploit information from multiple levels to improve the accuracy of our estimates. With this in mind, we can interpret utility-weighted sampling as a single species within a genus of estimators that uses the higher-level distribution to constrain estimates of the lower-level options. There are numerous other algorithms within this genus that have demonstrated the ability to outperform those that neglect the multilevel structure of the environment (Efron & Morris, 1973; Greenland, 2000). With recent advances in methodological tools and computing power, these multilevel approaches have entered into routine statistical practice on both sides of the Frequentist-Bayesian divide (Bates et al., 2014; Carpenter et al., 2017; R. Gelman & Gallistel, 2004; Snijders & Bosker, 2011).

The most common approach to multilevel estimation imposes shrinkage or regularisation on lower-level items that reflects a compromise between evaluating each item in

isolation and aggregating them to evaluate the higher-level context (A. Gelman, 2006; Heck & Thomas, 1999). In contrast with utility-weighted sampling where the same importance distribution is used for every option, the amount of shrinkage is influenced by the uncertainty for each option relative to the context. The influence of the context exists on a continuum. When there is perfect knowledge regarding the lower-level items, the context has no influence and when there is considerable uncertainty—such as when there is missing data—the lower-level estimates are determined almost entirely based on the upper-level distribution.

Given that uncertainty mediates the relationship between multiple levels in these models, this attribute might offer a clue to the question emphasised by utility-weighted sampling. Namely, the influence of the context on individual options might make sense when there is uncertainty regarding the lower level options. This is the prevailing situation in decisions from experience. Most decisions involve numerous options, and therefore, uncertainty arises from both insufficient experience and our limited cognitive capacities. Attentional and encoding mechanisms allow us to remember some options and outcomes but this comes at the cost of remembering others.

This raises the question of how we should prioritise outcomes in memory—is it possible to do better than forgetting at random? This could be achieved by prioritising outcomes that reduce either the *probability* that forgetting will contribute to a suboptimal action or the *consequences* of those actions. We often have some knowledge whether information will be useful in the future. For example, it makes more sense to pay attention to the name of the barista where you get coffee every morning than the one who served you in a cafe you visited on an overseas holiday. We can exploit the spatio-temporal autocorrelation in these situations by prioritising outcomes that we have experienced recently or that occur frequently in our environment (e.g., J. R. Anderson, 1991; Plonsky et al., 2015).

Throughout the remainder of this chapter, we will attempt to demonstrate that it is also possible to improve our choices by prioritising memory for extreme outcomes

within the higher-level context. This influence arises as a consequence of two attributes of decision-making: 1) similarly to the multilevel estimators described above, uncertainty means that the context becomes relevant when evaluating lower-level options and 2) the comparative nature of choice means that the position of an option within the context influences the consequences of uncertainty. The conjunction of these attributes entails that extreme outcomes have a greater influence on whether our choices are aligned with our goals.

We will examine these attributes in considerable detail but in the interim, the potential rationality of the influence of extreme outcomes can be grasped intuitively in the following example: Suppose that you were presented with a choice between a pair of options, one known and the other unknown. How likely are you to make the correct choice? On the one hand, if the known option is either the best or worst you have encountered within a given context, you can be relatively confident in either selecting or rejecting it. On the other hand, the probability of making the correct choice is equivalent to a coin toss when the known option is the median because there are an equal number of experienced options that are better and worse than the known option.

In this example, the higher-level context is relevant due to uncertainty regarding the lower-level options and the extreme outcomes provide more certainty when determining the better option. In other words, extreme outcomes are more informative than intermediate outcomes and forgetting them would have greater consequences. We will formalise this intuition into a rational explanation that consists of two closely related models: one that uses an ordinal criterion to explain the influence of extreme outcomes with judgements that can be either correct or incorrect and another that uses a continuous criterion to explain their influence when aiming to maximise expected utility.

2.3.1 Ordinal criterion

We will begin our exposition by putting some extra meat onto the bones of the intuition sketched above. Imagine an idealised scenario in which each option is associated with a

CHAPTER 2. RATIONAL EXPLANATIONS

single outcome from a set of possible outcomes from 1 (the worst outcome) to 10 (the best outcome). A decision-maker enters this context without any knowledge about these options and must learn from experience by making a series of choices and observing the outcome. Some of these options will be encountered again in the future and this will allow the agent to perform better than chance but their memory is limited so they are unable to recall all of the options. Importantly, they have an ordinal criterion where their goal is to make the correct choice on as many decisions as possible rather than maximising the expected value.

Imagine that they encounter a choice between the option associated with outcome 3 and the option associated with outcome 8. If they can access a memory associated with *both* of these options, they will always choose outcome 8. They care about making the correct choice but do not care how much the outcomes differ. In this case, the outcome was CORRECT. On the other hand, if they can access a memory associated with *neither* of these options because they have not previously encountered or have forgotten them, the decision-maker cannot do better than choosing an option at random. The outcome will be CORRECT half the time and INCORRECT the other half of the time.

What about when only one option is remembered? This is less straightforward than when both options were remembered or forgotten but we can derive an optimal decision-rule by observing the consequences associated with each choice. In a situation where the known option is associated with outcome 1 (the worst outcome), selecting this option will always be INCORRECT. The decision-maker should always select the unknown option. Similarly, when the known option is 10 (the best outcome), selecting the known option will always be CORRECT. In other words, when the known option is associated with the best or worst outcome in the context, whether an option is CORRECT or INCORRECT can be known with certainty.

In contrast, uncertainty is the rule when the known option is associated with one of the intermediate outcomes. For example, selecting an option associated with outcome 5 will be CORRECT when the unknown option is associated with the four outcomes

in $\{1, 2, 3, 4\}$ and INCORRECT when the unknown option is associated with the five outcomes in $\{6, 7, 8, 9, 10\}$. If we tentatively grant the assumption that each option occurs with equal frequency, there would be a $4/9$ probability that this option is CORRECT. Likewise, selecting an option associated with outcome 6 will be CORRECT when the unknown option is associated with the five outcomes in $\{1, 2, 3, 4, 5\}$ and INCORRECT when the unknown option is associated with the four outcomes in $\{7, 8, 9, 10\}$. In this case, there is a $5/9$ probability that this option will be CORRECT.

These examples demonstrate that the probability that a known option will be CORRECT depends entirely on its rank whenever the unknown options occur with equal frequency. Although these ranks are usually expressed as numbers from 1 to n , we can normalise them using

$$rank'_{option} = \frac{rank_{option} - rank_{min}}{rank_{max} - rank_{min}} \quad (2.4)$$

where $rank_{option}$ is the unnormalised rank of the option and $rank'_{option}$ is the normalised equivalent. This ensures that the options are evenly spaced across the range where the worst option is 0 and the best option is 1 and we can interpret the normalised ranks as probabilities.

Therefore, the probability that an option will be CORRECT is always greater than 0.5 whenever its rank is above the median and we can use this to derive an optimal decision rule: the known option should be selected when it is above the median and avoided when it is below. Admittedly, this requires knowledge of the median option but this can be estimated using a relatively small sample of previous outcomes (Rider, 1960). We can also generalise this rule to situations where the assumption of equiprobable options is violated by weighting each option by its probability of occurrence.

At this point, it is worth emphasising three elements of our exposition that determine the influence of extreme outcomes: Firstly, when either both options are known or both are unknown, the probability of a choice being CORRECT is not influenced by the rank

of the options. A choice is either made with certainty or an option is chosen randomly. Secondly, when one option is known and the other is unknown, the probability that the known option will be CORRECT is determined by its normalised rank. Thirdly, the agent should select the known option whenever its value is above the value of the median option. These elements are the three premises that underlie our rational explanation for the influence of extreme outcomes using an ordinal criterion.

The journey from these premises to the conclusion involves seven steps. First, the benefit associated with remembering an option is simply $P(\text{correct}|\text{known}) - P(\text{correct}|\text{unknown})$ for the decision-maker with an ordinal criterion. This is simply the difference between the probability of selecting the correct option when only that option is known and when both options are unknown. Second, given that the median option is defined as having an equal number of options on either side, the normalised rank for the median, $\text{rank}'_{\text{median}}$, is equal to 0.5. The probability that the correct option will be chosen when both options are unknown is also 0.5, and therefore $P(\text{correct}|\text{unknown})$ is equal to $\text{rank}'_{\text{median}}$.

Third, we can use the function in Equation 2.4 to transform the rank of the known option into a normalised rank, $\text{rank}'_{\text{known}}$, which corresponds to the probability that the known option will be CORRECT. Fourth, the known option is selected whenever its rank is higher than the median, and therefore, $P(\text{correct}|\text{known})$ is equal to $\text{rank}'_{\text{known}}$ above the median. Conversely, the known option is rejected whenever its rank is below the median, and given that one option being CORRECT implies that the other is INCORRECT, $P(\text{correct}|\text{known})$ is equal to $(1 - \text{rank}'_{\text{known}})$ below the median.

Fifth, using the equivalences from Step 2 and Step 4 and some basic arithmetic, we can demonstrate that $P(\text{correct}|\text{known}) - P(\text{correct}|\text{unknown})$ is $\text{rank}'_{\text{known}} - \text{rank}'_{\text{median}}$ when $\text{rank}'_{\text{known}}$ is higher than $\text{rank}'_{\text{median}}$ and is $\text{rank}'_{\text{median}} - \text{rank}'_{\text{known}}$ when $\text{rank}'_{\text{known}}$ is lower. In both of these cases, this corresponds to the difference between the normalised ranks for the known option and the median option but the subtrahend and minuend are mirrored when the known option is above or below the median. When the known

option is the same as the median option, this difference is equal to zero. Otherwise, the larger number is always subtracted from the smaller number and $P(\text{correct}|\text{known}) - P(\text{correct}|\text{unknown})$ is always positive.

Sixth, we can use the arithmetic equivalence between $|a-b|$ and $|b-a|$ to demonstrate that $|P(\text{correct}|\text{known}) - P(\text{correct}|\text{unknown})| = |\text{rank}'_{\text{known}} - \text{rank}'_{\text{median}}|$ regardless of the option that is chosen. Because $P(\text{correct}|\text{known}) - P(\text{correct}|\text{unknown})$ is always greater or equal to zero, we can simplify this to

$$P(\text{correct}|\text{known}) - P(\text{correct}|\text{unknown}) = |\text{rank}'_{\text{known}} - \text{rank}'_{\text{median}}|. \quad (2.5)$$

The final step is to recognise that the definition of “extreme outcomes” in this rational explanation is determined by the difference in rank between the known option and the median, $|\text{rank}'_{\text{known}} - \text{rank}'_{\text{median}}|$. The probability of choosing the correct option increases monotonically with this conception of extremity. Therefore, we can conclude that, *ceteris paribus*, it is rational for a decision-maker with an ordinal criterion to prioritise extreme outcomes in memory.

2.3.2 Continuous criterion

Imagine that a second decision-maker enters into the scenario described in the previous section. Everything about the context remains identical: there are still ten possible options associated with outcomes from 1 to 10 and knowledge of these options must be acquired through experience. The second decision-maker is identical to the first one except that their goal is to maximise the outcome of their choices. In other words, they have the same continuous criterion as the expectation-maximising *Homo economicus* but they possess an extremely limited memory rather than omniscience and omnipotence.

There are numerous similarities between the first and second decision-makers. When the second decision-maker can access a memory for *both* options, they will always select

the option that the first decision-maker considered CORRECT. Similarly, when they remember *neither* of the options, they cannot do better than selecting a random option. In this scenario—where the outcomes are evenly distributed from 1 to 10—the two decision-makers also make the same choice when one option is known and the other is unknown. Nonetheless, they choose this option for a slightly different reason. Namely, given that the second decision-maker has a continuous criterion, they care about the distance between outcomes and should choose the known option whenever its outcome is above the *mean* rather than the *median*.

We can use this information to construct a set of three premises that are analogous to the ones we used for the ordinal criterion. Firstly, the decision-maker selects a random option when neither option is known and selects the better option when *both* options are known. Secondly, when one option is known and the other is unknown, the expected value of the known option is simply the outcome associated with this option. Thirdly, the agent should select the known option whenever it is better than the average option. Once again, we will use these premises as the foundation of a rational explanation for the influence of extreme outcomes—this time using a continuous criterion.

The first premise introduces an additional element of complexity. When examining the ordinal criterion, we focused exclusively on the benefit of knowing one option relative to knowing neither. This was not an erroneous neglect of the second difference between knowing one option and knowing both. Instead, given that the probability of selecting the correct option is a constant value in both cases (0.5 or 1), the latter difference is a simple transformation of the former. The cost associated with forgetting an option when both are known depends on the rank of the option that remains, such that the cost is simply $|rank'_{median} - rank'_{known}|$. This value *decreases* monotonically as the ordinal extremity of the remembered option increases, and therefore, focusing on this difference would have led to exactly the same conclusion.

This is *not* the case using a continuous criterion. Neither of these choices is independent of the (potentially) known options, and therefore, we would reach different conclusions

depending on the comparison we chose. Examining the difference between decisions that involve one known option and two known options will reveal a broader pattern that will be relevant when we generalise from binary choices to those involving multiple options. This pattern is also relevant to the ordinal criterion so we will examine the difference between zero and one known option in this section and multiple known options in the following section.

In contrast with the obstacle lingering in the first premise, the difference between the second premise used for the ordinal and continuous criteria will make our task much easier. A considerable proportion of our examination of the ordinal criterion was spent justifying the transition between rank and probability. For the continuous criterion, the outcome of the known option is directly isomorphic with the quantity that the decision-maker is aiming to maximise—at least until we discuss risky options—and a similar transformation will be unnecessary. Specifically, this will allow us to bypass the third step in the derivation for the ordinal criterion leaving a path with six rather than seven steps.

First, the benefit associated with remembering an option is $EV_{known} - EV_{unknown}$ for the decision-maker with a continuous criterion. This is the difference between the expected value of the chosen option when only that option is known and the expected value when both options are unknown. Second, when neither option is known, the decision-maker selects a random option. In contrast with the ordinal criterion where the probability of selecting the correct option was always 0.5, the outcome of the potentially—but not actually—known option influences the expected value of this random choice. The potentially known option and the other unknown option have an equal probability of being selected. On average, the outcome of the other unknown option will be equivalent to the mean, and therefore, the expected value of this choice is $(outcome_{known} + outcome_{mean})/2$, where $outcome_{known}$ is the potentially known option.

As we mentioned above, the third step in the derivation for the ordinal criterion is not necessary for the continuous criterion but we will keep the numbering consistent to facilitate comparison between the two criteria. Fourth, when only one option is known, the

known option is selected whenever its outcome is above average, and therefore, EV_{known} is equal to $outcome_{known}$ above this value. Conversely, the unknown option is selected whenever the outcome associated with the known option is below average. On average, the outcome of the unknown option is the average outcome in the context, and therefore, EV_{known} is equal to $outcome_{mean}$ when the known option is below this value.

Fifth, we can once again use the equivalences from Step 2 and Step 4 and some basic arithmetic to calculate the expected utility gain associated with remembering an option. $EV_{known} - EV_{unknown}$ is $outcome_{known} - (outcome_{known} + outcome_{mean})/2$ when $outcome_{known}$ is above average and $outcome_{mean} - (outcome_{known} + outcome_{mean})/2$ when $outcome_{known}$ is below. Similarly to the ordinal criterion, the smaller number is always subtracted from the larger number, and therefore, $EV_{known} - EV_{unknown}$ will always be positive. This occurs because the average of the known outcome and the mean is always 1) smaller than the known outcome when the known outcome is larger than the mean and 2) smaller than the mean when the mean is larger than the known outcome.

Sixth, we can perform an identical sequence of arithmetic operations to those used to rearrange the equation for the ordinal criterion based on the equivalence between $|a - b|$ and $|b - a|$. For the continuous criterion, this simplifies to

$$EV_{known} - EV_{unknown} = \frac{|outcome_{known} - outcome_{mean}|}{2} \quad (2.6)$$

Once again, the final step is to recognise that the benefit of remembering an option corresponds to our definition of “extreme outcomes” in this rational explanation but we will define these outcomes differently for the continuous criterion. The numerator of the benefit, $|outcome_{known} - outcome_{mean}|$, is simply the continuous distance between the potentially known option and the average outcome in the higher-level context. The expected value associated with remembering an option increases monotonically with this conception of extremity. Therefore, we can conclude that, *ceteris paribus*, it is rational for a decision-maker with a continuous criterion to prioritise extreme outcomes in memory.

2.4 Chapter discussion

In the previous section, our use of examples was limited to a single scenario where the outcomes were amorously described as “1 (the worst outcome)” and “10 (the best outcome)” and it would be natural to question the generalisability of our analysis. Nonetheless, if you only remembered one of the two options, you would have a better chance of answering which movie is longer if you remembered the length of *Gone with the Wind* than if you remembered *The Wizard of Oz*, which element has a higher atomic number if you know lithium than if you know barium, which of two routes is quicker if you remember the road where the speed limit is 110km/h than if you remember the one where the limit is 80km/h, which holiday is earlier in the year if you know Hanukkah than if you know Halloween, which mountain is taller if you know Everest than if you know Kilimanjaro, which monarch reigned longer if you remember King Charles III than if you know King Edward II, and which job applicant was better if you remember the one who stumbled in 30 minutes late reeking of alcohol than if you remember the one who stumbled over a couple of questions.

Given that the economics literature is inundated with situations where the decision-maker aims to maximise the expected utility of their choices, it feels superfluous to construct a similar list for the continuous criterion. Our point is that our use of idealisation contributes to the generalisability of our analysis rather than detracting from it. This analysis used the ordinal and the continuous criteria to define the benefit associated with remembering one option compared with remembering neither and supplemented this criterion with three additional premises. Insofar as one is convinced that those premises and the subsequent derivation are correct, they also have reason to believe that prioritising extreme outcomes is rational within the domain to which our analysis pertains.

2.4.1 Multiple options

The mathematical approach used in our analysis is based on the structural relationship between items rather than their specific content. Similarly to the way that gravity applies to objects regardless of whether they are planets within a galaxy or apples falling from a tree, our model applies whenever items can be sensibly ordered along a dimension and are selected according to either an ordinal or continuous criterion. Nonetheless, there are also definite limitations on the applicability of our rational explanation that deserve further emphasis. Of these, perhaps the most prominent boundary is between decisions in which there are two available options and those in which there are three or more. In our analysis, we focused exclusively on the former and we must venture beyond this province before reaching any strong conclusions regarding multiple options.

There is a sense in which the benefits associated with remembering options across this broader domain is a separate question. Regardless of our conclusion, our analysis using two options would remain unscathed and we would still have reason to believe that there are benefits to remembering extreme outcomes within this limited domain. Be that as it may, it is simultaneously true that this independence would break down when using our rational explanation to constrain potential mechanisms. It is a plausible assumption that there would be some cost associated with having one mechanism for choices with two options and another for those with three or more. We might need to take this potential expense into consideration if prioritising extreme outcomes has a detrimental effect on choices with multiple options.

Whilst a proper formal treatment of this broader class would inevitably double the length of this chapter, we will attempt to provide a sketch of the implications for multiple options. For the decision-maker with an ordinal criterion, they should always select the best known option when at least one is above the median because it is not possible to reliably perform better when selecting an unknown option. The probability that this known option is better than one particular unknown option is $\max(RANK'_{known})$. Therefore, the probability that this option is better than *all* unknown options is

$\max(RANK'_{known})^{n_{unknown}}$, where $n_{unknown}$ is the number of unknown options.

Conversely, when none of the known options is above the median, they cannot perform better than randomly selecting one of the unknown options. Two things must be true for this randomly selected option to be CORRECT: it must be the best unknown option and it must be better than $\max(RANK'_{known})$. Given that $\max(RANK'_{known})$ is the probability of that the best *known* option is CORRECT, the conjunction of these events can be written as $(1/n_{unknown})(1 - \max(RANK'_{known}))$. As we would expect, this reduces to the equations in the *ordinal criterion* section when there is one known and one unknown option.

We can roughly summarise the behaviour of these equations as follows: 1) Forgetting any option—regardless of its position in the context—increases $n_{unknown}$, which decreases the probability that the choice is CORRECT. 2) When one or more known options is above the median, the probability of selecting the correct option increases monotonically as $\max(RANK'_{known})$ approaches the upper extreme. 3) When none of the known options is above the median, the probability of selecting the correct option increases monotonically as $\max(RANK'_{known})$ approaches the lower extreme. 4) Forgetting an option only changes these probabilities when $\max(RANK'_{known})$ changes.

Similarly, for the decision-maker with a continuous criterion, they should always select the best known option when at least one is above the mean. The expected value of this choice is simply the outcome associated with this option, $\max(OUTCOME_{known})$. Conversely, when none of the known options is above the mean, they cannot perform better than randomly selecting an unknown option and the expected value of this choice is $outcome_{mean}$.

When an option is forgotten, there are two possible consequences: the best known option will be selected when $\max(OUTCOME_{known})$ remains above the mean or a randomly selected option will be chosen. As we discussed in the *continuous criterion* section, we cannot merely ignore the potentially known option in the counterfactual where it has been forgotten. Therefore, the expected value when an option has been forgotten is a

weighted average between the potentially known option and the mean. The probability of selecting this option decreases as the number of unknown options increases so that the expected value is $(outcome_{known} + outcome_{mean} \times (n_{unknown} - 1))/n_{unknown}$.

These broader equations encompass the ones we used in the *continuous criterion* section and also capture the comparison between two and one known option that we consigned to a promissory note. Once again, we can summarise the behaviour of these equations as follows: 1) When at least one known option is above the mean, the expected value of the choice increases monotonically as $max(OUTCOME_{known})$ approaches the upper extreme. 2) When none of these options is above the mean, the cost associated with forgetting an option increases monotonically as $outcome_{known}$ approaches the lower extreme. 3) The influence of each potentially known option decreases as $n_{unknown}$ increases but this also ensures that there are more forgotten options. 4) When at least one known option is above the mean, the expected value only changes when forgetting an option influences $max(OUTCOME_{known})$.

We now possess the information we need to evaluate whether our explanation generalises to multiple options. Consistent with our analysis using two options, the benefit associated with remembering an option increases as it becomes more extreme but there is a second variable that moderates these benefits. In this scenario, the decision-maker must select one option and reject the others and this entails that remembering an option only influences their choice when it modifies the best known option. When this occurs, remembering extreme outcomes has advantages similar to those we ascertained using two options but there is no guarantee that a memory lapse will change the best known option.

This attribute has two main implications for our rational explanation. Firstly, the probability that forgetting an option will change the best known option decreases as the number of known options increases. As a consequence, the advantages of extreme outcomes are strongest when there is a single known option and decreases as this number increases. Secondly, the probability that an option will change the best known option depends on their rank within the distribution. The options that are located near the higher extreme

are more likely to change the best option and this means that remembering or forgetting them is more consequential. Therefore, increasing the number of known options mitigates the influence of extreme options with one hand whilst introducing an additional reason for remembering higher extreme options with the other.

2.4.2 Risky options

When introducing our discussion of risky options, it is useful to recall that our critique of the utility-weighted sampling explanation was based on their substitution of the expected value of the option with the average outcome in the context. We emphasised how Lieder and colleagues provided a rigorous mathematical proof for the minimum-variance distribution and then merely assumed that this capacity would generalise to the average outcome. This allowed the model to explain the behaviour of participants in the experiments by Ludvig and colleagues but it undermined the rational foundation of their explanation. We are bringing this up yet again because we suspect that if someone were to compose a similar critique of our alternative explanation it would be focused on the substitution discussed in this section—the section in which we will attempt to formulate that critique ourselves.

We constructed our rational explanation based on a scenario in which each option is associated with a single outcome. When there is only one possible outcome, the expected value of each *option* is isomorphic with the *outcome* for that option and this was the basis of our conclusion regarding extreme outcomes within this domain. Although it might be enticing to perform a similar substitution when there are multiple possible outcomes—for risky options—this would be illegitimate. The outcomes of risky options are partially determined by their expected value but they also fluctuate around it. Therefore, we must also consider the variance when determining whether extreme outcomes should be prioritised for risky options.

Before we attempt to do this, we should reflect on why we even want to generalise our explanation to these options. Can we not be satisfied with having explained the influence of

CHAPTER 2. RATIONAL EXPLANATIONS

extreme outcomes within a limited domain? The first reason is that prioritising extreme *options* is entirely unable to explain the pattern of decisions observed by Ludvig and colleagues. Their experiments employed safe and risky options with the same expected value and we are unable to explain their observations without generalising from extreme options to extreme outcomes. A second reason is that the way we learn about options is usually by observing outcomes and an encoding bias towards extreme outcomes might offer a more feasible strategy for prioritising extreme options. Therefore, it appears that we have found ourselves in the same position as Lieder and colleagues where we might be required to sacrifice the rational component of our model to explain the phenomena.

At least to some extent, the utility of our explanation depends on our justification for substituting options with outcomes but fortunately this only requires us to emphasise two attributes: 1) The benefit of remembering an option increases monotonically towards the edges so that remembering an even more extreme option is always preferable. 2) Even when the variance of an option is enormous, there is always a positive correlation between the expected value and the outcomes associated with the option. The strength of this relationship depends on the variance but prioritising extreme outcomes will never reduce the probability that you will remember an extreme option. Therefore, we can conclude that remembering extreme *outcomes* always produces some proportion of the benefits associated with remembering extreme *options*.

Have we achieved our aim? What we have demonstrated is that the relationship between extreme outcomes and options increases the probability that the correct risky option will be selected and the benefits associated with doing so. In contrast with utility-weighted sampling, this relationship is favourable across its entire domain of applicability instead of sacrificing performance in some common scenarios. Admittedly, this benefit is moderated by the number of known options and the variance of each option but we have ostensibly found an explanation that obeys the rationality equivalent of the Hippocratic Oath: at the very least, it never makes things worse.

Although this is an unambiguously desirable quality, we have not yet answered

whether it is rational to prioritise extreme outcomes in memory. This question remains up for debate because there are two interpretations of our claim to do no harm and only one of them is true. What we have demonstrated is that there are benefits associated with remembering extreme outcomes but we have not shown that it is necessarily beneficial to prioritise them. It is possible that there are negative consequences that outweigh the positives, and therefore, adopting a strategy where extreme outcomes are prioritised might give rise to worse decisions.

This is a lingering concern for every rational explanation because even when the analytical component includes an ironclad logical deduction, the empirical component can reveal neglected variables. There is some evidence, however, that our explanation is vulnerable to this possibility when extended to include a bias towards extreme outcomes. This bias in memory increases the probability of remembering informative options but contributes to the mean squared error when subsequently making decisions. Therefore, the viability of this explanation for the influence of extreme outcomes depends on whether the positive consequences outweigh the negative ones.

The bias towards extreme outcomes would ideally be present in memory but completely absent in choice and this is another place where we can take inspiration from Lieder and colleagues. Their model includes a bias-correction mechanism that weights the retrieved outcomes based on their probability of being included in the sample. Comparing the light and dark blue lines in Figure [2.3](#) and [2.4](#) highlights the capacity of this component to mitigate the consequences of biased sampling. It is plausible that a similar mechanism attenuates the disadvantages associated with biases in choice sufficiently to enable the advantages associated with biases in memory.

2.4.3 Conclusion

We began this chapter with the observation that extreme outcomes are disproportionately influential across numerous cognitive domains and raised the question whether this could be attributed to an underlying reason. We examined the explanation offered by

utility-weighted sampling but ultimately concluded that its variance-reduction strategy is only rational within a limited domain. This was the motivation for developing an alternative rational explanation based on the informativeness of extreme options relative to the intermediate options. We demonstrated that prioritising extreme options and outcomes increases the probability of selecting the correct option and the expected utility gained from these choices. This rational explanation applies across a broad range of judgements and decisions and might serve to unify several phenomena that have been assigned separate explanations.

Chapter 3

Levels of measurement

In the previous chapter, we used the experiments conducted by Ludvig et al. (2014) to critique utility-weighted sampling and motivate an alternative that applies across numerous domains. In this chapter, we will narrow our focus to extreme outcomes in decisions from experience and examine several mechanical explanations for this phenomenon. Ludvig, Madan, and colleagues were responsible for the majority of the empirical work in this area and our experiments in this section were influenced by their methods (Ludvig et al., 2018; Ludvig et al., 2014; Ludvig & Spetch, 2011; Madan et al., 2014, 2017). Therefore, although we have already discussed their experiments in the previous chapters, the following paragraphs offer another brief summary.

Participants in each of these experiments made numerous choices between pairs of options. They were initially unaware of the potential outcomes and probabilities associated with each option but were able to acquire information by observing the outcomes of their choices. In contrast with most decisions from experience tasks, these participants were presented with numerous options within the same context rather than a single pair (Wulff et al., 2018). For example, one experiment included a low-value safe option that always resulted in -25 points paired with a risky option that resulted in either -45 or -5 points and a high-value safe option that always resulted in 25 points paired with a risky option

that either resulted in 5 or 45 points (Ludvig et al., 2014, Experiment 2).

Participants' choices stemming from this seemingly minor alteration were quite different from those that are observed when either of these pairs is presented in isolation (Kahneman & Tversky, 1979).¹ Specifically, they selected the risky option more often when they chose between the high-value pair than when they chose between the low-value pair. Ludvig et al. (2014) explained this pattern by emphasising that the risky option results in *both* the best outcome *and* the worst outcome when these options are presented in isolation. In contrast, when more than one pair is encountered in the same context, the low-value risky option results in the worst outcome and the high-value risky option results in the best outcome. Therefore, a preference towards options that result in the best outcome and against options that result in the worst outcome could be explained by these extreme outcomes (e.g., -45 or 45 points) exerting greater influence on choices than intermediate outcomes (e.g., -5 or 5 points).

Ludvig et al. (2014) provided three additional pieces of evidence in support of this explanation. First, that a similar preference towards the risky option of the high-value pair and against the risky option of the low-value pair is observed when all outcomes are gains *or* losses (Ludvig et al., 2014, Experiment 4). Second, that the effect is weaker or entirely absent when options are not associated with the best or worst outcomes. And third, that extreme outcomes are over-represented when participants estimate the frequency of each outcome and when they report which outcome comes to mind first (Ludvig et al., 2018; Madan et al., 2014, 2017). This suggests that the observed pattern of choices is associated

¹Prospect theory implies that people are risk-seeking for losses and risk-averse for gains that involve moderate to high probabilities. Although there is considerable support for this conjecture when tasks involve written descriptions (e.g., Abdellaoui, 2000; Baucells & Villasís, 2010; Holt & Laury, 2005; Tversky & Kahneman, 1992), there is growing evidence for risk-neutrality across domains when decisions are based on experience (Erev et al., 2010; Erev et al., 2008; Ert & Haruvy, 2017). This difference has been partly attributed to the use of small magnitude outcomes in experience-based tasks (Erev et al., 2008; Konstantinidis et al., 2017; B. J. Weber & Chapman, 2005), and therefore, we might expect to observe relative indifference when these pairs of options are presented in isolation. Nonetheless, these theories remain unable to explain the effects of context observed by Ludvig, Madan, and colleagues.

with a bias towards remembering the best and worst outcomes, which is consistent with evidence that memory biases lead to extreme forecasts when a small number of outcomes are retrieved (Fredrickson, [2000]; Morewedge et al., [2005]; D. L. Thomas & Diener, [1990]; Wilson & Gilbert, [2003]; Wirtz et al., [2003]).

Based on this evidence, Ludvig et al. ([2014]) devised an *extreme-outcome rule* that the best and worst outcomes within a given context are over-represented in memory and exert a disproportionate influence on decisions from experience. This theory suggests that the best and worst outcomes differ *categorically* from intermediate outcomes. In this chapter we will examine whether other ways to conceptualise extreme outcomes are equally consistent with the existing evidence. If outcomes were instead over-represented in memory based on their *ordinal* or *continuous* extremity could we explain preferences towards the risky option of the high-value pair and against the risky option of the low-value pair? We suggest that the answer is “yes”. Furthermore, each of these alternatives corresponds to a distinct literature in which they have been used to explain experimental findings ranging from discrimination amongst perceptual stimuli to attentional phenomena regarding sequences of numbers.

3.1 Categorical extreme outcomes

As is evident in their nomenclature, the extreme-outcome rule was influenced by the peak-end rule, which describes the relationship between experiences as they unfold and evaluations using memory (Fredrickson, [2000]; Fredrickson & Kahneman, [1993]). We might expect these retrospective evaluations to correspond to the sum of the momentary experiences. Instead, people appear to evaluate events based on the single most intense moment (peak) and the final state (end) of an experience. Specifically, numerous experiments have demonstrated that the peaks of an experience are highly predictive of retrospective evaluations across a diverse range of modalities, including painful and aversive experiences (Ariely, [1998]; Chajut et al., [2014]; Kahneman et al., [1993]; Redelmeier & Kahneman, [1996]; Stone et al., [2000]), video clips (Baumgartner et al., [1997]; Fredrickson & Kahneman, [1993]),

mental or physical exertion (Hargreaves & Stych, 2013; Hsu et al., 2018), music (Rozin et al., 2004; Schäfer et al., 2014), prizes (Do et al., 2008), and odours (Scheibehenne & Coppin, 2020).

Why might this be the case? One reason is that categorical extreme outcomes are unique. There is considerable evidence going back to von Restorff (1933) that an item possessing a unique attribute is easier to remember, such as when one word in a list of black words is presented in red (for a review, see Schmidt, 1991). Categorical extreme outcomes differ from other outcomes in that they only have neighbours on one side whereas other outcomes have neighbours that are better and neighbours that are worse. Consequently, they are also the only outcomes that are either always preferred or always avoided no matter which outcome they are compared with and this uniqueness might increase their salience in memory.

These categorical extreme outcomes are also uniquely meaningful because they define the range of experienced outcomes. According to Fredrickson (2000), this “conveys the personal capacity necessary for achieving, enduring, or coping with that episode. In other words, the moment of peak affect intensity is the single moment that defines the personal capacity needed to face the experience again” (p. 590). Therefore, it might be advantageous to consider categorical extreme outcomes when making a decision because if the most extreme outcomes are not beyond your capacity, neither are the intermediate ones.

Despite this, there are also reasons to remain sceptical of the categorical peak-end rule and extreme-outcome rule. Although the peak-end rule emphasises the role of categorical extreme outcomes, the bulk of the evidence supporting this component of the rule depends on its ability to *predict* evaluations. Whilst this might initially appear innocuous, the assertion that categorical extreme outcomes are *disproportionately* influential is necessarily also dependent on the predictive ability of intermediate outcomes. Most early experiments included only the peak and end as predictors but some recent experiments have demonstrated that other aspects of the experience, such as the average, are similarly

capable of predicting retrospective evaluations. It has, therefore, become less clear whether the peaks are disproportionately influential or whether they merely capture equivalent information to the mean and median (Cojuharencu & Ryvkin, 2008; Ganzach & Yaor, 2019; Kemp et al., 2008; Miron-Shatz, 2009; Rozin et al., 2004; Schäfer et al., 2014; Seta et al., 2008; Steffens & Guastavino, 2015; Strijbosch et al., 2019).

Regarding the extreme-outcome rule, a recent experiment conducted by Ludvig and colleagues demonstrated that the preference towards the risky option in the high-value pair and against the risky option in the low-value pair can be observed with outcomes that are *adjacent* to the best and worst outcomes (Ludvig et al., 2018). A similar effect was observed in an experiment where the outcomes were drawn from a continuous distribution. Participants reported experiencing outcomes that were near the edges of the distribution that were not necessarily the best or worst outcome (Mason et al., 2020). Ludvig et al. (2018) accounted for their results by invoking an auxiliary assumption that “the psychological representation of the edges is fuzzy” (p. 1916). Although this modification allows the categorical extreme-outcome rule to accommodate these observations, they are more easily explained using ordinal or continuous extreme outcomes.

3.2 Ordinal extreme outcomes

A genuine choice *between* options requires us to select some options and forgo others. This action is necessarily comparative. The best options should be chosen regardless of the magnitude of the difference and this is reflected in the influence of ordinal comparisons on our judgements and decisions. For example, an outcome’s rank influences the relationship between financial status and satisfaction (Boyce et al., 2010; G. D. Brown, Gardner, et al., 2008; G. D. Brown et al., 2017; Osafo Hounkpatin et al., 2015), assessments regarding the morality of actions (Aldrovandi et al., 2013; Parducci, 1968), perceptions of whether behaviour is healthy (Maltby et al., 2012; Melrose et al., 2013; Moore et al., 2016; Wood et al., 2012), and manipulations that aim to improve health-related decisions (Aldrovandi et al., 2015; M. J. Taylor et al., 2015). Furthermore, ordinal representations of subjective

value have been identified in the brain (Mullett & Tunney, 2013; Tremblay & Schultz, 1999; Winston et al., 2014) and even incidental exposure to rank-based information is predictive of subsequent choices (Stewart, 2009; Stewart et al., 2015; Ungemach et al., 2011).

Similarly, the limited capacity of attention and memory causes us to neglect some outcomes and this implies that they are influenced by ordinal characteristics. This observation has given rise to the development of many rank-dependent models of selective attention (e.g., Birnbaum, 2008; Birnbaum & Chavez, 1997; Diecidue & Wakker, 2001; Jaffray, 1988; Lopes & Oden, 1999; Quiggin, 1982; Wakker, 2001; E. Weber & Kirsner, 1997). It has also contributed to several attentional theories that employ ordinal definitions of extreme outcomes (Pleskac et al., 2019; Tsetsos et al., 2012; Vanunu et al., 2020; Zeigenfuse et al., 2014). These researchers suggest that items capture more attention as their *rank* approaches the edges of the distribution and have demonstrated that extreme items are associated with attention-based phenomena such as the attentional blink (Kunar et al., 2017; Raymond et al., 1992).

In the previous chapter, we proposed an additional rank-based account of extreme outcomes. Our rational explanation was based on the consequences of neglecting or forgetting extreme options. To illustrate the role of rank in our account, suppose that you were presented with a choice between a pair of options, one known and the other unknown. How likely are you to make the correct choice? On the one hand, if the known option is the best that you have encountered within a given context, you can be confident in selecting it and if the known option is the worst that you have encountered, you can be confident in rejecting it. On the other hand, the probability of making the correct choice is equivalent to a coin toss for the median because there are an identical number of experienced options that were better or worse than the known option.

Other options fall between these two extremes and the amount of information lost by neglecting or forgetting an option is proportional to its ordinal distance from the median option. We also showed that a bias towards ordinal extreme *outcomes* can be used as

a proxy to minimise the information lost by neglecting extreme *options*. This ordinal definition is similar to the categorical extreme-outcome rule in that the best and worst outcomes are also most extreme based on rank. This allows ordinal theories to capture the observations of Ludvig et al. (2014) but the influence of ordinal extreme outcomes also extends to outcomes that are located near the edges. Therefore, in contrast with the extreme-outcome rule, ordinal theories can capture the findings of Ludvig et al. (2018) and Mason et al. (2020) without appealing to fuzzy representations.

3.3 Continuous extreme outcomes

Categorical extreme outcomes are always the best and worst and ordinal extremity always refers to an outcome's rank. In contrast, continuous extreme outcomes can be defined with reference to multiple aspects of the experienced distribution. An outcome can be considered extreme based on the extent to which it deviates from the *centre* or is located near one of the *edges*. Finally, extreme outcomes can refer to the distance between an outcome and its *neighbours* because this increases as outcomes approach the edges of a distribution. These *referents* to which extreme outcomes are defined are identical when the centre of an experienced distribution is equidistant from the edges but this equivalence breaks down when outcomes are asymmetrical. Therefore, in the following sections, we will briefly examine the evidence regarding continuous extreme outcomes with reference to the centre, edges, and neighbours.²

²It is also possible to conceptualise ordinal extreme outcomes with reference to the centre, edge, or neighbouring outcomes. Nonetheless, these three referents are functionally identical because rank ignores the magnitude of continuous distances. The ordinal distance between an outcome and the median or the edges is, therefore, complementary rather than providing independent information. There might be theoretical reasons to favour one of these referents over the others but our empirical work is silent regarding this distinction.

3.3.1 Centre of the distribution:

Although the distance from the centre is employed less often compared with the edge-based and neighbour-based accounts, we encountered two examples in the previous chapter. Recall that utility-weighted sampling explained the influence of extreme outcomes based on an optimal response to cognitive resource limitations. The idea that biases can be rational initially appears counter-intuitive but is based on the well-established bias-variance trade-off (S. Geman et al., 1992; Gigerenzer & Brighton, 2009; Hastie et al., 2009). The bias component decays faster than the variance component as the number of samples increases and the variance eventually approximates the overall estimation error. This principle suggests that the negative consequences of memory biases can be compensated by their ability to reduce the variance.

Lieder et al. (2018) demonstrated that the minimum-variance estimator is based on the *continuous distance* between each outcome and the expected value of the option. The problem is that the expected value is the parameter being estimated and will be unavailable whenever there is reason to use the estimator. Therefore, Lieder and colleagues suggested that the *average experienced outcome* should be used as a proxy for the expected value. This model is based on the continuous distance from the centre of the distribution and is able to capture the qualitative effects reported by Ludvig et al. (2014) using a single set of parameters. Insofar as the average outcome in the context is a reasonable substitute for the expected value of the option, utility-weighted sampling offers a centre-based explanation for why extreme outcomes are over-represented in memory.

In the second half of the previous chapter, we suggested that prioritising extreme options and outcomes 1) increases the probability of selecting the correct option and 2) increases the expected utility gained from these choices. The first claim was based on an ordinal definition of extremity whereas the second employed a centre-based continuous definition. Specifically, the expected utility of remembering an option is proportional to the continuous distance between the option and the average outcome in the higher-level context. We demonstrated that centre-based extreme *outcomes* can be used as an adequate

proxy for these extreme *options*. Therefore, similarly to utility-weighted sampling, this explanation suggests that recalling centre-based extreme outcomes can be advantageous.

3.3.2 Edges of the distribution:

Continuous extreme outcomes have primarily been used to explain the encoding of perceptual stimuli (e.g., Berliner et al., [1977]; Braida et al., [1984]; Marley & Cook, [1984]) and the effect of presentation order in short-term memory (e.g., Farrell & Lelièvre, [2009]; Henson, [1998]; Jou, [2010]). In each of these domains, a robust phenomenon has been observed in which items near the edges of a distribution (e.g., brightness, loudness, line length, pitch, area, weight, numerosity, or temporal order) are retrieved from memory with greater speed and accuracy than items near the centre (Berliner et al., [1977]; Bower, [1971]; Eriksen & Hake, [1957]; Lacouture & Marley, [2004]; Luce et al., [1982]; Murdock, [1960]; Neath et al., [2006]; D. L. Weber et al., [1977]). The serial position curve corresponding to the accuracy of these items is typically U-shaped, gradually increasing towards the edges, and therefore, either ordinal or continuous extreme outcomes are necessary to capture these phenomena.

Several explanations of these serial position curves have been proposed that suggest that items are either encoded or retrieved with reference to salient aspects of the distribution (Berliner et al., [1977]; Braida et al., [1984]; Farrell & Lelièvre, [2009]; Henson, [1998]; Jou, [2010]; Marley & Cook, [1984]). These accounts suggest that the U-shaped curve reflects a gradual increase in encoding or retrieval noise as items get further away from these salient reference points or anchors. One of the most well-known examples of a reference point underlies the diminishing sensitivity between outcomes in prospect theory as their distance increases from the salient zero-point separating gains from losses (Kahneman & Tversky, [1979]). This reference point is particularly salient but other aspects of the distribution, such as the best and worst outcomes, might serve as additional reference points and this could explain the U-shaped curves where accuracy diminishes as distance increases from the edges.

These theories can be divided into two broad classes: the first suggests that items

are imperfectly *encoded* based on their distance from salient anchors at the edges of the distribution (Berliner et al., 1977; Braida et al., 1984) and the second suggests that the edges are used as implicit reference points when items are *compared* with each other (e.g., Holyoak, 1978; Jou, 2010). In an experiment that teases apart these theories, Madan et al. (2021) presented the same participants with options across multiple contexts. Some options that were experienced in the first (encoding) context were later encountered in a second (choice) context. Participants' choices were relatively independent of the context in which these choices were made, and instead, were based on whether outcomes were extreme relative to the encoding context. This observation is consistent with the encoding theories—referred to as end-anchor theories—but is also compatible with ordinal accounts that are based on attention.

3.3.3 Neighbouring outcomes:

An alternative explanation for the U-shaped serial position curves in the previous section is that items stored in memory tend to interfere with other similar memory traces, and therefore, items are easier to remember when they are distinctive (M. C. Anderson & Neely, 1996; Murdock, 1960; Schmidt, 1991). The characteristic U-shape arises because the average distance between an item and the other remembered items is greater for those that are located towards the edges of the distribution. You can convince yourself of this claim either by remembering that the median minimises the sum of absolute deviations or noticing that the distance between the highest and lowest items corresponds to the entire range of the distribution whereas the furthest item from the midpoint is half this distance.

Neath et al. (2006) proposed an influential neighbour-based model (SIMPLE) in which the distinctiveness of an item is primarily influenced by its immediate neighbours (also see, G. D. Brown et al., 2007; Neath & Brown, 2006). In addition to explaining the U-shaped serial position curves described above, this *local similarity* model is also able to explain a phenomenon that arises when outcomes are irregularly spaced throughout the distribution. Specifically, their model correctly predicts that central items with relatively

few near neighbours are remembered more accurately than peripheral items with many near neighbours (Bower, 1971; G. D. Brown et al., 2007; Neath et al., 2006). This observation cannot be explained using global similarity models that are based on the unweighted average distance between items or using the accounts described in the previous sections.

The SIMPLE model was designed to explain discrimination performance based on perceptual attributes. It offers a promising explanation of these phenomena but Ludvig et al. (2018) argued that it cannot be extended to preferences in decisions from experience. Similarly to their earlier experiment that was described in the opening section of this chapter, Ludvig and colleagues presented participants with pairs of low-value options (safe: 25 points; risky: 5 or 45 points) and high-value options (safe: 75 points; risky: 55 or 95 points). This combination of options would likely give rise to the now familiar pattern of preferences towards the high-value risky option and against the low-value risky option. In this experiment, however, they presented participants with two additional risky options that were contingent on the condition they were allocated to.

In the *near neighbours* condition, these options resulted in outcomes that were separated from the other risky options by a single point (low-value: 6 or 44 points; high-value: 56 or 94 points) whereas in the *far neighbours* condition these options were positioned adjacent to the safe options (low-value: 24 or 26 points; high-value: 74 or 76 points). The SIMPLE model stipulates that introducing additional near neighbours reduces the distinctiveness of an outcome. Therefore, the best and worst outcomes should have been less influential in the near-neighbours condition but Ludvig and colleagues observed a negligible difference between their conditions.

How can we explain this observation? There were at least two large differences between the experiments conducted by Neath et al. (2006) using perceptual attributes and Ludvig et al. (2018) in decisions from experience. Firstly, the outcomes of Ludvig and colleagues were separated from their near neighbours by a single point and it is quite plausible that these outcomes were chunked in memory (Miller, 1956; Shiffrin & Nosofsky, 1994). In previous serial position experiments, spontaneous grouping of items led to a

step-like “scaloping” pattern where near neighbours were treated as equivalent and an increase in interposition errors within groups was observed (G. D. Brown et al., 2007; Henson, 1998; Hitch et al., 1996; Ryan, 1969). These patterns are compatible with the choices and memory responses reported by Ludvig et al. (2018).

Secondly, the additional risky options were designed to make the best and worst outcomes less distinctive but they also resulted in outcomes that were one point away from the *non-extreme* outcomes. It is possible that participants perceived the extreme outcome (e.g., 95 points) that was adjacent to its near neighbour (94 points) as more distinctive than the non-extreme outcome (55 points), which was similar to multiple outcomes clustered at the centre of the distribution (44, 45, and 54 points). To the extent that this argument is persuasive, it remains plausible that the observations of Ludvig et al. (2018) are compatible with a neighbour-based account of extreme outcomes.³

3.4 Overview of experiments

In the previous sections, we described numerous theories that employed extreme outcomes and defined them using different levels of measurement. These theories cited their own phenomena and mechanisms. Categorical extreme outcomes were used to explain memory-based evaluations of affective experiences, ordinal extreme outcomes were used to explain the allocation of attention, and continuous extreme outcomes were used to explain the U-shaped serial position curve in short-term memory. Based on the existing evidence, however, it is unclear where the experience-based choices observed by Ludvig, Madan, and colleagues fit into this picture. Therefore, the aim of this chapter is to tease apart

³A similar argument was made by Brown and colleagues (2007): “The forgetting curve in SIMPLE is closely approximated by exponential forgetting in the short term and power-law forgetting over longer time periods, but the form of the best fitting function was found to depend to a large (and perhaps intuitively surprising) extent on parameter values that, from a theoretical point of view, seem rather peripheral to the core assumptions of the model. We therefore suggest that the search for ‘the’ forgetting function may be misguided” (p. 566).

these explanations by manipulating attributes such as the value and variance of outcomes and the skewness of the contexts in which they are experienced.

Across three experiments, we presented participants with a task that involved repeatedly making choices between pairs of options represented by coloured squares (see Figure 3.1). Similarly to the experiments conducted by Ludvig et al. (2014), these involved *safe options* that always resulted in a fixed outcome and *risky options* that had an equal probability of resulting in an outcome that was better or worse than the outcome of the safe option. For example, in the first experiment, participants were presented with a safe option that always resulted in 50 points and a risky option that resulted in 40 or 60 points. Although a single choice never involved more than two options, each experiment consisted of numerous options that were presented in an interspersed fashion. This allowed us to examine the influence of extreme outcomes on participants' choices.

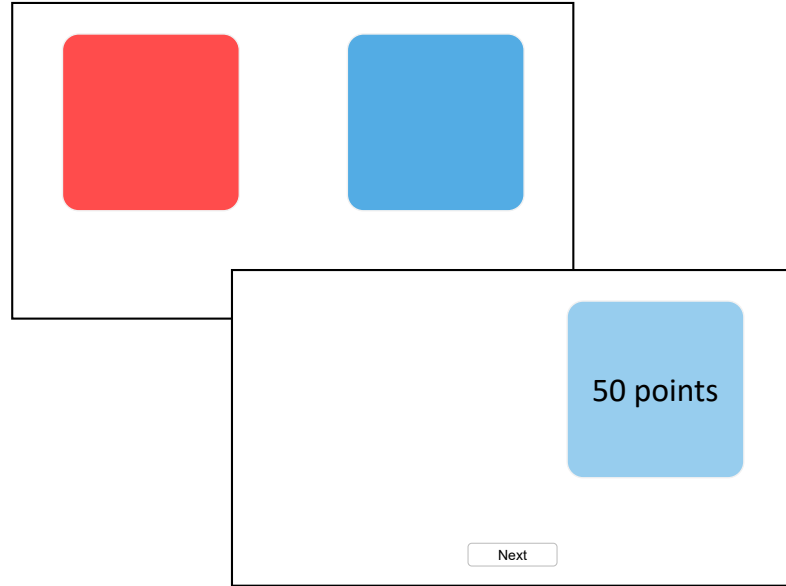


Figure 3.1: **Screenshot of a decision trial between two options.** In this example, the participant selects the blue option and receives 50 points. Feedback is presented until they click the ‘Next’ button.

3.4.1 Terminology

One of the hurdles that we encountered when describing the experiments in this chapter was that there were multiple manipulations involving “value” that could be described using very similar labels. The inadequacy of language was particularly salient when describing the high- and low-value options, the high- and low-value outcomes associated with these options, the highest or lowest values within a given context, and the context manipulations that resulted in high or low average values. In order to mitigate the inevitable confusion that would arise from using nearly identical terminology, we aimed to consistently describe these manipulations using the conventions described below, which are also presented visually in Figure 3.2 with reference to the outcomes used in Experiment 1.

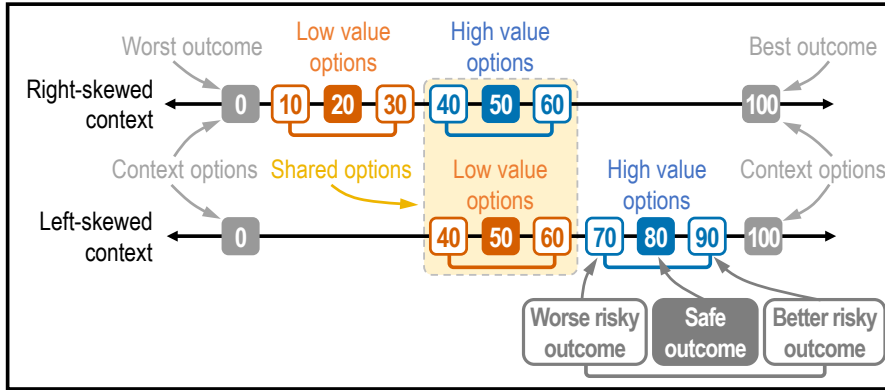


Figure 3.2: **Terminology use to describe each manipulation with reference to the outcomes used in Experiment 1.** Each square represents an outcome associated with an option. Outcomes connected with a $\sqcup$ were associated with the same risky option.

In each experiment, participants encountered multiple pairs of safe and risky options. These often included one pair that had a higher or lower expected value than the other options. We will refer to these pairs as *high-value* and *low-value options* (or *higher* and *lower value options*). For example, the first experiment involved a pair of high-value options that had an expected value of 50 points and a low-value pair that had an expected value of 20 points. The risky option within each pair was associated with one outcome that was better than the value of the safe option and another that was worse (e.g., safe: 50 points; risky: 40 or 60 points). Therefore, we will refer to these as the *better* and *worse*

outcomes associated with the risky option.

The extreme-outcome rule suggests that the single most extreme outcomes within a given context are over-represented in memory. We will refer to these categorical extreme outcomes as the *best* and *worst outcomes*. Finally, we manipulated the average value of the outcomes that participants encountered in each condition to examine the role of centre-based continuous extreme outcomes. To avoid confusion with the high- and low-value options, we will refer to these distributions based on skewness. *Left-skewed contexts* are those that have outliers on the left tail and an overall average value that is above the midpoint of the distribution. In contrast, *right-skewed contexts* have outliers on the right tail and an average below the midpoint.

3.4.2 Manipulation 1: Value

The extreme-outcome rule proposed by Ludvig et al. (2014) entails that the best and worst outcomes are uniquely influential. We evaluated this conjecture by presenting participants with low-value and high-value pairs of options that were not associated with either of these categorical extreme outcomes. To provide a concrete example, participants in our second experiment were presented with a low-value pair (safe: 25 points; risky: 10 or 40 points), a high-value pair (safe: 75 points; risky: 60 or 90 points), and an *extreme* pair that resulted in the best and worst outcomes (safe: 50 points; risky: 0 or 100 points). According to the ordinal and continuous accounts, the better outcome of the high-value risky option and the worse outcome of the low-value risky option are more extreme than the other outcomes associated with those options. Over-representation of these ordinal and continuous extreme outcomes in memory should, therefore, generate a preference towards the high-value risky option and against the low-value risky option. Given that these options never resulted in the best or worst outcomes, the categorical accounts would be entirely unable to explain these preferences.

3.4.3 Manipulation 2: Variance

The defining characteristic of ordinal extreme outcomes is that they disregard the *magnitude* of continuous differences between outcomes. This allowed us to empirically disambiguate the ordinal and continuous accounts by manipulating the variance of the risky options. Again providing an example from Experiment 2, participants in the *high-variance* condition were presented with the options described in the preceding paragraph whereas participants in the *low-variance* condition were presented with a low-value pair (safe: 25 points; risky: 20 or 30 points) and a high-value pair (safe: 75 points; risky: 70 or 80 points) in which the standard deviations of the risky options were three times smaller. Therefore, the continuous accounts predict that the influence of extreme outcomes should be stronger in the high-variance condition than the low-variance condition and the ordinal accounts suggest identical preferences across conditions because the outcomes' rankings are unaffected by the variance manipulation.

3.4.4 Manipulation 3: Skewness

Whereas every outcome exerts an influence on the average outcome, the edges of a distribution are only influenced by the best and worst outcomes. This implies that manipulating the skewness of intermediate outcomes should have a different impact on extreme outcome theories depending on whether they are defined with reference to the centre or the edges. For example, in our first experiment (depicted in Figure 3.2), participants were presented with a pair of options (safe: 50 points; risky: 40 or 60 points) that was located at the midpoint between the best outcome (100 points) and the worst outcome (0 points). This *shared* pair was identical across conditions.

In the right-skewed condition, participants also encountered a pair of options that had a lower expected value (safe: 20 points; risky: 10 or 30 points). These options shifted the average outcome below the midpoint, and therefore, the better outcome of the shared risky option (60 points) was further from the average than the worse outcome

Table 3.1: Summary of the predictions associated with each level of measurement and referent (the aspect of the distribution to which ‘extremity’ refers).

Level	Referent	Predicts greater risk-seeking for...				
		Higher value option?			Right-skewed context?	
		Exp. 1	Exp. 2	Exp. 3	Exp. 1	Exp. 3
Categorical	Edge	No	No	No	No	No
Ordinal	Equivalent	Yes	Yes [†]	Yes [‡]	Yes	No
Continuous	Edge	Yes	Yes [*]	Yes	No	No
Continuous	Centre	Yes	Yes [*]	Yes	Yes	Yes
Continuous	Neighbours	?	?	?	Yes	Yes

^{*} These theories predict that greater risk-seeking for higher value option will be stronger in the high variance condition.

[†] These theories predict that greater risk-seeking for higher value option will be equal across variance condition.

[‡] These theories predict a weaker effect in this experiment.

(40 points). This was reversed in the left-skewed condition where the additional pair had a higher expected value than the shared pair (safe: 80 points; risky: 70 or 90 points). For this reason, the centre-based accounts entail a preference towards the risky option of the shared pair in the right-skewed condition and against this option in the left-skewed condition whereas the edge-based accounts suggest similar preferences across conditions.

A summary of the hypotheses associated with each level of measurement are provided in Table [3.1](#).

3.5 General method

3.5.1 Participants

All of the experiments described in this chapter were conducted using undergraduate psychology students enrolled at UNSW Sydney. 80 students participated in Experiment 3a and 130 students participated in each of the other experiments. Skewness and variance were manipulated between-subjects and participants were randomly allocated into condi-

Table 3.2: Summary of demographics for each experiment in Chapter 3.

Experiment	N	Age		Gender		Bonus	
		Mean	SD	Female	Male	Mean	SD
1 Skewed (non-extreme)	130	20.0	2.4	84	45	\$5.95	\$2.85
2 Variance	130	19.7	4.0	78	52	\$4.00	\$0.00
3a Skewed (extreme)	80	19.6	4.4	60	20	\$5.33	\$2.47
3b Skewed (extreme)	130	19.6	2.8	87	43	\$4.97	\$2.88

Note:

The standard deviation of payments for Experiment 2 appears as \$0.00 because participants were paid based on the total number of points earned instead of based on the outcome of a randomly selected choice and because payments were rounded to the nearest dollar.

tions with balanced sample sizes. Further demographic information for each experiment is presented in Table 3.2.

3.5.2 Design and procedure

Upon entering the laboratory, we informed participants—in groups no larger than four—that the experiment involved a computer-based task in which they could earn real money based on their choices. Participants completed this task in individual rooms where detailed instructions were presented on the screen. These instructions emphasised that their objective was to earn as many points as possible and explained how those points would be converted into dollars. Following this, participants made a series of choices between pairs of options (the choice task) and were then asked which outcome came to mind first and the percentage of choices that resulted in each outcome (the memory tasks). These tasks were closely analogous to those used by Ludvig et al. (2014). Each experiment was programmed in MATLAB using PsychToolbox (Brainard, 1997; Kleiner et al., 2007; Pelli, 1997) and the code used for each experiment can be accessed at <https://github.com/joelholwerda>.

Choice task

Participants encountered pairs of options in a task that interspersed between four and six options. The options that participants encountered are displayed in Figure 3.3 and will be described within the sections associated with each experiment. Although we never explicitly distinguished them, each choice was either a *decision trial* or a *catch trial*. Decision trials involved pairs of safe and risky options that had *identical* expected values. These choices were used to examine our primary hypotheses regarding value, variance, and skewness. Conversely, catch trials involved pairs that had *different* expected values and allowed us to ensure that participants were adequately engaging with the task. Following the criterion proposed by Ludvig et al. (2014), we excluded data from participants who selected the better option on less than 60% of choices.

Participants were not given written descriptions of each option and were required to learn about options by selecting them and receiving feedback on the number of points earned. No feedback was given for the option that was not chosen. In addition to the decision and catch trials, we also presented participants with *single-option trials* in which they had to click the presented option to continue. These trials were included to mitigate the possibility that participants would completely avoid an option that was initially unfavourable, and therefore, not have an opportunity to learn that it also produces better outcomes (Denrell & March, 2001).

Following the experiment, we converted the points that participants earned into real money. In Experiment 2, we paid them \$1 for every 5000 points accumulated in the choice task for comparability with Madan et al. (2014). These payments were based on participants' total scores, which incorporated the outcomes of 360 individual choices, and therefore, the average outcome associated with each risky option was very similar to its expected value. We aimed to address this issue in Experiment 1 and 3 by randomly selecting a choice and paying participants \$1 for every 10 points earned on that one choice. This ensured that selecting risky options was associated with greater variability in the amount of money that the participant earned.

CHAPTER 3. LEVELS OF MEASUREMENT

There was no time limit when making choices and feedback was presented until participants selected the “next” button, which was positioned at the horizontal centre of the screen to ensure that the starting position of the cursor was always roughly equidistant from the two options. The colours associated with each option were randomised for each participant and the side of the screen on which options were presented was randomised for each trial. Choices were presented in a randomised order across five blocks and participants were encouraged to take a short break before commencing each subsequent block.

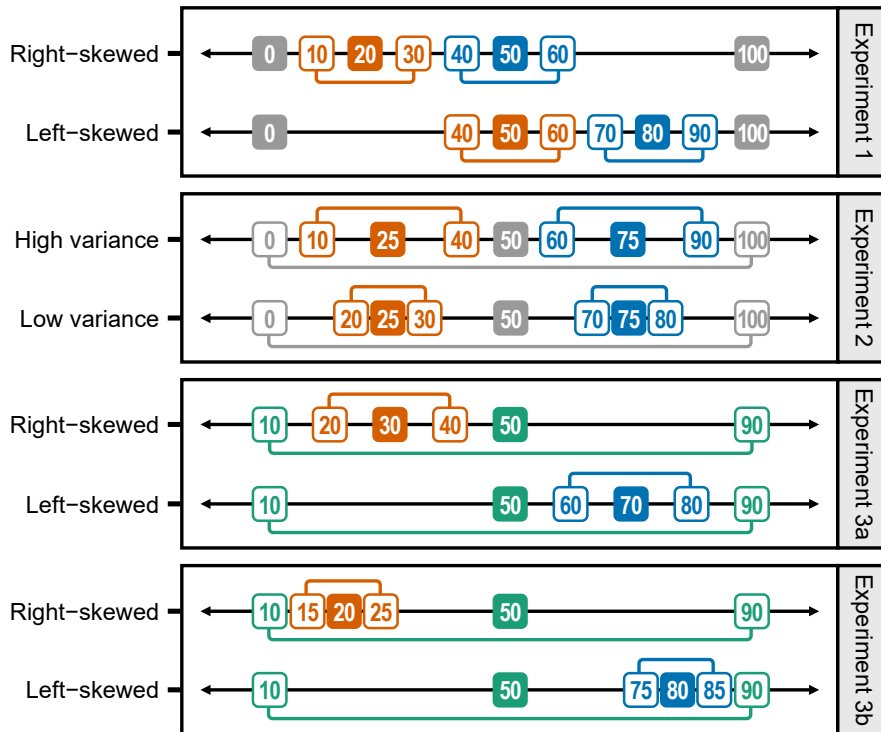


Figure 3.3: **Diagram of the designs of each experiment in Chapter 3.** Each square represents an outcome associated with an option. Outcomes connected with a $\sqcup$ were associated with the same risky option and occurred with equal probability. In the figures displayed in this chapter, we adopt a convention that the low-value options are depicted in orange, high-value options are depicted in blue, other options associated with a hypothesis are depicted in green, and other options (not associated with a hypothesis) are depicted in grey.

Memory tasks

After participants completed the choice task, we examined their memory for each of the experienced outcomes. We presented them with the coloured squares associated with each option and asked them to report “the outcome that comes to mind first”. This first question was designed to estimate the *availability* of each outcome in memory. Following this, we again presented them with each coloured square and asked them, “When you selected the option presented above, on what percentage of these choices did you experience each of the outcomes listed below?” All experienced outcomes were listed regardless of whether they were associated with the displayed option. This second question was designed to assess whether the *frequency* of experienced outcomes was distorted in memory. In both of these tasks, the coloured squares were presented sequentially in a semi-randomised order with the risky options—used to examine the influence of extreme outcomes—always presented before the safe options.

3.5.3 Analysis

We used Bayesian regression to analyse participants’ responses because it allowed us to flexibly model the hierarchical structure of our experimental tasks, incorporate regularisation, and examine degrees of credibility rather than dichotomous indicators of significance or non-significance. All posterior distributions reported in this section were determined by Hamiltonian Monte Carlo using the brms package in R (Bürkner, 2017). For each posterior, we report the probability that there was a difference between conditions in the predicted direction. This statistic provides information about whether our results can be attributed to sampling error but does not indicate the magnitude of the difference and cannot provide evidence for the absence of an effect (Makowski, Ben-Shachar, et al., 2019; Makowski, Ben-Shachar, et al., 2019). Consequently, we also report the highest density interval that contained 95% of the posterior (95% CI) and discuss whether these intervals

are consistent with the predictions of each theory.⁴

Weakly regularising priors were selected for each parameter. A student- $t(7, 0, 0.5)$ distribution was used for the slope, intercept, and threshold parameters. A half-student- $t(7, 0, 0.5)$ distribution was used for the standard deviation parameters in the hierarchical models and an LKJ(4) distribution was used for the correlation between intercept and slope parameters (Lewandowski et al., 2009). These priors were selected to conform with the Stan prior choice recommendations (Stan Development Team, 2020)⁵ and predictions from the prior distribution were inspected to ensure that they allowed the range of plausible observations (Gabry et al., 2019). The same prior distributions were used for the subsequent experiments with a small number of additions that will be described alongside the relevant models. Numerical predictors were standardised (mean = 0, SD = 1) and categorical variables were deviation coded (-1, 1) so that they were centred and their scale was comparable when setting priors (A. Gelman, 2008).

All parameters had bulk and tail effective sample sizes greater than 10000 and an $\hat{R} < 1.01$ suggesting adequate chain convergence (Vehtari et al., 2020). There were no divergent transitions and the other Stan diagnostics did not indicate issues with estimation. Rank histograms, posterior predictive distributions, and other diagnostic plots

⁴We chose to present these summaries of the posterior distribution because they are analogous to frequentist statistics that might be more familiar to some readers. The probability of direction corresponds roughly to the complement of a one-tailed p-value ($1 - p$) and the highest density interval corresponds roughly to a confidence interval. They should not be treated as equivalent because our analysis incorporates informative priors (Nalborczyk et al., 2019), but they both attempt to answer similar questions. The full posterior distributions can be accessed at <https://github.com/joelholwerda>.

⁵The Stan prior choice recommendations suggest using a student- t distribution with degrees of freedom between 3 and 7. We selected the latter to provide stronger protection against implausible parameter estimates whilst still allowing us to learn from the data. To examine the sensitivity of our conclusions to our choice of priors, we also assessed our hypotheses using two alternate sets of priors: one more informative set that used normal(0, 0.5) distributions and one less informative set that used student- $t(3, 0, 1)$ distributions. To examine the sensitivity of our conclusions to our choice of likelihood functions, we assessed each hypothesis using numerous alternate models that other researchers might have justifiably used to answer our research questions. The parameter estimates using these priors and likelihoods can be accessed at <https://github.com/joelholwerda>.

were examined for each model and can be accessed at <https://github.com/joelholwerda>. Preregistered hypotheses for each experiment can be accessed at <https://osf.io/d8pq3>.

3.6 Experiment 1: Value and skewness (non-extreme shared options)

To understand the design of the experiments in this chapter, it will be useful to consider an analogy with the rabbit-duck image that features in every undergraduate course on visual perception. This image demonstrates that the same component that, from one perspective, was interpreted as the ears of a rabbit can also be interpreted, from another, as the bill of a duck. Each experiment in this chapter is similar to the rabbit-duck in that each option is a component of a design that creates two distinct pictures when viewed from different perspectives. Consider the options used in our first experiment that are depicted in Figure 3.2 (also see the top panel of Figure 3.3):

The first perspective interprets these options as manipulating the *value* of outcomes (Manipulation 1 in the overview of experiments). This aspect consists of the horizontally-aligned low-value options (orange) and high-value options (blue). Additional context options (grey) were included to ensure that these options were never associated with the best or worst outcomes, and for that reason, the categorical theories assert that choices involving *only* these options (e.g., the decision trials) should be unaffected by the overrepresentation of extreme outcomes. In contrast, the ordinal and continuous theories assert that the worse outcome of the low-value risky options (right-skewed: 10 points; left-skewed: 40 points) and the better outcome of the high-value options (right-skewed: 60 points; left-skewed: 90 points) are extreme relative to the other outcomes associated with these options. Consequently, they predicted a preference towards the high-value risky option and against the low-value risky option whereas the categorical theories would be unable to explain this pattern.

The second perspective interprets these same options as manipulating the *skewness*

of the experienced distribution (Manipulation 3). This aspect consists of the vertically-aligned shared options (within the yellow-shaded region) that are identical across conditions and the flanking context options (right-skewed: orange; left-skewed: blue) that were used to manipulate the average experienced outcome. Additional context options (grey) were included to ensure that the range was identical across conditions, and for that reason, the distance between each shared outcome and the edges of the distribution was similarly identical. In contrast, the centre-based theories assert that the worst outcome of the shared risky option (40 points) is further from the average in the left-skewed condition and the better outcome of this option (60 points) is further from the average in the right-skewed condition. Consequently, they predicted a preference towards the shared risky option in the right-skewed condition and against this option in the left-skewed condition whereas the edge-based theories would be unable to explain this pattern.

Therefore, viewing the options in our first experiment from these two perspectives allowed us to simultaneously evaluate the adequacy of the categorical extreme-outcome rule and disambiguate the centre-based and edge-based theories. Participants encountered these options in a task that involved making a total of 330 choices across five blocks. Each block consisted of 24 decision trials, 30 catch trials, and 12 single-option trials (see the general method for more detail). The choice and memory data from one participant (Right-skewed condition) was excluded because they selected the better option on less than 60% of catch trials.

3.6.1 Results and discussion

Value manipulation

The proportion of decision trials in which participants selected the risky option is displayed in Figure 3.4a. The continuous and ordinal theories predicted that participants would select the risky option more often for the high-value pairs (blue) than the low-value pairs (orange) whereas the categorical theories predicted similar preferences across these

options. Safely ignoring the shared options (bottom panel) until we discuss the skewness manipulation, we observe a pattern consistent with the continuous and ordinal theories in both the right-skewed condition (top panel) and the left-skewed condition (middle panel).

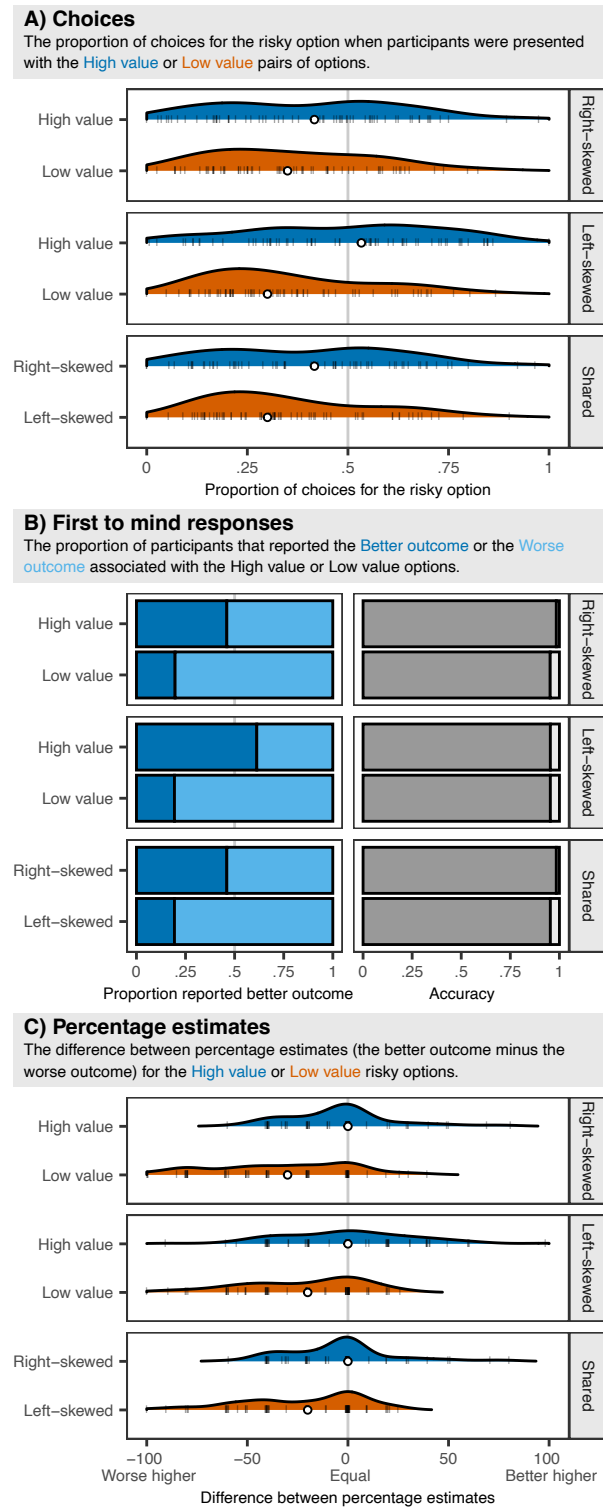
We examined this pattern further using Bayesian hierarchical logistic regression predicting the choice that participants made on each decision trial. We included the condition (right-skewed or left-skewed), option value (high-value or low-value), and their interaction as fixed predictor variables. The intercept and the slope for option value were allowed to vary for each participant. This model suggests that there is a 99.96% probability that participants were more likely to select the risky option for the high-value options than the low-value options. This difference corresponds to being roughly 1.50 times more likely to select the risky option for the high-value options (95% CI = [1.18, 1.89]) and is difficult to explain using the categorical extreme outcome theories.

Assuming that participants' preferences were influenced by the availability of extreme outcomes in memory, their responses to the first to mind questions (displayed in Figure 3.4b) should echo the choices displayed in Figure 3.4a. Specifically, we would expect that more participants would report the better outcome (dark blue) as coming to mind first in response to the high-value risky option compared with the low-value risky option. This is precisely what we observe in Figure 3.4b (again ignoring the bottom panel until we discuss the skewness manipulation).

We examined this pattern further using Bayesian multinomial regression and the same predictors as the model of participants' choices. Consistent with the increased availability of continuous or ordinal extreme outcomes, this model suggests that there is a greater than 99.99% probability that participants were more likely to report the better outcome associated with the risky option for the high-value options compared with the low-value options, a difference that corresponds to being roughly 5.27 times more likely to report the better outcome for the high-value options (95% CI = [2.71, 10.79]).

A similar pattern is echoed in participants' responses to the percentage estimate questions, which are displayed in Figure 3.4c. This figure—and the subsequent analysis—

Figure 3.4: **Choices and responses to the memory questions for Experiment 1.** The white dots represent the median response. The bottom (Shared) panels reproduce the responses to the high-value options in the right-skewed condition (top panel) and the low-value options in the left-skewed condition (middle panel), but emphasise that they were identical across conditions. Accuracy represents the proportion of participants' responses to the first to mind questions that corresponded to an experienced outcome.



depicts the difference between participants' percentage estimates for the better and worse outcomes of each option, thus allowing us to gauge the perceived *relative* frequency of these outcomes.⁶ These responses appear to demonstrate that the perceived relative frequency of the better outcome was higher for the high-value options (blue) compared with the low-value options (orange).

We examined them further using Bayesian linear regression and the same predictors as the previous models except that the varying slope was removed due to issues with model convergence. This model suggests that there is a greater than 99.99% probability that people are more likely to report the better outcome as occurring with greater frequency when presented with the high-value options compared with the low-value options, corresponding to a difference of roughly 0.73 standard deviations (95% CI = [0.51, 0.95]). Therefore, we should conclude that participants' choices and memory responses provide considerable evidence against the categorical extreme outcome theories.

Skewness manipulation

The bottom (Shared) panel in each sub-figure of Figure 3.4 displays participants' responses to the option that was shared across conditions. This panel reproduces the responses to the high-value options in the right-skewed condition (top panel) and the low-value options in the left-skewed condition (middle panel), but emphasises their alignment within the second perspective.

The centre-based theories predicted that participants would choose the risky option more often in the right-skewed condition than the left-skewed condition. Participants' choices in the bottom panel of Figure 3.4a appear to be broadly consistent with this

⁶Roughly 80% of percentage estimates included only the outcomes that were associated with the relevant option. For these responses, the estimates for the better and worse outcomes are complementary whereas the other (inaccurate) responses might not sum to 100% for these outcomes. Using the difference between them integrates information from both estimates whilst eliminating the problems that arise when including highly collinear variables in a regression analysis.

conjecture. We further examined participants' choices between the shared options using Bayesian hierarchical logistic regression. We included the condition (right-skewed or left-skewed) as a fixed predictor variable and the intercept was allowed to vary for each participant. This model suggests that there is an 85.03% probability that participants were more likely to select the risky option in the right-skewed condition than the left-skewed condition (Median odds = 1.23, 95% CI = [0.82, 1.81]). This estimate provides some evidence in favour of the centre-based theories but also assigns non-trivial probability to small differences that would be consistent with the edge-based theories.

The responses to the memory questions, however, provide considerably stronger evidence in favour of the centre-based theories (see the bottom panels of Figure 3.4b and 3.4c). Similarly to the first perspective described above, we expected that choosing the risky option would be echoed in both over-estimating the relative frequency of the better outcome and reporting it as the first outcome that comes to mind. We examined the first to mind responses using multinomial regression with the condition (right-skewed or left-skewed) as a fixed predictor. This model suggests that there is a 99.92% probability that participants were more likely to report the better outcome as coming to mind first in the right-skewed condition than the left-skewed condition. This corresponds to being roughly 2.99 times more likely to report the better outcome in the right-skewed condition (95% CI = [1.43, 5.95]).

Similarly, we examined the percentage estimates displayed in Figure 3.4c using Bayesian linear regression and the same predictor variable as the model of first to mind responses. Consistent with the centre-based theories, this model suggests that there is a 99.95% probability that participants in the right-skewed condition were more likely to provide a higher percentage estimate for the better outcome than those in the left-skewed condition. This difference corresponds to roughly 0.59 standard deviations (95% CI = [0.25, 0.92]).

Thus, to summarise the results of our first experiment, viewing the options from the first perspective provided strong evidence against the categorical theories and viewing

them from the second perspective provided some evidence in favour of the centre-based theories. It is worth noting, however, that the evidence of an effect of centre-based extreme outcomes was much stronger in participants' memory responses than it was on their choices. We will revisit this distinction between centre- and edge-based theories in Experiment 3.

3.7 Experiment 2: Value and variance

The first perspective in our second experiment interprets the options as a variation on the *value* manipulation that was introduced in the previous experiment (Manipulation 1). Looking now at the second panel of Figure 3.3, this aspect consists of the horizontally-aligned low-value options (orange) and high-value options (blue) that similarly never resulted in the best or worst outcomes. In accordance with the previous experiment, the ordinal and continuous theories predicted a preference towards the high-value risky option and against the low-value risky option whereas the categorical theories would be unable to explain this pattern.

So what additional value is offered by this variation? One possible rebuttal to our evidence against the categorical theories is that the psychological representation of the edges is fuzzy, and therefore, outcomes that are close to the edges (e.g., 10 or 90 points in our previous experiment) might be confused or chunked with the best and worst outcomes (Ludvig et al., 2018). The variation used in our second experiment addresses this rebuttal by ensuring that the categorical extreme outcomes were separated from their nearest neighbours by a distance that would make fuzzy encoding implausible. Specifically, in the low variance condition, this distance was equivalent to 20% of the entire range of experienced outcomes and adopting this amount of fuzziness would render the categorical accounts utterly vacuous.

The second perspective interprets these same options as manipulating the *variance* of the outcomes associated with the risky options (Manipulation 3). This aspect consists

of the vertically-aligned options that have identical expected values. The rank order of these outcomes is identical across conditions, and for that reason, the ordinal theories assert that choices involving *only* these options should be unaffected by the disproportionate influence of extreme outcomes. Nonetheless, the continuous distance separating the outcomes of the risky options was three times greater in the high variance condition (30 points) compared with the low variance condition (10 points). This continuous distance corresponds to the magnitude of the *difference* in continuous extremity between these outcomes. Consequently, the continuous theories predicted that the preferences towards the high-value risky option and against the low-value option would be stronger in the high variance condition whereas the ordinal theories would be unable to explain this pattern.

Therefore, viewing the options in our second experiment from these perspectives allowed us to address the fuzzy encoding rebuttal whilst differentiating between the ordinal and continuous theories. Participants encountered these options in a task that involved making a total of 360 choices across five blocks. Each block consisted of 36 decision trials, 24 catch trials, and 12 single-option trials. All participants selected the better option on greater than 60% of catch trials.

3.7.1 Results and discussion

Value manipulation

Similarly to the first experiment, participants' choices (displayed in Figure 3.5a) appear to exhibit a pattern of preferences towards the high-value risky option (blue) and against the low-value risky option (orange). We examined this further using the same logistic regression model that was used in the previous experiment but replacing the skewness conditions with variance conditions. This model suggests that there is a greater than 99.99% probability that participants were more likely to select the risky option for the high-value options than the low-value options (Median odds = 1.59, 95% CI = [1.24, 2.05]). This provides additional evidence against the categorical extreme outcome theories.

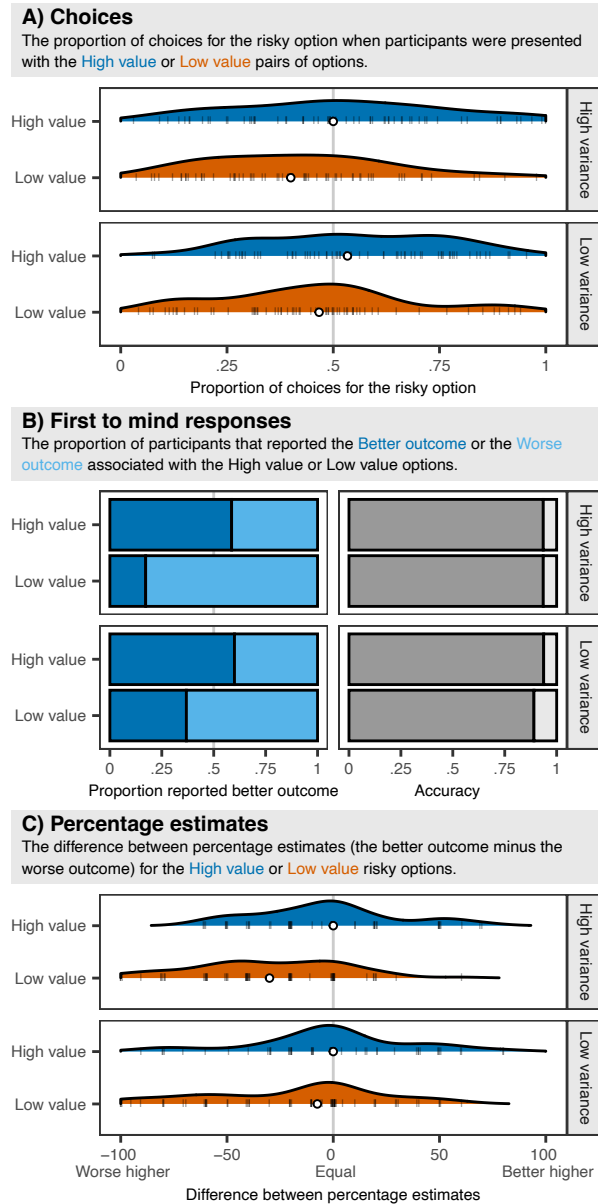


Figure 3.5: **Choices and responses to the memory questions for Experiment 2.** The white dots represent the median response. Accuracy represents the proportion of participants' responses to the first to mind questions that corresponded to an experienced outcome.

To evaluate the fuzzy encoding rebuttal, we used contrasts to focus on participants' choices in the low variance condition. These contrasts suggest that there is a 99.04% probability that there was a similar pattern of preferences in this condition, corresponding to participants being roughly 1.54 times more likely to select the risky option for the high-value options 95% CI = [1.09, 2.21]). To explain these choices, the fuzzy encoding rebuttal would be required to chunk outcomes together that were separated by a fifth of the overall range of the experienced distribution.

This pattern of preferences is echoed in participants responses to the memory questions. The proportion of participants that reported the better outcome (dark blue) as coming to mind first when presented with the risky options is displayed in Figure 3.5b. We examined these responses using the same model as the previous experiment. This model suggests that there is a greater than 99.99% probability that participants were more likely to report the better outcome as coming to mind first for high-value options compared with the low-value options (median odds = 5.25, 95% CI = [2.56, 11.15]). Focusing on the low variance condition, there is a 99.77% probability that this pattern was present in the low variance condition, corresponding to participants being 3.17 times more likely to report the better outcome for the high-value risky option (95% CI = [1.38, 8.08]).

A similar pattern is exhibited in participants' percentage estimates (displayed in Figure 3.5c). Based on the same model as the previous experiment, there is a greater than 99.99% probability that people are more likely to report that the better outcome occurs more often when presented with the high-value (blue) options compared with the low-value (orange) options (Median = 0.57, 95% CI = [0.37, 0.75]). Again focusing on the low variance condition, there is a 99.85% probability that participants in the right-skewed condition were more likely to provide a higher percentage estimate for the better outcome than those in the left-skewed condition. This difference corresponds to roughly 0.45 standard deviations (95% CI = [0.15, 0.74]).

Interpreted together, the choices that participants made and their responses to the memory questions in Experiment 2 provide further evidence against the categorical ex-

treme outcome theories. Although none of the outcomes were the best or worst, participants were still more likely to select the risky option when the better outcome associated with the risky option was further from the centre of the distribution than when the worse outcome was further from the centre. This evidence is particularly compelling because the same pattern of responses was observed for the lower variance condition where there was considerable distance between the outcomes associated with the risky options and the edges of the distribution.

Variance manipulation

The continuous extreme outcome theories assert that the effect of the value manipulation on choices (displayed in Figure 3.5a) should be stronger in the high variance condition (top panel) than the low variance condition (bottom panel) whereas the ordinal theories assert that preferences should be independent of variance. The evidence regarding these hypotheses is not convincing. Someone arguing in favour of the continuous theories might point out that the difference between the *median* of the low-value options (orange) and high-value options (blue) is roughly seven percent greater in the higher variance condition. In response to this assertion, someone arguing in favour of the ordinal theories might emphasise that the difference in the *means* is less than one percent. They might further declare that, using the model described in the previous section, there is only a 58.82% probability that the effect is stronger in the higher variance condition (Median odds = 1.03, 95% CI = [0.80, 1.32]). This is scarcely better than a coin toss!

Participants' responses to the memory questions offer some further evidence in favour of the continuous theories. The difference between the low-value options and high-value options appears to be larger in the high-variance condition for both the first to mind questions (Figure 3.5b) and percentage estimates (Figure 3.5c). We examined these patterns further using the same predictors that we employed for the value manipulation in the previous section. These models suggest that there is a 94.72% probability that the value manipulation had a stronger effect on participants' responses to the first to mind

questions in the higher variance condition (Median = 1.64, 95% CI = [0.90, 3.09]) and an 88.87% probability of a similar interaction for percentage estimates (Median = 0.12, 95% CI = [-0.07, 0.31]). This is hardly irrefutable evidence and further investigation will be required before reaching any strong conclusions regarding these theories.

Thus, to summarise the results of the second experiment, the first perspective provided compelling evidence against the categorical extreme outcome theories. Even if these theories were to evoke the auxiliary assumption that encoding is fuzzy, the amount of noise required to explain our results would leave them categorical in name only. Unfortunately, the second perspective was considerably less decisive. The memory responses provide some evidence in favour of the continuous theories but an advocate of the ordinal theories would undoubtedly maintain considerable scepticism. We will provide some further evidence regarding these theories in the following experiment whilst attempting to differentiate between the centre- and edge-based theories.

3.8 Experiment 3: Value and skewness (extreme shared options)

Our third experiment was conducted in two parts (3a and 3b) and these are displayed in the third and fourth panels in Figure [3.3](#). For simplicity, we will illustrate the broad patterns with reference to Experiment 3a and then highlight the distinct attributes of Experiment 3b. The first perspective in this experiment interprets the options as another variation on the *value* manipulation that was employed in the previous experiments (Manipulation 1). In contrast with these experiments, the low-value options (orange) and high-value options (blue) are aligned along the diagonal so that participants were presented with *either* low-value options (right-skewed condition) *or* high-value options (left-skewed condition). This allowed us to examine whether the pattern of preferences towards the high-value risky option and against the low-value risky option requires that these options are experienced by the same participant.

3.8. EXPERIMENT 3: VALUE AND SKEWNESS (EXTREME SHARED OPTIONS)

A broad distinction can be formed between two classes of theories based on their predictions regarding this manipulation. As we emphasised with reference to the skewness manipulation (Manipulation 3), the best and worst outcomes are the only other outcomes that influence the continuous distance between an outcome and the edges of the distribution. This independence from intermediate outcomes necessitates that the edge-based theories would predict similar preferences regardless of whether these options are presented within- or between-subjects. Conversely, every experienced outcome influences both the overall average and the rank of other outcomes, and therefore, introducing or eliminating an option might alter the predictions of the centre- and edge-based theories. This attribute offers these theories a method for explaining potential differences between our within- and between-subjects manipulations—a method that is unavailable to the edge-based theories.⁷

The second perspective in this experiment interprets options as a variation on the *skewness* manipulation introduced in the first experiment (Manipulation 3). This aspect consisted of the vertically-aligned shared options (green) and the context options that manipulated the average experienced outcome. In contrast with the first experiment where the shared option was located close to the midpoint of the distribution, the shared risky

⁷More specific predictions might be ascertained by examining the value manipulations from our second and third experiments. The average outcome was roughly 53 points in both conditions of our second experiment whereas in our third experiment the average was 41.3 points in the right-skewed condition and 61.5 points in the left-skewed condition. These averages are much closer to the non-extreme risky outcomes (40 and 60 points) than they were in our second experiment. Therefore, assuming that our choices are influenced by the *ratio* of the distances from the average outcome (e.g., utility-weighted sampling), reducing the distance from the non-extreme outcomes would entail even *stronger* preferences in the third experiment. Regarding edge-based theories, the extreme risky outcomes were separated from the closest edge by a single outcome in both experiments. In contrast, the non-extreme outcomes in our second experiment (40 and 60 points) were separated from the closest edge by three outcomes whereas those in our third experiment (also 40 and 60 points) were separated from the *opposite* edge by just two outcomes. Although these designs both entail a preference towards the high-value risky option and against the low-value risky option, the effect would arguably be *weaker* in our third experiment. Despite this, these predictions should probably be taken with a grain of salt because they are made based on specific instantiations of these broader classes. It is conceivable that using a different extremity function would alter—and even reverse—these predictions and the more useful distinction is with the edge-based theories that predict independence from the intermediate outcomes.

option in our third experiment resulted in both the best outcome (90 points) and the worst outcome (10 points). This ensured that the rank associated with these shared outcomes—and therefore the predictions of the ordinal and the edge-based theories—were identical across conditions. Nonetheless, the context options in the right-skewed condition (orange) were positioned further away from the better outcome of the shared risky option and the context options in the left-skewed condition (blue) were positioned further away from the worse outcome. This manipulated the average experienced outcome and ensured that the centre-based theories predicted a preference towards the shared risky option in the right-skewed condition and against this option in the left-skewed condition.

Therefore, viewing the options in our third experiment from these perspectives allowed us to simultaneously examine the impact of presenting options between-subjects and further disambiguate the centre- and edge-based theories. In Experiment 3b, we attempted to increase the strength of the skewness manipulation to provide a rebuttal to the objection that the difference between conditions in Experiment 3a was insufficient to expect an effect of continuous extreme outcomes. We achieved this by presenting the context options (i.e., the low-value and high-value options) twice as often and shifting their outcomes closer to the best or worst outcomes. The difference between the average outcome of the conditions was doubled from around 20 points in Experiment 3a to 40 points in Experiment 3b, which is equivalent to half the range of the experienced outcomes.

In Experiment 3a, participants encountered these options in a task that involved making a total of 240 choices across five blocks. Each block consisted of 24 decision trials, 16 catch trials, and 8 single option trials. In Experiment 3b the task involved making 300 choices and each block consisted of 36 decision trials, 16 catch trials, and 8 single option trials. The options selected on catch trials were not used to exclude data in this experiment. This was because, although the catch trial options had different expected values, risk-seeking and risk-averse participants may have genuinely preferred the lower value option based on its variance.

3.8.1 Results and discussion

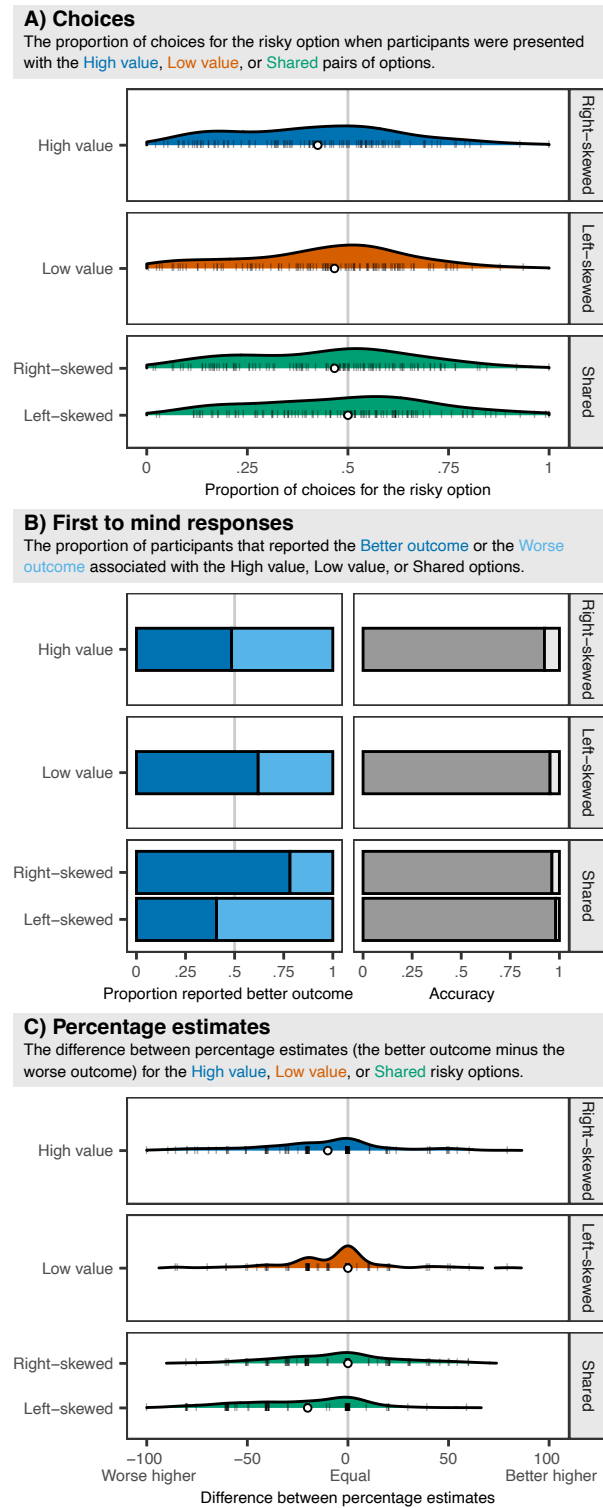
Value manipulation

The proportion of decision trials in which participants selected the risky option is displayed in Figure 3.6a. The previous experiments in this chapter provided us with strong evidence that participants were *more* likely to select the risky option when choosing between the high-value options than when choosing between the low-value options. Our third experiment examined whether a similar pattern can be observed when participants are presented with *either* low-value options (orange) *or* high-value options (blue). This modification generated a strikingly different configuration of preferences. Participants in this experiment, instead, selected the risky option *less* often for high-value options than low-value options.

We examined this pattern further using Bayesian logistic regression predicting the choice that participants made on each trial. We included the experiment (3a or 3b) and option value (low-value or high-value) as fixed predictors and allowed the intercept to vary for each participant. This model suggests that there is only a 34.72% probability that the effect of the between-subjects manipulation is even in the same direction as the within-subject manipulations from the previous experiments (Median = 0.94, 95% CI = [0.70, 1.27]).

A similar pattern can be observed in participants' responses to the first to mind questions (Figure 3.6b) and percentage estimates (Figure 3.6c). Regarding the first to mind questions, fewer participants reported the better outcome (dark blue) associated with the high-value option (top panel) than the low-value option (middle panel). We examined this further using Bayesian multinomial regression predicting the outcome reported in response to the first to mind questions. We included the experiment (3a or 3b) and option value (low-value or high-value) as fixed predictors. This model suggests that there is only a 6.83% probability that the effect was in the same direction as the within-subjects manipulations (Median = 0.65, 95% CI = [0.36, 1.16]).

Figure 3.6: **Choices and responses to the memory questions for Experiment 3.** The white dots represent the median response. The bottom (Shared) panels reproduce the responses to the high-value options in the right-skewed condition (top panel) and the low-value options in the left-skewed condition (middle panel), but emphasise that they were identical across conditions. Accuracy represents the proportion of participants' responses to the first to mind questions that corresponded to an experienced outcome.



Correspondingly, the average percentage estimate for the better outcome was lower for the high-value option (blue) than the low-value option (orange). We examined this pattern further using a linear regression model with the same predictors as the first to mind responses. This model suggests that there is only a 5.80% probability that the effect was in the same direction as the within-subjects manipulation (Median = -0.23, 95% CI = [-0.52, 0.06]). These responses produce substantial challenges for the edge-based theories whereas the ordinal and centre-based theories might be more well-equipped to explain the observed differences.

Skewness manipulation

Unlike the skewness manipulation in the first experiment, the bottom (Shared) panels in each subfigure of Figure 3.6 do not merely reproduce the results displayed in the other panels. This is because—viewed from the first perspective—these shared options were neither high-value nor low-value and were instead context options that ensured that none of the other outcomes was the best or worst. From the second perspective, locating the shared risky outcomes at the edges of the distribution ensured that the continuous theories predicted a preference towards the shared risky option in the right-skewed condition and against this option in the left-skewed condition whereas the ordinal and edge-based theories predicted similar responses across conditions.

Participants' choices involving the shared options (green) are displayed in the bottom panel of Figure 3.6a. The average proportion of risky choices was *lower* for the right-skewed condition than the left-skewed condition and this is evidently not consistent with the predictions of the continuous theories. We examined this pattern further using a Bayesian logistic regression model predicting participants' choices between the shared options. We included the experiment (3a or 3b) and condition (right-skewed or left-skewed) as fixed predictors and allowed the intercept to vary for each participant. This model suggests that there is only a 5.65% probability that the effect was even in the same direction that was predicted by the continuous theories—even the upper bound of the 95% credible interval

is scarcely compatible (Median = 0.78, 95% CI = [0.57, 1.04]).

This provides some evidence against the continuous theories and we might, therefore, expect that participants' preferences would also be reflected in their responses to the memory questions. Despite this, there was a considerable divergence between their choices and memory responses. Looking at participants' responses to the first to mind questions in the bottom panel of Figure 3.6b, their responses were, instead, broadly consistent with the continuous theories. We examined this further using a multinomial regression model with the same predictors as the choice model. This model suggests that there is a greater than 99.99% probability that participants were more likely to report the better outcome in the right-skewed compared with the left-skewed condition, a difference equivalent to being roughly 4.33 times more likely to report the better option in the right-skewed condition (95% CI = [2.34, 7.84]).

Participants' responses to the percentage estimate questions (displayed in the bottom panel of Figure 3.6c) were similarly consistent with the continuous theories. The average percentage estimate for the better outcome was higher in the right-skewed condition than the left-skewed condition. We examined this pattern using a Bayesian linear regression model with the same predictors as the choice model. This model suggests that there is a 99.97% probability that participants were more likely to report that the better outcome occurs more often in the right-skewed condition than the left-skewed condition. This difference is equivalent to roughly 0.48 standard deviations 95% CI = [0.22, 0.76]. Therefore, in contrast with participants' choices, their responses to the memory questions offer evidence that is compatible with the centre-based theories and is difficult to reconcile with the ordinal and edge-based theories.

To summarise the results of our third experiment, viewing the options from the first perspective provided evidence that the influence of extreme outcomes depends on whether options are presented within- or between-subjects. This was particularly troubling for the edge-based theories. Viewing them from the second perspective unveiled a clear divergence between participants' choices and memory. This presented us with a conundrum where

their choices seemed to provide evidence against the centre-based theories whereas their memory responses seemed to provide evidence in favour of these theories and against the ordinal and edge-based theories. We will examine potential resolutions to this conundrum in the following discussion.

3.9 Chapter discussion

Ludvig et al. (2014) observed that their participants consistently selected the options associated with the best outcome and avoided the options associated with the worst outcome and explained this pattern using a categorical extreme-outcome rule. Although this rule accurately captures their observations, it is not unique in this ability, and the same pattern is also compatible with participants choosing options associated with ordinal or continuous extreme outcomes. The experiments in this chapter, therefore, aimed to tease apart the predictions of these theories by introducing three novel experimental manipulations. In the following sections, we will employ these manipulations to demonstrate the inadequacy of the categorical theories and examine the viability of other plausible alternatives.

3.9.1 Value manipulation

Consistent with the ordinal and continuous extreme outcome theories, the participants in our first and second experiments were more likely to choose the risky options for the high-value options than the low-value options. Their responses to the memory questions were similarly compatible with the over-representation of either ordinal or continuous extreme outcomes. But what about the categorical theories? Given that none of these outcomes was the best or worst, are we able to reconcile these observations with their predictions? One possible approach is to amend the categorical theories with an auxiliary assumption regarding noisy encoding or chunking so that neighbouring outcomes are also considered extreme. This approach offered a reasonable explanation for a similar pattern of preferences when outcomes were separated from the edges by a single point (Ludvig

et al., 2018) and perhaps even when outcomes were drawn from a continuous distribution (Mason et al., 2020), but noisy encoding is unable to account for participants' choices in our experiments.

To demonstrate this, consider the design of the low-variance condition in our second experiment in which the outcomes of the low-value and high-value options differed from the best and worst outcomes by at least 20 points (see the second panel of Figure 3.3). To appreciate the magnitude of this difference, it can be restated in the following ways: 1) there was a difference equivalent to 20% of the overall range, 2) there was a difference of at least twice the distance between the better outcomes and worse outcomes associated with these options, and 3) there was a difference of only 5 points between these outcomes and the outcomes that were equidistant from the centre and the edge.

This amount of noise is simply not compatible with our experiments. For example, over 90% of our participants in the low-variance condition accurately reported an experienced outcome in response to the first to mind questions (see the right-hand column in Figure 3.5b). This creates a dilemma for the categorical theories where avoiding one horn inevitably results in being impaled by the other. They are compelled to propose an amount of noise that is simultaneously large enough to explain participants' biased choices and small enough to explain their accurate memory responses and—assuming this is even possible—doing so would fundamentally alter their categorical nature.

For that reason, we find ourselves searching for a viable alternative amongst the ordinal or continuous theories and our value manipulation also provides us with an initial bearing in this search. Specifically, our third experiment presented participants with *either* high-value options *or* low-value options and this modification led to substantial differences from our previous experiments where these options were presented within-subjects. Participants were no longer more likely to choose the risky option for the high-value options than the low-value options and this observation is problematic for the edge-based theories. These theories suggest that the evaluation of a particular option is completely blind to the presence or absence of other options unless they influence the range of experienced

outcomes. This is clearly contradicted by participants' responses in our third experiment when compared with those in our previous experiments.

The ordinal and centre-based theories are somewhat better equipped to explain these differences because their predictions are influenced by every single experienced outcome. Nonetheless, this does not necessitate that their predictions are consistent with our observations. The better high-value risky outcome and the worse low-value risky outcome were still further from the average outcome than their counterparts and their ranks were still closer to the edges. Therefore, at least some versions of these theories would predict that participants in our third experiment should have chosen the risky option more often for the high-value options than the low-value options. In contrast with the edge-based theories, however, these predictions depend on the specific extremity function (e.g., whether it is linear or non-linear) and our experiments did not conclusively rule out a smaller effect in the same direction as the previous experiments. Consequently, although it might be possible to reject specific instantiations (see Footnote [7](#)), it is much harder to justify a general claim regarding the ordinal and centre-based theories.

The implications of the value manipulations might, therefore, be summarised as providing evidence against three classes of theories with decreasing levels of certainty: Firstly, there was strong evidence against the categorical theories across two experiments that involved outcomes that were neither the best nor the worst. Secondly, there was considerable evidence against the edge-based theories because the effect seems to be modulated by the presence of intermediate outcomes. Thirdly, there was evidence against some specific instantiations of the ordinal and centre-based theories that were similarly inconsistent with our third experiment but there might be versions of these theories that could escape relatively unscathed, and consequently, the evidence regarding this third class should be considered merely suggestive.

3.9.2 Variance manipulation

The variance manipulation was designed to further differentiate between these ordinal and continuous theories. The continuous theories predicted that the effect described above would be stronger in the higher variance condition whereas the ordinal theories predicted that the effect would be similar across conditions. Although there was very little evidence for an effect of variance on the choices that participants made, the 95% credible interval contained values that were consistent with both the ordinal and continuous theories. Nonetheless, there was some evidence that variance influences memory for outcomes in a way that is consistent with the over-representation of continuous extreme outcomes.

One possible explanation for these ambiguous results is that the variance manipulation was insufficiently strong, which could easily be addressed by increasing the difference in the variance between the two conditions. This solution, however, might not be as straightforward as it seems. The variance manipulation was specifically designed to be as strong as possible without causing indifference between the safe and risky options or confusion between the outcomes associated with the risky options and the best or worst outcomes.

These constraints can be understood by comparing the design of our variance manipulation with a similar experiment conducted by Ludvig and colleagues (Experiment 2b, 2018). Although their goal was different, they also manipulated the variance of risky options in their experiment. In their high-variance condition, there were risky options that resulted in outcomes that were separated by 38% of the overall range of experienced outcomes and in their low-variance condition, there were risky options that resulted in outcomes separated by just 2% of the overall range. As would be predicted by the continuous-level theories, participants in the high variance condition were more likely to choose the high-value risky option than the low-value risky option but a similar pattern was not observed in the low-variance condition.

Given that outcomes in our low-variance condition were separated by 10% of the overall range, the inconclusive results in our second experiment might be interpreted as

reflecting an insufficiently strong manipulation. Whilst this might be the case, the problem with this conclusion is that the results of the experiment conducted by Ludvig and colleagues can also be interpreted as participants displaying indifference between the safe and risky options that were separated by a single point in their low-variance condition. Finding evidence of an effect of variance might, therefore, require a compromise between increasing the strength of the manipulation whilst mitigating the possibility that the difference between conditions results from attributes other than extremity, such as indifference or confusion.

3.9.3 Skewness manipulation

The skewness manipulations were designed to further tease apart the edge- and centre-based theories by keeping the distance from the edges constant while manipulating the intermediate outcomes to shift the average of the experienced distribution. They also provided some evidence regarding the ordinal-level theories because the rank-based extremity of the shared risky outcomes differed between conditions in the first experiment but remained constant in the third. There was some suggestive evidence from Experiment 1 that the skewness of the distribution influences choices in conformity with the continuous-level theories but the choices in Experiment 3 were not consistent with these theories. In contrast to this, there was strong evidence across both of these experiments that the distance between outcomes and the average of the distribution influenced participants' responses to the memory questions.

In previous experiments, such as those conducted by Madan et al. (2014), the choices that participants made and their outcome memory responses after completing the choice task were highly correlated. This was clearly not the case in our third experiment. As a result, our findings are not fully compatible with either the edge- or centre-based continuous-level theories or with the ordinal-level theories. On one hand, the edge-based and ordinal-level theories are able to explain the choices that participants made in the third experiment but are unable to explain why people over-reported the outcome further from the average

in the memory tasks. On the other hand, the centre-based theories struggle to explain the choices made in the third experiment but successfully predicted the results of the memory tasks.

So what are we to make of this dissociation between memory and choice? Similarly to the variance manipulation, one possible explanation is that the difference between conditions was strong enough to influence participants' memory but not their preferences. In favour of this explanation, there was greater variability in participants' choices than in their memory responses. In each of our experiments, there were some participants who almost exclusively chose the safe option and others who almost exclusively chose the risky option. The greater variability in their choices makes sense because—although choices are necessarily memory-dependent—they are influenced by numerous other idiosyncratic variables, such as risk-preferences.

Whilst a generic account of the relationship between memory and choice is not a very convincing explanation, there was at least one confounding variable in our third experiment that could further substantiate this account. The rank of the outcomes associated with the shared *risky* option were the same in both conditions. In contrast, the shared *safe* option was associated with the second best outcome (out of six outcomes) in the right-skewed condition and the second worst in the left-skewed condition. Therefore, to the extent to which there is evidence that our evaluation of outcomes is at least partially dependent on their rank within the experienced distribution (Parducci, 1968; Stewart, 2009), it seems plausible that this variable may have overwhelmed the effect of centre-based continuous extreme outcomes on choice.

This explanation might initially appear to rescue the centre-based theories from the seemingly contradictory evidence observed in Experiment 3 but it gives rise to another potential issue. The difference in the average outcome between conditions in Experiment 3b was equal to roughly half the range of experienced values and the difference between conditions in Experiment 1 was contrived to be as large as possible without leading to indifference or confusion between outcomes. While it is still possible that the manipulation

was insufficient to observe an effect of these extreme outcomes on participants' choices, this would place considerable limitations on the situations in which we should expect to observe the effect—perhaps even constraining it to situations in which there are non-overlapping low-value and high-value pairs of options.

3.9.4 Conclusion

The experiments presented in this chapter were designed to tease apart different extreme outcome theories by organising them according to their level of measurement and referent of extremity. The evidence was most compelling against the categorical theories, such as the extreme-outcome rule, but none of the other theories we examined was left entirely unscathed. As a result, further examination using contexts that deviate from those used in previous experiments will be required to evaluate the adequacy and generalisability of these theories—this was one of the primary aims of the experiments described in the next chapter.

Chapter 4

Types and tokens

The extreme-outcome effect observed by Ludvig and colleagues—that people choose the risky option more often for higher value options than lower value options—clearly demonstrates that the choices people make depend on the context in which they are made. In the previous chapter, we examined the relationship between value and the extreme-outcome effect as either categorical, ordinal, or continuous. The results of these experiments showed that the categorical conceptualisation of extremity was an inadequate explanation of the effect. Whilst there was also some evidence regarding the ordinal and continuous theories, this was less conclusive. In this chapter we continue to examine these levels of measurement but also investigate whether the frequency of occurrence of each outcome influences the extreme-outcome effect.

When representing a distribution of outcomes, it is possible to either only consider the value of outcomes or also include information about how often each outcome occurred. As an example of this, suppose that you were offered a role with a salary of \$75000 and were presented with the salaries of the ten existing employees. Seven of your potential colleagues received \$50000, one receives \$95000, and one receives \$100000. Considering that the position of your salary within the distribution affects satisfaction, it seems likely that you would compare the offered amount to those received by your peers (e.g., G. D.

Brown, Gardner, et al., 2008).

When doing so, whether you consider the frequency of each salary would play a large role on the conclusion reached from this comparison. On one hand, the only salary level lower than yours is \$50000 but there are two levels that are higher. Based on this frequency-independent assessment, you might conclude that your salary is at the lower end of the pay scale. On the other hand, your potential salary would be better than seven of your colleagues and only worse than two. If you include this frequency-based information, you would be more likely to conclude that the salary is on the higher end of the scale.

This example illustrates a distinction between *types* and *tokens* that applies to both continuous and ordinal interpretations of extremity. First proposed by Pierce (1906), types refer to distinct classes or categories and tokens refer to concrete instances of a type. Wetzel (2018) provides an illustrative example using the line “a rose is a rose is a rose” from the poem Sacred Emily by Gertrude Stein. This line includes three *types* of word (“is”, “a”, “rose”) but eight word instances or *tokens*; this same example includes six types of letter (“a” “e” “i” “o” “r” “s”) and 19 letter tokens.

This distinction is ubiquitous in natural language. As a more prosaic example, in order to comprehend a simple phrase such as “the two ladies were wearing the same dress”, it is essential that we recognise that “ladies” refers to two tokens and “dress” refers to a single type. The type-token distinction has been applied widely in linguistics (e.g., Richards, 1987; Templin, 1957) and philosophy (e.g., Fodor, 1974; Putnam, 1975; Quine, 1987). Although there is a sizeable literature discussing the nature of types and tokens (for a review, see Armstrong, 1989), for the purpose of this chapter, it will suffice that their distinguishing feature is that the number of tokens is sensitive to frequency of occurrence whereas the number of types is not sensitive to frequency.

In decision theory, the distinction between types and tokens is rarely explicitly discussed but is nonetheless inescapably present in the way theories represent the value of options. The multitude of theories that evaluate options based on some form of expected utility entail type-based representation because they represent value and probability as

separate dimensions, and therefore, the representation of value is independent of the information about their frequency of occurrence (e.g., Kahneman & Tversky, [1979]; von Neumann et al., [1944]). To give an example, when evaluating a gamble involving a coin toss, rather than recalling specific instances, these theories represent the outcome *types* associated with heads and tails (e.g., winning or losing a dollar) and weight these types by their associated probabilities.

This type-based representation allows these theories to explain the observation that people often act as if they are neglecting probabilities and making decisions entirely based on the value of the possible outcomes. As an illustration of this, participants in an experiment conducted by Rottenstreich and Hsee ([2001]) were asked to imagine a scenario in which they were able to pay money to avoid a painful electric shock. As such, in the described scenario, there were two outcome types: one in which the participant would receive a shock and another in which they would not receive a shock.

The median price that participants said they were willing to pay to avoid the shock with 100% certainty—thus removing the outcome type in which they would receive a shock—was around \$20. In contrast, when participants were told that they could pay to avoid the shock with either 1% or 99% probability—in other words, where both outcome types remained possible—the median amounts paid in these scenarios were \$7 and \$10, respectively. This clearly demonstrates that the outcome types that remained had a much larger effect on the price that participants were willing to pay than the probabilities associated with them.

Although historically type-based representation of value has dominated decision-theory, it is certainly not the only way to represent the value of outcomes. Several theories of decision-making have been proposed that assume that options are evaluated based on individual stored tokens (Dougherty et al., [1999]; Stewart et al., [2006]). For example, the decision by sampling model proposed by Stewart et al. ([2006]) renders type-based representation of outcomes unnecessary by sampling and comparing stored outcome tokens. This token-based approach has a number of advantages, including that it can parsimoniously

explain the characteristic shapes of psycho-economic functions—such as those described by prospect theory—based on the distribution of experienced outcomes and probabilities (Stewart et al., 2006; Stewart et al., 2015). It can also at least partially explain loss aversion (Walasek & Stewart, 2015, 2018), delay discounting (Stewart et al., 2015), and decoy effects (Noguchi & Stewart, 2014; Ronayne & Brown, 2017).

These token-based theories have the additional benefit that they are similar to token-based models of memory and categorisation. For example, the multiple trace theory of memory is able to explain an array of experimental observations using a model that suggests that tokens are encoded as separate memory traces that coexist in memory (Hintzman, 1976). Likewise, in categorisation research, Nosofsky (1988) compared type- and token-based versions of a categorisation model by changing the frequency of exemplars in a task that required people to categorise colours. In these experiments, the frequency-sensitive models provided a better explanation of participants' classification accuracy and typicality ratings, which has contributed to the development of numerous influential token-based models of categorisation (e.g., Kruschke, 1992; Nosofsky, 1986).

Given the success of the research programs that have formed based on the expected utility and sampling approaches, it seems that there is some empirical support for both type- and token-based representations of value. Therefore, instead of suggesting that people encode value using one *or* the other, another possibility is that people are able to use types *and* tokens. In support of this idea, there is evidence that people encode information about individual people (types) and instances in which they were encountered (tokens; Barsalou et al., 1998). Similarly, Brainerd and Reyna (1990) suggest that people encode events using multiple levels of abstraction and show a preference for the highest level that allows them to complete the required task. People also seem to use different strategies depending on the nature of the task, determining the frequency of outcomes using enumeration of instances, availability of memory traces, or direct retrieval of frequency-based information (N. R. Brown, 1997).

Applying the distinction between types and tokens to the experiments described in

the previous chapter, extreme-outcomes could be represented as either types or tokens. Despite this, given the design of our experiments, these different outcome representations remained confounded. This is because whenever an outcome was extreme based on types, it was also extreme based on tokens. As a concrete example, Experiment 1 manipulated the skewness of the experienced distribution so that the number of outcome types above the midpoint in the left-skewed condition was greater (60, 70, 80, 90, and 100 points) than below the midpoint (0 and 40 points). Each of these outcomes was experienced a similar number of times, and as such, the number of outcome tokens was manipulated along with the number of outcome types. Therefore, it remains unclear whether the extreme-outcome effect is driven by type- or token-based extremity or some combination of both forms of representation.

Some aspects of the experimental design used in our previous experiments may have influenced whether outcomes were represented as types or tokens. For example, the options in the choice task were represented by coloured squares that were identical each time they were presented, which reduces the ease of distinguishing between individual tokens, whereas the outcome types were given verbal labels (e.g., “10 points”), which increases the ease of representing outcomes as types (Waxman & Markow, 1995). Given the differences between the choice task and memory tasks, it is also possible that a token-based representation of value was used in one task and a type-based representation was used in the other. This might offer a potential explanation of the differences between choices and memory in the previous chapter.

In light of this, the experiments in this chapter aimed to distinguish between type-based and token-based theories of the extreme-outcome effect. In Experiment 4, we manipulate the skewness of the distribution based on tokens while holding the number of types constant. We swap these manipulated and controlled variables in Experiment 5 by manipulating the skewness of the distribution based on types while holding the number of tokens constant. Finally, these experiments also provide additional evidence regarding the ordinal and continuous levels of measurement that we examined in the previous chapter. In order to manipulate only types *or* tokens, the experiments in this chapter deviated fur-

ther from the designs used by Ludvig and colleagues, and therefore, allowed us to examine the generalisability of these theories regarding the influence of extreme outcomes.

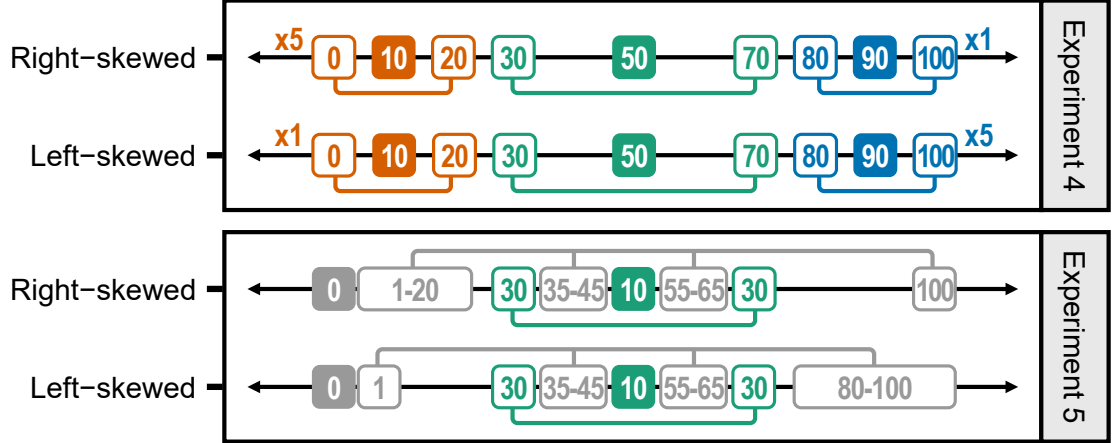


Figure 4.1: **Diagram of the designs of each experiment in Chapter 4.** Each square represents an outcome associated with an option. Rectangles that contain multiple numbers represent a discrete uniform distribution between these values. Outcomes (or sets of outcomes) connected with a $\sqcup$ were associated with the same risky option and (each set) occurred with equal probability. In the figures displayed in this chapter, we adopt a convention that the low-value options are depicted in orange, high-value options are depicted in blue, other options associated with a hypothesis are depicted in green, and other options (not associated with a hypothesis) are depicted in grey.

4.1 General method

4.1.1 Participants

Experiment 4 (skewed distributions - tokens) and 5 (skewed distributions - types) recruited participants using Amazon Mechanical Turk. A total of 102 participants signed up for Experiment 4 (Left-skewed: 51 participants, Right-skewed: 51 participants)¹. Seven participants did not complete the experiment (Left-skewed: 4; Right-skewed: 3) and three failed the colour-blindness test. 125 participants signed up for Experiment 5 (Left-skewed:

¹Unlike the experiments run in-person, allocating a precise number of participants to each condition was not feasible on Mechanical Turk, and therefore, the sample size varies slightly.

58 participants; Right-skewed: 67 participants). Five participants did not complete the task (Left-skewed: 4; Right-skewed: 1) and one failed the colour-blindness test. The average age was 33.4 years ($SD = 10.0$). 103 participants were female, 122 were male, and two were non-binary. Participants were paid US\$3.50 for participating and were able to earn an additional amount depending on their performance in the task ($M = US\$4.44$, $SD = US\$3.20$).

4.1.2 Design and procedure

Participants signed up for the experiment using Amazon Mechanical Turk and completed it on their own personal computers. At the beginning of the task, detailed instructions were presented on the screen including the conversion rate from points to dollars and that their objective within the task was to earn as many points as possible. Participants were required to correctly answer three multiple-choice questions about the instructions before continuing and completed a reCAPTCHA v2 challenge to prevent automated responding. They were also required to answer a question from the Ishihara colour-blindness test to mitigate the possibility that their data was included if they were unable to distinguish between the options. Following this, participants completed choice and memory tasks similar to those used in the previous chapter. These experiments were programmed in Javascript using jsPsych (de Leeuw, 2015) and the code used for each experiment is available at <https://github.com/joelholwerda>.

Choice task

In the choice task, participants repeatedly made choices involving safe and risky options that were represented by coloured squares. On most trials, participants were presented with two options but these choices were made in a context involving six options in Experiment 4 and four options in Experiment 5. Choices were either *decision trials* or *catch trials*—these trials were not explicitly differentiated for participants. Decision trials involved two options that had the same expected value. This allowed us to examine

participants' risk-preferences. Catch trials involved two options with different expected values and these trials allowed participants to earn points and acted as a measure of performance. In addition to these choices, on single-option trials, options were presented by themselves and participants had to click the option to continue. These options were included to mitigate the possibility of hot stove effects (Denrell & March, 2001).

Participants were not given explicit information about the distribution of outcomes associated with each option and were required to learn about options by selecting them and receiving feedback. No feedback was given for the option that was not chosen. There was no time limit for making choices and feedback was presented until the participant chose to begin the next trial—the next button was positioned in the horizontal centre of the screen to ensure that the cursor was equidistant from the options at the beginning of each trial. Choices were presented in a randomised order in five blocks and participants were encouraged to take a short break before starting the next block. The colours assigned to each option in each experiment and the side on which options were presented on each trial were also randomised to prevent the data being influenced by preferences for irrelevant attributes.

Participants were able to convert the points they earned in the choice task into real money following the experiment. Participants were given \$1 for every 10 points earned on a randomly selected choice. This method simplified the conversion from points to dollars and ensured that the difference between the safe and risky options influenced the amount of money earned.

Memory tasks

After completing the choice task, participants' memory for each outcome was assessed using two tasks. First, they were presented with the stimulus associated with each option and asked to report "the outcome that comes to mind first" when presented with each coloured square. The risky options were presented first in a randomised order and then the safe options were presented. Second, participants were presented with the coloured squares

and were asked “When you selected the option presented above, on what percentage of these choices did you experience each of the outcomes listed below?”. Every outcome was listed below the image in Experiment 4 but only the outcomes associated with the shared option were listed in Experiment 5. This modification was necessary because we used uniform distributions to manipulate type-based rank and participants experienced around 50 outcome types. To ensure that they understood that they were required to report percentages, their answers needed to sum to 100 before they could continue to the next section.

4.1.3 Analysis

Similarly to the experiments in the previous chapter, data was excluded from analysis for participants that selected the better option on less than 60% of choices where all possible outcomes of one option were better than the possible outcomes of the other. Participants were also excluded if they started the task but failed to complete every section of the experiment.

The posterior distributions were determined by Hamiltonian Monte Carlo using the `brms` package in R (Bürkner, 2017) and weakly regularising priors were selected for each parameter. These priors were selected using the Stan prior choice guidelines (Stan Development Team, 2020). Their plausibility was checked using prior predictive distributions. The prior for each analysis was identical to the ones used in the previous chapter with the addition that option value in Experiment 4 was modelled as a monotonic effect and a Dirichlet(3) distribution was used for the simplex parameter (Bürkner & Charpentier, 2020).

Once again, all parameters had bulk and tail effective sample sizes greater than 10000 and an $\hat{R} < 1.01$ suggesting adequate chain convergence (Vehtari et al., 2020). Rank histograms, posterior densities for each chain, and posterior predictive distributions were examined for each model. Sensitivity analyses were conducted to examine the robustness of our conclusions to our choice of exclusion criterion. The code for these diagnostic

analyses can be accessed at <https://github.com/joelholwerda>. Preregistered hypotheses for each experiment can be accessed at <https://osf.io/d8pq3>.

4.2 Experiment 4: Skewed distribution (token-based)

Similarly to the experiments described in the previous chapter, the options in the top panel of Figure 4.1 can be viewed from two different perspectives. From the first perspective, these options can be interpreted as another version of the *value* manipulation (Manipulation 1 in the overview of experiments). In our previous experiments, participants chose the risky option more often for the high-value pair than the low-value pair. This pattern is consistent with outcomes near the upper *and* lower edges being over-represented in memory but only one edge is necessary to explain this pattern with two options. If *only* the upper extreme outcomes were over-represented, the expected value of the high-value risky option would be over-estimated and it would be chosen more often *relative* to the low-value risky option—the converse applies to the lower extreme outcomes.

A single comparison between two options cannot disambiguate the contribution of the upper and lower extremes. Ludvig et al. (2014) addressed this issue by introducing a third risky option that was associated with *both* the upper and lower extreme outcomes. Assuming that the two extremes are equally influential, the influence of one would negate the other and the proportion of choices for this option would be half-way between the low-value and high-value risky options. In their experiment, Ludvig and colleagues observed that this proportion was intermediate to the other options but also observed that it was slightly closer to the proportion for the low-value options. This asymmetry suggests that the upper extreme might be more influential than the lower extreme.

This evidence is suggestive rather than conclusive but there are some theoretical reasons to believe that the influence of extreme outcomes might be asymmetrical. For example, Fredrickson (2000) argued that the most extreme outcome conveys the personal capacity necessary to endure an episode. Someone who can endure the worst outcome

can usually endure intermediate outcomes but this claim is less compelling regarding the best outcomes. In this manipulation, we aimed to further examine the possibility that the influence of extreme outcomes is asymmetrical by comparing low-value options (orange), medium-value options (green), and high-value options (blue). The medium-value options provide an alternative baseline to the both-extremes method used by Ludvig et al. (2014).

The second perspective interprets these same options as a token-based version of the *skewness* manipulation from our previous experiments (Manipulation 3). This aspect consists of the shared medium-value options (green) and the high- and low-value context options (blue and orange). Importantly, the outcome types (0, 10, 20, 30, 50, 70, 80, 90, and 100 points) were identical across conditions and the token-based skewness was manipulated by changing the number of times the context options were encountered (See Table 4.1).

In the left-skewed condition, the high-value context options were presented five times as often as the low-value context options. This shifted the token-based rank of the shared medium-value outcomes and the token-based mean so that the worse outcome of the shared risky option was more extreme than the better outcome. Conversely, in the right-skewed condition, the low-value context options were presented five times as often as the high-value context options. The better outcome was more extreme than the worse outcome on both token-based metrics. Consequently, the token-based theories predicted riskier choices for the shared options in the right-skewed condition whereas the type-based theories would be unable to explain this pattern.

Therefore, viewing the options in our fourth experiment from these perspectives allowed us to examine the contribution of the two extremes and disambiguate the type-based and token-based theories. Participants encountered these options in a task that involved making a total of 204 choices across five blocks. Each block consisted of 12 decision trials that involved the shared medium-value options. In the right-skewed condition, each block included 20 trials involving the high-value context options and four trials involving the low-value context options. These frequencies were swapped in the left-skewed condition.

Table 4.1: **The number of points associated with each option in Experiment 4.**

	Left-skewed		Right-skewed	
	Safe	Risky	Safe	Risky
Low value	10	0/20	10	0/20
Medium value	50	30/70	50	30/70
High value	90	80/100	90	80/100

Note:

Each of the outcomes associated with a risky option (separated by '/') occurred with equal probability. In the left-skewed condition, the high value options were presented five times as often as the low value options. In the right-skewed condition, the low value options were presented five times as often as the high value options.

We presented the participants with 24 catch trials after they had completed the other trials so we could assess their performance without sacrificing control over the number of context options that were chosen. The choice and memory data from seven participants was excluded because they selected the better option on less than 60% of catch trials. Five of these participants were in the Left-skewed condition and two were in the Right-skewed condition.²

4.2.1 Results and discussion

Value manipulation

The proportion of decision trials in which participants selected the risky option is displayed in Figure 4.2a. Recall that the low-value risky option (orange) resulted in the worst outcome but not the best outcome, the shared medium-value risky option (green) resulted

²Due to a coding error, the option chosen and feedback received was recorded but the foregone option was omitted from the data. This meant that we could not always determine whether the better option was chosen in the catch trials. Despite this, we were able to recover the proportion of choices for the better option with enough precision to establish whether the 60% inclusion criterion was met for all but four participants. This was possible because we were certain that the participant chose the better option if they chose the best possible option and were certain that they chose the worse option if they chose the worst possible option. Incidentally, this recovery process was the inspiration for our rational explanation in Chapter 2.

in neither the best nor worst outcomes, and the high-value risky option (blue) resulted in the best outcome but not the worse outcome. Therefore, the relative distance from the high-value and low-value options to the medium-value options allows us to estimate the relative weighting of the upper and lower extreme outcomes. These options appear to be roughly equidistant from the medium-value options in the right-skewed condition (top panel) but the low-value options are much more similar to the medium-value options in the left-skewed condition (middle panel).

We examined this pattern further using Bayesian hierarchical logistic regression predicting the choice that participants made on each decision trial. We included the condition (right-skewed or left-skewed), option value (high-value, medium-value, or low-value), and their interaction as fixed predictor variables. The intercept and the slope for option value were allowed to vary for each participant. Compared with the medium-value options, this model suggests that there was a greater than 99.9% probability that participants were *more* likely to select the risky option for the high-value options (Median = 1.73, 95% CI = [1.28, 2.37]) and *less* likely to select the risky option for the low-value options (Median = 0.67, 95% CI = [0.54, 0.83]). What about the difference between these distances? The model suggests that there is a 90.9% probability that the distance is larger for the high-value options than the low-value options (Median = 1.16, 95% CI = [0.94, 1.44]).

We expected that this pattern would be echoed in participants' responses to the first to mind task displayed in Figure 4.2b). The proportion of participants that reported the better outcome (dark blue) was highest for the high-value options and lowest for the low-value options in both skewness conditions. We examined this pattern further using Bayesian multinomial regression and the same predictors as the model of participants' choices. There was a greater than 99.9% probability that participants were *more* likely to report the better outcome for the high-value options (Median = 1.88, 95% CI = [1.28, 2.80]) and *less* likely to report the better outcome for the low-value options (Median = 0.48, 95% CI = [0.31, 0.73]). In contrast with their choices, however, there was only a 31.2% probability that this difference was larger for the high-value options (Median =

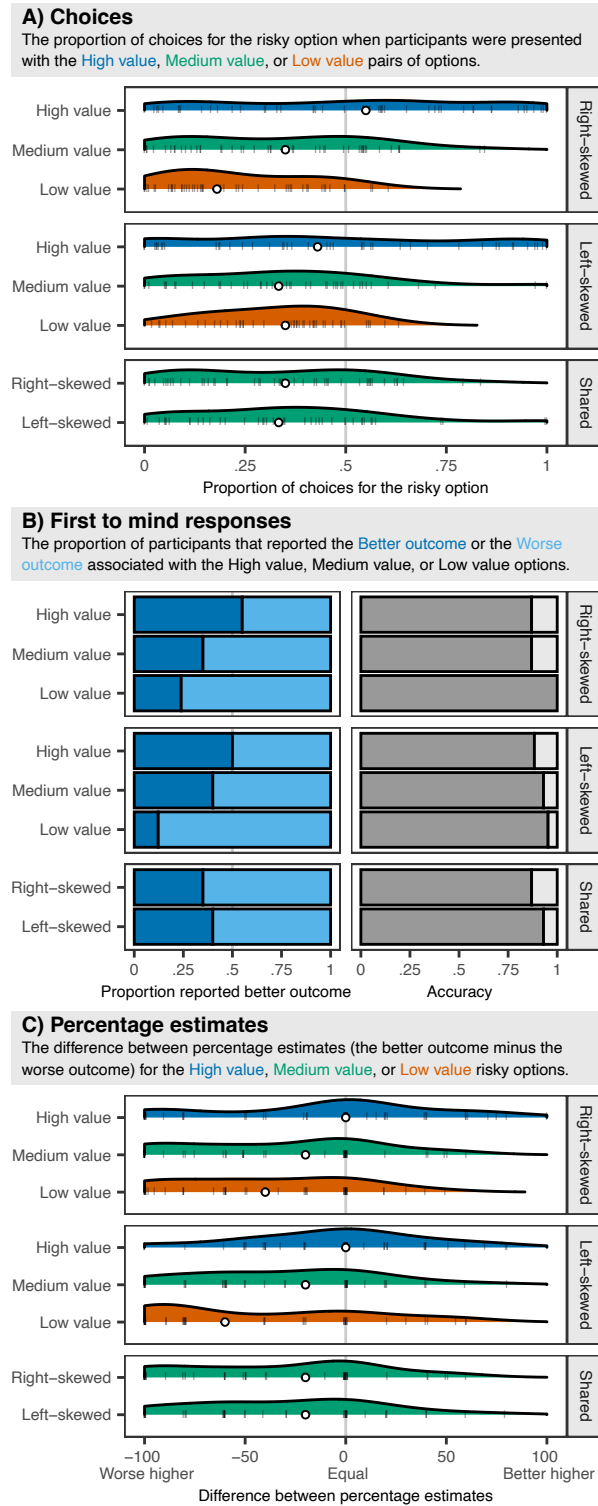


Figure 4.2: **Choices and responses to the memory questions for Experiment 4.** The white dots represent the median response. The bottom (Shared) panels reproduce the responses to the high-value options in the right-skewed condition (top panel) and the low-value options in the left-skewed condition (middle panel), but emphasise that they were identical across conditions. Accuracy represents the proportion of participants' responses to the first to mind questions that corresponded to an experienced outcome.

0.90, 95% CI = [0.59, 1.37]).

A similar pattern can be observed in participants' responses to the percentage estimate questions, which are displayed in Figure 4.2c. We examined them further using Bayesian linear regression and the same predictors as the previous models except that the varying slope was removed due to issues with model convergence. Compared with the medium-value options, there is a greater than 99.9% probability that people are *more* likely to report the better outcome as occurring with greater frequency for the high-value options (Median = 0.36, 95% CI = [0.21, 0.51]) and less likely for the low-value options (Median = -0.27, 95% CI = [-0.41, -0.12]). Similarly to their choices, there was an 88.7% probability that this difference was slightly larger for the high-value options (Median = 0.09, 95% CI = [-0.06, 0.22]).

Skewed distributions

Similarly to the previous chapter, the ordinal-level and centre-based extreme outcome theories suggested that participants were more likely to choose the shared medium-value risky option in the right-skewed condition. In contrast, however, we manipulated the frequency of outcomes whilst controlling type-based extremity so that only the token-based theories would predict an effect of skewness on choice. The bottom (Shared) panel in Figure 4.2a displays the proportion of risky choices for the shared medium-value option. This panel reproduces the responses to the high-value options in the right-skewed condition (top panel) and the low-value options in the left-skewed condition (middle panel), but emphasises their alignment within the second perspective.

Participants' choices in this experiment do not appear to have been influenced by the token-based skewness manipulation so we examined them further using Bayesian hierarchical logistic regression. We included the condition (right-skewed or left-skewed) as a fixed predictor variable and the intercept was allowed to vary for each participant. This model suggests that there is only a 27.94% probability that participants were more likely to select the risky option in the right-skewed condition than the left-skewed condition. The

median odds suggest that people were roughly 0.84 times as likely to select the risky option in the right-skewed condition than the left-skewed condition (95% CI = [0.47, 1.53]).

Participants' responses to the first to mind questions similarly do not appear to be consistent with an effect of token-based extreme-outcomes (see the bottom panels of Figure 4.2b). The better outcome was reported slightly more often in the left-skewed condition and we examined this pattern further using multinomial regression with the condition (right-skewed or left-skewed) as a fixed predictor. This model suggests that there is only a 32.77% probability that people are more likely to report the better outcome associated with the risky option as coming to mind first in the right-skewed compared with the left-skewed condition (Median = 0.83, 95% CI = [0.38, 1.94]).

This pattern was echoed in participants' percentage estimates displayed in Figure 4.2c. We examined these responses further using Bayesian linear regression with the condition (right-skewed or left-skewed) as a fixed predictor. This model suggests that there is only a 49.40% probability that participants were more likely to report that the better outcome occurred more often in the right-skewed condition (Median = -0.01, (95% CI = [-2.29, 2.29])). Once again, an effect in the direction predicted by the token-based theories was even *less* probable than an effect in the opposite direction. The credible intervals for the memory questions contain both values that are consistent and inconsistent with the frequency of occurrence influencing outcome memory. Nonetheless, interpreted as a whole, the results of this experiment do not provide evidence in favour of the token-based theories of extremity.

4.3 Experiment 5: Skewed distribution (type-based)

Our main goal in this experiment was to manipulate the type-based extremity of outcomes while controlling their token-based extremity. In order to make this manipulation as strong as possible and reduce the number of options that participants had to remember, the options in the bottom panel of Figure 4.1 focused exclusively on this manipulation

rather than offering two different perspectives. A single context option (grey) was used to manipulate the type-based extremity of the shared medium-value options (green). A fourth option that always resulted in the worst outcome (0 points) was paired with the context option. This option was only included to ensure that participants would select the context option on the majority of context trials (See Table 4.2).

There was an equal probability that the outcome of the context option would be above or below the outcomes of the shared medium-value options. Nonetheless, the context option included a combination of discrete uniform distributions and single outcomes that changed the type-based extremity of the shared medium-value outcomes. In the left-skewed condition, when the context option resulted in an outcome better than the shared medium-value options, the outcome was drawn from a discrete uniform distribution of numerous outcome types (80 to 100 points) whereas when the outcome was below the shared medium-value options, the result was always the same outcome type (1 point). This ensured that the worse outcome of the shared medium-value risky option was close to the edge of the distribution of outcome types and that the better outcome was near the centre.

In the right-skewed condition the uniform distribution and single outcome were switched. When the context option resulted in an outcome better than the shared medium-value options, the result was always the same outcome type (100 points) whereas when the outcome was below the shared medium-value options, the outcome was drawn from a discrete uniform distribution of numerous outcome types (1 to 20 points). This ensured that the better outcome of the shared medium-value risky option was close to the edge of the distribution of outcome types and that the worse outcome was near the centre.

To increase the difference in type-based rank between the better and worse outcomes of the shared medium-value risky option, the context option also resulted in outcomes that were located between the better and worse outcomes of the shared medium-value pair. These outcomes were drawn with equal probability from two discrete uniform distributions: one below the shared medium-value safe option (35 to 45 points) and one above the shared

4.3. EXPERIMENT 5: SKEWED DISTRIBUTION (TYPE-BASED)

medium-value safe option (55 to 65 points). The outcomes that were located between the better and worse outcomes of the shared medium-value risky option, the outcomes located above the better risky outcome, and those that were positioned below the worse outcome, occurred with equal probability.

This ensured that there was eventually a difference of 24 type-based ranks between the better and worse outcomes of the shared medium-value risky option in both conditions. This meant that in the right-skewed condition, the worse outcome of the shared medium-value risky option was the third worst outcome whereas the worse outcome in the right-skewed condition was 22nd worst out of 47 potential outcomes. Similarly, the better outcome of the shared medium-value risky option was the 20th best outcome of 47 possible outcomes in the right-skewed condition whereas it was the second best outcome in the right-skewed condition.

As trials progressed, the type-based rank differences approached those described above but in order to increase the speed at which that happened, twice as many choices between the context option and the worst option were presented than choices between the shared medium-value options. Additionally, the uniform distributions were rigged so that the same outcomes were not presented a second time until all other outcomes had been presented. As a result, for a participant that avoided the worst option on most trials, the difference between the rank of the outcomes associated with the shared medium-value risky option would be roughly seven type-based ranks after the first block and 15 after the second block.

Participants encountered these options in a task that involved making a total of 220 choices across five blocks. Each block consisted of 12 decision trials, 24 context trials, and 8 single-option trials. The catch trials were omitted from this task to ensure that the average outcome remained as similar as possible between conditions.

Table 4.2: The number of points associated with each option in Experiment 5.

	Left skewed		Right skewed	
	Safe	Risky	Safe	Risky
Shared	50	30/70	50	30/70
Context	0	1/35-45/55-65/80-100	0	1-20/35-45/55-65/100

Note:

Each of the outcomes associated with a risky option (separated by '/') occurred with equal probability. Outcomes separated by '-' were drawn from a uniform distribution where the two numbers are the min and max.

4.3.1 Results and discussion

In this experiment, the type-based continuous and ordinal theories of extremity predict that people should choose the riskier option more often for the higher value pair of options. The token-based mean and median were similar across conditions, and therefore, these theories did not predict a difference between conditions. Participants' choices in Figure 4.3a are not consistent with the type-based explanations for the extreme-outcome effect. On average participants chose the risky option *less* often in the right-skewed condition where there were fewer outcome types between the better risky outcome and the edge of the distribution than the worse risky outcome.

We further examined participants' choices between the shared options using Bayesian hierarchical logistic regression. We included the condition (right-skewed or left-skewed) as a fixed predictor variable and the intercept was allowed to vary for each participant. This model suggests that there is only a 9.81% probability that participants were more likely to select the risky option in the right-skewed condition than the left-skewed condition and the median odds of selecting the risky option is 1.41 times *lower* in the right-skewed condition (95% CI = [0.82, 2.44]).

Participants' responses to the memory questions were less conclusive (see Figure 4.3b and 4.3c). On one hand, more participants in the right-skewed condition than the left-skewed condition reported the better risky outcome as coming to mind first. We examined these responses further using multinomial regression with the condition (right-

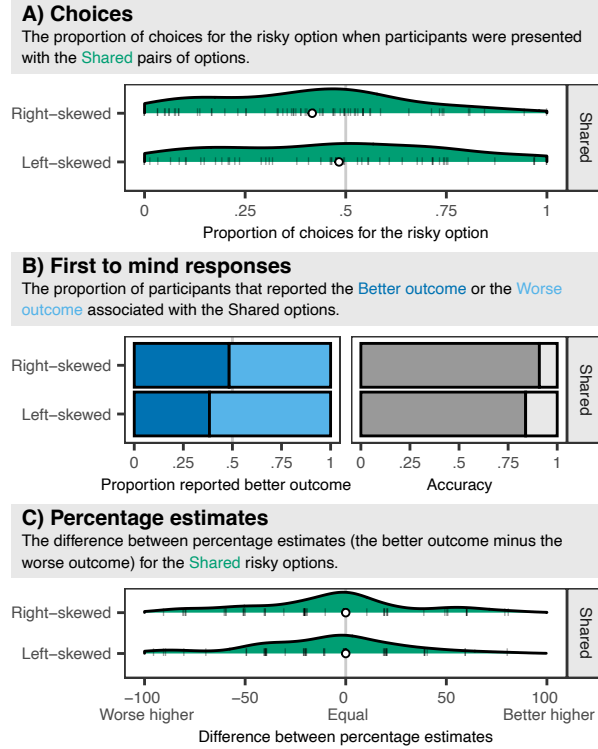


Figure 4.3: **Choices and responses to the memory questions for Experiment 5.** The white dots represent the median response. Accuracy represents the proportion of participants' responses to the first to mind questions that corresponded to an experienced outcome.

skewed or left-skewed) as a fixed predictor. This model suggests that there is an 85.52% probability that participants were more likely to report the better outcome in the right-skewed condition (Median = 1.48, 95% CI = [0.70, 2.93]). Even though they do not provide compelling evidence either way, these estimates are compatible with a modest effect of type-based extreme outcomes on memory.

On the other hand, participants' percentage estimates appear to be quite similar across conditions. We examined this further using Bayesian linear regression and the same predictor variable as the model of first to mind responses. According to this model, there is only a 65.10% probability that people are more likely to report that the better outcome occurs more often in the right-skewed condition than the left-skewed condition (Median = 0.07, 95% CI = [-0.27, 0.42]). The upper bound of the credible interval suggests that it is unlikely that there is a difference greater than 16% in the percentage estimates between the two conditions. Even if there is a difference, it is most likely fairly modest. Therefore, interpreted together, the results of this experiment are not consistent with the type-based theories of extremity.

4.4 Chapter discussion

The experiments discussed in this chapter were designed to follow up on some of the ambiguous findings from the previous chapter while also looking at the distinction between types and tokens. At the end of the previous chapter, we concluded that the categorical theories were inadequate to explain our results. Empirical challenges also arose for the ordinal and continuous theories in that shifting the centre of the distribution seemed to influence memory but there was conflicting evidence regarding its influence on choice. The challenges faced by the ordinal and continuous theories deepened further in the present chapter.

The influence of token-based extremity was examined in Experiment 4 by manipulating the frequency of outcomes while keeping the outcome types constant across conditions. This shifted both the token-based rank of the shared medium-value outcomes and their relative distance from the token-based average of the distribution. The token-based ordinal-level and continuous-level theories predicted that participants would select the shared risky option more often in the right-skewed condition than the left-skewed condition. This is not what we observed. Instead, for both choices and memory, this pattern was roughly half as likely as a pattern in the opposite direction to the predicted effect.

Similarly, the influence of type-based extremity was examined in Experiment 5 by manipulating the number of types above and below the shared options while keeping the distribution of tokens constant across conditions. This shifted the type-based rank of the shared medium-value outcomes and their relative distance from the average of the outcome types. As such, the type-based ordinal-level and centre-based theories predicted that the shared risky option would be chosen more often in the right-skewed than left-skewed conditions. Again this is not what we observed. The choices that participants made suggest that it is quite unlikely that type-based extremity had more than a small effect on choices in the predicted direction. In fact, the predicted effect was roughly ten times less likely than an effect in the opposite direction.

So what are we to make of these results? Neither types nor tokens generated the expected patterns, and therefore, it is difficult to draw strong conclusions about the role of these representations of extremity for an effect that we failed to observe in either experiment. In pursuit of exploring the implications of these results more fully, however, let us briefly consider a generous case for each of these representations. An obvious place to start—especially for the token-based manipulation in Experiment 4—is to emphasise that while the median estimate suggests an effect in the wrong direction, a non-trivial proportion of the posterior distribution was consistent with the predicted effect. As such, despite providing suggestive evidence, the jury is still out and we are simply not able to conclusively eliminate these representations of extremity.

One rationale for continuing to examine skewness in this chapter was that there was some evidence that manipulating type-based *and* token-based extremity influenced choices in the first experiment of the previous chapter. Specifically, the median odds suggested that participants were roughly 1.24 times more likely to select the shared risky option in the right-skewed condition compared to the left-skewed condition. Using this as a crude benchmark, although there was only a 2% probability for an effect of at least this magnitude for the type-based manipulation in Experiment 5, there was an 11% probability for the token-based manipulation in Experiment 4. Admittedly, the odds aren't great for either, but if you had to place a bet on one of them based solely on these experiments, you would be foolish to choose the type-based over the token-based representation of extremity.³

Before we completely rule out the type-based representation, however, it might be

³Using the median of the previous experiment as a benchmark may seem somewhat arbitrary but a similar conclusion would have been reached if we had instead looked at the overlap between the posterior estimate in Experiment 1 and posterior estimates for Experiment 4 (29%) or Experiment 5 (12%). Likewise, this conclusion would remain if we had looked at the proportion of the posterior distribution consistent with any effect in the predicted direction in Experiment 4 (30%) and Experiment 5 (10%). A more substantial critique is that this analysis assumes that the strength of the manipulation was similar across each experiment. This cannot be guaranteed and this analysis should be viewed as a rough approximation.

worth considering whether our manipulation even captures how people represent outcome types. In Experiment 5, we used uniform distributions to ensure there was a large difference between the number of types positioned above and below the shared pair of options (e.g., in the right-skewed condition, there were 21 types below and only one type above). On paper, this should have produced a stronger effect than the skewness manipulation in the first experiment where fewer types were introduced (e.g., in the right-skewed condition, there were four outcomes below and one above).

In reality, however, using discrete uniform distributions dramatically increased the number of outcome types that participants were presented with—they experienced up to 48 different outcomes. This is much larger than the number of outcomes that can easily be held in short-term memory (Miller, 1956). As such, it is possible that participants chunked each of the uniform distributions into a single chunked type (e.g., “outcomes between 80 to 100 points”) and this might have contributed to the observed results.

Now that we have considered the generous case for the type-based and token-based representations of extremity, we can move on to a more sceptical evaluation. While the evidence around types and tokens was inconclusive, this distinction only makes sense if there is an effect of rank or distance from the average in general. As we mentioned above, there was some suggestive evidence in favour of an effect of skewness in Experiment 1, but including the experiments in this chapter, we have manipulated skewness across four experiments and the median estimate was in the wrong direction for choices in three of those experiments. Although the evidence from each individual experiment was not conclusive, together they paint a picture where the rank of outcomes and their distance from the centre of the distribution might not have the predicted effect on choice.

These experiments also provided some further evidence regarding the connection between choice and memory. In Experiment 3 of the previous chapter, we observed that although participants’ memory reports were consistent with the extreme-outcome effect, their choices were not consistent. In the previous chapter, we offered the rank of the safe outcome as a potential explanation of this dissociation. Although certainly not conclusive,

there was some evidence of a similar effect for the first to mind responses in Experiment 5 but the rank of the safe outcome was no longer a confounding variable. This might give us reason to suspect the explanation that we provided in the previous chapter and further question the connection between memory and choice.

Finally, our results provide evidence regarding the influence of upper and lower extreme outcomes. There was a nontrivial probability that the distance between the high-value options and the medium-value options was larger than the low-value options for participants' choices. This is compatible with the observation by Ludvig et al. (2014) that the difference between participants' choices was slightly larger between the high-value options and the option that resulted in the best *and* worst outcomes. Whilst neither of these provides strong evidence, they both suggest that outcomes towards the upper extreme influence decisions from experience more than outcomes towards the lower extreme.

Nonetheless, there are some reasons that we might want to take these results with a grain of salt. Firstly, the difference between these options was observed in the left-skewed condition but there was little evidence of a difference in the right-skewed condition. Secondly, the evidence regarding participants' memory responses was less conclusive and the outcomes that they reported as coming to mind first provided some evidence in the opposite direction. Thirdly, there is some circumstantial evidence that the effect of extreme outcomes on memory is more consistent for lower extreme outcomes than upper extreme outcomes (Ganzach & Yaor, 2019; Hui et al., 2014; Kemp et al., 2008; Ludvig et al., 2018; Madan et al., 2014, 2017; Miron-Shatz, 2009; Rode et al., 2007). Fourthly, whilst encountering a lower extreme outcome that is beyond your capacity might propel you out of the gene pool, there is no equivalent consequence for upper extreme outcomes (Fredrickson, 2000).

Future experiments will be required to conclusively determine the relative weighting of these outcomes but there was one finding in these experiments that was unequivocal. There was strong evidence of a monotonic trend from the low-value to medium-value to high-value options in 1) the proportion of choices for the risky option, 2) the proportion

of participants that reported the better risky outcome, and 3) the percentage estimates for the better and worse outcomes. Regardless of whether extreme outcomes are weighted equally, the effect almost certainly influences both the upper *and* the lower ends of the distribution.

4.4.1 Conclusion

The evidence from these experiments progressed our understanding of the extreme-outcome effect in two seemingly opposite directions. As our doubts regarding the ability of ordinal and continuous theories to explain the results of the skewness manipulation grew, we became more sure that the effect influences both high-value and low-value options. This has been a consistent pattern throughout the experiments that we have presented in this section. Namely, there seemed to be strong evidence for the extreme-outcome effect whenever there were pairs of higher and lower value options in the same context, but whenever we deviated from this design, it was exceedingly difficult to determine the criteria required to observe the effect. The scope of the effect may be narrower than previously assumed and we might not be able to make broad statements about extremity beyond the narrow bounds of the manipulations that have been used in previous experiments.

Chapter 5

Temporal and distributional

In the previous chapters we examined numerous explanations for the influence of extreme outcomes. These theories defined extreme outcomes using categorical, ordinal, or continuous levels of measurement, identified them with reference to the centre, edges, or neighbouring outcomes, and employed either type-based or token-based representations. Despite their myriad differences, these theories shared one common attribute: they were based on the *distribution* of outcomes rather than *temporal* relationships between them. In the distributional theories, the best and worst outcomes, the edges of the distribution, and the rank of outcomes can all change as someone encounters new outcomes. Nonetheless, the order in which outcomes are experienced is disregarded when identifying outcomes as extreme.

The origin of the distributional approach can be traced to a common interpretation of the peak-end rule (Fredrickson, [2000](#); Fredrickson & Kahneman, [1993](#)). This rule describes how people evaluate *continuous* affective experiences, such as painful medical procedures or an exciting trip abroad. These experiences do not contain discrete outcomes that can be aggregated by simply calculating their average (Langer et al., [2005](#)). Instead, continuous experiences consist of an infinite number of moments and this gives rise to a more challenging computational task. How then do people evaluate these experiences?

One possibility is that we harness a small number of salient aspects that correlate with the overall quality of the experience (Ariely & Carmon, 2000).

Most theories suggest that these aspects are distributional components, such as the most intense moment, but this is not the only possibility. There are numerous other aspects that might offer an efficient method of aggregating the experience as a series of discrete points. An experience could be evaluated based on the intensity of *temporal peaks* where there is a change in the direction of a trend (e.g., things stop getting worse and start getting better). This aggregation method captures the most intense moment because the distributional maximum is always a temporal peak. It does this without needing to continuously monitor whether the current experience is more intense than the previous maximum and this arguably has computational advantages.

A second possible temporal explanation for the influence of extreme outcomes is that an outcome's subjective value is influenced by preceding outcomes due to adaptation, loss aversion, or extrapolation (Ariely, 1998; Hsee & Abelson, 1991; Hsee et al., 1991; Loewenstein & Prelec, 1993). This theory does not ascribe special importance to extreme outcomes or temporal peaks. Instead, it suggests that previous outcomes determine the reference point against which the current outcome is evaluated. An outcome is more likely to induce disappointment when the previous outcome was better and excitement when the previous outcome was worse. This could heighten the intensity of extreme outcomes because the best outcome is always preceded by a worse (or equal) outcome and the worst outcome is always preceded by a better (or equal) outcome.

As such, in this section, we have described two related temporal theories: the first based on salient temporal peaks and the second based on whether an outcome is better or worse than the previous outcome. In favour of these explanations, people prefer improving rather than deteriorating sequences (Loewenstein & Prelec, 1993), they avoid options that result in a higher number of peak losses (Langer et al., 2005), and satisfaction is influenced by differential partitioning, the spacing of outcomes, and whether events are perceived as a sequence (Ariely & Zauberman, 2000, 2003). The influence of extreme outcomes was

less pronounced when images were presented simultaneously rather than as a sequence (D. Thomas et al., 2018) and is not observed in decisions from description where temporal peaks are absent (Madan et al., 2017).

Even though our focus was on distributional attributes, our previous experiments also provided some evidence regarding temporal theories. Specifically, when options are presented in a random order, the probability of each outcome being a temporal peak and following a better or worse outcome is determined by the token-based rank of the outcome. When an outcome is close to the edge of the distribution, the preceding and following outcomes will seldom be one of the few outcomes closer to the edge and will usually be one of the many intermediate outcomes. As a consequence of this correlation, the temporal explanations for the influence of extreme outcomes rise and fall with their token-based ordinal-level counterparts.

5.1 Experiment 6: Temporal extremity

In the previous chapter, there was some ambiguity regarding the adequacy of the ordinal-level theories, and therefore, in this chapter, we aimed to directly examine the temporal theories by manipulating the order of outcomes. To do this, we used the same outcomes in both conditions so that the distributional forms of extremity were constant across conditions. Similarly to the experiments by Ludvig et al. (2014), we used a low-value pair that resulted in the worst outcome and high-value pair that resulted in the best outcome (see Table 5.2). The order in which these outcomes were presented was manipulated between two conditions.

The *random order* condition was similar to our previous experiments in that options were presented in a randomised order. This is depicted in the top half of Table 5.1 where the best outcome in the sequence (90 points) is always preceded by a worse outcome and is always a temporal peak because the direction of change reverses. The worst outcome (-40 points) holds an analogous position for losses—it is always preceded by a better outcome

and is always a temporal trough. On the other hand, the intermediate outcomes can form part of either an upward or downward trend and can be either a peak or trough.

For example, 60 points is a peak on the third trial because the outcomes before and after it are both worse and is part of an upward trend on the sixth trial because the outcome before it is worse and the outcome after it is better. The median outcome has an equal chance of following a better or worse outcome. As outcomes get closer to the edge of the distribution, there are fewer outcomes that are more extreme, and therefore, they are more likely to be a peak or trough.

In the *alternating order* condition, in contrast, trials cycled between one choice involving the high-value pair and one choice involving the low-value pair in an alternating fashion. In the bottom half of Table 5.1, every outcome associated with the high-value pair that results in gains is surrounded on either side by a loss, and therefore, is always a local peak. Similarly, every outcome of the low-value pair that results in losses is surrounded by gains and is a local trough. Because the outcomes of the pair do not overlap, this was always the case regardless of whether the participant chose the safe or risky option. As a result, extreme outcomes were no longer more likely to result in a temporal peak relative to non-extreme outcomes, and therefore, observing the extreme-outcome effect in the alternating condition would provide evidence against these temporal theories.

In order to possess full control of the order of outcomes we did not include catch trials for this experiment. As a result, choices were only made among pairs of options that each had the same expected value. All of the previous experiments included choices that mixed the low- and high-value pairs and excluding these choices allowed us to examine whether this is required to observe the extreme-outcome effect. Participants encountered these options in a task that involved making a total of 200 choices across five blocks. Each block consisted of 16 high-value decision trials, 16 low-value decision trials, and 8 single-option trials.

Table 5.1: **Example of sequences presented in a random order or alternating between gains and losses.**

Options presented in a random order								
Trial	1	2	3	4	5	6	7	8
Outcome	-10	-40	75	-25	-40	60	90	-25
Position	Start	Trough	Peak	Down	Trough	Up	Peak	End
Domain	Losses	Losses	Gains	Losses	Losses	Gains	Gains	Losses
Choice	Risky	Risky	Safe	Safe	Risky	Risky	Risky	Risky

Options presented as alternating gains and losses								
Trial	1	2	3	4	5	6	7	8
Outcome	75	-25	90	-25	60	-40	75	-10
Position	Start	Trough	Peak	Trough	Peak	Trough	Peak	End
Domain	Gains	Losses	Gains	Losses	Gains	Losses	Gains	Losses
Choice	Safe	Safe	Risky	Safe	Risky	Risky	Safe	Risky

Note:

When outcomes are presented in a random sequence, each outcome can be either a peak, trough, or an intermediate outcome in an upward or downward trend. When options alternate between gains and losses, all losses are troughs and all gains are peaks regardless of whether the safe or risky option is chosen.

Table 5.2: **The number of points associated with each option in Experiment 6.**

	Random		Alternating	
	Safe	Risky	Safe	Risky
Low-value	-25	-40/-10	-25	-40/-10
High-value	75	60/90	75	60/90

Note:

Each of the outcomes associated with a risky option (separated by '/') occurred with equal probability. In the random condition, the order of choices was randomised. In the alternating condition, choices between high-value and low-value options alternated in an ABAB fashion.

5.1.1 Method

Participants

Experiment 6 (temporal extremity) was conducted using 130 undergraduate psychology students enrolled at UNSW Sydney. These participants were randomly allocated into either the Random or Alternating condition with balanced sample sizes. The average age was 19.7 years ($SD = 3.3$). 411 participants were female and 188 were male. In addition to receiving course credit, participants were able to earn \$1 for every 1000 points that they earned in the choice task ($M = AU\$5.02$, $SD = AU\$2.18$).

Design and procedure

Similarly to the experiments in Chapter 3, participants in this experiment were given verbal instructions upon entering the laboratory that they could complete a computer-based task in which they could earn real money based on their choices. They completed this task in individual rooms where detailed instructions were presented on the screen. These instructions emphasised that their objective was to earn as many points as possible and explained how those points would be converted into dollars. Following this, they completed a version of the choice task and memory tasks introduced in Chapter 3, in which they encountered the outcomes in Table 5.2 in either a random or alternating sequence. The experiment was programmed in MATLAB using PsychToolbox (Brainard, 1997; Kleiner et al., 2007; Pelli, 1997) and the code is available at <https://github.com/joelholwerda>.

5.1.2 Analysis

The posterior distributions were determined by Hamiltonian Monte Carlo using the brms package in R (Bürkner, 2017) and weakly regularising priors were selected for each parameter. These priors were selected using the Stan prior choice guidelines (Stan Development Team, 2020). Their plausibility was checked using prior predictive distributions. The

prior for each analysis was identical to the ones used in the previous chapters except where specified below.

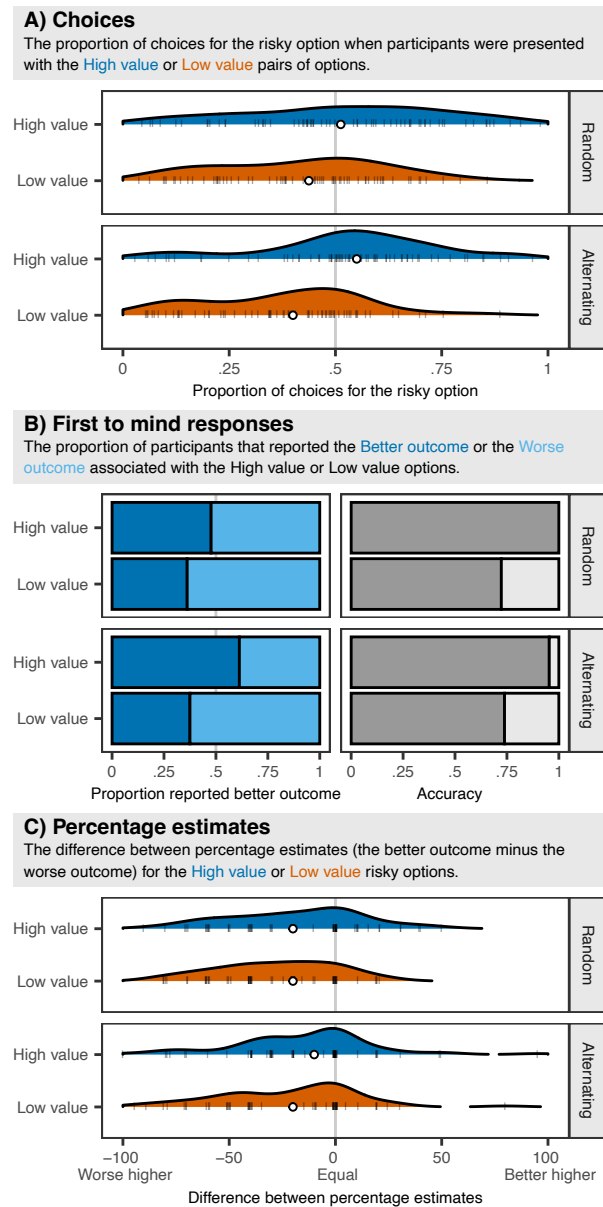
Once again, all parameters had bulk and tail effective sample sizes greater than 10000 and an $\hat{R} < 1.01$ suggesting adequate chain convergence (Vehtari et al., 2020). Rank histograms, posterior densities for each chain, and posterior predictive distributions were examined for each model and can be accessed at <https://github.com/joelholwerda>. Preregistered hypotheses for each experiment can be accessed at <https://osf.io/d8pq3>.

5.1.3 Results and discussion

The temporal extreme outcome theories predicted that the high-value risky option would be selected more often than the low-value risky option in the random condition but not in the alternating condition. This is not compatible with participants' choices in Figure 5.1a. In the alternating condition, there is a greater than 99.99% probability that participants were more likely to select the risky option for the higher value pair of options than the lower value pair (Median = 2.01, 95% CI = [1.55, 2.66]). Additionally, there is only a 13.12% probability that the extreme-outcome effect is stronger in the random condition where outcomes closer towards the edges of the distribution were more likely to be temporal peaks (Median = 0.90, 95% CI = [0.73, 1.07]). This provides fairly conclusive evidence that the explanations based on temporal peaks are unable to explain the influence of extreme outcomes.

A similar pattern was also observed in participants' responses to the memory tasks that are displayed in Figure 5.1b and Figure 5.1b. There is a probability of 99.43% that people are more likely to report the better outcome associated with the risky option as coming to mind first in the alternating condition. The median odds of reporting the better outcome for the alternating sequences is 2.83 times greater for the higher value options compared with the lower value options (95% CI = [1.27, 6.64]). Similarly to the choices, there is only a 21.22% probability that the effect is stronger in the random order condition (Median = 0.79, 95% CI = [0.44, 1.38]).

Figure 5.1: **Choices and responses to the memory questions for Experiment 6.** The white dots represent the median response. Accuracy represents the proportion of participants' responses to the first to mind questions that corresponded to an experienced outcome.



Finally, there is a probability of 90.10% that people in the alternating condition are more likely to report that the better outcome occurs more often when presented with the higher value options compared with the lower value options and the median response is 0.22 percent higher for the higher value options (95% CI = [-0.12, 0.52]). Although this evidence is not quite as strong as was provided by the choices and responses to the first to mind questions, the percentage estimates were quite similar across conditions and there is only a 62.77% probability that the effect is stronger in the random order condition (Median = 0.03, 95% CI = [-0.17, 0.24]).

Interpreting the choices and memory responses together, the results of this experiment provide strong evidence against the extreme-outcome theories that are based on temporal peaks or whether the previous outcome is better or worse. This conclusion is particularly compelling in combination with the results of Experiment 4 in the previous chapter. The results of this experiment were not consistent with the ordinal distributional theories, and therefore, also provide some evidence against the ordinal temporal theories. Although the results of Experiment 4 alone were not entirely conclusive, these complementary pieces of evidence allow us to reject these temporal accounts as candidate explanations of the extreme-outcome effect with some degree of confidence.

Having said that, these theories are not an exhaustive set of all possible temporal theories of the influence of extreme outcomes. An analogy can be established between these temporal theories and the levels of measurement that we examined in Chapter 3. The temporal theories examined in this chapter were analogous to the ordinal-level theories because the number of temporal peaks and the probability that the previous outcome was better or worse was correlated with the token-based rank of outcomes. Similarly to the way that distance from the centre or edge of the distribution could be measured ordinally or continuously, it is possible to conceive of a temporal theory that over-weights outcomes not only based on *whether* the previous outcome was better or worse, but the *amount* that it was better or worse.

In essence, this theory is a temporal counterpart to the continuous-level theories

based on the distance between neighbouring outcomes (e.g., Murdock, 1960; Neath et al., 2006). The difference between these theories is that the temporal version only compares the present outcome with the previous outcome, but when outcomes are presented in a random order, the temporal and distributional theories eventually end up predicting identical results. As such, the continuous-level temporal theories rise and fall with the continuous-level distributional theories that are based on similarity between outcomes.¹

There is some evidence in favour of these continuous interpretations, for example, Hsee and colleagues demonstrated that the rate of change from one outcome to the next influences evaluations of hypothetical scenarios that are described or presented graphically (Hsee & Abelson, 1991; Hsee et al., 1991). They are also able to account for choices and memory responses described in this chapter. Although the high-value options in the alternating condition were always presented following a worse outcome and the low-value options were always presented following a better outcome, the rate of change was greater for outcomes towards the edges. For example, if the previous outcome was -25 points, the difference between this and the intermediate high-value risky outcome (60 points) is 85 points whereas the change for the high-value risky outcome at the edge of the distribution (90 points) is 115 points. As a result, this continuous-level temporal theory might remain a viable explanation of the extreme-outcome effect.

As we mentioned in the introduction to this chapter, this experiment also provides some evidence regarding whether making choices between the high-value and low-value pairs of options is required to observe the extreme-outcome effect. Recent experiments have demonstrated that the effect is influenced by whether outcomes are considered as

¹Although the evidence regarding the distributional theories applies to the temporal theories, the applicability of evidence is asymmetrical. For example, the evidence provided in this chapter by manipulating the temporal ordering of outcomes does not rule out that we might observe a token-based distributional effect driven by non-temporal attributes. Another issue that might influence whether evidence from one variant is applicable to the other is that the theories of the extreme-outcome effect that rely on distance from neighbouring outcomes can produce quite different predictions depending on the function that governs the relationship between outcomes. For the accounts based on similarity, there is evidence in favour of local rather than global similarity, but there is little reason why the same function should govern theories based on adaptation, loss aversion, or extrapolation.

belonging to the same context (Madan et al., 2021). The previous experiments each provided participants with an opportunity to make choices between high-value and low-value outcomes and this might emphasise that both pairs should be considered as belonging to a single context. The experiment in this chapter only included choices within pairs but we still observed the extreme-outcome effect. This clearly demonstrates that choices between pairs are not required to observe the effect.

Consequently, as was the case in the previous chapter, our theoretical understanding of the extreme-outcome effect has made progress in two entirely different directions. Whilst we provided considerable evidence that the ordinal-level temporal theories are inadequate to explain the effect, we also provided evidence that increases our confidence that the effect is still observed regardless of whether people are able to make choices *between* high-value and low-value options. In Chapter 7, we will attempt to integrate the results of the four chapters in this thesis that examined extreme outcomes. Before doing so, however, we will shift direction to explore how people interpret uncertainty in decisions from experience.

Part II

Uncertainty

Chapter 6

Epistemic and aleatory uncertainty

Suppose that you were a contestant on a television quiz show and the host asks you the following multiple-choice question: “What was John Lennon’s middle name? A) Alfred or B) Winston”.¹ After you answer this question, the host pulls a coin out of their pocket and asks “What will be the outcome when I toss this coin? A) Heads or B) Tails”. Unless you happen to be an aficionado of Beatles’ trivia, these questions would both entail some degree of uncertainty but your perception of that uncertainty would likely differ considerably. Even if you had never even heard of John Lennon and exclaimed to the host that you might as well toss a coin, you would still interpret your uncertainty regarding the first question as arising from a *lack of knowledge*. The question is about a specific instance for which the truth is knowable, in principle. In contrast, when presented with the actual coin toss on the next round, your uncertainty would likely be interpreted as resulting from an *inherently stochastic process*.

This duality in uncertainty has been proposed by numerous philosophers, statisticians, and scientists. Poisson (1837) was the first to discuss these two forms of uncertainty

¹For those playing at home, the answer is Winston.

in print where he distinguished between “probability” and “chance”, a distinction that was further elaborated by Cournot (1843) using the terms “subjective probability” and “objective possibility”. Carnap (1945) distinguished between “probability₁” and “probability₂”, Russell (1948) between “credibility” and “probability”, and Savage (1954) between “personalistic” and “objectivistic” probability. Likewise, in psychology, Kahneman and Tversky (1982) differentiated “internal” and “external” uncertainty and similar distinctions have subsequently been proposed numerous times (e.g., Ascough et al., 2008; Bedford & Cooke, 2001; Fox & Ülkümen, 2011; Frisch & Baron, 1988; Gillies, 2000; Helton et al., 2010; Hoffman & Hammonds, 1994; Keren, 1991; Lawson, 1988; Peterson & Pitz, 1988; Robinson et al., 2006; Thunnissen, 2003; W. Walker et al., 2003). Although these forms of uncertainty have been described using almost as many distinct labels as the number of people that have written about them, each suggests a duality between uncertainty that arises due to insufficient knowledge and probability that subsists in things themselves independently of knowledge.

Perhaps the most detailed examination of these concepts was conducted by Hacking (1975) who described them as *epistemic uncertainty* (derived from the Greek word for “knowledge”) and *aleatory uncertainty* (derived from the Latin word for “dice”). Using the work of Pascal as a representative example, he traces the duality back to the earliest applications of probability theory in the seventeenth century. On one hand, the wager that Pascal (1670) made regarding the existence of god is epistemic in nature—God is, or He is not—was based on the inadequacy of our knowledge to rule out the existence of god. On the other hand, the famous correspondence between Pascal and Fermat that is widely credited as the origin of mathematical probability discussed interrupted games of chance that can be perceived as essentially aleatory in nature. This duality between epistemic and aleatory uncertainty traverses the history of probability and lies at the heart of the recently reignited debates between the subjectivist-Bayesian approaches to statistics (e.g., Jeffreys, Keynes, Savage, Ramsey, and de Finetti) and the objectivist-frequentist approaches (e.g., von Mises, Reichenbach, Kolmogorov, Neyman, Pearson, and Fisher).²

²Although Hacking suggests that the distinction between epistemic and aleatory uncertainty was implicitly evident in the writings of the early probablists, it is unlikely that

6.1 Uncertainty as a psychological concept

Hacking (1975) emphasised the curiously autonomous nature of epistemic and aleatory uncertainty so that “the same idea crops up everywhere, on the pens of people who have never heard of each other” (p. 16). And yet despite this, he also pointed out that the vast majority of people that employ probability seem oblivious to the distinction. This raises the question of whether epistemic and aleatory uncertainty are only discernible after someone has spent too long contemplating the foundations of probability or whether the distinction is reflected in our daily lives. Several authors have suggested that the way uncertainty is communicated in natural language might offer an answer to this question (Hacking, 1975; Kahneman & Tversky, 1982; Teigen, 1988). Ülkümen et al. (2016) examined a two-year corpus of New York Times articles and conducted laboratory experiments based on the conjecture that people tend to use confidence statements (“sure”, “certain”) to refer to epistemic uncertainty and likelihood statements (“chance”, “probability”) to refer to aleatory uncertainty. They observed that these statements differ across a wide range of attributes including whether they refer to the past or future, the propensity to quantify uncertainty numerically, and the perceived level of control (also see Olson & Budescu, 1997). Likewise, experiments have demonstrated that first person (epistemic) language was perceived as less informative, more indicative of the opinion of the speaker, and was used more often to describe high probabilities than third person (aleatory) language (Juanchich et al., 2017; Løhre & Teigen, 2016). Assuming that linguistic differences either shape or reflect people’s concept of uncertainty, these differences observed in natural language might indicate that the reach of epistemic and aleatory uncertainty extends far beyond the deliberations of the philosopher or statistician.

Several other lines of research have provided additional support that people intuitively distinguish between epistemic and aleatory uncertainty. For example, numerous

Pascal and his contemporaries recognised the distinction. As argued by Daston (1995), it is more likely that the associationist psychology that was prevalent during the European enlightenment obscured the distinction by implying a necessary connection between degrees of belief and relative frequencies.

experiments have demonstrated that adults and children distinguish between situations that involve inadequate knowledge about an event that has already occurred and situations that involve a stochastic process that will occur in the future (e.g., Beck et al., [2011]; Chua Chow & Sarin, [2002]; A. J. Harris et al., [2011]; Heath & Tversky, [1991]; Robinson et al., [2009]; Robinson et al., [2006]). People tend to make more extreme probability judgments when they interpret events as entailing more epistemic uncertainty and less aleatory uncertainty (Tannenbaum et al., [2017]). Investors are more willing to pay a financial advisor when they interpret their uncertainty regarding the stock market as resulting from a lack of knowledge (Walters et al., [2022]). Beliefs about the nature of uncertainty are associated with political ideology such that liberals tend to attribute higher aleatory uncertainty to outcomes regarding financial well-being (Krijnen et al., [2020]). Finally, brain imaging studies have demonstrated distinct activation patterns associated with tasks that suggest inadequate knowledge and those that involve stochastically determined outcomes (Volz et al., [2004], [2005]).

6.2 What is the nature of this duality?

In philosophy and statistics, aleatory and epistemic interpretations of probability have been understood as answering the question what *is* probability? Within the context of the apparent success of the deterministic laws of Newtonian mechanics, proponents of the duality have often faced a criticism, articulated by Boole ([1854]), that “a perfect acquaintance with all the circumstances affecting the occurrence of an event would change expectation into certainty, and leave neither room nor demand for a theory of probabilities” (p. 188). More recently, the emergence of quantum indeterminacy has revived the discussion regarding the ontological status of probability but the uncertainty that we encounter in our lives is generally macroscopic and the difference between insufficient knowledge and processes that are truly stochastic can be thought of as a “distinction without a difference” (Taleb, [2008], p. 319).

If the philosophical question has no bearing on our lives, another question arises

as to the origin of the distinction that has been observed in natural language. Fox and colleagues have recently suggested an approach that examines the cognitive differences that characterise aleatory and epistemic uncertainty (Fox & Ülkümen, 2011). This psychological approach no longer runs into the question of whether god plays dice but instead distinguishes aleatory uncertainty as attributed to outcomes that “for practical purposes cannot be predicted and are therefore treated as stochastic” (Fox & Ülkümen, 2011, p. 26). The philosophical question is transformed into a pragmatic choice about whether to treat outcomes as members of a class or whether it is worth attempting to determine their internal structure.³

This psychological approach has at least two important consequences: Firstly, it is no longer necessary to classify uncertainty as *either* epistemic *or* aleatory. Instead, uncertainty can be interpreted as consisting of *degrees* of each form of uncertainty—the weather tomorrow might be considered as more epistemic than the weather in two weeks, but these events would also be associated with a considerable degree of aleatory uncertainty. Secondly, the interpretation of uncertainty is subjective, and therefore, different people might experience different forms of uncertainty regarding the same event. Just as over-confidence can arise from blissful ignorance, someone might be convinced that a problem is easily conquered and continue to strive in vain. In contrast, someone might treat uncertainty as insurmountable when it would have been easily resolved. The psychological distinction between epistemic and aleatory uncertainty exists only within the mind of the person who doubts. Whilst resolving the challenge posed by determinism, this approach runs into a challenge of its own in the necessity to define “practical purposes” and why this should matter. But as we shall see, although the psychological approach makes the duality of uncertainty subjective, it does not make it entirely arbitrary and is constrained

³The psychological distinction between epistemic and aleatory uncertainty that was proposed by Fox and colleagues is closely related to the psychological concept of ambiguity when interpreted as partial knowledge. For example, Frisch and Baron define ambiguity as “the subjective experience of missing information relevant to a prediction”. Similarly, Camerer and Weber (1992) emphasise the pragmatic nature of the distinction, defining ambiguity as “uncertainty about probability, created by missing information that is relevant and could be known”.

by a number of factors, most notably, exploration and competition.

6.2.1 Exploration

Homo sapiens is a species of informavores, compelled by an epistemic hunger to improve our grasp on the world (Miller, 1983). We are constantly striving to reduce our uncertainty but the probabilistic representation of that uncertainty alone is not sufficient to dictate whether seeking information will prove beneficial. Analogous to the evolution of food-foraging strategies, the success of the informavore is determined by the amount of valuable information they acquire per unit cost. Pirolli and Card (1999) describes the objective of information foraging as follows:

If profitability of prey is defined as the energy returned per unit of handling time, then clearly less profitable prey should be ignored if they would prevent the predator from the opportunity to pursue a more profitable prey. For example, a predator that relentlessly pursued small hard-to-catch prey while large easy-to-catch prey were equally available would have a suboptimal diet. It has been noted in biology that predators will often ignore potential low-profitability prey in order to seek out higher-profitability prey (p. 11).

Acquiring information is rarely free, and therefore, distinguishing between epistemic situations in which seeking useful information will bear fruit and aleatory situations in which such attempts would prove futile is often beneficial. This suggests that our choices should be influenced by an evaluation of whether—for practical purposes—it seems possible to improve on our current understanding. Consistent with this, Walters et al. (2022) demonstrated that stock market investors who interpret uncertainty regarding the market as predominantly epistemic in nature are more likely to seek information by paying for financial advice whereas people who view uncertainty as aleatory are more likely to diversify their portfolio. Likewise, A. R. Walker et al. (2021) demonstrated that people increase the amount that they explore options after experiencing outcomes that are difficult

to integrate into their understanding of the outcome distribution, and therefore, suggest that their knowledge is incomplete.

Finally, in an experiment that demonstrated the effect of uncertainty on exploration, Goodnow (1955) presented participants with either a problem-solving task that involved matching geometric patterns (emphasising epistemic uncertainty) or a gambling task (emphasising aleatory uncertainty). Importantly, the problem-solving task was insoluble. Perfect accuracy could not be attained using the geometric patterns, but instead, the outcomes were determined stochastically using the same probabilities used in the gambling task (e.g., one pattern was correct 70% of the time, the other was correct 30% of the time). Therefore, the problem-solving and gambling tasks involved options that were identical with respect to their underlying outcome probabilities, but participants' behaviour differed depending on the task framing. Specifically, participants in the problem-solving condition were considerably more likely than those in the gambling condition to select the option that was associated with the lower probability of success. This observation might reflect participants sacrificing short-run accuracy in order to seek information that would improve their performance in the long-run.

6.2.2 Competition

The belief that it is possible to reduce our uncertainty, whilst providing a motivation to seek additional information, also implies the possibility that others are more knowledgeable. This is consequential because the outcomes of our choices are often influenced by the knowledge and choices of others—either through cooperation or competition. To provide an illustration, imagine that you are in the market to purchase a used car. Ideally, you want to find the best car for the lowest possible amount. The person selling the car, on the other hand, is seemingly hoping to sell you a lemon worth a small fraction of the price. The problem that you face is that you are unsure whether the car is one of the lemons whereas the seller has more experience with the car and might know something that you

are missing⁴

As such, the ability to estimate the degree to which additional information would reduce your uncertainty might help you avoid exploitation by a more knowledgeable adversary, either by attempting to improve your own knowledge or avoiding options associated with epistemic uncertainty. Consistent with this, people are more willing to bet on uncertain events—such as stock prices, football matches, and die rolls—that will be resolved in the (indeterminate) future compared with the same events that have already occurred in the (potentially knowable) past (Brun & Teigen, 1990; Heath & Tversky, 1991; Rothbart & Snyder, 1970). Likewise, Ellsberg (1961) demonstrated that people generally prefer to gamble on events where probabilities are explicitly presented rather than events where the probabilities are unknown to them (also see Becker & Brownson, 1964; Curley & Yates, 1989; Einhorn & Hogarth, 1985; Keren & Gerritsen, 1999; Sarin & Weber, 1993). Using an extended version of the Ellsberg task, Fox et al. (2021) showed that these preferences cannot be attributed, as is often suggested (e.g., Halevy, 2007; Segal, 1987), to an aversion to compound lotteries and instead reflects a preference for gambles that are associated with lower epistemic uncertainty.⁵

Providing further evidence for competition as a basis of the psychological distinction between epistemic and aleatory uncertainty, numerous experiments suggest that people make decisions based on the amount they know relative to the amount that they believe is knowable. Chua Chow and Sarin (2002) found that people are more averse to ambiguous

⁴This example was explored by George Akerlof (1970) in a paper for which he was eventually awarded the Nobel Prize in economics. He describes how asymmetrical information can lead to complete market failure. This can occur because the amount that the buyer is willing to pay takes into account their uncertainty regarding the quality of the car and the sellers who have reason to believe their product is not a lemon are unwilling to accept this lower price. Many of the solutions to this problem—such as signalling (Spence, 1973), screening (Stiglitz, 1975), and fair disclosure legislation—involve reducing epistemic uncertainty so that predominantly aleatory uncertainty remains.

⁵Ellsberg (1961) emphasises that ambiguity aversion cannot merely be attributed to a “mistake” that is corrected upon further reflection. He notes that even L. J. Savage, upon realising that his choice in the task violated his own axioms, decided to persist with the offending choice rather than following the axioms.

events when it is plausible that other people might possess more information than when that information is not available, for practical purposes, to anyone. Heath and Tversky (1991) showed that people avoid epistemic uncertainty when they perceive their own lack of knowledge or competence within a specific domain but that they seek epistemic uncertainty when they are confident in their knowledge (also see, Graham et al., 2009; Hadar et al., 2013). Likewise, people avoid betting on options when more knowledgeable individuals are evaluating the same bet and are more sensitive to their relative competence when playing a competitive game (Fox & Tversky, 1995; Fox & Weber, 2002). As such, the basis of the duality of uncertainty appears, at least in part, to arise from the necessity to gauge one's own knowledge relative to the knowledge of others.⁶

6.3 An example with two dice

Given the psychological approach proposed by Fox and colleagues, we are left with the possibility that people might disagree whether uncertainty is epistemic or aleatory. But are there some general features that influence the interpretation of uncertainty? A possible answer to this question might be derived from an examination of the elements that are logically required in order to reduce uncertainty. In essence, uncertainty arises from the inability to map outcome states onto observable states. Therefore, in order to always perfectly predict the outcome of some mechanism, there must be *at least* as much possible observable variability as there is outcome variability. As an example, imagine a game in which you attempt to guess the outcome of two unbiased six-sided dice that are rolled sequentially—one die and then the other—so that there are 36 possible outcomes (6 sides x 6 sides). Before the first die is rolled, it seems reasonable to believe that, *for practical purposes*, there is no observable variability that can be mapped onto those 36 outcome

⁶A third suggested influence on whether uncertainty is interpreted as epistemic or aleatory is the psychological consequences that result from self-evaluation or evaluation by others (Curley et al., 1986; Fox et al., 2021; K. A. Taylor, 1995). According to this account, attributing negative outcomes to chance rather than ignorance can mitigate the regret and blame that is experienced whereas attributing positive outcomes to one's knowledge or skill can allow credit for making the choice.

states. In this situation, most people would interpret their uncertainty as almost entirely aleatory.⁷

The dealer rolls the first die. Each of the six observable faces of this die could be mapped onto a different set of the 36 possible outcome states so that observing the first die roll would reduce your uncertainty so that only six outcome states remain plausible. Suppose that instead of openly rolling the first die on the table, the dealer rolls the die in secret and it remains hidden beneath a cup until you decide whether to place a bet. Mapping an observable state onto the set of outcome states certainly seems more plausible than before the first die was rolled. Indeed, both adults and children treat gambles differently before and after they have been resolved (Robinson et al., 2009; Robinson et al., 2006). But in this case, observing the state of the first die seems implausible within the constraints of the game. Thus, even if you are *aware* of possible variability that might be mapped onto the outcome states, your assessment of “practical purposes”, and therefore your interpretation of uncertainty, might also depend on your perceived ability to *acquire* this information.

To complicate the game further, suppose that instead of numbers on the six faces of the first die, there are six shapes that the dealer assures you will be converted into numbers after the second die is rolled. In this version of the game, you might observe the state of the first die and yet still remain uncertain about the outcome unless you are able to correctly map the outcome state to the variability in the observable state. As such, your definition of “practical purposes” will also depend on your perceived ability to *understand* this mapping. To the degree that you possess information without understanding the mapping, it remains possible that variability that you currently believe is related to the outcome states, in fact, provides no information. When this is the case, we will refer to the observable variability as *surface variability* in contrast with *structural variability* that

⁷Some epistemic uncertainty might enter the mix as demonstrated by the gambler’s fallacy (Jarvik, 1951) and the superstitious beliefs—wearing red, itchy hands, and lucky charms—that commonly arise among casino patrons. In these cases, people are mapping their uncertainty regarding the outcomes onto surface variability, a concept that will be discussed below.

can be successfully mapped onto the outcome states in order to reduce our uncertainty.

As such, an evaluation of “practical purposes” is based on your *awareness* of potentially observable variability, your perceived ability to *acquire* information regarding that variability, and your perceived ability to *understand* the mapping between the observable variability and the variability in the outcome states. In general, as your interpretation of uncertainty increases along each of these three dimensions, the tendency to interpret uncertainty as epistemic rather than aleatory should also increase. Having said that, this framework does not provide a comprehensive explanation and there are possible approaches to evaluating the solubility of uncertainty that might not include these three dimensions. Instead, in some situations, it might be reasonable to differentiate between epistemic and aleatory uncertainty based on an assessment of whether it has been possible to reduce uncertainty in previous similar situations or infer the solubility of uncertainty based on the existence of experts within a given domain. Nonetheless, this three dimensional framework might provide us with an approach to understand some of the differences in the interpretation of uncertainty across domains and individuals.⁸

⁸There are two additional observations that—whilst not central to the current investigation—are demonstrated by our game of two dice. Firstly, this example illustrates that uncertainty can be interpreted as both epistemic and aleatory with varying degrees. In this case, epistemic uncertainty was associated with the die that had already been rolled and aleatory uncertainty was associated with the die that would be rolled on the next round. Their combination arose due to a conjunction of separate events but it is equally possible to experience second-order (epistemic) uncertainty regarding (aleatory) propensities or dispositions that are associated with a single event. This latter possibility underlies the uncertainty in the Ellsberg urn task and is expressed in statements such as “the probability is 50% [$\pm 20\%$]”. Secondly, we suggested that epistemic uncertainty arises when another person could be more knowledgeable but your response to this possibility should also differ markedly depending on *who* possesses that knowledge. If the dealer is able to observe the die under the cup, you might still be willing to place a bet. It is unlikely that would still be the case if the dealer also showed the die to the person betting against you on the other side of the table.

6.4 Decisions from experience

The three dimensional account suggests that observable variability plays an essential role in estimating the degree to which uncertainty is perceived as epistemic or aleatory but most experiments that examine decisions based on experience involve options that are identical each time they are encountered (for a review, see Wulff et al., 2018). Adapting the words of George Loewenstein (2007), we treat options in our experiments “like the characters ‘thing one’ and ‘thing two’ in Dr Suess’ *Cat in the Hat*...and all other information that might make the situation familiar and provide a clue about how to behave is removed” (p.155). It is no mere accident that observable variability has been neglected, and instead, it has been intentionally expunged from experimental designs because the dominant theories of decision-making consider it to be irrelevant. These theories (e.g., Kahneman & Tversky, 1979) represent uncertainty as subjective probabilities over classes of events—as aleatory uncertainty—and therefore, there is no reason to differentiate between specific instances in which an option is experienced. There is no reason to suggest that a coloured square presented on a computer screen would differ from the options that we encounter in our daily lives.

This conclusion, however, falls apart when you consider even a mundane example such as buying apples from a supermarket. Similarly to the options presented in experiments, apples come in distinct varieties (Granny Smith, Royal Gala) and there is some degree of uncertainty regarding the value (the sweetness of the apple) that will result from selecting each of these options. On average, a Royal Gala apple will be sweeter than a Granny Smith but there is also considerable variability within each class. So far, so good. But in contrast with the experiments, there is also observable variability *within* each of the distinct classes. The presence of this variability suggests that it might be possible to map the variability in the apples’ sweetness onto their colour, shape, or size. You might attempt to acquire additional information by examining their aroma or firmness. Although there are some similarities with the options presented in decision-making experiments, the way that people interact with options in the real world and the reasons they select one

option over another might differ considerably.

Even in previous experiments in which options were visually identical each time they were presented, there is evidence that people, nonetheless, interpret uncertainty as being somewhat epistemic. The choices that participants make in these experiments often exhibit sequential dependencies that suggest they are not treating the outcomes as independent and identically distributed (for a review, see D. Cohen & Erev, 2021). They search for patterns even when there are none. This also seems to be the case in the few experiments that have used options that were not visually identical each time they were presented. As described above in the section on exploration, Goodnow (1955) observed that when each presentation of an option was individuated using unique geometric patterns, people were more likely to attempt to discern underlying patterns in the outcomes, as demonstrated by a tendency to split their choices between alternatives (probability matching) rather than always selecting the option with the greater probability of success (probability maximising).

In the three experiments described in this chapter, we examine whether introducing *surface variability* to a decisions from experience task—the bandit task—influences the degree to which people interpret uncertainty as epistemic or aleatory. Learning about options in this task requires participants to choose that option and experience the consequences, and as such, epistemic uncertainty might prove to be a double-edged sword. On the forward edge, it might allow people to increase their knowledge over time, but on the reverse edge, it might suggest that other people are more knowledgeable. This presents participants with a dilemma not encountered in previous experiments where they could learn about options without selecting them (e.g., by paying for advice in Walters et al., 2022) or where they were not able to learn from experience (e.g., when presented with a one-shot decision in Ellsberg, 1961). We examined which of the two horns of this dilemma people tend to select and whether their choices are markedly different from the standard bandit task that omits observable variability.

6.5 Experiment 7: Partial-feedback

The first experiment in this chapter aimed to examine whether introducing observable surface variability influences the way people interpret uncertainty in a bandit task. Similarly to most decisions from experience tasks, participants in the *identical images* condition were repeatedly presented with a pair of options that were represented by the same image each time the option was encountered. In contrast, in the *unique images* condition, participants were presented with pairs of options that were easily distinguished by their colour and location on the screen (similarly to the options in the identical images condition), but each time an option was encountered, it was represented by an image that was subtly differentiated from the other images used to represent the option (see Figure 6.1). This observable variability was not predictive of the experienced outcomes. It was surface variability rather than structural variability, and therefore, all participants had access to the same amount of useful information.⁹

Despite this, to the degree that participants in the unique images condition believed it was possible to understand the mapping between this surface variability and the variability in the outcomes, we predicted that they would interpret their uncertainty as being more epistemic. We also predicted that the surface variability manipulation would have an effect on the amount of aleatory uncertainty that participants reported. Assuming that the amount of *total uncertainty* remains constant across conditions, an increase in epistemic uncertainty would lead to a decrease in aleatory uncertainty. This might be further decreased if participants believed that they actually possessed an accurate mapping between the surface variability and the outcome variability. In this case, their total uncertainty would decrease (even though this would not be accompanied by an increase in

⁹We expected that introducing observable variability would influence participants' choices and were interested in the pathway mediated by their interpretation of uncertainty. Using surface variability was essential for examining this relationship because structural variability would allow the participants to improve their performance across trials. This might cause participants to preference riskier options based on the increased expected value rather than differences in their interpretation of uncertainty. Using surface variability thus allowed us to decouple the influence of performance and uncertainty.

their performance) and this might lead to a consequent further reduction in their aleatory uncertainty.

In both conditions, there was a *safe option* that always resulted in similar outcomes and a *risky option* that resulted in outcomes that were either better or worse—with equal probability—than the outcomes associated with the safe option. We expected that participants would experience relatively little epistemic and aleatory uncertainty associated with the safe option. In contrast, we expected that participants would experience a considerable amount of uncertainty associated with the risky option and that its interpretation would depend on the presence or absence of observable variability, as described above. Therefore, we suggest that differences in participants’ choices between conditions might reflect similar differences in their interpretation of uncertainty. Specifically, we expected the risky option to be associated with greater epistemic uncertainty in the unique images condition, and therefore, a stronger preference for the risky option in this condition might suggest epistemic uncertainty seeking and a stronger preference for the safe option might suggest epistemic uncertainty aversion.

The results of Experiment 7a suggested that participants might be responding with reference to their uncertainty regarding the appearance of the options rather than uncertainty regarding the outcomes. Experiment 7b aimed to minimise this possibility by emphasising that the questions referred to the outcome of a specific future choice. It also aimed to increase the likelihood that participants were correctly differentiating between the options by paying them based on their performance and making the outcome distributions more distinct. These and other relevant differences between Experiment 7a and 7b are discussed below.¹⁰

¹⁰In addition to examining the influence of surface variability on participants’ interpretation of uncertainty and their choices, we also examined two supplementary questions: one theoretical and the other methodological. The theoretical question aimed to ascertain whether there is a relationship between one-shot measures of epistemic uncertainty preferences (Ellsberg tasks) and responses to epistemic uncertainty when uncertainty might be resolved through repeated experience (bandit tasks). The methodological question examined whether it is possible to reduce within-condition variability by controlling for the specific sequence that participants experienced throughout the task. Given their periph-

6.5.1 Method

Participants

A total of 240 undergraduate psychology students from UNSW Sydney participated in Experiment 7 (120 each in Experiment 7a and 7b). The average age was 19.3 years ($SD = 2.4$) and 178 participants were female. In addition to receiving course credit, participants in Experiment 7b were able to earn a small amount of money depending on their performance in the task ($M = \text{AU\$}5.67$, $SD = \text{AU\$}1.48$).

Design and procedure

Bandit task. Participants completed the experiment in individual testing booths. At the beginning of the task, they were presented with written instructions that the task involved repeatedly making choices between pairs of options presented on a screen and that they should try to earn as many points as possible. They were not given information about the distributions of points; instead they were required to learn about options by receiving feedback about the number of points that resulted from selecting an option. Participants did not receive feedback on options unless they selected them, and therefore, being exposed to the consequences of choosing an option (in this case, gaining a specific number of points) was necessary in order to learn about it.

Each participant made 110 choices between a safe and a risky option that both had the same expected value. In Experiment 7a the outcomes of both options were drawn from a Gaussian distribution with a mean of 50 points. The standard deviation of the safe option was 1 point and the risky option was 20 points—this distribution was truncated so that all outcomes were two-digit numbers between 10 and 90. The outcome distribution for the risky option in Experiment 7a was centred on the same mean as the safe condition, and therefore, 50 points was the most likely outcome for both options. To accentuate the

eral relationship with the main research questions in this chapter, these investigations are described in Appendix [D](#) and [E](#), respectively.

different levels of risk associated with each of the options in Experiment 7b, a bimodal distribution that had peaks at 30 and 70 points was used for the risky option.

Participants were randomly allocated into either the identical images or unique images condition with balanced sample sizes. In the identical images condition, the safe and risky options were differentiated by their colour (red or blue) and by their position on the screen (left or right) and the images used to represent them remained identical across trials. In the unique image condition, the options were differentiated by colour and position but a slightly different image was used to represent these options on every trial. Importantly, the amount of observable variability was identical for the safe and risky options and although each image was unique, the outcomes were still drawn from the same distribution as the options in the constant image condition.

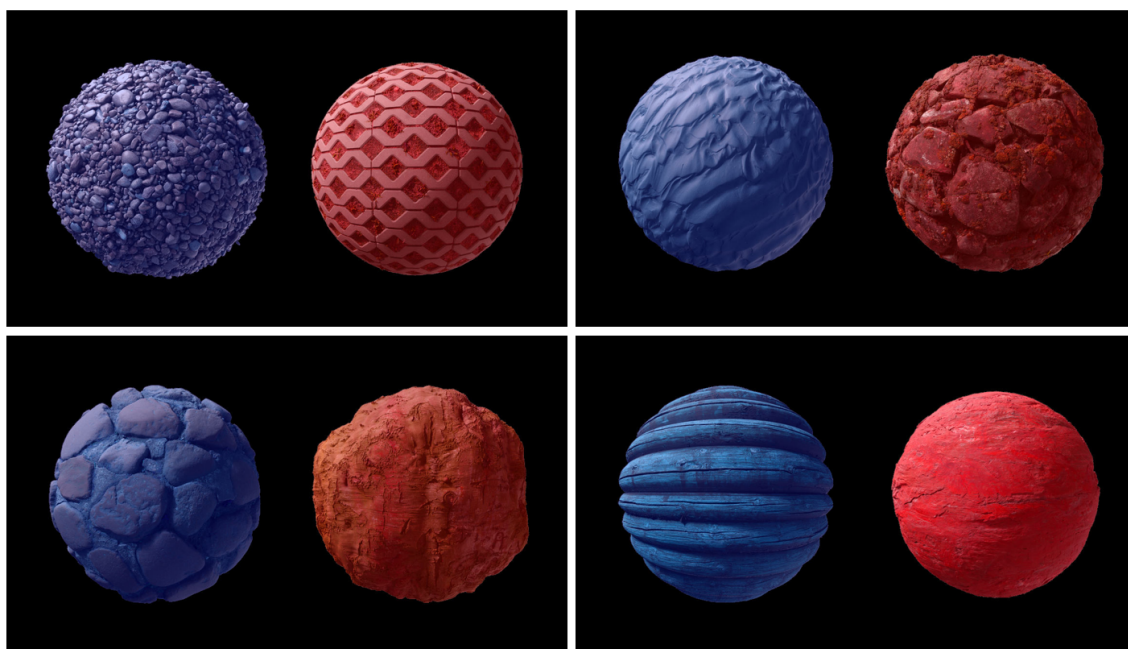


Figure 6.1: **Examples of the images used to represent options in Chapter 6.** Each quadrant depicts a choice between a pair of options as they were presented on the screen to participants. They were required to learn about each option by selecting it and observing the outcome.

Epistemic and Aleatory Rating Scale. Following this decision-making task, participants completed the ten-item Epistemic and Aleatory Rating Scale (EARS) that

was developed by Ülkümen et al. (2016) to examine the perceived amount of epistemic and aleatory uncertainty associated with an event—in this case, the outcome of selecting the safe or risky option. In Experiment 7a, participants were asked to think about the outcomes (numbers of points) that they received in the decision-making task when they selected a red or blue option.¹¹ They responded to items on a seven-point Likert-scale that either indicated an epistemic interpretation (e.g., “The outcomes were in principle knowable in advance”) or an aleatory interpretation (e.g., “The outcomes could play out in different ways on similar occasions”). The EARS items, instructions, and reliability estimates are available in Appendix B.

The EARS in Experiment 7a was phrased with reference to aggregated choices in the decision-making task, referring collectively to the outcomes of blue or red options instead of uncertainty regarding a specific future choice. Based on the results of this experiment, it was plausible that referring to aggregated past choices caused some participants to respond with reference to their uncertainty regarding the observable variability in the appearance of the options. Experiment 7b aimed to address this issue by presenting the EARS with reference to a future choice regarding a specific instance of each option. Participants were asked to imagine that they were going to select an option that was displayed on the screen and were presented with the EARS with reference to the “outcome (number of points)” that would result from that specific choice.

Analysis

We used Bayesian regression to analyse participants’ responses because it allowed us to flexibly model the hierarchical structure of our experimental tasks, incorporate regularisation, and examine degrees of credibility rather than dichotomous indicators of significance or non-significance. All posterior distributions reported in this chapter were determined by Hamiltonian Monte Carlo using the `brms` package in R (Bürkner, 2017). For each

¹¹The “safe” and “risky” options were never explicitly described to participants using those labels.

posterior, we report the probability that there was a difference between conditions in the predicted direction. This statistic provides information about whether our results can be attributed to sampling error but does not indicate the magnitude of the difference and cannot provide evidence for the absence of an effect (Makowski, Ben-Shachar, et al., 2019; Makowski, Ben-Shachar, et al., 2019). Consequently, we also report the highest density interval that contained 95% of the posterior (95% CI) and discuss whether these intervals are consistent with the predictions of each theory. We chose to present these summaries of the posterior distribution because they are analogous to frequentist statistics that might be more familiar to some readers. The probability of direction corresponds roughly to the complement of a one-tailed p-value ($1 - p$) and the highest density interval corresponds roughly to a confidence interval. They should not be treated as equivalent because our analysis incorporates informative priors (Nalborczyk et al., 2019), but they both attempt to answer similar questions.¹² The full posterior distributions can be accessed at <https://github.com/joelholwerda>.

Weakly regularising priors were selected for each parameter. A student-t(7, 0, 0.5) distribution was used for the slope, intercept, and threshold parameters. A half-student-t(7, 0, 0.5) distribution was used for the standard deviation parameters in the hierarchical models and an LKJ(4) distribution was used for the correlation between intercept and slope parameters (Lewandowski et al., 2009). These priors were selected to conform with the Stan prior choice recommendations (Stan Development Team, 2020)¹³ and predictions

¹²One possible source of confusion when interpreting these two posterior summaries is that the probability that there is an effect in the specified direction excludes probability from a single tail whereas the highest density interval excludes probability from both tails. They answer slightly different questions. This is not a issue but the reader should remember that the probability of an effect in the specific direction can be slightly greater than 95% when the 95% highest density interval includes the possibility of a small effect in the opposite direction.

¹³The Stan prior choice recommendations suggest using a student-t distribution with degrees of freedom between 3 and 7. We selected the latter to provide stronger protection against implausible parameter estimates whilst allowing us to learn from the data. To examine the sensitivity of our conclusions to our choice of priors, we also assessed our hypotheses using two alternate sets of priors: one more informative set that used normal(0, 0.5) distributions and one less informative set that used student-

from the prior distribution were inspected to ensure that they allow the range of plausible observations (Gabry et al., 2019). The same prior distributions were used for the subsequent experiments with a small number of additions that will be described alongside the relevant models. Numerical predictors were standardised (mean = 0, SD = 1) and categorical variables were deviation coded (-1, 1) so that they were centred and their scale was comparable when setting priors (A. Gelman, 2008).

All parameters had bulk and tail effective sample sizes greater than 10000 and an $\hat{R} < 1.01$ suggesting adequate chain convergence (Vehtari et al., 2020). There were no divergent transitions and the other Stan diagnostics did not indicate issues with estimation. Rank histograms, posterior predictive distributions, and other diagnostic plots were examined for each model and can be accessed at <https://github.com/joelholwerda>. Preregistered hypotheses for each experiment can be accessed at <https://osf.io/d8pq3>.

6.5.2 Results and discussion

Does surface variability influence the interpretation of uncertainty? Participants' responses to the EARS for the risky option are shown in Figure 6.2a. We hypothesised that participants who were presented with unique images for each choice would report higher epistemic and lower aleatory uncertainty associated with the risky option compared with participants who were always presented with the same images. To assess this hypothesis, we predicted responses to the EARS questionnaire for the risky option using hierarchical Bayesian ordinal (probit) regression. We included the experiment (7a or 7b), condition (unique images or identical images), option type (safe or risky), uncertainty type (epistemic or aleatory), and their interactions as fixed predictor variables. The correlated intercept and slope parameters for option type and uncertainty type were allowed to vary for participants within each condition. The correlated intercept and slope parameters for experiment, option type and condition were allowed to vary for each item

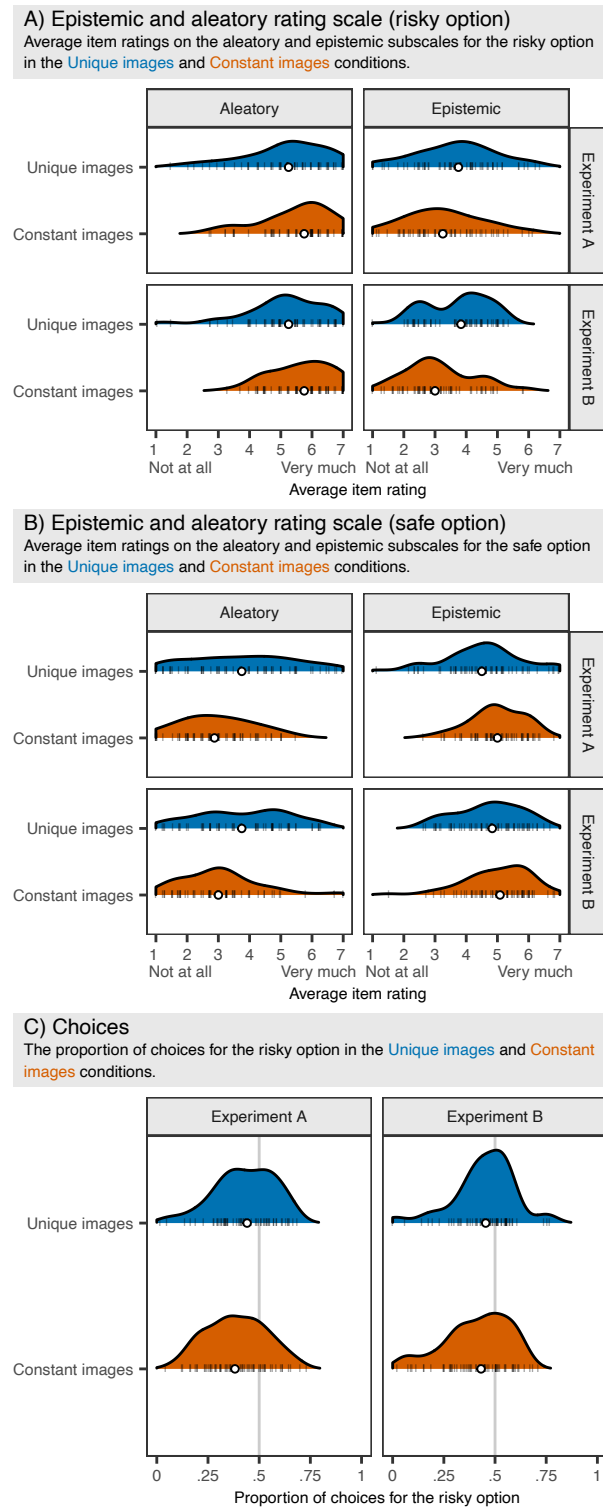
$t(3, 0, 1)$ distributions. The parameter estimates using these priors can be accessed at <https://github.com/joelholwerda>.

of the EARS within each uncertainty type.

Based on this model, there is a 99.32% probability that presenting participants with unique or identical images influenced the difference between their responses to the epistemic and aleatory items of the EARS (median = 0.43, 95% CI = [0.09, 0.77]). This is consistent with our first hypothesis that introducing surface variability influences the interpretation of uncertainty. We followed up this analysis by interrogating the posterior distribution using contrasts to examine the epistemic and aleatory items. These contrasts suggest that there is a 98.72% probability that participants in the unique images condition were more likely to report *higher* epistemic uncertainty than participants in the identical images condition and this corresponds to a difference of roughly 0.40 standard deviations on the latent scale (95% CI = [0.05, 0.75]). Conversely, there is a 97.22% probability that these participants were more likely to report *lower* aleatory uncertainty and this difference was around -0.46 standard deviations on the latent scale (95% CI = [-0.92, 0.03]). These observations are consistent with our hypothesis that people interpret uncertainty as more epistemic when risky options are associated with greater surface variability. Given that surface variability cannot be used to improve performance, we assumed that total uncertainty would be comparable between conditions, and therefore, the reduction in aleatory uncertainty could be explained on the basis of its complementary relationship with epistemic uncertainty.

EARS responses for the safe option. Participants' responses to the EARS for the safe option are shown in Figure 6.2b. Although our primary hypothesis concerned the risky option, we also analysed participants' responses to the EARS with reference to the safe option. The outcomes that resulted from selecting the safe option had a standard deviation that was 20 times smaller than the risky option, and therefore, we predicted that participants would experience considerably less uncertainty regarding this option. We did not have strong hypotheses regarding the interpretation of this uncertainty for the safe option but assumed that the proportion of epistemic and aleatory uncertainty might coincide with the risky option. In contrast, however, participants' responses appear to differ markedly between the safe and risky option. Using the same model that was applied

Figure 6.2: Responses to the EARS questionnaire and choices for Experiment 7. The white dots represent the median response. Average item ratings reflect the average response for each participant on the epistemic or the aleatory subscales.



to the risky option, there is a 99.39% that responses to the safe option showed the *opposite* pattern to the median estimate for the risky option (median = -0.46, 95% CI = [-0.77, -0.11]).

It is debatable how much emphasis should be placed on this observation. On one hand, it seems plausible that participants merely provided confused answers when asked a confusing question. We asked them to evaluate their uncertainty regarding outcomes for which they had minimal uncertainty and they may have responded to this absurdity with reference instead to the visual appearance of the options rather than their outcomes. We aimed to address this possibility in Experiment 7b by emphasising that the EARS items referred to a specific outcome of a concrete choice associated with a specific image but the results were nearly identical to Experiment 7a (the instructions for the EARS can be found in Appendix [B](#)). In the chapter discussion, we will propose a more substantial explanation for this pattern that is based on the relationship between observable variability and outcome variability.

Is there an effect on choices? Participants' choices in the bandit task are shown in Figure [6.2c](#). We hypothesised that epistemic interpretations might influence participants' choices, either by promoting exploration or avoidance of the risky option. Given that participants probably interpreted uncertainty as more epistemic in the unique images condition—based on their responses to the EARS—we were interested in whether there was a resultant influence of surface variability on their choices. We examined this using Bayesian hierarchical logistic regression predicting the choice participants' made on each trial. The experiment (7a or 7b), condition (unique images or identical images), and their interaction were included as fixed predictor variables. The intercept was allowed to vary for each participant. This model suggests that there is a 95.74% probability that participants in the unique images condition were more likely to select the risky option than participants in the identical images condition. This evidence is more consistent with participants responding to epistemic uncertainty with exploration than responding with avoidance and the median odds ratio suggests that participants in the unique images condition were roughly 1.18 times more likely to select the risky option (95% CI = [0.97, 1.41]). This

posterior distribution is consistent with a moderate effect of surface variability but also assigns non-trivial credibility to small differences that we would not consider meaningful. We will attempt to differentiate between these possibilities with greater precision in the subsequent experiments.

To further examine the role of uncertainty, we also examined the relationship between participants' epistemic uncertainty ratings and their choices. We used a similar model to the one described in the previous paragraph but replaced the condition predictor with participants' ratings on the epistemic items of the EARS. Based on this model, there is an 93.55% probability that reporting greater epistemic uncertainty is associated with choosing the risky option more often, but similarly to the previous model, our results are consistent with participants displaying a relatively modest preference (median odds-ratio = 1.07, 95% CI = [0.98, 1.17]). Even if the true parameter value is situated near the upper 95% credible interval, scoring one standard deviation higher on the epistemic items would be associated with being merely 1.2 times more likely to select the risky option. This suggests that, in the context of this experiment, epistemic uncertainty plays a minor role in determining the influence of observable variability on participants' choices.

To examine this relationship more directly, we conducted a mediation analysis using a Bayesian multivariate linear regression with the mediator variable (participants' average response to the epistemic EARS items) predicted by the condition (unique images or constant images) and the outcome variable (the proportion of choices for the risky option) predicted by the condition and mediator variable. We then used the `bayestestR` package (Makowski, Ben-Shachar, et al., 2019) to calculate the direct effect (the posterior distribution for the condition variable in the model predicting the outcome variable) and the indirect effect (the multiplication of the posterior distribution for the mediator variable in the model predicting the outcome variable and the condition variable in the model predicting the mediator variable). Based on this analysis, the median estimate of the direct effect was 0.016 (95% CI = [-0.003, 0.035] and the median estimate of the indirect effect was 0.002 (95% CI = [-0.001, 0.007]). Consistent with our conclusions based on the relationship between epistemic uncertainty and choice, the indirect effect suggests

that epistemic uncertainty does not play a large role in the relationship between surface variability and choice. The estimated proportion of this relationship that was mediated by epistemic uncertainty was roughly 12%. In the chapter discussion, we will propose an alternative pathway that appears to offer a better explanation of the relationship between surface variability and participants' choices. At this point, however, we will merely mention the possibility that participants' choices reflect an attempt to balance the positive and negative consequences associated with epistemic uncertainty. Determining the plausibility of this account would require us to tease apart the influence of exploration and epistemic uncertainty aversion and this was our primary objective in the next experiment.

6.6 Experiment 8: Information or reward

The first experiment in this chapter manipulated observable surface variability in a bandit task where participants could learn about options by selecting them and receiving feedback. The information acquired about an option was necessarily dependent on the experienced outcome and we suggested that this creates a potential dilemma. Interpreting uncertainty as epistemic suggests that you might improve your future performance by acquiring additional information but people are often averse to choosing options associated with epistemic uncertainty.

We observed some evidence that, on average, people resolve this dilemma by increasing exploration rather than avoiding epistemic uncertainty (but the effect size was small). In this experiment, we aimed to further examine the contribution of these two aspects of the dilemma by separating the ability to acquire information about an option from the outcomes that might be obtained. Based on a task designed by Tversky and Edwards (1966), each time a participant selected an option, they were presented with a subsequent choice between acquiring information *or* reward from the option (also see, Navarro et al., 2016; Rakow et al., 2010). They could choose to *observe* the outcome but not have the points added to their final score or *claim* the outcome but not find out the number of points until after the experiment. This task required all participants to complete the same

number of trials, thus eliminating a potential confound that is present in free sampling tasks that people might engage for longer in the unique images condition because the novel images make the task more interesting.

There was no longer a dilemma between approaching and avoiding options associated with epistemic uncertainty, but there was still a trade-off between exploring and exploiting their existing knowledge. Observing an outcome was inherently costly because it precluded the opportunity to claim the points associated with the outcome. Therefore, the number of times that participants chose to observe an option might reflect whether they believed that acquiring additional information would be beneficial. In other words, it might reflect the degree to which they perceived uncertainty as soluble or epistemic. As such, we hypothesised that participants would observe options more often when they rated them on the EARS as associated with greater epistemic uncertainty. We predicted that the risky option in the unique images condition would be associated with greater epistemic uncertainty, and therefore, that participants would choose to observe more outcomes in that condition, particularly for the risky option.

Similarly to one-shot experiments based on Ellsberg's (1961) urn task, claiming an outcome associated with an option did not provide the participant with information that would improve their performance. Consequently, to the degree to which people display an aversion to epistemic uncertainty, we expected that participants would avoid claiming options that were associated with greater epistemic uncertainty. Separating information and reward, therefore, allowed us to individually examine the two horns of the dilemma that participants, supposedly, faced in the first experiment in this chapter.

6.6.1 Method

Participants

137 undergraduate psychology students from UNSW Sydney participated in Experiment 8. The average age was 19.2 years ($SD = 2.3$) and 100 participants were female. In addition to

receiving course credit, participants were able to earn a small amount of money depending on their performance in the task ($M = \text{AU\$}4.59$, $SD = \text{AU\$}0.68$).

Design and procedure

Participants completed a decision-making task similar to the one used in the previous experiment in which they were repeatedly presented with safe and risky options represented by either identical or unique images. The images and outcome distributions were identical to Experiment 7b. The key difference was that instead of having the points both displayed on the screen and added to their total score, participants were required to decide whether they wanted to *observe* or *claim* the points associated with the option.

If they chose to observe, the number of points associated with the option was displayed on the screen but not added to their final score. If on the other hand, they chose to claim the outcome, the number of points was not displayed but was added to their final score and “points added to total” was presented on the screen. The distinction between observing and claiming was explained in detail prior to beginning the task and participants were told that they would make a total of 100 choices. They were also required to obtain a perfect score on a short multiple-choice questionnaire designed to ensure adequate knowledge of the task.

Participants were not given explicit information about the outcome distributions and the only way to learn about options was to choose to observe the outcome. Doing so required them to forego the opportunity of claiming the points associated with the option, and therefore, it only made sense to observe an outcome if they believed it would provide information that could be exploited in future choices. Following this decision-making task, participants were presented with the EARS that was used in Experiment 7b.

6.6.2 Results and discussion

Does surface variability influence the interpretation of uncertainty? Participants' responses to the EARS for the risky option are shown in Figure 6.3a. We again hypothesised that participants in the unique images condition would report higher epistemic uncertainty and lower aleatory uncertainty than the participants in the identical images condition. To address this hypothesis, we used the same ordinal regression model that was used in the previous experiment with the exception that—for obvious reasons—the experiment predictor (7a or 7b) was omitted. Based on this model, there is a 96.22% probability that presenting participants with unique or identical images influences the difference between participants' responses to the epistemic and aleatory items (median = 0.36, 95% CI = [-0.04, 0.74]). This is consistent with the results of the previous experiment and provides further evidence that surface variability can influence the interpretation of uncertainty.

Once again, we used contrasts to examine the influence of surface variability for the epistemic and aleatory items of the EARS. These contrasts suggested that there is a 78.52% probability that people report experiencing *more* epistemic uncertainty (median = 0.14, 95% CI = [-0.22, 0.50]) and a 97.14% probability that people report experiencing *less* aleatory uncertainty (median = -0.57, 95% CI = [-1.14, 0.02]) when considering options that are represented by unique images relative to options represented by identical images. Whilst both of these estimates are *compatible* with our hypotheses and the results of the previous experiment, there is a notable difference in the strength of the evidence. The first estimate provides some weak evidence that surface variability *increases* epistemic uncertainty but is also compatible with a modest *decrease* in epistemic uncertainty. In contrast, the second estimate provides considerable evidence that surface variability decreases aleatory uncertainty.

Given that our hypotheses primarily concerned epistemic uncertainty and its effects on exploration and choice, what are we to make of these findings? We hypothesised that surface variability would *indirectly* impact aleatory uncertainty. This was based on

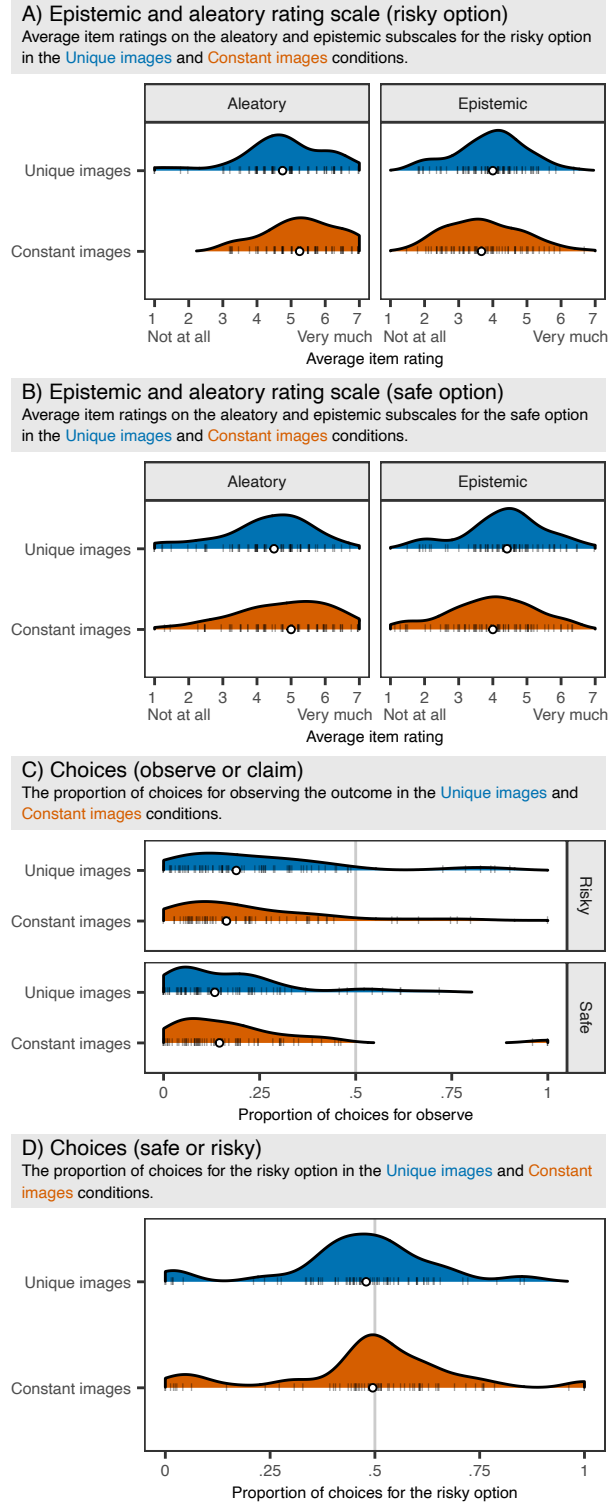


Figure 6.3: **Responses to the EARS questionnaire and choices for Experiment 8.** The white dots represent the median response. Average item ratings reflect the average response for each participant on the epistemic or the aleatory subscales.

the assumption that the amount of total uncertainty would be roughly equivalent across conditions because surface variability cannot be used to improve performance. Therefore, increasing epistemic uncertainty would indirectly cause a decrease in its complement, aleatory uncertainty. Our assumption that total uncertainty would remain constant would be erroneous, however, if participants believed that they possessed an accurate mapping between the surface variability and the outcome variability. This spurious association might decrease their total uncertainty (even though this would not be accompanied by an increase in performance) and this would lead to a consequent reduction in both their epistemic and aleatory uncertainty. The separation between observing and claiming outcomes in this experiment might have played a role in participants' illusory mapping between surface and outcome variability because most participants chose to observe a small number of outcomes, and therefore, were not presented with disconfirmatory evidence that would cause them to doubt their illusions. This could explain the potential differences with the results of the previous experiment.

EARS responses for the safe option. Participants' responses to the EARS for the safe option are shown in Figure 6.3b. In the first experiment in this chapter, we unexpectedly found that the EARS questionnaire referring to the safe option produced the opposite pattern of responses to the risky option. Based on our model, there is only a 2.77% probability that we observed this effect again in this experiment (median = 0.44, 95% CI = [-0.01, 0.85]). How can we explain this salient contrast between our experiments? As was the case with the risky option, one plausible explanation emphasises that the median number of observations for the safe option was ten times lower in this experiment than the number of observations in the previous experiment. The plausibility of this account will be discussed further in the chapter discussion where we will attempt to integrate the evidence regarding the safe option into a coherent explanation.

Is there an effect on exploration? The second main question we aimed to address was whether epistemic uncertainty would lead to higher exploration when there was no longer a necessary connection with the consequences of choosing an option. The proportion of trials in which participants chose to observe an outcome rather than having

its value added to their score is shown in Figure 6.3c. We predicted that participants would choose to observe more outcomes in the unique images condition, especially for the risky option, which we predicted would be associated with greater epistemic uncertainty. We examined this hypothesis using Bayesian hierarchical logistic regression predicting whether participants observed or claimed the outcome on each trial. The condition (unique images or identical images), the outcome variance associated with the option type (safe or risky), and their interaction were included as fixed predictor variables. The intercept and slope for the outcome type was allowed to vary for each participant. This model suggests that there is a 75.11% probability that participants were more likely to observe outcomes in the unique images (median odds-ratio = 1.15, (95% CI = [0.75, 1.73])). Regarding our prediction that the differences in the amount of exploration would be particularly salient for the risky option, the model suggests that there is a 76.34% probability associated with this interaction (median odds-ratio = 1.07, (95% CI = [0.88, 1.30])). These estimates do not provide meaningful evidence in favour of the hypothesised effect of surface variability on information seeking, which is not entirely surprising given that there was considerable uncertainty regarding the effect of surface variability on the epistemic items of the EARS.

We also directly modelled the relationship between responses to the epistemic items and whether participants observed or claimed the outcome on each trial. We used the same model as described in the previous paragraph but replaced the condition predictor with the participants' ratings on the epistemic items of the EARS. According to this model, there is only a 33.39% probability that epistemic uncertainty is associated with choosing to observe more outcomes (median odds-ratio = 0.96, (95% CI = [0.76, 1.17])). This is not consistent with our hypothesis that epistemic uncertainty would lead participants to observe more outcomes. It is possible that we failed to observe an effect of epistemic uncertainty on exploration because participants observed such a small number of outcomes. A single additional observation would be equivalent to a 17% increase for the median participant and this might suggest that our experiment lacked adequate precision. We attempt to address this potential issue in the next experiment.¹⁴

¹⁴Before selecting the information or reward task used in this experiment, we piloted another task in which there was an explicit cost associated with acquiring observations.

Is there an effect on choice? The option (safe or risky) that participants chose when they subsequently claimed the points associated with the option are shown in Figure 6.3d. Given that participants would not learn the outcome of these choices, and therefore, there was no possibility to reduce their uncertainty, we hypothesised that they would avoid options associated with greater epistemic uncertainty. We predicted that this would result in fewer choices for the risky option in the unique images condition. We examined this using Bayesian hierarchical logistic regression predicting the choice participants' made each time they chose to claim the outcome. The condition (unique images or identical images) was included as a fixed predictor variable and the intercept was allowed to vary for each participant. This model suggests that there is a 78.01% probability that participants in the unique images condition were more likely to avoid the risky option (median odds-ratio = 0.85, (95% CI = [0.56, 1.28])). This does not provide strong evidence for an effect of surface variability on choice, and similarly to the effect on exploration, this is not entirely surprising given our uncertainty regarding the effect of surface variability on the epistemic items of the EARS.

Finally, we modelled the relationship between participants' responses to the epistemic items for the risky option and whether they chose to claim the safe or risky option. We used the same model that was described in the preceding paragraph but replaced the condition predictor with the average response to the epistemic items for the risky option. Similarly to our hypothesis regarding surface variability, there was no real evidence that epistemic uncertainty influences information seeking. Instead, our model suggests that there is only a 39.50% probability that participants in the unique images condition were more likely to avoid claiming the risky option (median odds-ratio = 1.01, (95% CI = [0.94, 1.08])). Similarly to the previous experiment, these findings do not provide evidence that

Participants could pay a small amount (less than one cent) to observe outcomes from each option before making a consequential choice. Most participants chose to observe less than two outcomes and some even opted to make a choice without observing a single outcome. Based on this pilot, we concluded that participants were more willing to observe in the "information or reward" task used in this chapter in which the cost of observing was the lost opportunity to claim than in the task where there were (very small) direct monetary costs.

epistemic uncertainty was associated with information seeking in our tasks, even when exploration was decoupled from choice. We will discuss these observations further in the chapter discussion after we have described our final attempt to investigate the relationship between observable variability, epistemic uncertainty, information seeking, and choice.

6.7 Experiment 9: Free sampling

In the first experiment of this chapter, exploration was necessarily yoked with choice and entailed the possibility of exploitation by others that were more knowledgeable. In the second experiment, exploration was decoupled from choice but precluded the ability of participants to exploit their own knowledge by claiming the outcome. In this final experiment, we removed both of these trade-offs by allowing participants to freely observe outcomes until they decided that they were ready to make a single consequential choice. We hoped that this would mitigate the possibility that participants were responding to these trade-offs in a way that was masking their preference to explore options associated with greater epistemic uncertainty. Specifically, participants in Experiment 8 chose to claim on most trials and observed a relatively small number of outcomes. As a consequence, it is quite plausible that using the number of observations as a measure of exploration was insufficiently sensitive to capture the effect of epistemic uncertainty and we removed the constraints on sampling to mitigate this issue.

6.7.1 Method

Participants

130 undergraduate psychology students from UNSW Sydney participated in Experiment 9. The average age was 19.1 ($SD = 1.9$) and 84 participants were female. In addition to receiving course credit, participants were able to earn a small amount of money depending on their performance in the task ($M = \text{AU\$}4.47$, $SD = \text{AU\$}0.41$).

Design and procedure

Participants completed a decision-making task involving six rounds. In each round, there was an *observation phase* where they were able to learn about a risky option that would be encountered in the subsequent *choice phase*. These six rounds alternated between two sampling types. In *fixed sampling* rounds, participants were allowed to observe a predetermined number of outcomes associated with the option and were then required to progress to the choice phase regardless of whether they believed that they possessed enough information. Specifically, they were able to observe 5, 10, or 20 outcomes with the order of these restrictions randomised across the three fixed sampling rounds. In *free sampling* rounds, on the other hand, participants were informed that they should continue observing outcomes until they were comfortable progressing to the choice phase. There was no limit on the number of outcomes that they could observe.¹⁵

In order to encourage sampling in each of the six rounds, we varied the mean associated with the risky option. Each outcome was drawn from a normal distribution with a standard deviation of 20 points but the mean in each round was drawn from one of six uniform distributions: 10-20, 30-40, 50-60, 70-80, 90-100, or 110-120 points. The order of these distributions was randomised for each participant. Using this method allowed us to analyse the influence of the mean outcome on the number of outcomes that participants chose to observe while ensuring initial uncertainty regarding the distribution of outcomes in each round.

Participants were allocated to either the identical images or unique images condition. The same images were used as the previous experiments but they were differentiated across

¹⁵The free sampling rounds were included to examine whether surface variability influences the amount of information that people choose to acquire. In contrast, the fixed sampling rounds were used to provide a measure of epistemic and aleatory uncertainty independent of potential differences in the number of samples taken. For example, if participants in the identical images condition chose to observe 10 outcomes and those in the unique images condition chose to observe 25 outcomes, any differences in their responses on the EARS might reflect the outcomes sampled rather than the presence of surface variability.

rounds using six colours (red, orange, yellow, green, blue, and purple). We predicted that participants in the unique images condition would report greater epistemic uncertainty and choose to observe more outcomes in the free sampling rounds than those in the identical images condition.

During the choice phases, participants were presented with a decision between the risky option they had sampled in the preceding observation phase and a specified number of points that was displayed on the screen. Depending on which option they selected, a draw from the risky option or the specified number of points was added to their final score, which was converted into real money at the end of the experiment. They were paid AU\$1 for every 100 points earned during the task.

The specified number of points served the same purpose as the experienced safe options in the previous experiments. It provided an alternative to selecting the risky option, and therefore, we predicted that participants would select this option more often in the unique images condition because the risky option would be associated with greater epistemic uncertainty. In contrast, however, the expected value of the safe option was not always equal to the risky alternative. The safe outcome was either 20 points better than, equal to, or 20 points worse than the mean of the risky option. The order of these differences was shuffled within each sampling type (free or fixed).

After completing the choice phase in each round, participants were presented with a 4-item version of the EARS to examine their uncertainty regarding the risky options.

6.7.2 Results and discussion

Does surface variability influence the interpretation of uncertainty? Participants' responses to the EARS are shown in Figure 6.4a. We again predicted that participants in the unique images condition would report higher epistemic uncertainty and lower aleatory uncertainty than participants in the identical images condition. We used the same ordinal regression model that was used in the previous experiment. Consistent with the

predicted effect of surface variability on uncertainty, this model suggested that there is a 95.81% probability that presenting participants with unique or identical images influenced the difference between their responses to the epistemic and aleatory items (median = 0.59, 95% CI = [-0.13, 1.20]).

Using contrasts to examine epistemic and aleatory items separately suggested that there is an 82.46% probability that people report experiencing *more* epistemic uncertainty and this corresponds to a difference of roughly 0.40 between these groups on the latent scale (95% CI = [-0.76, 1.38])). This only provides weak evidence that surface variability increases epistemic uncertainty. The evidence regarding aleatory uncertainty was more clear-cut. There was a 96.67% probability that experiencing unique images resulted in reporting *less* aleatory uncertainty and that the difference between groups was roughly -0.81 standard deviations on the latent scale (95% CI = [-1.54, 0.05]). These estimates echo the results of Experiment 8. Although we predicted that surface variability would predominantly influence epistemic uncertainty, the evidence regarding this relationship has been somewhat inconclusive. In contrast, there has been considerable evidence across all three experiments that surface variability influences aleatory uncertainty.

Is there an effect on exploration? The number of outcomes that participants chose to observe in free sampling blocks is displayed in Figure 6.4b. We predicted that greater epistemic uncertainty regarding the risky option in the unique images condition would give rise to increased exploration. Given that we had some reason to doubt the relationship between surface variability and epistemic uncertainty, we also have reason to doubt this hypothesis. Nonetheless, we evaluated this hypothesis using Bayesian hierarchical negative binomial regression predicting the number of outcomes that participants chose to observe before making a choice.¹⁶ The condition (unique images or identical images), block (1-6), and their interaction were included as fixed predictor variables. The intercept was allowed to vary for each participant. This model suggests that there is a

¹⁶We apply the student-t distribution to the inverse of the shape parameter so that the mode of the prior distribution corresponds to the simpler Poisson distribution rather than a negative binomial with large amounts of over-dispersion. Taking the square root places the shape parameter on a similar scale to the other parameters (Simpson, 2018).

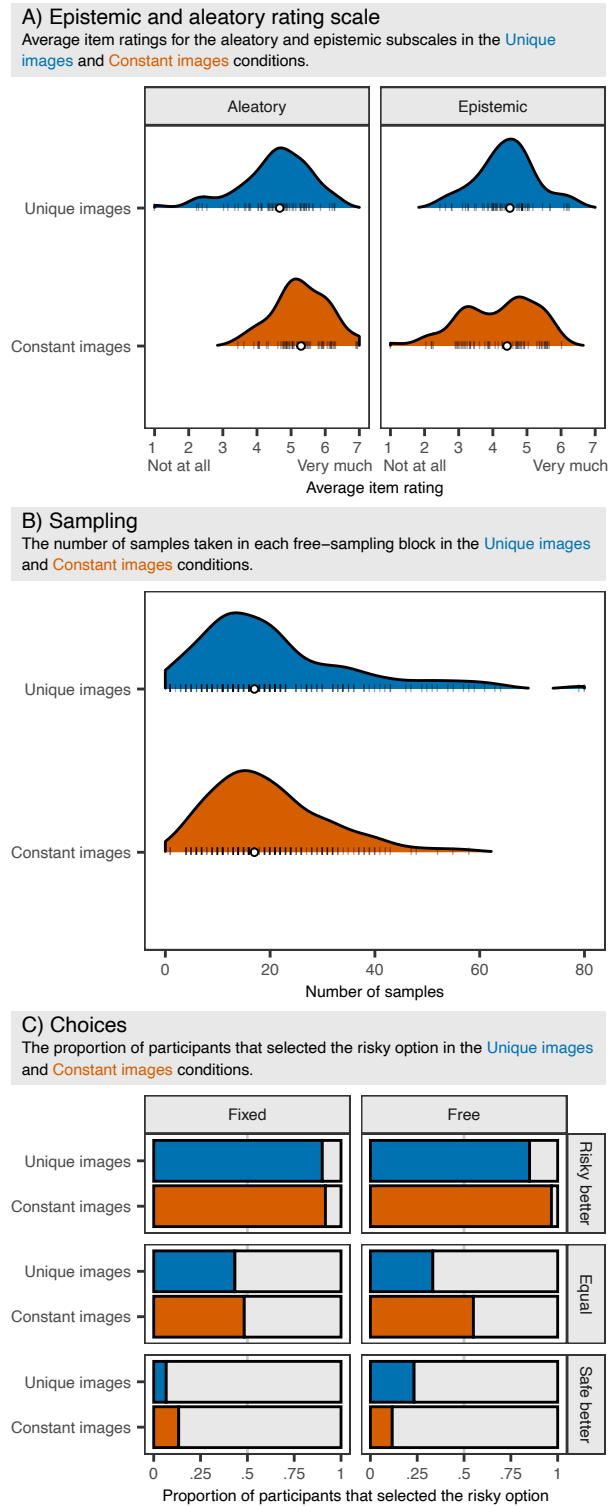


Figure 6.4: Responses to the EARS questionnaire, sampling, and choices for Experiment 9. The white dots represent the median response. Average item ratings reflect the average response for each participant on the epistemic or the aleatory subscales.

66.20% probability that participants were more likely to observe outcomes in the unique images condition (median odds ratio = 1.05, (95% CI = [0.83, 1.32])). This does not provide any meaningful evidence that surface variability increases exploration and there is even a considerable probability that the effect could fall in the opposite direction. To examine the relationship between epistemic uncertainty and exploration more directly, we also modelled the relationship between participants' responses to the epistemic items of the EARS and the number of outcomes they chose to observe. This model suggested that there is a only a 19.81% probability that there is a positive relationship between epistemic uncertainty and observing more outcomes (median odds-ratio = 0.96, (95% CI = [0.86, 1.06])). These results provide additional problems for our hypothesis regarding epistemic uncertainty and information seeking.

Is there an effect on choice? Participants' choices between the safe and risky options are shown in Figure 6.4c. We hypothesised that participants would avoid options associated with epistemic uncertainty, and consequently, that participants in the unique image condition would select the risky option less often than participants in the identical images condition. We examined this hypothesis using Bayesian hierarchical logistic regression predicting the choices that participants made in each round. The condition (unique images or identical images) was included as a fixed predictor variable and the intercept was allowed to vary for each participant. Based on this model, there is an 88.71% probability that participants were more likely to avoid the risky option in the unique images condition (median odds-ratio = 0.56, (95% CI = [0.21, 1.47])). We again examined the relationship between epistemic uncertainty and participants' choices by replacing the condition predictor in the model described above with responses to the epistemic items of the EARS. According to this model, there is only a 40.45% probability that participants in the unique images condition were more likely to avoid the risky option (median odds-ratio = 1.06, (95% CI = [0.66, 1.79])). Similarly to our findings regarding information seeking, these results place doubt on our hypothesis regarding epistemic uncertainty aversion and we will examine both of these findings in the chapter discussion.

6.8 Chapter discussion

The main purpose of these experiments was to determine whether introducing observable variability would influence the interpretation of uncertainty when people could learn about options through experience. This research question was motivated by the simple observation that the observable characteristics of real-world options vary each time they are encountered. In contrast, experimental methods, such as the classic bandit task, represent options using images that remain constant throughout the task, and therefore, uncertainty cannot be resolved by mapping outcome variability onto observable variability. Consequently, we expected that people would experience greater epistemic uncertainty when instances of options were differentiated by varying the images used to represent each instance. We found that observable variability influenced how uncertainty was interpreted but our results were not wholly consistent with our expectations. Namely, the evidence that surface variability influences aleatory uncertainty was stronger than for epistemic uncertainty and we observed little evidence of an impact on either exploration or choice.

6.8.1 Epistemic and aleatory uncertainty

In this chapter, we examined the influence of surface variability on epistemic uncertainty across three experiments. One reasonable approach to aggregating this evidence would be to include the experiments in a single multilevel model and this combined analysis suggests that there is a 97.88% probability that surface variability increases epistemic uncertainty (median = 0.26, 95% CI = [0.01, 0.52]). This model ostensibly suggests that surface variability influences epistemic uncertainty but this might not tell the whole story. Examining each of these experiments separately, there was considerable evidence of an effect on epistemic uncertainty in the first experiment in this chapter but the second and third experiments were less conclusive. On one hand, if this difference between these experiments is merely attributable to sampling error, the combined model would provide the best estimate and we might conclude that, on average, surface variability increases epistemic uncertainty. On the other hand, if this difference is attributable to dissimilarities

between the experiments, it would call into question the robustness of the effect and further investigation would be required to determine the extent to which we can generalise beyond the context of the first experiment.

In contrast, the evidence that surface variability influenced participants' responses to the aleatory items was relatively unequivocal. The combined model suggests that there is a 99.49% probability that surface variability influenced responses to the aleatory items (median = -0.49, 95% CI = [-0.78, -0.17]) and the effect was consistent across the three experiments. We expected that ratings of aleatory uncertainty would decrease but mainly as a consequence of increasing epistemic uncertainty. Assuming that the total amount of uncertainty experienced in each condition was similar, an increase in epistemic uncertainty would lead to a decrease in its complement, aleatory uncertainty, but were we justified in making this assumption regarding total uncertainty? Although surface variability cannot be used to improve their performance, this does not preclude participants from believing they had successfully mapped surface variability onto the outcome variability. This would lead to a decrease in total uncertainty, which is consistent with participants reporting that the outcomes were less "random" or "unpredictable." This decrease in total uncertainty would compound with the decrease in the absolute amount of aleatory uncertainty, but even if there was an increase in the proportion of epistemic uncertainty *relative* to aleatory uncertainty, a decrease in total uncertainty would act in the opposite direction to decrease the *absolute* amount of epistemic uncertainty.

We will revisit this distinction between relative and absolute uncertainty later when we discuss measuring uncertainty using the EARS. First, let us consider the possible scenario that—at least in the second and third experiments in this chapter—there might have been a reduction in reported aleatory uncertainty without a corresponding increase in reported epistemic uncertainty. This scenario seems to suggest two possible alternatives: either at some point participants interpreted their uncertainty as epistemic but we failed to accurately measure this or their aleatory uncertainty decreased without ever affecting their epistemic uncertainty. The first alternative might plausibly have resulted from inadequate measurement precision for epistemic uncertainty or because we presented

the EARS at the end of the experiment. Participants might have figured out that there was no mapping between observable and outcome variability, and therefore, presenting the EARS at different points throughout the task would allow us to examine the time course of uncertainty. The second alternative—that observable variability does not influence epistemic uncertainty—might be asserted because people instead use heuristics, such as the performance of others, to evaluate whether uncertainty is resolvable. This alternative must contend, however, with the evidence of an effect on epistemic uncertainty in the first experiment and the apparent absurdity, but not impossibility, of reducing uncertainty without ever believing that it was possible.

6.8.2 Outcome and observable variability

In Experiment 7a, we observed an unexpected pattern in participants' responses to the EARS. In contrast with the risky option, participants rated the safe option as *higher* in aleatory uncertainty and *lower* in epistemic uncertainty when unique images were presented on each trial. We were initially sceptical of this observation so replicated the experiment to reduce the possibility that we were grasping illusory patterns. Given that we observed the same pattern again in Experiment 7b, we are confident that this was not a statistical anomaly but it is less clear whether this pattern is theoretically meaningful. In response to this latter concern, we examined the amount of observable and outcome variability associated with each option. This relationship is presented visually in Figure [6.5](#). Below the diagonal in Figure [6.5](#), there is less observable variability than outcome variability and this divergence increases with distance along the y-axis from the diagonal. For this reason, we predicted that the risky option in the constant images condition would be primarily associated with aleatory uncertainty. For options positioned along the diagonal, such as the risky option in the unique images condition, the observable variability is proportional to the outcome variability. Therefore, to the degree that the person believes they can map this observable variability onto the outcome variability, we expect that people would predominantly experience epistemic uncertainty.

The safe option in the unique images condition is positioned in the vast space above the diagonal where there is more observable variability than is expected in the outcomes. Although the space above the diagonal initially seems unusual, options could be located there whenever there are factors that influence observable variability but not variability in the outcomes—in other words, whenever there is a considerable amount of surface variability. Returning to participants' responses to the EARS, the location above the diagonal of the safe option in the unique images condition might have highlighted that there were other factors present in the environment that might influence the outcome of selecting the option. Consequently, participants might have perceived the option as less predictable. This is consistent with responses to the EARS, and if this offers an adequate explanation, we might also expect to observe an effect on participants' choices. Given that participants were generally risk averse, increasing uncertainty regarding the safe option would decrease its attractiveness relative to the risky option. As such, we would expect that the preference for the safe option would decrease in the unique images condition. Consistent with this prediction, there is evidence that responses to the aleatory EARS items for the safe option mediated roughly 42% of the effect of the surface variability manipulation on participants' choices in Experiment 7.¹⁷ Contrast this with responses to the epistemic items for the risky option—the subject of our primary hypothesis—which only mediated roughly 12% of this effect.

Despite this evidence, there are two forceful rebuttals against our interpretation of these results. Firstly, it seems plausible that participants were merely responding in a strange way to a strange question. We asked them to describe their uncertainty regarding outcomes for which they had little uncertainty, and therefore, they might have responded with reference to their uncertainty regarding the appearance of the options rather than the outcomes. To address this concern, we presented participants in Experiment 7b with a specific instance of each option and emphasised that the EARS questions referred to the

¹⁷The median estimate of the direct effect was 0.011 (95% CI = [-0.009, 0.031] and the median estimate of the indirect effect was 0.008 (95% CI = [0.002, 0.015]). The estimates for the risky option and further details on the mediation analysis are presented in the results section of Experiment 7.

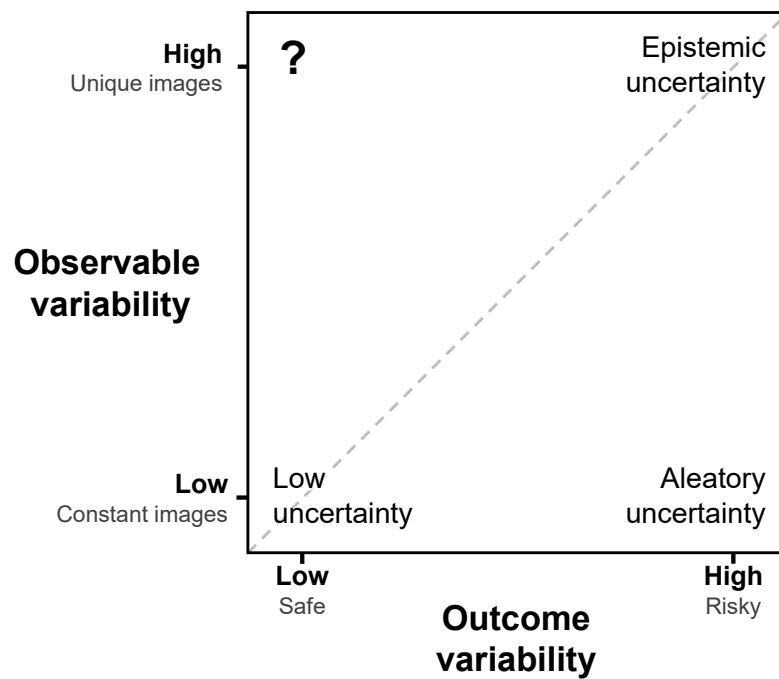


Figure 6.5: **The relationship between observable and outcome variability.** There is a straightforward interpretation of three corners of the square: 1) When observable and outcome variability are low, there is minimal uncertainty, 2) When observable and outcome variability are high, uncertainty is interpreted as epistemic, 3) when observable variability is low and outcome variability is high, uncertainty is interpreted as aleatory. The interpretation of the final corner is less obvious.

outcome that would result from selecting this instance. Their responses in this version were similar to Experiment 7a. As an additional piece of evidence that their responses to the safe option were—at the very least—not complete nonsense, we conducted exploratory and confirmatory factor analyses, which suggested that participants were interpreting items in a way that was broadly similar to their interpretation of items for the risky option (see Appendix [C](#) for more details).

The second rebuttal is more difficult to address using the available evidence. This rebuttal asserts that, although we replicated the effect in Experiment 7b, the same pattern was not observed in Experiment 8. Prior to conducting the latter experiment, we had no reason to believe that there would be a difference in participants’ responses to the EARS. Therefore, in order to convincingly argue that the pattern is meaningful, we would need to describe a region that includes the manipulation in the first experiment and excludes the manipulation in the second. One plausible difference between these experiments is that participants in the second experiment observed ten times fewer outcomes (the median was six outcomes for the safe option) and it is plausible that the contrast between the observable and outcome variability was less pronounced than the first experiment where the median was 61.5 outcomes. Further investigation would be required to determine whether this is an adequate defence against the second rebuttal. Given that the current experiments focused primarily on the region below the diagonal in Figure [6.5](#), this investigation might provide a more detailed understanding of the role of observable variability in decision-making.

6.8.3 Measuring uncertainty using the EARS

In these experiments we aimed to differentiate between two distinct concepts: the *amount* of uncertainty (total uncertainty) and the *type* of uncertainty (epistemic or aleatory uncertainty). This distinction can be captured using two equivalent conceptualisations provided that epistemic and aleatory uncertainty are complementary: either by measuring the *absolute* amount of epistemic and aleatory, the sum of which corresponds to the total amount

of uncertainty, or measuring the total uncertainty and the *relative* amount (proportion) of epistemic and aleatory uncertainty.^[18] As an example, when comparing EARS responses for options involving different amounts of total uncertainty, if the EARS is an absolute measure, the safe options should be rated as involving both lower epistemic and aleatory uncertainty than the risky options. If the EARS is a relative measure, the proportion of epistemic and aleatory uncertainty might instead be similar for comparable safe and risky options. In our first experiment, there was a greater than 99.9% probability that whether an option was safe or risky influenced ratings of both epistemic (median = 1.17, 95% CI = [0.63, 1.67]) and aleatory uncertainty (median = -1.72, 95% CI = [-2.35, -0.97]).^[19] This *decrease* in aleatory uncertainty for the safe option might be consistent with the EARS measuring the absolute amount of each uncertainty type but the large *increase* in epistemic uncertainty is difficult to reconcile with either conceptualisation and suggests that the scale might be confounding the amount and type of uncertainty.

To examine the origin of this issue, let us examine the framework proposed by Fox and Ülkümen (2011). They propose that “aleatory uncertainty can be measured by entropy”, which maps neatly onto EARS items such as, “... is unpredictable”. Likewise, they suggest that “as subjective knowledge decreases, epistemic uncertainty increases”,

¹⁸Although not explicitly discussed, the absolute conceptualisation can be recognised in the separate predictions associated with the epistemic and aleatory subscales in Ülkümen et al. (2016) and relative conceptualisation in the “epistemicness” index that averaged epistemic items and reverse-coded aleatory items in Tannenbaum et al. (2017). In this latter article, the authors justify their use of the relative index by asserting that they obtained “qualitatively identical results” using the separate scales. One exception to this was in their final experiment where they used a modified bandit task to prime the interpretation of uncertainty in a subsequent task. Similarly to the second and third experiments in this chapter, there was a difference between conditions of the aleatory subscale but no reliable difference on the epistemic subscale. Due to the evident similarities, the discussion of epistemic uncertainty regarding our experiments might equally apply to the experiment described by Tannenbaum et al. (2017).

¹⁹A similar effect was observed in the second experiment but the effect size was diminished (median for epistemic uncertainty = 0.28, 95% CI = [0.00, 0.57]; median for aleatory uncertainty = -0.48, 95% CI = [-0.88, -0.04]). This was most likely because—as discussed above in the section regarding outcome and observable variability—the small number of observations in the second experiment reduced the distinctiveness of the two options.

which maps onto items such as “... is knowable in advance, given enough information”. With the benefit of additional experimental evidence and hindsight, the authors of this chapter recognise that both of these conceptualisations confound the amount and type of uncertainty. Entropy (or total uncertainty) only determines aleatory uncertainty when the absolute amount of epistemic uncertainty is held constant and subjective knowledge (or the proportion of uncertainty that is attributed to inadequate knowledge) only determines epistemic uncertainty to the extent that total uncertainty is held constant. Therefore, although the EARS has demonstrated convergent validity as a measure of epistemic and aleatory uncertainty (e.g., Tannenbaum et al., 2017; Ülkümen et al., 2016; Walters et al., 2022), our experiments seem to suggest that it does not adequately discriminate between the amount and type of uncertainty. Further refinement of the EARS items to improve its discriminative validity would offer a valuable contribution to our ability to investigate epistemic and aleatory uncertainty.

6.8.4 Information seeking

We predicted that observable variability would influence information seeking by increasing epistemic interpretations of uncertainty. The conclusions drawn in this section are consequently less definitive given the inconclusiveness of the evidence regarding this potential mediator—the absence of an effect of surface variability on exploration would be wholly unsurprising if surface variability does not influence epistemic uncertainty. Despite this, let us briefly consider the implications of the alternative possibility in which—at least at some point—surface variability caused participants to interpret uncertainty as more epistemic and we merely failed to capture this effect using the EARS. Based on our results, would we expect a difference in the number of outcomes that participants observed?

Our prediction that surface variability would cause an increase in exploration relied on an assumption that participants would not be able to reduce their total uncertainty because surface variability was uncorrelated with the outcomes. Given that this appears to have been violated, an increase in epistemic uncertainty might coincide with an increased

impression that each observation contains useful information. Based on a simple model in which people seek information until the disadvantages associated with exploration (e.g., opportunity cost) no longer outweigh the perceived benefits of further reducing uncertainty, a steeper decline in uncertainty caused by illusory patterns in surface variability could mean that people more rapidly approach the threshold at which further exploration no longer appears beneficial. This perceived difference in the informativeness of each observation could offer a potential explanation for the similar number of outcomes that participants observed in the constant images and unique images conditions.

6.8.5 Competition and risk

The decisiveness of our conclusions based on participants' choices between safe and risky options are equally subject to the inconclusiveness of our evidence regarding epistemic uncertainty. As such, we will allow the reader to draw their own conclusions based on each experiment, and instead, focus in this section on a reasonable criticism of our hypothesis regarding risk. One of the reviewers of an earlier conference paper based on the first two experiments in this chapter (Holwerda & Newell, 2021) responded to our claim that epistemic uncertainty implies that adversaries might possess more knowledge with the question, "Who is the opponent?" Indeed, many of the explanations for epistemic uncertainty aversion hinge on social characteristics (e.g., information asymmetry, potential blame, or embarrassment) so why should it be surprising that participants in private testing booths are indifferent to epistemic uncertainty? Good question. Nonetheless, we conjectured that we would observe a difference in participants' choices because epistemic uncertainty aversion was founded using equally sterile conditions (Camerer & Weber, 1992). For example, in his thought experiments, Ellsberg (2001) emphasised that the game is set up by an observer that has "no other interest in your choices, nor in the outcome of a gamble, than to learn your opinions; the prize money at risk is, in his eyes cheap payment for this information" (p. 131). Although there is also evidence that this effect is stronger when people perceive their own comparative ignorance (Curley et al., 1986; Fox & Weber, 2002), our hypothesis is also consistent with an expansive heuristic that it is generally

beneficial to avoid unnecessary epistemic uncertainty (Al-Najjar & Weinstein, 2009).

Although we argue that our hypothesis is robust to this particular criticism, the question posed by this reviewer implicitly draws attention to a deeper issue that threatens the parsimonious duality of epistemic and aleatory uncertainty. Although our concept of epistemic uncertainty—uncertainty attributed to lack of knowledge—can be used to estimate the benefits of avoiding certain options, we could potentially improve our accuracy by evaluating our knowledge relative to our opponent and weighting this estimate using the predicted consequences of each piece of asymmetrical information. In contrast, we could more accurately evaluate the benefits associated with exploration by estimating the cost of acquiring more information and the probability that this would improve our performance. Furthermore, this approach might consider each potential source of information rather than our uncertainty regarding each option. Therefore, our concept of uncertainty attributed to lack of knowledge might be too crude to explain some meaningful patterns in behaviour. Notably our current concept of epistemic uncertainty does not capture the differential impact of who lacks knowledge or distinguish between inadequate information or understanding.²⁰ These distinctions might not make a meaningful difference on the way people interpret uncertainty but it is equally plausible that we may need to refine our understanding of epistemic and aleatory uncertainty.

6.8.6 Conclusion

The different ways that people interpret uncertainty has often been neglected when examining decisions based on experience. As is often the case when traversing largely unexplored

²⁰One potential example of sensitivity to who lacks knowledge was exhibited in participants' responses to the ninth question of the EARS in our experiments. This question asks whether “consulting an expert” would improve prediction, and unsurprisingly, this question made less sense to participants in our experiment than it would when discussing sports or the stock market. The standardised loading of this item onto the epistemic uncertainty factor in our confirmatory factor analysis reached as low as .13 for the safe option in the first experiment. Although our conclusions were robust to the removal of this question, the responses to this item demonstrate that the relevance of a particular agent's knowledge state might differ between contexts.

territory, the current investigation has raised almost as many questions as it answered. The treatment of epistemic and aleatory uncertainty as a psychological distinction is still in its infancy but there is already evidence that people differentiate between them in several aspects of natural language (Ülkümen et al., 2016) and the consequences of this distinction are slowly becoming apparent in areas such as the extremity of judgements (Tannenbaum et al., 2017) and investment behaviour (Walters et al., 2022). In our experiments, we did not observe strong evidence that surface variability influences risk preferences. This might provide some reassurance that standard bandit tasks generalise to decisions where the options involve observable variability but we also recognise the somewhat artificial nature of surface variability in our experiments. Future experiments could strengthen this conclusion by deviating further from the standard task and incorporating observable variability that is semantically meaningful to the participant. Finally, we have discussed several areas in which we could refine our theoretical understanding and measurement of epistemic and aleatory uncertainty. Implementing these suggestions—such as discriminating between the amount and type of uncertainty—will be essential in further investigations into variants of uncertainty.

Chapter 7

General discussion

The nine experiments in this thesis examined the influence of context and uncertainty on decisions from experience. The first section focused on extreme outcomes. We began our investigation in Chapter 2 by addressing the question of why these outcomes appear not only in decisions from experience but also throughout numerous other cognitive domains. Although utility-weighted sampling offered a persuasive answer to this question, a closer examination revealed multiple issues with both its mechanical and rational components. In response to this conclusion, we devised an alternative rational explanation based on the informativeness of extreme outcomes. We demonstrated that prioritising these outcomes increases both the probability of selecting the best option and the expected utility gained from these choices.

In Chapter 3, we introduced a framework that categorises the many plausible mechanical explanations according to whether: 1) their level of measurement is categorical, ordinal, or continuous, 2) extreme outcomes refer to the centre, the edges, or neighbouring outcomes, 3) outcomes are represented as types or tokens, and 4) peaks are identified using temporal or distributional characteristics. We then described three experiments that manipulated the expected value, the variance, and the skewness of the experienced outcomes. These experiments provide considerable evidence against both the categorical and

edge-based explanations but also hinted at potential issues for many of the alternatives.

The two experiments in Chapter 4 were designed to follow up on some of the anomalies from the previous chapter whilst investigating the distinction between types and tokens. Once again, participants were more likely to select the risky option in the high-value pair than the risky option in the low-value pair. Nonetheless, we also noted that whenever we deviated from the design used by Ludvig et al. (2014), we were unable to predict whether we would observe this phenomenon. Finally, in Chapter 5, we manipulated the order in which outcomes were experienced so that the temporal and distributional peaks were no longer correlated. The results of this experiment were highly incompatible with the theories that are based on ordinal-level temporal peaks.

The second section in this thesis focused on epistemic and aleatory uncertainty. The vast majority of previous decisions from experience tasks have represented each option using a single image that remained constant across trials. Therefore, in Chapter 6 we examined whether introducing variability to the appearance of each option would elicit an epistemic interpretation of uncertainty. Across three experiments, we demonstrated that observable variability influences how uncertainty is interpreted but the impact on aleatory uncertainty was stronger than epistemic uncertainty and we observed little evidence of an impact on either exploration or choice.

Given that the discussion section in the previous chapter integrated all three experiments that examined uncertainty, we will not subject the reader to this content twice in almost as many pages. We have not yet performed a similar integration of the six experiments in the extreme outcome chapters and we will rectify this deficit in the following section. We will then discuss a series of concepts that were important to our examination of both context and uncertainty. Finally, we will conclude our discussion by revisiting the relationship between theory and experiment that we encountered at the beginning of this thesis.

7.1 Context and extreme outcomes

The extreme-outcome effect described by Ludvig et al. (2014) has been replicated dozens of times (e.g., Madan et al., 2017; Madan et al., 2015; Madan et al., 2021), and therefore, our aim was not to question the existence of the *effect* but instead to reconcile two simple observations. On one hand, the *explanations* for this effect, such as the extreme-outcome rule, assume a broad definition of extremity that is applicable across a range of contexts. On the other hand, the experiments used to assess these theories have always operationalised extreme outcomes using symmetrical high- and low-value pairs of options. Several theories were consistent with the existing evidence and it remained unclear when it was legitimate to generalise to situations that deviated from these experiments.

Thus, in contrast with our predecessors, we exposed these theories to scenarios where none of the relevant outcomes were the best or worst. We manipulated the expected value, the variance, the skewness, and the order of the experienced outcomes. We observed that the broad definitions of extremity that appeared promising in the previous experiments crumbled under these new conditions. Some of these manipulations were present across multiple experiments and chapters. In this section, we will integrate them to assess each of the attributes that comprise the framework described in Chapter 3.

7.1.1 Categorical extreme outcomes

Perhaps the strongest evidence in the context section concerns the categorical definition of extremity. The extreme-outcome rule established by Ludvig et al. (2014) is a member of this class and given that their theory was based explicitly on decisions from experience, we should consider the rationale for its categorical definition. It consists primarily of two pieces of evidence: Firstly, when participants chose between two options with the same expected value, they were more likely to select the risky option associated with the best outcome than the risky option associated with the worst outcome. Secondly, there was no significant difference between participants' choices for high-value and low-value options

that were not associated with the best or worst outcomes (Ludvig et al., 2014, Experiment 3).

The prominence of the best and worst outcomes in this description naturally suggests a categorical definition but these observations could also be explained as suggesting that the effect is *weaker* rather than *absent* for outcomes near the centre of the distribution. This interpretation is compatible with the ordinal and continuous definitions. In contrast, the categorical class of theories is unable to explain the results of Experiment 1 and 2 of this thesis. These experiments found considerable evidence that participants were more likely to choose the risky option for high-value pairs relative to low-value pairs even when none of the relevant outcomes was the best or worst in the context of the experiment.

This evidence is consistent with two recent experiments that each suggest that extreme outcomes must extend—at the very least—to those located close to the best and worst outcomes (Ludvig et al., 2018; Mason et al., 2020). The deviations from the categorical extreme-outcome rule that were observed in these experiments were attributed to encoding noise but this is inadequate to explain our results. Specifically, participants in the low-variance condition of our second experiment chose the risky option more often for the high-value than low-value pair. We observed this pattern even though these outcomes were separated from the best and worst outcomes by 40% of the distance between the centre and the edges of the distribution.

Unless you are willing to accept a gargantuan amount of encoding noise, we suggest that even the modified version of the categorical definition is unable to explain the influence of extreme outcomes in decisions from experience. The only avenue that remains for these categorical theories would be to argue that the phenomenon in our experiments was not—in fact—an instance of the effect observed by Ludvig et al. (2014). Instead, our results should be attributed to a different bias that influences intermediate outcomes. This is possible but offers a much less parsimonious approach than adopting a definition that extends beyond the best and worst outcomes.

7.1.2 Continuous extreme outcomes

The most straightforward alternative to the categorical extreme-outcome rule is the class of continuous theories in which memory deteriorates with increasing distance between an outcome and the edges of the distribution (e.g., Berliner et al., 1977; Braida et al., 1984). The best and worst outcomes retain their unique status in these edge-based theories. Rather than acquiring their importance through their influence on estimates, however, they serve as anchor or reference points. The influence of extreme outcomes decreases gradually and this allows these theories to accommodate the results of our second experiment that were so disastrous for the categorical-level theories.

The attribute that distinguishes these edge-based definitions from the other continuous-level theories is that the role of intermediate outcomes is *passive*. In other words, memory is only influenced by the distance between each outcome and the edge of the distribution, and therefore, intermediate outcomes do not—even collectively—influence memory for other outcomes. This attribute entails that options can be evaluated quasi-independently, only taking into account the best and worst outcomes and remaining blind to the presence or absence of intermediate outcomes. Therefore, edge-based theories can easily account for why manipulating the skewness of intermediate outcomes in our third, fourth, and fifth experiments had little influence on choices.

Although the passive role of intermediate outcomes allows the edge-based theories to explain the relationship between skewness and choice in these experiments, it necessarily disqualified them from explaining two other observations. Firstly, there was evidence that skewness influenced the memory responses for the shared options in Experiment 1, 3, and 5. And secondly, in contrast with other experiments using similar options, participants did not select the risky option more often for the high-value pair than the low-value pair when they were presented separately in Experiment 3. The best and worst outcomes were held constant across these conditions, and therefore, these observations cannot be explained by assuming a *passive* relationship between intermediate outcomes.

In contrast, centre-based continuous-level theories, such as utility-weighted sam-

pling, posit an active relationship between intermediate outcomes. Every outcome influences the location of the centre of the distribution and the skewness manipulation would have influenced the weighting of each outcome associated with the shared risky options. Therefore, these theories can easily explain the influence of skewness on memory but—as was the case with its passive counterpart—this attribute cuts both ways. The centre-based theories struggle to explain why the predicted effect of skewness on choices was not observed in our third, fourth, and fifth experiments.

In summary, the most consequential difference between the edge-based and centre-based theories is whether there is a passive or active relationship between intermediate outcomes. They both struggle to explain some of our observations but they struggle with opposite aspects of the data. Whereas one class was unable to explain the *presence* of a phenomenon that *was not* predicted, the other class of theories was unable to explain the apparent *absence* of a phenomenon that *was* predicted. We will consider the implications of this distinction further in the *evidence and rebuttals* section of this discussion.

In contrast with the edge-based and centre-based theories, the final class of continuous-level theories do not involve a direct relationship between extreme outcomes and memory. Instead, their influence in neighbour-based theories arises because each outcome stored in memory interferes with the retrieval of other similar outcomes (Murdock, 1960; Neath et al., 2006). Items that are located near the edges of a distribution usually have fewer immediate neighbours than those located near the centre. Therefore, extreme outcomes are usually—though not necessarily—easier to retrieve from memory than intermediate outcomes.

One consequence of this indirect relationship is that the predictions regarding extreme outcomes are quite sensitive to the attributes of the similarity function. We provided an example of this in the *levels of measurement* chapter. Ludvig et al. (2018) demonstrated that the best and worst outcomes were still disproportionately influential when additional outcomes were included in the context that were separated from them by a single point. This design also included outcomes near the non-extreme risky outcomes so that there were

four highly similar outcomes in the centre of the distribution. Although the exponential similarity function employed by Neath et al. (2006) predicts a diminished extreme-outcome effect in this experiment, other distinctiveness models that assign less weight to almost identical outcomes would diverge from this prediction.

This creates a challenge when interpreting the value and variance manipulations in our experiments because there is no consensus amongst the neighbour-based theories. The skewness manipulation, however, produces unambiguous predictions for most plausible similarity functions. These predictions were broadly similar to the centre-based theories. The skewness manipulation introduced three additional outcomes near either the better or worse outcome of the shared option. Therefore, the neighbour-based theories assert that the shared outcome with fewer neighbours would be more influential on both memory responses and choices. In other words, these neighbour-based theories can explain the effect of skewness on memory but not the seeming absence of an effect on choice.

7.1.3 Ordinal extreme outcomes

Ordinal-level theories offer an alternative explanation for the difference between the high-value and low-value options when they were not associated with the categorical extreme outcomes. They assume that the influence of extreme outcomes diminishes gradually, but in contrast with the continuous-level theories, the value of these outcomes is not situated on an external scale. Instead, an outcome's rank depends on a direct comparison with the other outcomes and the distance between adjacent outcomes cannot vary—the distance is always exactly one rank. This attribute offers a parsimonious account for the invariance of extreme-outcome phenomena to both the scale and location of the distribution (e.g., Ludvig et al., 2014; Neath & Brown, 2006). Continuous-level theories resort to normalising outcomes (e.g., Lieder et al., 2018) but this is unnecessary for the ordinal theories due to the absence of an external scale.

This attribute also means that the rank of the median outcome is always equidistant from the edges and there is no difference between centre-based and edge-based ordinal-level

theories. Having said that, ordinal theories are more similar to centre-based or neighbour-based continuous-level theories than edge-based continuous-level theories. Introducing an intermediate outcome actively influences the extremity of other outcomes by changing their rank. Therefore, similarly to the centre-based theories, ordinal-level theories are able to explain the influence of skewness on memory but struggle to explain the inconsistency of its influence on choice.

We aimed to disentangle the ordinal-level and continuous-level theories in our second experiment by manipulating the variance of the risky options. This changed the distance between each outcome and the centre or edges of the distribution whilst keeping their rank constant. Unfortunately, the results of this manipulation were somewhat ambiguous. There was some evidence that variance influenced memory for the outcomes associated with low-value options but this was not observed for the high-value options. Therefore, distinguishing between these levels of measurement requires further experimentation.

7.1.4 Temporal extreme outcomes

In the final chapter of the context section, we examined the possibility that the extreme-outcome effect is driven by the temporal relationship between outcomes. Specifically, we examined two closely related ordinal-level temporal theories: the first suggests that the effect can be attributed to temporal peaks and the second suggests that the subjective value of each experienced outcome depends on whether the previous outcome was better or worse. We acquired strong evidence that neither of these theories is sufficient to explain the extreme-outcome effect. Although this was the case, our evidence does not rule out a continuous-level theory in which outcomes are evaluated based on the *amount* that the previous outcome was better or worse. This theory could be examined in future experiments using a similar approach in which presentation order is used to eliminate the correlation between temporal and distributional peaks.

7.1.5 Evidence and rebuttals

So far, in this discussion section, we have examined several candidate explanations for the extreme-outcome effect. Not one of those explanations was left entirely unscathed in the light of our results. So where does that leave us? If each of the existing explanations is refuted, we would need to develop new theories but before accepting this as necessary, we will examine four possible rebuttals that might be offered as a defence of the existing theories. These consist of the following: 1) attributing the results to sampling variability, 2) questioning the relationship between the manipulation and the theory, 3) questioning the relationship between the observations and the effect, and 4) emphasising that additional variables exist beyond the scope of the theory.

These rebuttals were not selected arbitrarily and instead reflect limitations associated with three fundamental scientific assertions. First, the *randomisation* assertion states that experiments can assess the adequacy of a theory by randomly allocating participants to conditions. This process ensures that the difference between conditions reflects either the manipulation or sampling variability. Although the latter component diminishes as the number of participants increases, the first rebuttal emphasises that our conclusions might be contingent on our sample. The plausibility of this rebuttal is quantified in the statistical analyses described in the previous chapters.

This rebuttal is not very compelling regarding the categorical-level or temporal theories but might gain some traction with the edge-based and centre-based theories evaluated using our skewness manipulation. Although there was considerable evidence that skewness influenced memory in our first and third experiments, there was little evidence that it influenced participants' behaviour in our fourth and fifth experiments. The edge-based theories imply a passive relationship between intermediate outcomes and the effect of skewness on memory was difficult to reconcile with these theories. Given that there were two experiments that found an effect and two that did not, this appears to legitimise the rebuttal that our evidence against the edge-based theories was based on sampling variability.

To explore this possibility, we combined the four experiments that manipulated skewness and analysed their results in a single multilevel Bayesian model. Although the same number of experiments were consistent and inconsistent with the edge-based theories, the combined model implies that there is a greater than 99.9% probability that there was an effect of skewness on memory. This analysis suggests that attributing the alleged skewness effect to sampling variability does not offer the edge-based theories a plausible rebuttal. Neither does it necessitate, however, that the effect of skewness was present in every single experiment and potential differences between them remain an open question.

What about the effect of skewness on choices predicted by the centre-based, neighbour-based, and ordinal theories? A sizeable proportion of the posterior distribution was consistent with an effect of skewness on choices and the sampling variability rebuttal might be more successful. This ranged from 9% in Experiment 3 to 86% in Experiment 1, and therefore, once again, we combined them into a single multilevel model. This analysis suggests that there is a 16% probability that there is any effect of skewness on choices in the predicted direction. This is roughly equivalent to predicting the outcome when rolling a single die—thus within the realm of possibility—but employing the sampling variability rebuttal is not without consequences. It subjects these theories to the dilemma of limiting the predicted strength of the effect or making predictions that are increasingly improbable given the empirical evidence.

The *mapping* assertion states that one of the primary aims of scientific theories is to explain phenomena using a function that maps input variables onto an output state that encompasses the phenomenon to be explained. This assertion features prominently in the deductive-nomological model of Hempel and Oppenheim (1948). According to their model, an explanation comprises a sentence describing the phenomenon that can be deduced from a class of sentences describing *antecedent conditions* and *general laws* (Hempel & Oppenheim, 1948). Although the rational and mechanical explanations described in this thesis are usually presented as fundamentally opposed to this account, they similarly depend on a version of the mapping assertion (Cartwright, 1983; Cartwright et al., 2020). This also arguably applies to other recent accounts of explanation such as counterfactual

(Pearl, [2000](#); J. Woodward & Woodward, [2003](#)), statistical relevance (Salmon, [1971](#)), probabilistic (Suppes, [1970](#)), unificationist (Kitcher, [1981](#)), pragmatist (Mitchell, [1997](#)), and explanatory virtue theories (Keas, [2018](#)).

To determine the adequacy of a theory, we conduct experiments that examine whether a given set of input variables produces an output state that is compatible with the mapping specified by the theory. This requires us to manipulate the input variables and the second rebuttal emphasises the potential mismatch between the variables in the theory and those that were actually manipulated. In the case of extreme outcomes, this rebuttal challenges the correspondence between the definition of extremity and the way we operationalised it in our experiments. This rebuttal is most frequently encountered when a predicted effect is not observed, and therefore, the researcher questions whether the manipulation was strong enough to observe the effect.

We discussed this possibility in Chapter 3 regarding the apparent absence of an effect of the skewness and variance manipulations. We conceded that a manipulation that affects memory might be inadequate to impact choice because the selected option is influenced by variables other than memory. Nonetheless, we also argued that each manipulation was designed to produce the strongest possible effect within the contrived experimental task. For example, the difference between the average outcomes in the skewness conditions in Experiment 3b was roughly equal to half the overall range of the experienced outcomes. It is possible that this manipulation was inadequate but adopting this rebuttal would severely limit the scenarios in which we should expect to observe the effect.

Similarly to the way that the second rebuttal questions the relationship between the input variables and the manipulation, the third rebuttal questions the correspondence between the output variables and our observations. It usually concedes that there was a legitimate difference between conditions but asserts that the difference does not share a common explanation with the phenomenon accounted for by the theory. We mentioned earlier that this might be one of the few remaining avenues for the categorical-level theories because there was strong evidence for a difference between conditions when the best

and worst outcomes were identical. Its persuasiveness depends on whether positing two separate causes gives the theory more explanatory power than the parsimonious explanation that attributes them to a single cause. There is some precedence for this approach in the distinct theories that account for the effect of extreme outcomes in different domains (e.g., Fredrickson & Kahneman, 1993; Neath et al., 2006), but further evidence would be required to make this a compelling case.

The *ceteris paribus* assertion renounces the claim that genuine scientific explanations must appeal to universal laws (e.g., Hempel & Oppenheim, 1948). This nomothetic account *might* be compatible with the theories in fundamental physics but is incompatible with those in the special sciences, such as psychology or economics. Instead, our explanations hold only within a limited range and there are often deviations and exceptions even within that range (Cartwright, 1999; Fodor, 1974; Wimsatt & Wimsatt, 2007; J. Woodward, 2000). Based on this assertion, the fourth rebuttal emphasises that even limited theories can be useful and that there are always potential disruptive factors that were not explicitly specified in the theory. An effect that was predicted might not be observed because there are multiple conditions that must be satisfied (e.g., water and sunlight are *both* required for a plant to grow) or there are other factors acting in the opposite direction (e.g., a feather falls slower than a bowling ball because of air resistance).

We employed this rebuttal to explain why skewness only affected memory responses in our third experiment. Specifically, we noted that the context variables were changing the rank of the safe option and this might be suppressing the effect of skewness on choice. Once again, however, employing this rebuttal is not without potential consequences. The usefulness of a theory depends on its degree of invariance across the domain in which it is applied. Lange (2000) describes this as follows: “Suppose someone says ‘I can run a four-minute mile’ but with each failure reveals a proviso that she had not stated earlier: ‘except on this track’, ‘except on sunny Tuesdays in march’ and so on. It quickly becomes apparent that this person will not acknowledge having committed herself to any claim by asserting ‘I can run a four-minute mile’” (p. 172). This problem is so pervasive in cognitive science that it has prompted some to propose that theories should come with

a toll-free number that the reader can call for advice on whether its predictions can be applied to a given scenario (Erev & Greiner, 2015).

These four rebuttals are not necessarily exhaustive but offer a systematic approach to examining many of the possible defences for the existing theories. In summary, there is considerable evidence that categorical and temporal extremity will be unnecessary in a parsimonious account of extreme outcomes. Not one of the rebuttals examined in this section offered a compelling argument for clinging to these theories. As for the remaining theories, the rebuttals available to the centre-based, neighbour-based, and ordinal-level theories appear to be somewhat more plausible than those available to the edge-based theories. Specifically, the absence of an effect of skewness on choice could be attributed to sampling variability, the strength of our manipulation, or disruptive factors acting in the opposite direction.

Finally, the conflicting predictions regarding skewness emphasise two possible ways that experimental results can contradict a theory: either a predicted difference is not observed or a difference is observed that was not predicted. When a difference between conditions is predicted but not observed—as was the case for skewness on choices—it challenges the validity of conditional statements, such as “if an outcome is extreme, then it will exert a disproportionate influence on choice”. It demonstrates that the conditions in the theory are not *sufficient* to observe the effect. Conversely, when an observed difference is not predicted—as was the case for skewness on memory—it shows that the conditions in the theory are not *necessary* to observe the effect. In other words, at the very least, the theory leaves some aspects of the phenomenon unexplained.

7.1.6 Future directions

If the existing theories cannot overcome the empirical challenges discussed in this thesis, new theories will be required. Some of these might be minor modifications of the existing theories but given the challenges faced by each of the theories we examined, a radically different conceptualisation of the effect might be required. A similar argument has been

made regarding the peak-end effect. As more experiments have been conducted, it has become less clear that peaks are used in retrospective evaluation or whether they merely capture similar information to other variables, such as the mean and median (Cojuharenco & Ryvkin, [2008]; Ganzach & Yaor, [2019]; Kemp et al., [2008]; Miron-Shatz, [2009]; Rozin et al., [2004]; Schäfer et al., [2014]; Seta et al., [2008]; Steffens & Guastavino, [2015]; Strijbosch et al., [2019]). It is indisputable that participants in the experiments conducted by Ludvig and colleagues selected the risky option more often when it was associated with the best outcome. As we have demonstrated, however, this might not be the only way to parse the difference between their conditions. If nothing else, our experiments provide a compelling case that we cannot rely on a narrow set of manipulations if we hope to understand the extreme-outcome effect.

Beyond this, one of the major contributions of this section was to provide an overarching framework in which theories regarding extremity can be categorised. We examined the level of measurement, the referent against which extremity is measured, whether it is represented using types or tokens, and whether it is temporal or distributional. Despite the breadth of this framework, there are some notable attributes that were absent. For example, most theories suggest that extreme outcomes are more *influential* than other experienced outcomes (e.g., Madan et al., [2014]) but it is also possible that their subjective *value* is systematically shifted. The influence-based accounts suggest that safe options are unbiased because influence is a relative concept and reweighting a single outcome has no impact. There is some evidence, however, that there is a similar effect of extreme outcomes for safe options (Wispinski et al., [2017]) and this could be explained by value-based theories (e.g., Chanales et al., [2020]; Favila et al., [2016]; Hulbert & Norman, [2015]).

As mentioned in the introduction, our motivation for conducting the experiments in the context section was simultaneously narrow and broad. The narrower aim was to examine the adequacy of the existing theories for the extreme-outcome effect in risky choice and we successfully eliminated some of the candidate theories. The broader aim was to begin looking at extremity in theories beyond just risky choice. Extreme outcomes can be found in theories across numerous domains but the relationship between them has

not received much attention. Although different explanations have been influential within each domain, the observed effects are often compatible with multiple explanations (for one exception, see Neath et al., 2006). We have made some progress in refining the explanation of the extreme-outcome effect in risky choice but similar efforts will be required in other domains.

In this section, we have demonstrated the importance of using a wide variety of manipulations and have provided a framework to organise the many definitions of extremity. When we examined the extreme-outcome effect in contexts that diverged from those used in the original experiments by Ludvig et al. (2014), we observed multiple anomalies that have improved our understanding of the effect. We suspect that examining the nature of extremity across other domains will lead to similar observations. This process will refine our understanding of each cognitive domain as well as the influence of extreme outcomes throughout cognition. In the few pages remaining in this thesis, we will shift our attention to concepts that appeared across multiple chapters—not only in the first section on context but also in the second section on uncertainty.

7.2 Representation

The concept of representation identifies the semantic content of mental states and has been foundational in many areas of philosophy, psychology, and artificial intelligence (Bringsjord & Govindarajulu, 2022; Pitt, 2022; Rescorla, 2020; Schlosser, 2019; Shapiro & Spaulding, 2021; Siegel, 2021; Thagard, 2020). In the early days of experimental psychology, a fierce disagreement erupted between Wundt and members of the Würzburg School about whether our representations necessarily include sensory content (Humphrey, 1951). This question of *how* we represent aspects of the environment acquired further significance following the emergence of the computational theory of mind (e.g., Johnson-Laird, 1983; Kosslyn, 1980; Pylyshyn, 1973). This approach emphasises that representational systems influence how efficiently certain computations can be performed. For example, the same number can be represented using Arabic or Roman numerals and these are equivalent for some purposes.

Nonetheless, calculating 42 times 96 is much easier than XLII times XCVI because Arabic numerals represent value using positional notation.

Mental representation remains a major topic of debate in many areas of cognitive science (e.g., Brette, [2019](#); Felin et al., [2017](#); Hebart et al., [2020](#); Szollosi et al., [2022](#)). It was similarly essential across multiple chapters in this thesis. Our discussion of context focused on how extreme outcomes are represented (e.g., as types or tokens) and the distinction between epistemic and aleatory uncertainty describes how people represent uncertainty. These chapters concentrated on *how* certain aspects of the environment are represented but also considered *whether* the reasons that explain our behaviour are necessarily *represented by* us. In Chapter 2, we developed a rational model that included the mean or median outcomes in the context and this inevitably raised the question of whether people have access to this information. In Chapter 6, we noted that the difference between epistemic and aleatory uncertainty can be observed in the behaviour of people who have never even heard of the distinction.

This suggestion that there are unrepresented reasons for our behaviour seems out of place in psychological science. It might be interpreted as regressing to the dark ages of behaviourism when mental representations were renounced as being irrelevant to psychological research (Watson, [1925](#)). Nonetheless, explanations without representations are commonplace in other fields such as biology and economics. For example, the density of leaves around a tree can be explained by claiming that each leaf behaves in a way that maximises the amount of sunlight it receives “as if it knew the physical laws” (Friedman, [1953](#) p. 19). There are myriad similarities between this explanation and those offered in psychology but scarcely anyone would claim that components of this theory are represented by the tree.

The reasons for an action can even differ from those represented by an organism. For example, a mother Black lace-weaver spider drums on her web in a way that triggers her offspring to attack and consume her body. Her behaviour is explained by arachnologists with reference to the reproductive benefits associated with transferring her body-mass to

her offspring compared with the alternative of producing a second clutch (Kim & Horel, 1998; Kim et al., 2000). The spider might have her own reasons for doing this but she never recognises the evolutionary rationale as the reason for her behaviour—neither did her mother or her grandmother. In fact, until it was represented by a human scientist, this reason was only recognised by Mother Nature herself (Dennett, 1983).

These unrepresented reasons are not exclusive to non-human animals and similar explanations can be given for why we prefer ice cream over Brussels sprouts or why we pursue romantic partners with certain physical characteristics. The reason for these preferences can be attributed to natural selection but what about clearly intentional behaviour with no obvious evolutionary rationale? In their classic experiment, Nisbett and Wilson (1977) presented customers in a shopping centre with four identical articles of clothing and asked them to indicate which one was the highest quality. The participants in this experiment showed a strong tendency to select the right-most item in the array but when they were asked to explain their decision, no one mentioned the position of the chosen item. Instead, they reported that “it was the knit, weave, sheerness, elasticity, or workmanship that they felt to be superior” (Wilson & Nisbett, 1978, p. 124).

Newell and Shanks (2014) attributed this phenomenon to the order in which options were evaluated. They suggested that people usually appraise options from left to right and compare subsequent items with the best one they have encountered so far. Each choice is based on the perceived attributes of the available options but the options in this experiment were identical so there should be no systematic quality-based preference for any single option. Nonetheless, assuming that the attributes of the current option are more salient than the previous items, which are stored in memory, this might increase the probability of selecting the current option on each choice. This would result in a tendency to select options that were evaluated later, and therefore, the observed pattern might reflect a temporal *process* rather than a spatial *preference*.

As such, there are multiple explanations for the choices that participants made in this task. Some reasons were represented by the participants and may correspond to their

verbal reports. Similarly to those of the mother spider, however, there were also reasons that were never represented by anyone until they were recognised by cognitive scientists. This account paints a very different picture from the claim that people are unable to accurately explain their own behaviour because they “have little or no introspective access to higher-order cognitive processes” (Wilson & Nisbett, 1978, p. 188). Instead, although we recognise that there are often unrepresented reasons that underlie our behaviour, it remains plausible that people accurately report *their* reasons—the reasons that are represented by them.¹

We can use this concept to interpret the explanations in this thesis that contain information that is not available to the people whom they describe. On one hand, people often have represented reasons for their behaviour. Someone might select an option because it gave rise to positive outcomes on previous occasions, because they are curious, or because the alternative is riskier. There are countless possible ways in which people might represent their environment and there is no guarantee that this will contain information such as the mean or median outcome. On the other hand, our rational model can be interpreted as an unrepresented reason that encompasses the broader system in which the choice is situated and the selective pressures that are exerted when people prioritise certain outcomes.

Which of these is the correct explanation? The answer to this question is the one given by the Cheshire Cat when Alice asked them which path she should follow: “That depends a good deal on where you want to get to” (Carroll, 1896, p. 90). Cognitive scientists are often looking for computational or psychological explanations and this destination must be reached by taking the representationist route. Other scientists travel down the anti-representationist route to discover interesting behavioural, ecological, neurological, or

¹An early precursor of this idea appears in the emphasis that Hegel (1807) bestows on the things that his readers recognise about the consciousness being discussed that are not explicitly available to the consciousness itself. It is also evident in his distinction between what Brandom (2019) describes as “a broadly behaviorist, externalist view, which identifies and individuates actions according to what is actually done... and an intentionalist, internalist view, which identifies and individuates actions by the agent’s intention or purpose in undertaking them” (p. 384).

dynamical explanations (e.g., Beer, [2000]; Chemero, [2009]; Gibson, [1979]; Skinner, [1953]). Many questions can be answered using either strategy but the chosen path is not always an ideological decision. Some problems involving absent, non-existent, or counterfactual states are “representation-hungry” and are not easily amenable to non-representational explanation (Adams & Aizawa, [2008]; Clark & Toribio, [1994]; Wheeler, [2005]). Other problems involve causal webs that can be explained as dynamical systems but are too strongly connected to decouple representational units from their environment (Brooks et al., [1991]; Chemero, [2009]; Clark, [1998]).

In this thesis, we have attempted to follow a middle path that recognises both represented and unrepresented reasons and that allows “many theoretical flowers to bloom” (Chemero, [2009], p. 16). This pluralist approach allowed us to answer a broader range of questions than either of the individual routes described above. Furthermore, as we have discussed in this section, recognising that there are unrepresented reasons for our behaviour functions as a prophylactic against several inadequate arguments that people lack insight into their own mental processes (for a review, see Newell & Shanks, [2014]). These concepts can also be used to refine our understanding of the explanations offered by rational models. When they are recognised as offering unrepresented reasons, concepts such as probability have a much broader scope than the relatively minor role they play when we only consider represented reasons (Szollosi et al., [2022]).

7.3 Idealisation

In the opening paragraphs of the general introduction, we emphasised that there have been countless refutations of expected utility theory as a true description of human behaviour. When we were developing our rational model in Chapter 2, we discussed two hypothetical agents that possessed limited memory but that were only slightly more plausible than the already refuted *Homo economicus*. Across each of the experimental chapters, participants made repeated choices between options that would never be encountered outside the psychology laboratory. Each of these theoretical and experimental models has been falsified

in the Popperian sense but we suggested that rejecting them on this basis would be as misguided as debunking the frictionless planes of Galileo.²

This assertion is incompatible with the widespread belief that the proper aim of science is to produce a veridical description of nature. Surely, we should always prefer models that conform to reality as closely as possible but the undesirable consequences of this position were sketched out by Borges (1998) in his short story, *On exactitude in Science*:

...In that Empire, the Art of Cartography attained such Perfection that the map of a single Province occupied the entirety of a City, and the map of the Empire, the entirety of a Province. In time, those Unconscionable Maps no longer satisfied, and the Cartographers Guilds struck a Map of the Empire whose size was that of the Empire, and which coincided point for point with it. The following Generations, who were not so fond of the Study of Cartography as their Forebears had been, saw that that vast Map was Useless, and not without some Pitilessness was it, that they delivered it up to the Inclemencies of Sun and Winters. In the Deserts of the West, still today, there are Tattered Ruins of that Map, inhabited by Animals and Beggars; in all the Land there is no other Relic of the Disciplines of Geography. (p. 325)

The moral of the story is that building models is pointless unless we can use them. Cartographers unapologetically make maps that deviate from reality in their spatial scale and are conveniently made using paper and ink rather than soil, rocks, trees, and water.

²Experimental methods are transparently non-propositional and the reader might object that they possess no truth value. Although less obvious, the same objection also applies to the *Homo economicus* and our other examples. They should be interpreted as models, which Morgan (2012) describes as follows: models are “either real objects, or pen-and-paper objects that are diagrammatic, algebraic, or arithmetic in form. Despite their variations in form, these objects share recognisable characteristics: each depicts, renders, denotes, or in some way provides, some kind of representation of ideas about some aspects of [nature]” (p. 13). Thus, our claim that these models were falsified refers to the proposition that they represent—they are isomorphic with—their target system.

These inaccuracies are tolerated because maps allow us to achieve our aims. Their usefulness depends not only on whether they preserve the necessary structure of the landscape but also on the attributes of the limited beings who use them. This is the reason that the map is not the territory (Korzybski, 1958), that all models are wrong (Box, 1976), that *ceci n'est pas une pipe* (Magritte, 1929) and that the menu is not the meal (Watts, 1957). Otherwise, they would not serve their intended purposes.

Psychological theories disregard countless attributes of their target systems. They scarcely ever mention neurons or glial cells and even fewer incorporate quarks or bosons. This is a strength rather than a weakness. To illustrate this, consider the two explanations offered by Putnam (1979) for why a square peg with one-inch sides will not fit through a round hole that has a one-inch diameter. The first explanation recognises that the peg and the board that surrounds the hole are both systems of particles. It uses equations from quantum mechanics to deduce that the peg system does not pass through the hole region. In contrast, the second explanation simply identifies that the peg and the board are inelastic and that the cross-section of the peg is larger than the diameter of the hole.

Which of these explanations is better? Although the geometric explanation omits the ultimate constituents of the target system, it emphasises the features that are relevant to the person seeking an explanation. The relationship between the size and shape of rigid objects applies to items made of wood, plastic or steel whereas the quantum explanation only applies to systems that have the same arrangement of particles. Thus, the geometric explanation applies to a more interesting class of systems than the quantum explanation. This is precisely *because* it abstracts away distinctions between microstates that realise equivalent macroscopic structures or functions.³

³The concept of bounded rationality would be incomprehensible to someone who only acknowledges the quantum explanation. According to this perspective, even if there is indeterminacy at the quantum-level, our behaviour is more or less deterministic. This is not compatible with rational explanation because ascribing rationality to an action implies that the agent can perform the behaviour *and* that they can do otherwise. Rational behaviour is influenced by the constraints that we choose to ignore and is not simply acquired through “some metaphysical hotline to the nature of absolute truth” (Binmore, 2008, p. 2).

Upon recognising that models are influenced by the purposes and capacities of their creators, we are compelled to relinquish the assumption that science aims to discover a single unified model that would render the others obsolete. This is not the only notion regarding the aims of science that is called into question. Instead, the compromise between preserving relevant structures and ensuring tractability opens the door to theories that deliberately maintain false assumptions. The costs incurred by using a more accurate model can outweigh the benefits of eliminating falsehoods. As a result, achieving our scientific aims might require us to employ idealisation (Batterman, [2009](#); Elgin, [2017](#); Levins, [1966](#); Potochnik, [2017](#); Wimsatt & Wimsatt, [2007](#)).

In contrast with the notion that science always pursues the truth, idealised models are known to be false but they are neither removed nor consigned to the periphery of our theories (Potochnik, [2017](#)). Scientists routinely discuss physical systems without any friction or intermolecular forces, economic agents that have perfect knowledge, and evolutionary processes in an infinite population. These systems cannot be found anywhere in nature. Nonetheless, the relevant structure is preserved and idealised models support counterfactuals that are “true enough” for their intended purpose (Elgin, [2017](#)). They highlight relevant attributes that are concealed in non-idealised models, can explain why a target phenomenon is observed across a heterogeneous class of systems, and can be used to indicate whether an explanation is invariant to certain assumptions (Elgin, [2007](#); Rice, [2015](#); Wimsatt & Wimsatt, [2007](#)).

The concept of idealisation is essential for understanding the role of probability in this thesis. An enormous catalogue of inconsistencies between human behaviour and the predictions of expected utility theory have accumulated over the past 50 years. Our aim was not to falsify these theories or to suggest that they should be abandoned. “The art of model-building is the exclusion of real but irrelevant parts of the problem” (P. Anderson, [1977](#), p. 381). But this exposes us to the danger that we will neglect an aspect that is genuinely important to achieving our aims. Thus, the experiments in this thesis examined whether context and how people interpret uncertainty are essential in at least *some* of our models of decisions from experience.

7.4 Theory and experiment

The reciprocal relationship between theory and experiment was the glue that held the two main sections of this thesis together. One side of this relationship is that theories influence the questions that we pursue and the methods that we use. This occurs because scientists employ tasks that manipulate the important variables in their theories and avoid those that include extraneous variables. The other side is that experiments influence our theories both in their capacity to provide evidence and as a consequence of pragmatic constraints. Scientists abandon theories that are incompatible with their observations and neglect those that are not amenable to available experimental methods.

In decisions from experience, we argued that probabilistic theories are well-suited to bandit task experiments because the researcher can manipulate the probability of experiencing each outcome. These experiments return the favour by precluding evidence that would be incompatible with the probabilistic assumptions underlying the theories. In particular, there are pragmatic constraints regarding the number of options that can be experienced in a single experiment and this might obscure the impact that context has on our choices. Furthermore, bandit task experiments were inspired by gambling tasks that intentionally eliminate observable variability and are unable to differentiate between alternative interpretations of uncertainty.

The reciprocal nature of this relationship suggests that many probabilistic theories are impervious to refutation but the significance of their persistence depends on whether there is something hiding in the shadows. In the general introduction, we emphasised that we could not know whether there was something lurking there until *after* we conducted our experiments. This is a challenging epistemic position that is analogous to realising that some conspiracy theories turn out to be true. As summarised by Sunstein and Vermeule (2008):

The Watergate hotel room used by Democratic National Committee was, in fact, bugged by Republican officials, operating at the behest of the White

House. In the 1950s, the Central Intelligence Agency did, in fact, administer LSD and related drugs under Project MKULTRA, in an effort to investigate the possibility of “mind control.” Operation Northwoods, a rumored plan by the Department of Defense to simulate acts of terrorism and to blame them on Cuba, really was proposed by high-level officials (though the plan never went into effect). (p. 4)

The difficulty is that we are required to reason with censored evidence. Suppose that Neil Armstrong’s “great leap for mankind” was fabricated by a government desperate to beat the Russians in the space race, that Covid-19 escaped from the Wuhan Institute of Virology, or that the September 11 attacks were orchestrated as pretence for an oil-driven invasion. There is a plausible motivation for covering up each of these operations. There is little evidence but this is exactly what we would expect if there were a successful concealment effort by powerful individuals or government agencies.⁴

Admittedly, drawing parallels between your position and those of tin-foil-wearing flat-earther crackpots is a terrible rhetorical move so let us attempt to recover at least some of our credibility by considering a Bayesian perspective. Scientists are active participants in the creation of knowledge. They approach experiments with a *prior* probability $P(h)$ that conducting them will achieve their epistemic aims and some theory that specifies the *likelihood* $P(d|h)$ of observing data d given different degrees to which the experiment would accomplish those aims.

Where does censored evidence fit into this perspective? Suppose that a scientist is considering a research project that would establish the existence of an alien spacecraft on an airbase in the Nevada desert. Their search through the existing literature reveals

⁴Some people suggest that the implied existence of competent government agencies is enough to rule out most conspiracy theories. Similarly, Popper¹⁹⁴⁵ criticises the “conspiracy theory of society” that whatever happens “is the result of direct design by some powerful individuals and groups” (p. 306). It is worth noting that—analogously to our discussion of reasons without representation—the relationship between theory and experiment gives rise to censorship without conspiracy.

numerous studies that failed to produce any photographic evidence. In the absence of censorship, this observation is more *likely* given a hypothesis that there are only tumbleweeds in the desert than that there is an alien spacecraft. The *relative likelihood* of there being no photographic evidence becomes more similar, however, when we assume the existence of a classified military operation. In this case, we would not expect there to be positive evidence either way.⁵

Therefore, the question regarding alien spacecraft remains unanswered and the scientist might decide that their proposal to storm Area 51 is aligned with their epistemic aims. We presented a similar argument regarding the influence of bandit tasks on decisions from experience. How does this argument hold up in the light of our results? Our experiments clearly demonstrate that context influences how people make decisions. We evaluated whether that influence is driven by extreme outcomes. We also demonstrated that surface variability influences how people interpret uncertainty.

These observations are consistent with theory and experiment concealing the role of context and uncertainty. Should we consider ourselves vindicated? One compelling rebuttal is that we are ignoring existing research programmes in each of these areas. Multiple experiments have examined context in decisions from experience (Ert & Lejarraaga, 2018; Hadar et al., 2018; Ludvig et al., 2014; Spektor et al., 2018) and there is a long history associated with epistemic and aleatory uncertainty (Fox & Ülkümen, 2011; Goodnow, 1955; Hacking, 1975; Kahneman & Tversky, 1982). Does this render our claim simply an elaborate rationalisation used to unify the unrelated sections of this thesis? There could be some truth to this rebuttal, but even so, the relationship between theory and experiment might offer a valuable heuristic for guiding scientific discovery (Gigerenzer, 1991; Reichenbach, 1978).⁶

⁵Less dramatic forms of censored evidence influence the likelihood of observing data in many scientific domains from biased sampling in election forecasts (Keeter, 2006) to the mesh size of fishing nets in marine surveys (Hamley, 1975). Censorship also impacts the way people acquire and generalise concepts (Hayes et al., 2019; Ransom et al., 2022; Shafto et al., 2014; Tenenbaum & Griffiths, 2001).

⁶Numerous heuristics to develop and investigate scientific theories have been proposed.

To understand the nature of this heuristic, it will be useful to compare our account of theory and experiment with that of Hacking (1992). He posits a self-vindicating system in which scientists create apparatus to assess their theories and judge the correctness of their apparatus by its alignment with their theories. In contrast with our account, Hacking applies this relationship exclusively to what he calls “mature laboratory science”. The defining feature of these sciences is that the phenomena under study seldom or never occur before they are created in the laboratory. The phenomena cannot be separated from the apparatus used to generate them, and therefore, the theories are tested with reference to the created phenomena rather than the untamed world.

Notably, he emphasises that “although there is plenty of experimentation in sociology, psychology, and economics, not much of it is what I call laboratory science, not even when there is a university building called the psychology laboratory” (Hacking, 1992, p. 34). Were we mistaken in applying this concept to decisions from experience? It is possible that Hacking was simply unaware of how contingent most psychological phenomena are on contrived laboratory conditions. Indeed, until bandit task phenomena were created in the laboratory, “nowhere in the universe, so far as we know”, were undergraduate students instructed to repeatedly choose between coloured squares in order to earn points.

Regardless of whether psychological phenomena can be created in the laboratory, the narrower scope of Hacking’s concept reveals a meaningful distinction between our accounts. His description of self-vindication is based on the claim by Heisenberg (1948) that Newtonian mechanics and quantum physics are internally consistent “closed systems” of knowledge. Their consistency ensures stability because there is no conflict within and thus no pressure for revision. They “hold for all time” wherever phenomena can be described using the concepts of the theory, and although this domain of applicability is limited, the

They include using tools from statistics and computer science as metaphors for the mind (Gigerenzer, 1991) and examining mechanisms that are adaptive (Matsumoto, 2021; Relihan, 2012) or rational (J. R. Anderson, 1991; Lieder & Griffiths, 2019). These heuristics are an essential element of scientific practice. As recognised by Tooby and Cosmides (1998), “being guided towards hypotheses that are more likely to be true is critical: the difference between a living science and an inert one is whether practising scientists have good heuristic principles guiding their research” (p. 197).

components are so interconnected that they are not susceptible to further modification (Bokulich, 2006; Kuhn, 2012).

The exclusion of the social sciences makes sense when mature laboratory science is viewed as a closed system. Laboratory sciences create non-contextual systems of theories, experiments, and observations that are independent from the external world. In contrast, experiments in the social sciences examine people who change in response to historical and cultural factors. They describe complex systems that seldom allow the degree of internal consistency required of a closed system. For example, the concept of risk preferences is central to many theories of decision-making. Nonetheless, there are discrepancies across elicitation methods and this might prompt scientists to amend their theories or experiments (L. R. Anderson & Mellor, 2009; Pedroni et al., 2017).

The existence of research programmes examining context and uncertainty is not compatible with the “closed system” account of probabilistic theories and bandit task experiments. This account describes stability that emerges from the absence of opposition rather than stability that endures in its presence. Our explanation for the robustness of these systems was based on feedback loops. This is compatible with the closed systems account but there are numerous other descriptions of this concept. Feedback loops have been examined across numerous domains and they have accumulated myriad different labels including self-organisation, recursivity, complexity, homeostasis, reflexivity, nonlinearity, autopoiesis, and autocatalysis (e.g., Arthur, 2014; Chomsky, 1995; Godfrey-Smith, 1998; Hofstadter, 2007; Kauffman, 1993; Maturana & Varela, 1991).

The concept of niche construction in evolutionary biology offers a better analogy for the reciprocal relationship between theory and experiment (Lewontin, 1983; Odling-Smee et al., 2013). This concept diverges from the traditional interpretation of natural selection as a one-way process in which “[t]he organism proposes and the environment disposes” (Lewontin, 2000, p. 43). Instead, organisms are moulded by the selective pressures in their environments but are also constantly modifying aspects of those environments. They can contribute to the creation or destruction of their own and others’ ecological niches and the

environment should be viewed as evolving alongside the organisms that live there.

In the same way, the relationship between theory and experiment is not a one-way process of conjecture and refutation. They each contribute to the preservation or destruction of the other. They can form self-maintaining systems but these are not the hermetically sealed units mentioned in the closed systems account. They exist within a dynamic ecosystem that also includes creatures such as funding bodies, research ethics committees, academic journals, scientists' aspirations to advance their careers, historical or cultural influences on the phenomena of interest, and the desire to use scientific knowledge to change aspects of the external world.

Novel theories and experiments can be introduced to this system but when the existing occupants are well-adapted, the introduced species encounter a hostile environment. In decisions from experience, probabilistic theories and bandit task experiments have created a niche in which they mutually benefit from each other. Introduced theories must contend with the inability of bandit tasks to contradict the assumptions underlying probabilistic theories. Introduced experiments must justify the inclusion of variables that seem irrelevant based on these assumptions. This gives probabilistic theories a competitive edge over potential alternatives and could explain their stability despite experiments that have attempted to overthrow the established order (e.g., Ert & Lejarraaga, 2018; Ludvig et al., 2014).

This raises the question of what we have achieved in this thesis. Although we have made numerous valuable empirical contributions, our experiments could be interpreted as yet another invasive species attempting to enter a hostile niche. We have not resolved the pragmatic constraints associated with examining multiple options and future experiments will encounter the same challenges when attempting to do so. Nonetheless, the discovery heuristic identified in this section has the capacity to make the ecosystem inhospitable to theories that fail to serve our scientific aims. That is to say, recognising experimental methods that censor incompatible evidence can enable us to weed out theories and

experiments that only prevail as self-maintaining systems.⁷

7.5 Summary and conclusion

In this thesis, we have made numerous theoretical and empirical contributions to the understanding of context and uncertainty in decisions from experience. On the theoretical side, we clarified the distinction between rational and mechanical explanations, evaluated the contribution of the utility-weighted sampling model, established an alternative rational explanation, developed a framework to organise the numerous theories regarding extreme-outcome phenomena, illuminated the relationship between observable variability and different forms of uncertainty, and illustrated the role of numerous philosophical elements, such as representation, idealisation, and the relationship between theory and experiment.

On the empirical side, we described three experiments that provided compelling evidence against the categorical and edge-based theories of extreme outcomes. We described two experiments that examined whether extremity is determined using a type-based or token-based representation. We then described an experiment that essentially ruled out any extreme-outcome theory based on ordinal-level temporal peaks. Finally, we described three experiments that examined how observable variability influences the way that people interpret uncertainty and whether they seek additional information. Therefore, in all of these different ways, the theoretical and empirical components of this thesis have advanced our knowledge regarding decisions from experience and the experimental methods we use to study them.

⁷As we emphasised in our discussion of conspiracy theories, recognising the existence of censorship can either uncover genuine information or entrench prior beliefs. To avoid the latter, we should consider the influence that censorship has on the likelihood function but also examine whether there is other evidence, avoid strong conclusions in the absence of compelling evidence, and where possible, conduct further experiments.

References

- Abdellaoui, M. (2000). Parameter-free elicitation of utility and probability weighting functions. *Management Science*, 46(11), 1497–1512. <https://doi.org/10.1287/mnsc.46.11.1497.12080>
- Adams, F., & Aizawa, K. (2008). *The Bounds of Cognition*. John Wiley & Sons.
- Akerlof, G. A. (1970). The Market for "Lemons": Quality Uncertainty and the Market Mechanism. *The Quarterly Journal of Economics*, 84(3), 488–500. <https://doi.org/10.2307/1879431>
- Aldrovandi, S., Brown, G. D., & Wood, A. M. (2015). Social norms and rank-based nudging: Changing willingness to pay for healthy food. *Journal of Experimental Psychology: Applied*, 21(3), 242–254. <https://doi.org/10.1037/xap0000048>
- Aldrovandi, S., Wood, A. M., & Brown, G. D. (2013). Sentencing, severity, and social norms: A rank-based model of contextual influence on judgments of crimes and punishments. *Acta Psychologica*, 144(3), 538–547. <https://doi.org/10.1016/j.actpsy.2013.09.007>
- Allais, M. (1953). Le comportement de l'homme rationnel devant le risque: Critique des postulats et axiomes de l'ecole americaine. *Econometrica*, 21(4), 503. <https://doi.org/10.2307/1907921>
- Al-Najjar, N. I., & Weinstein, J. (2009). The ambiguity aversion literature: A critical assessment. *Economics and Philosophy*, 25(3), 249–284. <https://doi.org/10.1017/S026626710999023X>

REFERENCES

- Anderson, J. R. (1991). Is human cognition adaptive? *Behavioral and Brain Sciences*, 14(3), 471–485. <https://doi.org/10.1017/S0140525X00070801>
- Anderson, J. R., Bothell, D., Byrne, M. D., Douglass, S., Lebiere, C., & Qin, Y. (2004). An Integrated Theory of the Mind. *Psychological Review*, 111, 1036–1060. <https://doi.org/10.1037/0033-295X.111.4.1036>
- Anderson, J. R. (1990). *The Adaptive Character of Thought*. Psychology Press.
- Anderson, L. R., & Mellor, J. M. (2009). Are risk preferences stable? Comparing an experimental measure with a validated survey-based measure. *Journal of Risk and Uncertainty*, 39(2), 137–160. <https://doi.org/10.1007/s11166-009-9075-z>
- Anderson, M. C., & Neely, J. H. (1996). Interference and inhibition in memory retrieval. *Memory* (pp. 237–313). Academic Press. <https://doi.org/10.1016/b978-012102570-0/50010-0>
- Anderson, P. (1977). Local Moments and Localized States.
- Ariely, D. (1998). Combining experiences over time: The effects of duration, intensity changes and on-line measurements on retrospective pain evaluations. *Journal of Behavioral Decision Making*, 11(1), 19–45. [https://doi.org/10.1002/\(SICI\)1099-0771\(199803\)11:1<19::AID-BDM277>3.0.CO;2-B](https://doi.org/10.1002/(SICI)1099-0771(199803)11:1<19::AID-BDM277>3.0.CO;2-B)
- Ariely, D., & Carmon, Z. (2000). Gestalt characteristics of experiences: The defining features of summarized events. *Journal of Behavioral Decision Making*, 13(2), 191–201. [https://doi.org/10.1002/\(SICI\)1099-0771\(200004/06\)13:2<191::AID-BDM330>3.0.CO;2-A](https://doi.org/10.1002/(SICI)1099-0771(200004/06)13:2<191::AID-BDM330>3.0.CO;2-A)
- Ariely, D., & Zauberman, G. (2000). On the making of an experience: The effects of breaking and combining experiences on their overall evaluation. *Journal of Behavioral Decision Making*, 13(2), 219–232. [https://doi.org/10.1002/\(SICI\)1099-0771\(200004/06\)13:2<219::AID-BDM331>3.0.CO;2-P](https://doi.org/10.1002/(SICI)1099-0771(200004/06)13:2<219::AID-BDM331>3.0.CO;2-P)
- Ariely, D., & Zauberman, G. (2003). Differential partitioning of extended experiences. *Organizational Behavior and Human Decision Processes*, 91(2), 128–139. [https://doi.org/10.1016/S0749-5978\(03\)00061-X](https://doi.org/10.1016/S0749-5978(03)00061-X)
- Armstrong, D. M. (1989). *Universals: An opinionated introduction*. Taylor & Francis.
- Arrow, K. (1951). *Social choice and individual values*. Wiley: New York.

- Arthur, W. B. (2014). *Complexity and the Economy*. Oxford University Press.
- Ascough, J., Maier, H., Ravalico, J., & Strudley, M. (2008). Future research challenges for incorporation of uncertainty in environmental and ecological decision-making. *Ecological Modelling*, 219(3-4), 383–399. <https://doi.org/10.1016/j.ecolmodel.2008.07.015>
- Barrett, D. (2014). Functional analysis and mechanistic explanation. *Synthese*, 191(12), 2695–2714. <https://doi.org/10.1007/s11229-014-0410-9>
- Barron, G., & Erev, I. (2003). Small feedback-based decisions and their limited correspondence to description-based decisions. *Journal of Behavioral Decision Making*, 16(3), 215–233. <https://doi.org/10.1002/bdm.443>
- Barsalou, L. W., Huttenlocher, J., & Lamberts, K. (1998). Basing categorization on individuals and events. *Cognitive Psychology*, 36(3), 203–272. <https://doi.org/10.1006/cogp.1998.0687>
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2014). Fitting Linear Mixed-Effects Models using lme4. <https://doi.org/10.48550/arXiv.1406.5823>
- Batterman, R. W. (2009). Idealization and modeling. *Synthese*, 169(3), 427–446. <https://doi.org/10.1007/s11229-008-9436-1>
- Baucells, M., & Villasís, A. (2010). Stability of risk preferences and the reflection effect of prospect theory. *Theory and Decision*, 68(1-2), 193–211. <https://doi.org/10.1007/s11238-009-9153-3>
- Baumgartner, H., Sujan, M., & Padgett, D. (1997). Patterns of affective reactions to advertisements: The integration of moment-to-moment responses into overall judgments. *Journal of Marketing Research*, 34(2), 219–232.
- Bechara, A., Damasio, A. R., Damasio, H., & Anderson, S. W. (1994). Insensitivity to future consequences following damage to human prefrontal cortex. *Cognition*, 50(1-3), 7–15. [https://doi.org/10.1016/0010-0277\(94\)90018-3](https://doi.org/10.1016/0010-0277(94)90018-3)
- Bechtel, W., & Abrahamsen, A. (2005). Explanation: A mechanist alternative. *Studies in History and Philosophy of Science Part C: Studies in History and Philosophy of Biological and Biomedical Sciences*, 36(2), 421–441. <https://doi.org/10.1016/j.shpsc.2005.03.010>

REFERENCES

- Bechtel, W., & Shagrir, O. (2015). The Non-Redundant Contributions of Marr's Three Levels of Analysis for Explaining Information-Processing Mechanisms. *Topics in Cognitive Science*, 7(2), 312–322. <https://doi.org/10.1111/tops.12141>
- Beck, S. R., McColgan, K. L., Robinson, E. J., & Rowley, M. G. (2011). Imagining what might be: Why children underestimate uncertainty. *Journal of Experimental Child Psychology*, 110(4), 603–610. <https://doi.org/10.1016/j.jecp.2011.06.010>
- Becker, S. W., & Brownson, F. O. (1964). What price ambiguity? or the role of ambiguity in decision-making. *Journal of Political Economy*, 72(1), 62–73. <https://doi.org/10.1086/258854>
- Bedford, T., & Cooke, R. (2001). *Probabilistic risk analysis: Foundations and methods*. Cambridge University Press. <https://doi.org/10.1017/CBO9780511813597>
- Beer, R. D. (2000). Dynamical approaches to cognitive science. *Trends in Cognitive Sciences*, 4(3), 91–99. [https://doi.org/10.1016/S1364-6613\(99\)01440-0](https://doi.org/10.1016/S1364-6613(99)01440-0)
- Beirlant, J., Goegebeur, Y., Teugels, J., & Segers, J. (2004). *Statistics of Extremes: Theory and Applications*. John Wiley & Sons, Ltd. <https://doi.org/10.1002/0470012382>
- Berliner, J. E., Durlach, N. I., & Braida, L. D. (1977). Intensity perception. VII. Further data on roving-level discrimination and the resolution and bias edge effects. *The Journal of the Acoustical Society of America*, 61(6), 1577–1585. <https://doi.org/10.1121/1.381471>
- Binmore, K. (2008). *Rational Decisions*. Princeton University Press.
- Birnbaum, M. H. (2008). New paradoxes of risky decision making. *Psychological Review*, 115(2), 463–501. <https://doi.org/10.1037/0033-295X.115.2.463>
- Birnbaum, M. H., & Chavez, A. (1997). Tests of theories of decision making: Violations of branch independence and distribution independence. *Organizational Behavior and Human Decision Processes*, 71(2), 161–194. <https://doi.org/10.1006/obhd.1997.2721>
- Bogacz, R., Brown, E., Moehlis, J., Holmes, P., & Cohen, J. D. (2006). The physics of optimal decision making: A formal analysis of models of performance in two-alternative forced-choice tasks. *Psychological Review*, 113(4), 700–765. <https://doi.org/10.1037/0033-295X.113.4.700>

- Bokulich, A. (2006). Heisenberg meets kuhn: Closed theories and paradigms. *Philosophy of Science*, 73(1), 90–107. <https://doi.org/10.1086/510176>
- Boole, G. (1854). *An investigation of the laws of thought: On which are founded the mathematical theories of logic and probabilities*. Dover.
- Borges, J. L. (1998). *Collected Fictions*. Viking.
- Botev, Z., & Ridder, A. (2017). Variance Reduction. *Wiley StatsRef: Statistics Reference Online* (pp. 1–6). John Wiley & Sons, Ltd. <https://doi.org/10.1002/9781118445112.stat07975>
- Bower, G. H. (1971). Adaptation-level coding of stimuli and serial position effects. *Adaptation-level theory* (pp. 175–2010).
- Box, G. E. P. (1976). Science and Statistics. *Journal of the American Statistical Association*, 71(356), 791–799. <https://doi.org/10.1080/01621459.1976.10480949>
- Boyce, C. J., Brown, G. D., & Moore, S. C. (2010). Money and happiness: Rank of income, not income, affects life satisfaction. *Psychological Science*, 21(4), 471–475. <https://doi.org/10.1177/0956797610362671>
- Braida, L. D., Lim, J. S., Berliner, J. E., Durlach, N. I., Rabinowitz, W. M., & Purks, S. R. (1984). Intensity perception. XIII. Perceptual anchor model of context-coding. *The Journal of the Acoustical Society of America*, 76(3), 722–731. <https://doi.org/10.1121/1.391258>
- Brainard, D. H. (1997). The psychophysics toolbox. *Spatial Vision*, 10(4), 433–436. <https://doi.org/10.1163/156856897X00357>
- Brainerd, C. J., & Reyna, V. F. (1990). Gist is the grist: Fuzzy-trace theory and the new intuitionism. *Developmental Review*, 10(1), 3–47. [https://doi.org/10.1016/0273-2297\(90\)90003-M](https://doi.org/10.1016/0273-2297(90)90003-M)
- Brandom, R. B. (2019). *A Spirit of Trust: A Reading of Hegel's Phenomenology*. Harvard University Press.
- Brette, R. (2019). Is coding a relevant metaphor for the brain? *Behavioral and Brain Sciences*, 42, 215. <https://doi.org/10.1017/S0140525X19000049>

REFERENCES

- Bringsjord, S., & Govindarajulu, N. S. (2022). Artificial Intelligence. In E. N. Zalta & U. Nodelman (Eds.), *The Stanford Encyclopedia of Philosophy* (Fall 2022). Metaphysics Research Lab, Stanford University.
- Brooks, L. R., Norman, G. R., & Allen, S. W. (1991). Role of specific similarity in a medical diagnostic task. *Journal of Experimental Psychology: General*, 120(3), 278–287. <https://doi.org/10.1037/0096-3445.120.3.278>
- Brown, G. D., Chater, N., & Neath, I. (2008). Serial and free recall: Common effects and common mechanisms? A reply to murdock (2008). *Psychological Review*, 115(3), 781–785. <https://doi.org/10.1037/a0012563>
- Brown, G. D., Gardner, J., Oswald, A. J., & Qian, J. (2008). Does wage rank affect employees' well-being? *Industrial Relations*, 47(3), 355–389. <https://doi.org/10.1111/j.1468-232X.2008.00525.x>
- Brown, G. D., Gathergood, J., & Weber, J. (2017). Relative rank and life satisfaction: Evidence from US households. *SSRN Electronic Journal*, (February), 1–39. <https://doi.org/10.2139/ssrn.2912892>
- Brown, G. D., Neath, I., & Chater, N. (2002). A ratio model of scale invariant memory and identification. *Memory Lab Technical Report*, 1–94.
- Brown, G. D., Neath, I., & Chater, N. (2007). A temporal ratio model of memory. *Psychological Review*, 114(3), 539–576. <https://doi.org/10.1037/0033-295X.114.3.539>
- Brown, N. R. (1997). Context memory and the selection of frequency estimation strategies. *Journal of Experimental Psychology: Learning Memory and Cognition*, 23(4), 898–914. <https://doi.org/10.1037/0278-7393.23.4.898>
- Brun, W., & Teigen, K. H. (1990). Prediction and postdiction preferences in guessing. *Journal of Behavioral Decision Making*, 3(1), 17–28. <https://doi.org/10.1002/bdm.3960030103>
- Brunswik, E. (1939). Probability as a determiner of rat behavior. *Journal of Experimental Psychology*, 25(2), 175–197. <https://doi.org/10.1037/h0061204>
- Bürkner, P. C. (2017). Brms: An R package for bayesian multilevel models using stan. *Journal of Statistical Software*, 80(1). <https://doi.org/10.18637/jss.v080.i01>

- Bürkner, P. C., & Charpentier, E. (2020). Modelling monotonic effects of ordinal predictors in bayesian regression models. *British Journal of Mathematical and Statistical Psychology*, 73(3), 420–451. <https://doi.org/10.1111/bmsp.12195>
- Bush, R. R., & Mosteller, F. (1955). Symmetric choice problems. *Stochastic models for learning*. (pp. 274–309). John Wiley & Sons Inc. <https://doi.org/10.1037/14496-013>
- Butler, J. (1736). *The Analogy of Religion, Natural and Revealed, to the Constitution and Course of Nature*. Creative Media Partners, LLC.
- Camerer, C., & Weber, M. (1992). Recent developments in modeling preferences: Uncertainty and ambiguity. *Journal of Risk and Uncertainty*, 5(4), 325–370. <https://doi.org/10.1007/BF00122575>
- Capitani, E., Della Sala, S., Logie, R. H., & Spinnler, H. (1992). Recency, primacy, and memory: Reappraising and standardising the serial position curve. *Cortex*, 28(3), 315–342. [https://doi.org/10.1016/S0010-9452\(13\)80143-8](https://doi.org/10.1016/S0010-9452(13)80143-8)
- Carandini, M., & Heeger, D. J. (2012). Normalization as a canonical neural computation. *Nature Reviews Neuroscience*, 13(1), 51–62. <https://doi.org/10.1038/nrn3136>
- Carnap, R. (1945). The two concepts of probability: The problem of probability. *Philosophy and Phenomenological Research*, 5(4), 513. <https://doi.org/10.2307/2102817>
- Carpenter, B., Gelman, A., Hoffman, M. D., Lee, D., Goodrich, B., Betancourt, M., Brubaker, M., Guo, J., Li, P., & Riddell, A. (2017). Stan : A Probabilistic Programming Language. *Journal of Statistical Software*, 76(1). <https://doi.org/10.18637/jss.v076.i01>
- Carroll, L. (1896). *Alice’s Adventures in Wonderland*. H.M. Caldwell.
- Cartwright, N. (1983). *How the Laws of Physics Lie*. <https://doi.org/10.1093/0198247044.001.0001>
- Cartwright, N. (1999). The dappled world : A study of the boundaries of science. <https://doi.org/10.1017/cbo9781139167093>
- Cartwright, N., Pemberton, J., & Wieten, S. (2020). Mechanisms, laws and explanation. *European Journal for Philosophy of Science*, 10(3), 1–19. <https://doi.org/10.1007/s13194-020-00284-y>

REFERENCES

- Chajut, E., Caspi, A., Chen, R., Hod, M., & Ariely, D. (2014). In pain thou shalt bring forth children: The peak-and-end rule in recall of labor pain. *Psychological Science*, 25(12), 2266–2271. <https://doi.org/10.1177/0956797614551004>
- Chanales, A. J., Tremblay-McGaw, A. G., & Kuhl, B. A. (2020). Adaptive repulsion of long-term memory representations is triggered by event similarity. *bioRxiv*. <https://doi.org/10.1101/2020.01.14.900381>
- Charnov, E. L., & Orians, G. H. (1973). *Optimal foraging: Some theoretical explorations*.
- Chater, N., & Brown, G. D. (2008). From Universal Laws of Cognition to Specific Cognitive Models. *Cognitive Science*, 32(1), 36–67. <https://doi.org/10.1080/03640210701801941>
- Chemero, A. (2009). *Radical Embodied Cognitive Science*. MIT Press.
- Chomsky, N. (1995). *The Minimalist Program*. MIT Press.
- Chua Chow, C., & Sarin, R. K. (2002). Known, unknown, and unknowable uncertainties. *Theory and Decision*, 52(2), 127–138. <https://doi.org/10.1023/A:1015544715608>
- Clark, A. (1998). *Being There: Putting Brain, Body, and World Together Again*. MIT Press.
- Clark, A., & Toribio, J. (1994). Doing without representing? *Synthese*, 101(3), 401–431. <https://doi.org/10.1007/BF01063896>
- Cohen, D., & Erev, I. (2021). Over and under commitment to a course of action in decisions from experience. *Journal of Experimental Psychology: General*, 150(12), 2455–2471. <https://doi.org/10.1037/xge0001066>
- Cohen, L. J. (1981). Can human irrationality be experimentally demonstrated? *Behavioral and Brain Sciences*, 4(3), 317–370. <https://doi.org/10.1017/s0140525x00009092>
- Cojuharencu, I., & Ryvkin, D. (2008). Peak-end rule versus average utility: How utility aggregation affects evaluations of experiences. *Journal of Mathematical Psychology*, 52(5), 326–335. <https://doi.org/10.1016/j.jmp.2008.05.004>
- Cournot, A. A. (1843). *Exposition de la théorie des chances et des probabilités*.
- Craver, C. F. (2006). When mechanistic models explain. *Synthese*, 153(3), 355–376. <https://doi.org/10.1007/s11229-006-9097-x>

- Cummings, R. (2000). "How does it work?" vs. "What are the laws?" Two conceptions of psychological explanation. In F. C. Keil & R. A. Wilson (Eds.), *Explanation and cognition* (pp. 117–144). MIT Press.
- Curley, S. P., & Yates, J. (1989). An empirical evaluation of descriptive models of ambiguity reactions in choice situations. *Journal of Mathematical Psychology*, 33(4), 397–427. [https://doi.org/10.1016/0022-2496\(89\)90019-9](https://doi.org/10.1016/0022-2496(89)90019-9)
- Curley, S. P., Yates, J., & Abrams, R. A. (1986). Psychological sources of ambiguity avoidance. *Organizational Behavior and Human Decision Processes*, 38(2), 230–256. [https://doi.org/10.1016/0749-5978\(86\)90018-X](https://doi.org/10.1016/0749-5978(86)90018-X)
- Daston, L. (1995). *Classical Probability in the Enlightenment*. Princeton University Press.
- Daw, N. D., O'Doherty, J. P., Dayan, P., Seymour, B., & Dolan, R. J. (2006). Cortical substrates for exploratory decisions in humans. *Nature*, 441(7095), 876–879. <https://doi.org/10.1038/nature04766>
- Dawes, R. M., & Mulford, M. (1996). The False Consensus Effect and Overconfidence: Flaws in Judgment or Flaws in How We Study Judgment? *Organizational Behavior and Human Decision Processes*, 65(3), 201–211. <https://doi.org/10.1006/obhd.1996.0020>
- de Leeuw, J. R. (2015). jsPsych: A JavaScript library for creating behavioral experiments in a web browser. *Behavior Research Methods*, 47(1), 1–12. <https://doi.org/10.3758/s13428-014-0458-y>
- de Beauvoir, S. (1949). *The Second Sex*. Random House.
- Denison, S., Bonawitz, E., Gopnik, A., & Griffiths, T. L. (2013). Rational variability in children's causal inferences: The Sampling Hypothesis. *Cognition*, 126(2), 285–300. <https://doi.org/10.1016/j.cognition.2012.10.010>
- Dennett, D. C. (1983). Intentional systems in cognitive ethology: The “panglossian paradigm” defended. *Behavioral and Brain Sciences*, 6(3), 343–355. <https://doi.org/10.1017/S0140525X00016393>
- Denrell, J., & March, J. G. (2001). Adaptation as information restriction: The hot stove effect. *Organization Science*, 12(5), 523–538. <https://doi.org/10.1287/orsc.12.5.523.10092>

REFERENCES

- Deutsch, D. (2011). *The Beginning of Infinity: Explanations that Transform The World*. Penguin UK.
- Diaconis, P., & Efron, B. (1983). Computer-Intensive Methods in Statistics. *Scientific American*, 248(5), 116–130. <https://doi.org/10.1038/scientificamerican0583-116>
- Diecidue, E., & Wakker, P. P. (2001). On the intuition of rank-dependent utility. *Journal of Risk and Uncertainty*, 23(3), 281–298. <https://doi.org/10.1023/A:1011877808366>
- Do, A. M., Rupert, A. V., & Wolford, G. (2008). Evaluations of pleasurable experiences: The peak-end rule. *Psychonomic Bulletin and Review*, 15(1), 96–98. <https://doi.org/10.3758/PBR.15.1.96>
- Dougherty, M. R., Gettys, C. F., & Ogden, E. E. (1999). MINERVA-DM: A memory processes model for judgments of likelihood. *Psychological Review*, 106(1), 180–209. <https://doi.org/10.1037/0033-295X.106.1.180>
- Duane, S., Kennedy, A. D., Pendleton, B. J., & Roweth, D. (1987). Hybrid Monte Carlo. *Physics Letters B*, 195(2), 216–222. [https://doi.org/10.1016/0370-2693\(87\)91197-X](https://doi.org/10.1016/0370-2693(87)91197-X)
- Duhem, P. M. M. (1991). *The Aim and Structure of Physical Theory*. Princeton University Press.
- Dyson, F. (2004). A meeting with Enrico Fermi. *Nature*, 427(6972), 297–297. <https://doi.org/10.1038/427297a>
- Edwards, W. (1956). Reward probability, amount, and information as determiners of sequential two-alternative decisions. *Journal of Experimental Psychology*, 52(3), 177–188. <https://doi.org/10.1037/h0047727>
- Efron, B., & Morris, C. (1973). Stein's Estimation Rule and Its Competitors—An Empirical Bayes Approach. *Journal of the American Statistical Association*, 68(341), 117–130. <https://doi.org/10.2307/2284155>
- Einhorn, H. J., & Hogarth, R. M. (1985). Ambiguity and uncertainty in probabilistic inference. *Psychological Review*, 92(4), 433–461. <https://doi.org/10.1037/0033-295X.92.4.433>
- Elgin, C. (2007). Understanding and the facts. *Philosophical Studies*, 132(1), 33–42. <https://doi.org/10.1007/s11098-006-9054-z>

- Elgin, C. (2017). *True Enough*. MIT Press.
- Ellsberg, D. (1961). Risk, Ambiguity, and the Savage Axioms. *The Quarterly Journal of Economics*, 75(4), 643–669. <https://doi.org/10.2307/1884324>
- Ellsberg, D. (2001). *Risk, Ambiguity, and Decision*. Taylor & Francis.
- Elqayam, S., & Evans, J. S. B. (2011). Subtracting ought from is: Descriptivism versus normativism in the study of human thinking. *Behavioral and Brain Sciences*, 34(5), 233–248. <https://doi.org/10.1017/S0140525X1100001X>
- Erev, I., Ert, E., Roth, A. E., Haruvy, E., Herzog, S. M., Hau, R., Hertwig, R., Stewart, T., West, R., & Lebiere, C. (2010). A choice prediction competition: Choices from experience and from description. *Journal of Behavioral Decision Making*, 23(1), 15–47. <https://doi.org/10.1002/bdm.683>
- Erev, I., Ert, E., & Yechiam, E. (2008). Loss aversion, diminishing sensitivity, and the effect of experience on repeated decisions. *Journal of Behavioral Decision Making*, 21(5), 575–597. <https://doi.org/10.1002/bdm.602>
- Erev, I., & Greiner, B. (2015). The 1-800 critique, counterexamples, and the future of behavioral economics. In G. R. Fréchet & A. Schotter (Eds.), *Handbook of Experimental Economic Methodology* (pp. 151–165). Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780195328325.003.0010>
- Eriksen, C. W., & Hake, H. W. (1957). Anchor effects in absolute judgments. *Journal of Experimental Psychology*, 53(2), 132–138. <https://doi.org/10.1037/h0047421>
- Ert, E., & Haruvy, E. (2017). Revisiting risk aversion: Can risk preferences change with experience? *Economics Letters*, 151, 91–95. <https://doi.org/10.1016/j.econlet.2016.12.008>
- Ert, E., & Lejarraaga, T. (2018). The effect of experience on context-dependent decisions. *Journal of Behavioral Decision Making*. <https://doi.org/10.1002/bdm.2064>
- Farrell, S., & Lelièvre, A. (2009). End anchoring in short-term order memory. *Journal of Memory and Language*, 60(1), 209–227. <https://doi.org/10.1016/j.jml.2008.09.004>
- Favila, S. E., Chanals, A. J., & Kuhl, B. A. (2016). Experience-dependent hippocampal pattern differentiation prevents interference during subsequent learning. *Nature Communications*, 7. <https://doi.org/10.1038/ncomms11066>

REFERENCES

- Fawcett, T. W., Fallenstein, B., Higginson, A. D., Houston, A. I., Mallpress, D. E., Trimmer, P. C., & McNamara, J. M. (2014). The evolution of decision rules in complex environments. *Trends in Cognitive Sciences*, 18(3), 153–161. <https://doi.org/10.1016/j.tics.2013.12.012>
- Felin, T., Koenderink, J., & Krueger, J. I. (2017). Rationality, perception, and the all-seeing eye. *Psychonomic Bulletin and Review*, 24(4), 1040–1059. <https://doi.org/10.3758/s13423-016-1198-z>
- Fiedler, K. (2000). Beware of samples! A cognitive-ecological sampling approach to judgment biases. *Psychological Review*, 107(4), 659–676. <https://doi.org/10.1037/0033-295X.107.4.659>
- Fishman, G. (2013). *Monte Carlo: Concepts, Algorithms, and Applications*. Springer Science & Business Media.
- Fodor, J. A. (1974). Special sciences (or: The disunity of science as a working hypothesis). *Synthese*, 28(2), 97–115. <https://doi.org/10.1007/BF00485230>
- Fodor, J. A. (1981). *Representations: Philosophical Essays on the Foundations of Cognitive Science*. MIT Press.
- Fox, C. R., Goedde-Menke, M., & Tannenbaum, D. (2021). Ambiguity aversion and epistemic uncertainty. *SSRN Electronic Journal*, 1–43. <https://doi.org/10.2139/ssrn.3922716>
- Fox, C. R., & Tversky, A. (1995). Ambiguity aversion and comparative ignorance. *The Quarterly Journal of Economics*, 110(3), 585–603. <https://doi.org/10.2307/2946693>
- Fox, C. R., & Ülkümen, G. (2011). Distinguishing two dimensions of uncertainty. *Perspectives on Thinking, Judging, and Decision Making* (pp. 21–35).
- Fox, C. R., & Weber, M. (2002). Ambiguity aversion, comparative ignorance, and decision context. *Organizational Behavior and Human Decision Processes*, 88(1), 476–498. <https://doi.org/10.1006/obhd.2001.2990>
- Fredrickson, B. L. (2000). Extracting meaning from past affective experiences: The importance of peaks, ends, and specific emotions. *Cognition and Emotion*, 14(4), 577–606. <https://doi.org/10.1080/026999300402808>

- Fredrickson, B. L., & Kahneman, D. (1993). Duration neglect in retrospective evaluations of affective episodes. *Journal of Personality and Social Psychology*, 65(1), 45–55. <https://doi.org/10.1037/0022-3514.65.1.45>
- Frey, R., Mata, R., & Hertwig, R. (2015). The role of cognitive abilities in decisions from experience: Age differences emerge as a function of choice set size. *Cognition*, 142, 60–80. <https://doi.org/10.1016/j.cognition.2015.05.004>
- Frey, R., Pedroni, A., Mata, R., Rieskamp, J., & Hertwig, R. (2017). Risk preference shares the psychometric structure of major psychological traits. *Science Advances*, 3(10), e1701381. <https://doi.org/10.1126/sciadv.1701381>
- Friedman, M. (1953). *Essays in Positive Economics*. University of Chicago Press.
- Frisch, D., & Baron, J. (1988). Ambiguity and rationality. *Journal of Behavioral Decision Making*, 1(3), 149–157. <https://doi.org/10.1002/bdm.3960010303>
- Gabry, J., Simpson, D., Vehtari, A., Betancourt, M., & Gelman, A. (2019). Visualization in bayesian workflow. *Journal of the Royal Statistical Society. Series A: Statistics in Society*, 182(2), 389–402. <https://doi.org/10.1111/rssa.12378>
- Ganzach, Y., & Yaor, E. (2019). The retrospective evaluation of positive and negative affect. *Personality and Social Psychology Bulletin*, 45(1), 93–104. <https://doi.org/10.1177/0146167218780695>
- Geisler, W. S. (2011). Contributions of ideal observer theory to vision research. *Vision Research*, 51(7), 771–781. <https://doi.org/10.1016/j.visres.2010.09.027>
- Gelman, A. (2006). Prior distributions for variance parameters in hierarchical models (comment on article by browne and draper). *Bayesian Analysis*, 1(3), 515–534. <https://doi.org/10.1214/06-BA117A>
- Gelman, A. (2008). Scaling regression inputs by dividing by two standard deviations. *Statistics in Medicine*, 27(15), 2865–2873. <https://doi.org/10.1002/sim.3107>
- Gelman, A. (2018). What is probability?
- Gelman, R., & Gallistel, C. R. (2004). Language and the origin of numerical concepts. *Science*, 306(5695), 441–443. <https://doi.org/10.1126/science.1105144>

REFERENCES

- Geman, S., Bienenstock, E., & Doursat, R. (1992). Neural networks and the bias/variance dilemma. *Neural Computation*, 4(1), 1–58. <https://doi.org/10.1162/neco.1992.4.1.1>
- Geman, S., & Geman, D. (1987). Stochastic Relaxation, Gibbs Distributions, and the Bayesian Restoration of Images. In M. A. Fischler & O. Firschein (Eds.), *Readings in Computer Vision* (pp. 564–584). Morgan Kaufmann. <https://doi.org/10.1016/B978-0-08-051581-6.50057-X>
- Gibson, J. J. (1979). *The Ecological Approach to Visual Perception*. Psychology Press.
- Giere, R. N. (1999). *Science Without Laws*. University of Chicago Press.
- Gigerenzer, G. (1991). From tools to theories: A heuristic of discovery in cognitive psychology. *Psychological Review*, 98(2), 254–267. <https://doi.org/10.1037/0033-295X.98.2.254>
- Gigerenzer, G., & Brighton, H. (2009). Homo heuristics: Why biased minds make better inferences. *Topics in Cognitive Science*, 1(1), 107–143. <https://doi.org/10.1111/j.1756-8765.2008.01006.x>
- Gigerenzer, G., & Gaissmaier, W. (2011). Heuristic Decision Making. *Annual Review of Psychology*, 62(1), 451–482. <https://doi.org/10.1146/annurev-psych-120709-145346>
- Gigerenzer, G., Swijtink, Z., Porter, T., Daston, L., & Kruger, L. (1989). *The Empire of Chance: How Probability Changed Science and Everyday Life*. Cambridge University Press.
- Gigerenzer, G., & Todd, P. M. (1999). *Simple heuristics that make us smart*. Oxford University Press.
- Gillies, D. (2000). *Philosophical theories of probability*. Routledge.
- Glennan, S. (1996). Mechanisms and the nature of causation. *Erkenntnis*, 44(1), 49–71. <https://doi.org/10.1007/BF00172853>
- Glennan, S. (2017). *The New Mechanical Philosophy*. Oxford University Press.
- Godfrey-Smith, P. (1998). *Complexity and the Function of Mind in Nature*. Cambridge University Press.

- Goodnow, J. J. (1955). Determinants of choice-distribution in two-choice situations. *The American Journal of Psychology*, 68(1), 106. <https://doi.org/10.2307/1418393>
- Gould, S. J., & Lewontin, R. C. (1979). The Spandrels of San Marco and the Panglossian Paradigm: A Critique of the Adaptationist Programme. *Proceedings of the Royal Society of London. Series B, Biological Sciences*, 205(1161), 581–598.
- Graham, J. R., Harvey, C. R., & Huang, H. (2009). Investor competence, trading frequency, and home bias. *Management Science*, 55(7), 1094–1106. <https://doi.org/10.1287/mnsc.1090.1009>
- Green, S. B., & Yang, Y. (2009). Reliability of summed item scores using structural equation modeling: An alternative to coefficient alpha. *Psychometrika*, 74(1), 155–167. <https://doi.org/10.1007/s11336-008-9099-3>
- Greenland, S. (2000). Principles of multilevel modelling. *International Journal of Epidemiology*, 29(1), 158–167. <https://doi.org/10.1093/ije/29.1.158>
- Hacking, I. (1975). *The emergence of probability: A philosophical study of early ideas about probability, induction and statistical inference*. Cambridge University Press. <https://doi.org/10.1017/CBO9780511817557>
- Hacking, I. (1992). The Self-Vindication of the Laboratory Sciences. In A. Pickering (Ed.), *Science as Practice and Culture* (pp. 29–64). University of Chicago Press.
- Hacking, I., Hacking, E. U. P. I., & Hacking, T. (1990). *The Taming of Chance*. Cambridge University Press.
- Hadar, L., Danziger, S., & Hertwig, R. (2018). The attraction effect in experience-based decisions. *Journal of Behavioral Decision Making*, 31(3), 461–468. <https://doi.org/10.1002/bdm.2058>
- Hadar, L., Sood, S., & Fox, C. R. (2013). Subjective knowledge in consumer financial decisions. *Journal of Marketing Research*, 50(3), 303–316. <https://doi.org/10.1509/jmr.10.0518>
- Halevy, Y. (2007). Ellsberg revisited: An experimental study. *Econometrica*, 75(2), 503–536. <https://doi.org/10.1111/j.1468-0262.2006.00755.x>
- Hamley, J. M. (1975). Review of Gillnet Selectivity. *Journal of the Fisheries Research Board of Canada*, 32(11), 1943–1969. <https://doi.org/10.1139/f75-233>

REFERENCES

- Hargreaves, E. A., & Stych, K. (2013). Exploring the peak and end rule of past affective episodes within the exercise context. *Psychology of Sport and Exercise*, 14(2), 169–178. <https://doi.org/10.1016/j.psychsport.2012.10.003>
- Harris, A. J., Rowley, M. G., Beck, S. R., Robinson, E. J., & McColgan, K. L. (2011). Agency affects adults', but not children's, guessing preferences in a game of chance. *Quarterly Journal of Experimental Psychology*, 64(9), 1772–1787. <https://doi.org/10.1080/17470218.2011.582126>
- Harris, J. F., & Lippman, D. (2020). Trump drops out. Biden gets sick. Pence is fired. What if 2020 gets really crazy?
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). *The elements of statistical learning: Data mining, inference, and prediction, second edition* (Vol. Second edi). Springer.
- Hastings, W. K. (1970). Monte Carlo sampling methods using Markov chains and their applications. *Biometrika*, 57(1), 97–109. <https://doi.org/10.1093/biomet/57.1.97>
- Hau, R., Pleskac, T. J., & Hertwig, R. (2010). Decisions from experience and statistical probabilities: Why they trigger different choices than a priori probabilities. *Journal of Behavioral Decision Making*, 23(1), 48–68. <https://doi.org/10.1002/bdm.665>
- Hayes, B. K., Banner, S., Forrester, S., & Navarro, D. J. (2019). Selective sampling and inductive inference: Drawing inferences based on observed and missing evidence. *Cognitive Psychology*, 113, 101221. <https://doi.org/10.1016/j.cogpsych.2019.05.003>
- Heath, C., & Tversky, A. (1991). Preference and belief: Ambiguity and competence in choice under uncertainty. *Journal of Risk and Uncertainty*, 4(1), 5–28. <https://doi.org/10.1007/BF00057884>
- Hebart, M. N., Zheng, C. Y., Pereira, F., & Baker, C. I. (2020). Revealing the multidimensional mental representations of natural objects underlying human similarity judgements. *Nature Human Behaviour*, 4(11), 1173–1185. <https://doi.org/10.1038/s41562-020-00951-3>
- Heck, R., & Thomas, S. L. (1999). *An Introduction to Multilevel Modeling Techniques: MLM and SEM Approaches* (Fourth). Routledge. <https://doi.org/10.4324/9780429060274>

- Heeger, D. J. (1992). Normalization of cell responses in cat striate cortex. *Visual Neuroscience*, 9(2), 181–197. <https://doi.org/10.1017/S0952523800009640>
- Hegel, G. W. F. (1807). *The Phenomenology of Spirit*. University of Notre Dame Press.
- Heisenberg, W. (1948). *Across the frontiers*. Harper & Row; 1st edition.
- Helton, J. C., Johnson, J. D., Oberkamp, W. L., & Sallaberry, C. J. (2010). Representation of analysis results involving aleatory and epistemic uncertainty. *International Journal of General Systems*, 39(6), 605–646. <https://doi.org/10.1080/03081079.2010.486664>
- Hempel, C. G., & Oppenheim, P. (1948). Studies in the Logic of Explanation. *Philosophy of Science*, 15(2), 135–175.
- Henson, R. N. A. (1998). Short-term memory for serial order. *Cognitive Psychology*, 36(2), 73–137.
- Hertwig, R., Barron, G., Weber, E. U., & Erev, I. (2004). Decisions from experience and the effect of rare events in risky choice. *Psychological Science*, 15(8), 534–539. <https://doi.org/10.1111/j.0956-7976.2004.00715.x>
- Hertwig, R., & Gigerenzer, G. (1999). The ‘conjunction fallacy’ revisited: How intelligent inferences look like reasoning errors. *Journal of Behavioral Decision Making*, 12(4), 275–305. [https://doi.org/10.1002/\(SICI\)1099-0771\(199912\)12:4<275::AID-BDM323>3.0.CO;2-M](https://doi.org/10.1002/(SICI)1099-0771(199912)12:4<275::AID-BDM323>3.0.CO;2-M)
- Hertwig, R., Pachur, T., & Kurzenhäuser, S. (2005). Judgments of Risk Frequencies: Tests of Possible Cognitive Mechanisms. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 31(4), 621–642. <https://doi.org/10.1037/0278-7393.31.4.621>
- Hintzman, D. L. (1976). Repetition and memory. *Psychology of Learning and Motivation - Advances in Research and Theory*, 10(100), 47–91. [https://doi.org/10.1016/S0079-7421\(08\)60464-8](https://doi.org/10.1016/S0079-7421(08)60464-8)
- Hitch, G. J., Burgess, N., Towse, J. N., & Culpin, V. (1996). Temporal grouping effects in immediate recall: A working memory analysis. *Quarterly Journal of Experimental Psychology Section A: Human Experimental Psychology*, 49(1), 116–139. <https://doi.org/10.1080/713755609>

REFERENCES

- Hoffman, F. O., & Hammonds, J. S. (1994). Propagation of uncertainty in risk assessments: The need to distinguish between uncertainty due to lack of knowledge and uncertainty due to variability. *Risk Analysis*, 14(5), 707–712. <https://doi.org/10.1111/j.1539-6924.1994.tb00281.x>
- Hofstadter, D. R. (2007). *I Am a Strange Loop*. Hachette UK.
- Holt, C. A., & Laury, S. K. (2005). Risk aversion and incentive effects: New data without order effects. *American Economic Review*, 95(3), 902–904. <https://doi.org/10.1257/0002828054201459>
- Holyoak, K. J. (1978). Comparative judgments with numerical reference points. *Cognitive Psychology*, 10(2), 203–243. [https://doi.org/10.1016/0010-0285\(78\)90014-2](https://doi.org/10.1016/0010-0285(78)90014-2)
- Holzmeister, F., & Stefan, M. (2021). The risk elicitation puzzle revisited: Across-methods (in)consistency? *Experimental Economics*, 24(2), 593–616. <https://doi.org/10.1007/s10683-020-09674-8>
- Hsee, C. K., & Abelson, R. P. (1991). "Velocity relation: Satisfaction as a function of the first derivative of outcome over time": Correction. *Journal of Personality and Social Psychology*, 60(6), 951–951. <https://doi.org/10.1037/0022-3514.60.6.951>
- Hsee, C. K., Abelson, R. P., & Salovey, P. (1991). The relative weighting of position and velocity in satisfaction. *Psychological Science*, 2(4), 263–266. <https://doi.org/10.1111/j.1467-9280.1991.tb00146.x>
- Hsu, C. F., Propp, L., Panetta, L., Martin, S., Dentakos, S., Toplak, M. E., & Eastwood, J. D. (2018). Mental effort and discomfort: Testing the peak-end effect during a cognitively demanding task. *PLoS ONE*, 13(2), 1–17. <https://doi.org/10.1371/journal.pone.0191479>
- Huber, J., Payne, J. W., & Puto, C. (1982). Adding asymmetrically dominated alternatives: Violations of regularity and the similarity hypothesis. *Journal of Consumer Research*, 9(1), 90. <https://doi.org/10.1086/208899>
- Hui, S. K., Meyvis, T., & Assael, H. (2014). Analyzing moment-to-moment data using a bayesian functional linear model: Application to TV show pilot testing. *Marketing Science*, 33(2), 222–240. <https://doi.org/10.1287/mksc.2013.0835>

- Hulbert, J. C., & Norman, K. A. (2015). Neural differentiation tracks improved recall of competing memories following interleaved study and retrieval practice. *Cerebral Cortex*, 25(10), 3994–4008. <https://doi.org/10.1093/cercor/bhu284>
- Hume, D. (1748). *An Inquiry Concerning Human Understanding*. J.B. Bebbington.
- Humphrey, G. (1951). *Thinking - An Introduction to Its Experimental Psychology*. Read Books.
- Hunt, R. R. (1995). The subtlety of distinctiveness: What von Restorff really did. *Psychonomic Bulletin & Review*, 2(1), 105–112. <https://doi.org/10.3758/BF03214414>
- Jaffray, J. Y. (1988). Choice under risk and the security factor: An axiomatic model. *Theory and Decision*, 24(2), 169–200. <https://doi.org/10.1007/BF00132460>
- Jarvik, M. E. (1951). Probability learning and a negative recency effect in the serial anticipation of alternative symbols. *Journal of Experimental Psychology*, 41(4), 291–297. <https://doi.org/10.1037/h0056878>
- Johnson-Laird, P. N. (1983). *Mental Models: Towards a Cognitive Science of Language, Inference, and Consciousness*. Harvard University Press.
- Jones, M., & Love, B. C. (2011). Bayesian fundamentalism or enlightenment? On the explanatory status and theoretical contributions of bayesian models of cognition. *Behavioral and Brain Sciences*, 34(4), 169–188. <https://doi.org/10.1017/S0140525X10003134>
- Jou, J. (2010). The serial position, distance, and congruity effects of reference point setting in comparative judgments. *American Journal of Psychology*, 123(2), 127–136. <https://doi.org/10.5406/amerjpsyc.123.2.0127>
- Juanchich, M., Gourdon-Kanhukamwe, A., & Sirota, M. (2017). “I am uncertain” vs “it is uncertain”: How linguistic markers of the uncertainty source affect uncertainty communication. *Judgment and Decision Making*, 12(5), 445–465.
- Kahneman, D., Fredrickson, B. L., Schreiber, C. A., & Redelmeier, D. A. (1993). When more pain is preferred to less: Adding a better end. *Psychological Science*, 4(6), 401–405. <https://doi.org/10.1111/j.1467-9280.1993.tb00589.x>
- Kahneman, D., & Tversky, A. (1979). Prospect theory: An analysis of decision under risk. *Econometrica*, 47(2), 263. <https://doi.org/10.2307/1914185>

REFERENCES

- Kahneman, D., & Tversky, A. (1982). Variants of uncertainty. *Judgment under Uncertainty* (pp. 509–520). Cambridge University Press. <https://doi.org/10.1017/CBO9780511809477.036>
- Kant, I. (1996). *Critique of Pure Reason*. Hackett Publishing.
- Kaplan, D. M. (2017). *Explanation and Integration in Mind and Brain Science*. Oxford University Press.
- Kaplan, D. M., & Craver, C. F. (2011). The explanatory force of dynamical and mathematical models in neuroscience: A mechanistic perspective. *Philosophy of Science*, 78(4), 601–627. <https://doi.org/10.1086/661755>
- Kareev, Y. (1995). Through a narrow window: Working memory capacity and the detection of covariation. *Cognition*, 56(3), 263–269. [https://doi.org/10.1016/0010-0277\(95\)92814-G](https://doi.org/10.1016/0010-0277(95)92814-G)
- Kauffman, S. (1993). *The Origins of Order: Self-organization and Selection in Evolution*. Oxford University Press.
- Kauffman, S., & Levin, S. (1987). Towards a general theory of adaptive walks on rugged landscapes. *New York*, 128(1), 11–45. [https://doi.org/10.1016/S0022-5193\(87\)80029-2](https://doi.org/10.1016/S0022-5193(87)80029-2)
- Keas, M. N. (2018). Systematizing the theoretical virtues. *Synthese*, 195(6), 2761–2793. <https://doi.org/10.1007/s11229-017-1355-6>
- Keeter, S. (2006). The Impact of Cell Phone Noncoverage Bias on Polling in the 2004 Presidential Election. *Public Opinion Quarterly*, 70(1), 88–98. <https://doi.org/10.1093/poq/nfj008>
- Kelley, K. (2007). Methods for the Behavioral, Educational, and Social Sciences: An R package. *Behavior Research Methods*, 39(4), 979–984. <https://doi.org/10.3758/BF03192993>
- Kelley, K., & Pornprasertmanit, S. (2016). Confidence intervals for population reliability coefficients: Evaluation of methods, recommendations, and software for composite measures. *Psychological Methods*, 21(1), 69–92. <https://doi.org/10.1037/a0040086>

- Kelley, M. R., Neath, I., & Surprenant, A. M. (2015). Serial position functions in general knowledge. *Journal of Experimental Psychology: Learning Memory and Cognition*, 41(6), 1715–1726. <https://doi.org/10.1037/xlm0000141>
- Kemp, S., Burt, C. D., & Furneaux, L. (2008). A test of the peak-end rule with extended autobiographical events. *Memory and Cognition*, 36(1), 132–138. <https://doi.org/10.3758/MC.36.1.132>
- Keren, G. (1991). Calibration and probability judgements: Conceptual and methodological issues. *Acta Psychologica*, 77(3), 217–273. [https://doi.org/10.1016/0001-6918\(91\)90036-Y](https://doi.org/10.1016/0001-6918(91)90036-Y)
- Keren, G., & Gerritsen, L. E. (1999). On the robustness and possible accounts of ambiguity aversion. *Acta Psychologica*, 103(1-2), 149–172. [https://doi.org/10.1016/s0001-6918\(99\)00034-7](https://doi.org/10.1016/s0001-6918(99)00034-7)
- Kesner, R. P., & Novak, J. M. (1982). Serial position curve in rats: Role of the dorsal hippocampus. *Science*, 218(4568), 173–175. <https://doi.org/10.1126/science.7123228>
- Keynes, J. M. (1921). *A Treatise on Probability*. Courier Corporation.
- Kim, K. W., & Horel, A. (1998). Matrophagy in the spider *amaurobius ferox* (araneidae, amaurobiidae): An example of mother-offspring interactions. *Ethology*, 104(12), 1021–1037. <https://doi.org/10.1111/j.1439-0310.1998.tb00050.x>
- Kim, K. W., Roland, C., & Horel, A. (2000). Functional value of matrophagy in the spider *amaurobius ferox*. *Ethology*, 106(8), 729–742. <https://doi.org/10.1046/j.1439-0310.2000.00585.x>
- Kitcher, P. (1981). Explanatory Unification. *Philosophy of Science*, 48(4), 507–531.
- Kleiner, M., Brainard, D., & Pelli, D. (2007). What's new in psychtoolbox-3? *Perception* 36 ECVF Abstract Supplement.
- Knight, F. H. (H. (1921). *Risk, uncertainty and profit*. Boston, New York, Houghton Mifflin Company.
- Konstantinidis, E., Taylor, R. T., & Newell, B. R. (2017). Magnitude and incentives: Revisiting the overweighting of extreme events in risky decisions from experience.

REFERENCES

- Psychonomic Bulletin and Review*, 1–9. <https://doi.org/10.3758/s13423-017-1383-8>
- Koppenol-Gonzalez, G. V., Bouwmeester, S., & Vermunt, J. K. (2014). Short-term memory development: Differences in serial position curves between age groups and latent classes. *Journal of Experimental Child Psychology*, 126, 138–151. <https://doi.org/10.1016/j.jecp.2014.04.002>
- Korzybski, A. (1958). *Science and Sanity: An Introduction to Non-Aristotelian Systems and General Semantics*. Institute of GS.
- Kosslyn, S. M. (1980). *Image and Mind*. Harvard University Press.
- Krijnen, J., Ulkumen, G., Bogard, J., & Fox, C. R. (2020). Lay beliefs about changes in financial well-being predict political and policy message preferences. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.3695322>
- Kroese, D. P., Taimre, T., & Botev, Z. I. (2013). *Handbook of Monte Carlo Methods*. John Wiley & Sons.
- Kruschke, J. K. (1992). ALCOVE: An exemplar-based connectionist model of category learning. *Psychological Review*, 99(1), 22–44. <https://doi.org/10.1037/0033-295X.99.1.22>
- Kuhn, T. S. (2012). *The Structure of Scientific Revolutions: 50th Anniversary Edition*. University of Chicago Press.
- Kunar, M. A., Watson, D. G., Tsetsos, K., & Chater, N. (2017). The influence of attention on value integration. *Attention, Perception, and Psychophysics*, 79(6), 1615–1627. <https://doi.org/10.3758/s13414-017-1340-7>
- Kyburg, H. E. (1983). Rational belief. *Behavioral and Brain Sciences*, 6(2), 231–245. <https://doi.org/10.1017/S0140525X00015661>
- Lacouture, Y., & Marley, A. A. (2004). Choice and response time processes in the identification and categorization of unidimensional stimuli. *Perception and Psychophysics*, 66(7), 1206–1226. <https://doi.org/10.3758/BF03196847>
- Lakatos, I. (1978). *The Methodology of Scientific Research Programmes: Philosophical Papers* (J. Worrall & G. Currie, Eds.; Vol. 1). Cambridge University Press. <https://doi.org/10.1017/CBO9780511621123>

- Lange, M. (2000). *Natural Laws in Scientific Practice*. Oxford University Press.
- Langer, T., Sarin, R., & Weber, M. (2005). The retrospective evaluation of payment sequences: Duration neglect and peak-and-end effects. *Journal of Economic Behavior and Organization*, 58(1), 157–175. <https://doi.org/10.1016/j.jebo.2004.01.001>
- Laplace, P.-S. (1814). *A Philosophical Essay on Probabilities*. Courier Corporation.
- Lawson, T. (1988). Probability and uncertainty in economic analysis. *Journal of Post Keynesian Economics*, 11(1), 38–65. <https://doi.org/10.1080/01603477.1988.11489724>
- Levins, R. (1966). The Strategy of Model Building in Population Biology. *American Scientist*, 54(4), 421–431.
- Lewandowski, D., Kurowicka, D., & Joe, H. (2009). Generating random correlation matrices based on vines and extended onion method. *Journal of Multivariate Analysis*, 100(9), 1989–2001. <https://doi.org/10.1016/j.jmva.2009.04.008>
- Lewontin, R. C. (1983). Gene, organism, and environment. In D. S. Bendall (Ed.), *Evolution From Molecules to Men* (pp. 273–286). CUP Archive.
- Lewontin, R. C. (2000). *The Triple Helix: Gene, Organism, and Environment*. Harvard University Press.
- Li, S.-C., & Lewandowsky, S. (1995). Forward and backward recall: Different retrieval processes. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21(4), 837–847. <https://doi.org/10.1037/0278-7393.21.4.837>
- Lichtenstein, S., & Slovic, P. (1971). Reversals of preference between bids and choices in gambling decisions. *Journal of Experimental Psychology*, 89(1), 46–55. <https://doi.org/10.1037/h0031207>
- Lichtenstein, S., Slovic, P., Fischhoff, B., Layman, M., & Combs, B. (1978). Judged frequency of lethal events. *Journal of Experimental Psychology: Human Learning and Memory*, 4(6), 551–578. <https://doi.org/10.1037/0278-7393.4.6.551>
- Lieder, F., & Griffiths, T. L. (2019). Resource-rational analysis: Understanding human cognition as the optimal use of limited computational resources. *Behavioral and Brain Sciences*. <https://doi.org/10.1017/S0140525X1900061X>

REFERENCES

- Lieder, F., Griffiths, T. L., & Hsu, M. (2018). Overrepresentation of extreme events in decision making reflects rational use of cognitive resources. *Psychological Review*, 125(1), 1–32. <https://doi.org/10.1037/rev0000074>
- Loewenstein, G. (2007). *Exotic preferences: Behavioral economics and human motivation*. OUP Oxford.
- Loewenstein, G., & Prelec, D. (1993). Preferences for sequences of outcomes. *Psychological Review*, 100(1), 91–108. <https://doi.org/10.1037/0033-295X.100.1.91>
- Løhre, E., & Teigen, K. H. (2016). There is a 60% probability, but I am 70% certain: Communicative consequences of external and internal expressions of uncertainty. *Thinking & Reasoning*, 22(4), 369–396. <https://doi.org/10.1080/13546783.2015.1069758>
- Lopes, L. L. (1983). Some thoughts on the psychological concept of risk. *Journal of Experimental Psychology: Human Perception and Performance*, 9(1), 137–144. <https://doi.org/10.1037/0096-1523.9.1.137>
- Lopes, L. L., & Oden, G. C. (1999). The role of aspiration level in risky choice: A comparison of cumulative prospect theory and SP/A theory. *Journal of Mathematical Psychology*, 43(2), 286–313. <https://doi.org/10.1006/jmps.1999.1259>
- Louie, K., Khaw, M. W., & Glimcher, P. W. (2013). Normalization is a general neural mechanism for context-dependent decision making. *Proceedings of the National Academy of Sciences*, 110(15), 6139–6144. <https://doi.org/10.1073/pnas.1217854110>
- Luce, R. D. (1959). *Individual choice behavior*. John Wiley.
- Luce, R. D. (1977). The choice axiom after twenty years. *Journal of Mathematical Psychology*, 15(3), 215–233. [https://doi.org/10.1016/0022-2496\(77\)90032-3](https://doi.org/10.1016/0022-2496(77)90032-3)
- Luce, R. D., Green, D. M., & Weber, D. L. (1976). Attention bands in absolute identification. *Perception & Psychophysics*, 20(1), 49–54. <https://doi.org/10.3758/BF03198705>
- Luce, R. D., Nosofsky, R. M., Green, D. M., & Smith, A. F. (1982). The bow and sequential effects in absolute identification. *Perception & Psychophysics*, 32(5), 397–408. <https://doi.org/10.3758/BF03202769>

- Luce, R. D., & Suppes, P. (1956). Preference, utility, and subjective probability. In R. D. Luce, R. R. Bush, & E. Galanter (Eds.), *Handbook of mathematical psychology* (pp. 249–410). Wiley.
- Luce, R. D., & Suppes, P. (1965). Preference, utility, and subjective probability. *Handbook of Mathematical Psychology* (pp. 249–410).
- Ludvig, E. A., Madan, C. R., McMillan, N., Xu, Y., & Spetch, M. L. (2018). Living near the edge: How extreme outcomes and their neighbors drive risky choice. *Journal of Experimental Psychology: General*, 147(12), 1905–1918. <https://doi.org/10.1037/xge0000414>
- Ludvig, E. A., Madan, C. R., & Spetch, M. L. (2014). Extreme outcomes sway risky decisions from experience. *Journal of Behavioral Decision Making*, 27(2), 146–156. <https://doi.org/10.1002/bdm.1792>
- Ludvig, E. A., & Spetch, M. L. (2011). Of black swans and tossed coins: Is the description-experience gap in risky choice limited to rare events? *PLoS ONE*, 6(6). <https://doi.org/10.1371/journal.pone.0020262>
- Luengo, D., Martino, L., Bugallo, M., Elvira, V., & Särkkä, S. (2020). A survey of Monte Carlo methods for parameter estimation. *EURASIP Journal on Advances in Signal Processing*, 2020(1), 25. <https://doi.org/10.1186/s13634-020-00675-6>
- Maccrimmon, K. R. (1968). Descriptive and normative implications of the decision-theory postulates. *Risk and Uncertainty*, (102), 3–32. https://doi.org/10.1007/978-1-349-15248-3_1
- Machamer, P., Darden, L., & Craver, C. F. (2000). Thinking about Mechanisms. *Philosophy of Science*, 67(1), 1–25.
- Madan, C. R., Ludvig, E. A., & Spetch, M. L. (2014). Remembering the best and worst of times: Memories for extreme outcomes bias risky decisions. *Psychonomic Bulletin & Review*, 21(3), 629–636. <https://doi.org/10.3758/s13423-013-0542-9>
- Madan, C. R., Ludvig, E. A., & Spetch, M. L. (2017). The role of memory in distinguishing risky decisions from experience and description. *Quarterly Journal of Experimental Psychology*, 70(10), 2048–2059. <https://doi.org/10.1080/17470218.2016.1220608>

REFERENCES

- Madan, C. R., Spetch, M. L., & Ludvig, E. A. (2015). Rapid makes risky: Time pressure increases risk seeking in decisions from experience. *Journal of Cognitive Psychology*, 27(8), 921–928. <https://doi.org/10.1080/20445911.2015.1055274>
- Madan, C. R., Spetch, M. L., Machado, F. M., Mason, A., & Ludvig, E. A. (2021). Encoding context determines risky choice. *Psychological Science*, 32(5), 743–754. <https://doi.org/10.1177/0956797620977516>
- Magritte, R. (1929). The Treachery of Images (This is Not a Pipe).
- Makowski, D., Ben-Shachar, M., & Lüdecke, D. (2019). bayestestR: Describing effects and their uncertainty, existence and significance within the bayesian framework. *Journal of Open Source Software*, 4(40), 1541. <https://doi.org/10.21105/joss.01541>
- Makowski, D., Ben-Shachar, M. S., Chen, S. H., & Lüdecke, D. (2019). Indices of effect existence and significance in the bayesian framework. *Frontiers in Psychology*, 10(December), 1–14. <https://doi.org/10.3389/fpsyg.2019.02767>
- Maltby, J., Wood, A. M., Vlaev, I., Taylor, M. J., & Brown, G. D. (2012). Contextual effects on the perceived health benefits of exercise: The exercise rank hypothesis. *Journal of Sport and Exercise Psychology*, 34(6), 828–841. <https://doi.org/10.1123/jsep.34.6.828>
- Marcus, G. (2008). *Kluge: The Haphazard Construction of the Human Mind*. Faber & Faber.
- Marley, A. A., & Cook, V. T. (1984). A fixed rehearsal capacity interpretation of limits on absolute identification performance. *British Journal of Mathematical and Statistical Psychology*, 37(2), 136–151. <https://doi.org/10.1111/j.2044-8317.1984.tb00797.x>
- Marr, D. (1982). *Vision: A Computational Investigation into the Human Representation and Processing of Visual Information*. MIT Press.
- Mason, A., Madan, C. R., Ludvig, E. A., & Spetch, M. L. (2020). Biased confabulation in risky choice. <https://doi.org/10.31234/osf.io/vphgc>
- Matsumoto, S. (2021). Making sense of the relationship between adaptive thinking and heuristics in evolutionary psychology. *Biological Theory*, 16(1), 16–29. <https://doi.org/10.1007/s13752-020-00369-0>

- Maturana, H. R., & Varela, F. J. (1991). *Autopoiesis and Cognition: The Realization of the Living*. Springer Science & Business Media.
- Mayo, D. G. (2018). *Statistical Inference as Severe Testing: How to Get Beyond the Statistics Wars*. Cambridge University Press.
- McMullin, E. (1985). Galilean idealization. *Studies in History and Philosophy of Science Part A*, 16(3), 247–273. [https://doi.org/10.1016/0039-3681\(85\)90003-2](https://doi.org/10.1016/0039-3681(85)90003-2)
- Melrose, K. L., Brown, G. D., & Wood, A. M. (2013). Am I abnormal? Relative rank and social norm effects in judgments of anxiety and depression symptom severity. *Journal of Behavioral Decision Making*, 26(2), 174–184. <https://doi.org/10.1002/bdm.1754>
- Metropolis, N., Rosenbluth, A. W., Rosenbluth, M. N., Teller, A. H., & Teller, E. (1953). Equation of State Calculations by Fast Computing Machines. *The Journal of Chemical Physics*, 21(6), 1087–1092. <https://doi.org/10.1063/1.1699114>
- Milkowski, M. (2013). *Explaining the Computational Mind*. MIT Press.
- Miller, G. A. (1956). The magical number seven, plus or minus two: Some limits on our capacity for processing information. *Psychological Review*, 63(2), 81–97. <https://doi.org/10.1037/h0043158>
- Miller, G. A. (1983). Informavores. In F. Machlup & U. Mansfield (Eds.), *The Study of Information: Interdisciplinary Messages* (pp. 111–113). Wiley.
- Miron-Shatz, T. (2009). Evaluating multiepisode events: Boundary conditions for the peak-end rule. *Emotion*, 9(2), 206–213. <https://doi.org/10.1037/a0015295>
- Mitchell, S. D. (1997). Pragmatic Laws. *Philosophy of Science*, 64(S4), S468–S479. <https://doi.org/10.1086/392623>
- Molière, & Laun, H. V. (1673). *The Imaginary Invalid*. Courier Corporation.
- Moore, S. C., Wood, A. M., Moore, L., Shepherd, J., Murphy, S., & Brown, G. D. (2016). A rank based social norms model of how people judge their levels of drunkenness whilst intoxicated. *BMC Public Health*, 16(1), 1–8. <https://doi.org/10.1186/s12889-016-3469-z>

REFERENCES

- Morewedge, C. K., Gilbert, D. T., & Wilson, T. D. (2005). The least likely of times: How remembering the past biases forecasts of the future. *Psychological Science*, 16(8), 626–630. <https://doi.org/10.1111/j.1467-9280.2005.01585.x>
- Morgan, M. S. (2012). *The World in the Model: How Economists Work and Think*. Cambridge University Press.
- Morin, C., Brown, G. D., & Lewandowsky, S. (2010). Temporal isolation effects in recognition and serial recall. *Memory & Cognition*, 38(7), 849–859. <https://doi.org/10.3758/MC.38.7.849>
- Mullett, T. L., & Tunney, R. J. (2013). Value representations by rank order in a distributed network of varying context dependency. *Brain and Cognition*, 82(1), 76–83. <https://doi.org/10.1016/j.bandc.2013.02.010>
- Murdock, B. B. J. (1960). The distinctiveness of stimuli. *Psychological Review*, 67(1), 16–31. <https://doi.org/10.1037/h0042382>
- Murdock, B. B. J. (1962). The serial position effect of free recall. *Journal of Experimental Psychology*, 64(5), 482–488. <https://doi.org/10.1037/h0045106>
- Nairne, J. S., & Dutta, A. (1992). Spatial and temporal uncertainty in long-term memory. *Journal of Memory and Language*, 31, 396–407. [https://doi.org/10.1016/0749-596X\(92\)90020-X](https://doi.org/10.1016/0749-596X(92)90020-X)
- Nalborczyk, L., Bürkner, P.-C., & Williams, D. R. (2019). Pragmatism should not be a substitute for statistical literacy, a commentary on albers, kiers, and van ravenswaaij (2018) (V. Savalei & V. Savalei, Eds.). *Collabra: Psychology*, 5(1). <https://doi.org/10.1525/collabra.197>
- Nash, J. F. (1950). Equilibrium points in n-person games. *Proceedings of the National Academy of Sciences*, 36(1), 48–49. <https://doi.org/10.1073/pnas.36.1.48>
- Navarro, D. J., Newell, B. R., & Schulze, C. (2016). Learning and choosing in an uncertain world: An investigation of the explore-exploit dilemma in static and dynamic environments. *Cognitive Psychology*, 85, 43–77. <https://doi.org/10.1016/j.cogpsych.2016.01.001>
- Neal, R. M. (2003). Slice sampling. *The Annals of Statistics*, 31(3), 705–767. <https://doi.org/10.1214/aos/1056562461>

- Neander, K. (1991). The teleological notion of 'function'. *Australasian Journal of Philosophy*, 69(4), 454–468. <https://doi.org/10.1080/00048409112344881>
- Neath, I. (1993). Distinctiveness and serial position effects in recognition. *Memory & Cognition*, 21(5), 689–698. <https://doi.org/10.3758/BF03197199>
- Neath, I., & Brown, G. D. (2006). SIMPLE: Further applications of A local distinctiveness model of memory. *Psychology of Learning and Motivation - Advances in Research and Theory*, 46(06), 201–243. [https://doi.org/10.1016/S0079-7421\(06\)46006-0](https://doi.org/10.1016/S0079-7421(06)46006-0)
- Neath, I., Brown, G. D., McCormack, T., Chater, N., & Freeman, R. (2006). Distinctiveness models of memory and absolute identification: Evidence for local, not global, effects. *Quarterly Journal of Experimental Psychology*, 59(1), 121–135. <https://doi.org/10.1080/17470210500162086>
- Newell, B. R., & Shanks, D. R. (2014). Unconscious influences on decision making: A critical review. *Behavioral and Brain Sciences*, 38(01), 1–19. <https://doi.org/10.1017/S0140525X12003214>
- Nisbett, R. E., & Wilson, T. D. (1977). Telling more than we can know: Verbal reports on mental processes. *Psychological Review*, 84(3), 231–259. <https://doi.org/10.1037/0033-295X.84.3.231>
- Noguchi, T., & Stewart, N. (2014). In the attraction, compromise, and similarity effects, alternatives are repeatedly compared in pairs on single dimensions. *Cognition*, 132(1), 44–56. <https://doi.org/10.1016/j.cognition.2014.03.006>
- Nosofsky, R. M. (1986). Attention, similarity, and the identification-categorization relationship. *Journal of Experimental Psychology: General*, 115(1), 39–57. <https://doi.org/10.1037/0096-3445.115.1.39>
- Nosofsky, R. M. (1988). Similarity, frequency, and category representations. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 14(1), 54–65. <https://doi.org/10.1037/0278-7393.14.1.54>
- Oaksford, M., & Chater, N. (2009). Précis of of bayesian rationality: The probabilistic approach to human reasoning. *Behavioral and Brain Sciences*, 32(1), 69–120. <https://doi.org/10.1017/S0140525X09000284>

REFERENCES

- Odling-Smee, F. J., Laland, K. N., & Feldman, M. W. (2013). *Niche Construction: The Neglected Process in Evolution*. Princeton University Press.
- Olson, M. J., & Budescu, D. V. (1997). Patterns of preference for numerical and verbal probabilities. *Journal of Behavioral Decision Making*, 10(2), 117–131. [https://doi.org/10.1002/\(sici\)1099-0771\(199706\)10:2<117::aid-bdm251>3.3.co;2-z](https://doi.org/10.1002/(sici)1099-0771(199706)10:2<117::aid-bdm251>3.3.co;2-z)
- Osafo Hounkpatin, H., Wood, A. M., Brown, G. D., & Dunn, G. (2015). Why does income relate to depressive symptoms? Testing the income rank hypothesis longitudinally. *Social Indicators Research*, 124(2), 637–655. <https://doi.org/10.1007/s11205-014-0795-3>
- Owen, A., & Zhou, Y. (2000). Safe and Effective Importance Sampling. *Journal of the American Statistical Association*, 95(449), 135–143. <https://doi.org/10.2307/2669533>
- Parducci, A. (1968). The relativism of absolute judgments. *Scientific American*, 219(6), 84–93. <https://doi.org/10.1038/scientificamerican1268-84>
- Parducci, A. (2011). Utility versus pleasure: The grand paradox. *Frontiers in Psychology*, 2(NOV), 1–2. <https://doi.org/10.3389/fpsyg.2011.00296>
- Parker, G. A., & Smith, J. M. (1990). Optimality theory in evolutionary biology. *Nature*, 348(6296), 27–33. <https://doi.org/10.1038/348027a0>
- Pascal, B. (1670). *Pensees*. Dover Publications.
- Payzan-LeNestour, E., & Woodford, M. (2020). 'Outlier Blindness': Efficient Coding Generates an Inability to Represent Extreme Values. *Behavioral & Experimental Finance eJournal*. <https://doi.org/10.2139/ssrn.3152166>
- Pearl, J. (2000). *Causality: Models, Reasoning, and Inference*. Cambridge University Press.
- Pedroni, A., Frey, R., Bruhin, A., Dutilh, G., Hertwig, R., & Rieskamp, J. (2017). The risk elicitation puzzle. *Nature Human Behaviour*, 1(11), 803–809. <https://doi.org/10.1038/s41562-017-0219-x>
- Pelli, D. G. (1997). The VideoToolbox software for visual psychophysics: Transforming numbers into movies. *Spatial Vision*, 10(4), 437–442. <https://doi.org/10.1163/156856897X00366>

- Peterson, D. K., & Pitz, G. F. (1988). Confidence, uncertainty, and the use of information. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 14(1), 85–92. <https://doi.org/10.1037/0278-7393.14.1.85>
- Piccinini, G., & Craver, C. (2011). Integrating psychology and neuroscience: Functional analyses as mechanism sketches. *Synthese*, 183(3), 283–311. <https://doi.org/10.1007/s11229-011-9898-4>
- Pierce, C. S. (1906). Prolegomena to an apology for pragmatism. *Collected papers of Charles Sanders Pierce* (pp. 530–572). Harvard University Press.
- Pirolli, P., & Card, S. (1999). Information foraging. *Psychological Review*, 106(4), 643–675. <https://doi.org/10.1037/0033-295X.106.4.643>
- Pitt, D. (2022). Mental Representation. In E. N. Zalta & U. Nodelman (Eds.), *The Stanford Encyclopedia of Philosophy* (Fall 2022). Metaphysics Research Lab, Stanford University.
- Pleskac, T. J., Yu, S., Hopwood, C., & Liu, T. (2019). Mechanisms of deliberation during preferential choice: Perspectives from computational modeling and individual differences. *Decision*, 6(1), 77–107. <https://doi.org/10.1037/dec0000092>
- Plonsky, O., Teodorescu, K., & Erev, I. (2015). Reliance on small samples, the wavy recency effect, and similarity-based learning. *Psychological Review*, 122(4), 621–647. <https://doi.org/10.1037/a0039413>
- Poisson, S. D. (1837). *Recherches sur la probabilité des jugements en matière criminelle et en matière civile*.
- Potochnik, A. (2017). *Idealization and the Aims of Science*. University of Chicago Press.
- Putnam, H. (1979). *Philosophical Papers: Volume 2, Mind, Language and Reality*. Cambridge University Press.
- Putnam, H. (1975). The meaning of 'meaning'. *Minnesota Studies in the Philosophy of Science*, 7, 131–193.
- Pylyshyn, Z. W. (1973). What the mind's eye tells the mind's brain: A critique of mental imagery. *Psychological Bulletin*, 80, 1–24. <https://doi.org/10.1037/h0034650>
- Quiggin, J. (1982). A theory of anticipated utility. *Journal of Economic Behavior and Organization*, 3(4), 323–343. [https://doi.org/10.1016/0167-2681\(82\)90008-7](https://doi.org/10.1016/0167-2681(82)90008-7)

REFERENCES

- Quine, W. V. (1987). *Quiddities: An intermittently philosophical dictionary* (Vol. 51). Belknap Press of Harvard University Press.
- Rakow, T., Newell, B. R., & Zougkou, K. (2010). The role of working memory in information acquisition and decision making: Lessons from the binary prediction task. *Quarterly Journal of Experimental Psychology*, 63(7), 1335–1360. <https://doi.org/10.1080/17470210903357945>
- Rangel, A., & Clithero, J. A. (2012). Value normalization in decision making: Theory and evidence. *Current Opinion in Neurobiology*, 22(6), 970–981. <https://doi.org/10.1016/j.conb.2012.07.011>
- Ransom, K. J., Perfors, A., Hayes, B. K., & Connor Desai, S. (2022). What do our sampling assumptions affect: How we encode data or how we reason from it? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, No Pagination Specified–No Pagination Specified. <https://doi.org/10.1037/xlm0001149>
- Raymond, J. E., Shapiro, K. L., & Arnell, K. M. (1992). Temporary suppression of visual processing in an RSVP task: An attentional blink? *Journal of Experimental Psychology: Human Perception and Performance*, 18(3), 849–860. <https://doi.org/10.1037/0096-1523.18.3.849>
- Redelmeier, D. A., & Kahneman, D. (1996). Patients’ memories of painful medical treatments: Real-time and retrospective evaluations of two minimally invasive procedures. *Pain*, 66(1), 3–8. [https://doi.org/10.1016/0304-3959\(96\)02994-6](https://doi.org/10.1016/0304-3959(96)02994-6)
- Reichenbach, H. (1978). Induction and probability. Remarks on Karl Popper’s The logic of scientific discovery. *Selected writings*, 2, 372–387.
- Rellihan, M. (2012). Adaptationism and adaptive thinking in evolutionary psychology. *Philosophical Psychology*, 25(2), 245–277. <https://doi.org/10.1080/09515089.2011.579416>
- Rescorla, M. (2020). The Computational Theory of Mind. In E. N. Zalta (Ed.), *The Stanford Encyclopedia of Philosophy* (Fall 2020). Metaphysics Research Lab, Stanford University.
- Rice, C. (2015). Moving beyond causes: Optimality models and scientific explanation. *Nous*, 49(3), 589–615. <https://doi.org/10.1111/nous.12042>

- Richards, B. (1987). Type/token ratios: What do they really tell us? *Journal of Child Language*, 14(2), 201–209. <https://doi.org/10.1017/S0305000900012885>
- Rider, P. R. (1960). Variance of the Median of Small Samples from Several Special Populations. *Journal of the American Statistical Association*, 55(289), 148–150. <https://doi.org/10.1080/01621459.1960.10482056>
- Robert, C. P., & Casella, G. (2004). *Monte Carlo Statistical Methods*. Springer New York. <https://doi.org/10.1007/978-1-4757-4145-2>
- Roberts, S., & Pashler, H. (2000). How persuasive is a good fit? A comment on theory testing. *Psychological Review*, 107(2), 358–367. <https://doi.org/10.1037//0033-295X.107.2.358>
- Robinson, E. J., Pendle, J. E. C., Rowley, M. G., Beck, S. R., & McColgan, K. L. T. (2009). Guessing imagined and live chance events: Adults behave like children with live events. *British Journal of Psychology*, 100(4), 645–659. <https://doi.org/10.1348/000712608X386810>
- Robinson, E. J., Rowley, M. G., Beck, S. R., Carroll, D. J., & Apperly, I. A. (2006). Children's sensitivity to their own relative ignorance: Handling of possibilities under epistemic and physical uncertainty. *Child Development*, 77(6), 1642–1655. <https://doi.org/10.1111/j.1467-8624.2006.00964.x>
- Rode, E., Rozin, P., & Durlach, P. (2007). Experienced and remembered pleasure for meals: Duration neglect but minimal peak, end (recency) or primacy effects. *Appetite*, 49(1), 18–29. <https://doi.org/10.1016/j.appet.2006.09.006>
- Ronayne, D., & Brown, G. D. (2017). Multi-attribute decision by sampling: An account of the attraction, compromise and similarity effects. *Journal of Mathematical Psychology*, 81, 11–27. <https://doi.org/10.1016/j.jmp.2017.08.005>
- Rosseel, Y. (2012). Llavaan: An R package for structural equation modeling. *Journal of Statistical Software*, 48(2). <https://doi.org/10.18637/jss.v048.i02>
- Rothbart, M., & Snyder, M. (1970). Confidence in the prediction and postdiction of an uncertain outcome. *Canadian Journal of Behavioural Science/Revue canadienne des sciences du comportement*, 2(1), 38–43. <https://doi.org/10.1037/h0082709>

REFERENCES

- Rottenstreich, Y., & Hsee, C. K. (2001). Money, kisses, and electric shocks: On the affective psychology of risk. *Psychological Science*, 12(3), 185–190. <https://doi.org/10.1111/1467-9280.00334>
- Rozin, A., Rozin, P., & Goldberg, E. (2004). The feeling of music past: How listeners remember musical affect. *Music Perception*, 22(1), 15–39. <https://doi.org/10.1525/mp.2004.22.1.15>
- Rubino, G., & Tuffin, B. (2009). *Rare Event Simulation using Monte Carlo Methods*. John Wiley & Sons.
- Rubinstein, R. Y., & Kroese, D. P. (2016). *Simulation and the Monte Carlo Method*. John Wiley & Sons.
- Russell, B. A. W. (1948). *Human knowledge: Its scope and limits* (Vol. 48). Simon and Schuster.
- Ryan, J. (1969). Grouping and short-term memory: Different means and patterns of grouping. *The Quarterly journal of experimental psychology*, 21(2), 137–147. <https://doi.org/10.1080/14640746908400206>
- Salmon, W. C. (1971). *Statistical Explanation and Statistical Relevance*. University of Pittsburgh Pre.
- Sanborn, A. N., Griffiths, T. L., & Navarro, D. J. (2010). Rational approximations to rational models: Alternative algorithms for category learning. *Psychological Review*, 117(4), 1144–1167. <https://doi.org/10.1037/a0020511>
- Sands, S. F., & Wright, A. A. (1980). Primate memory: Retention of serial list items by a rhesus monkey. *Science*, 209(4459), 938–940. <https://doi.org/10.1126/science.6773143>
- Sarin, R. K., & Weber, M. (1993). Effects of ambiguity in market experiments. *Management Science*, 39(5), 602–615. <https://doi.org/10.1287/mnsc.39.5.602>
- Savage, L. J. (1954). *The foundations of statistics*. John Wiley & Sons.
- Schäfer, T., Zimmermann, D., & Sedlmeier, P. (2014). How we remember the emotional intensity of past musical experiences. *Frontiers in Psychology*, 5(AUG), 1–10. <https://doi.org/10.3389/fpsyg.2014.00911>

- Scheibehenne, B., & Coppin, G. (2020). How does the peak-end rule smell? Tracing hedonic experience with odours. *Cognition and Emotion*, 34(4), 713–727. <https://doi.org/10.1080/02699931.2019.1675599>
- Schlosser, M. (2019). Agency. In E. N. Zalta (Ed.), *The Stanford Encyclopedia of Philosophy* (Winter 2019). Metaphysics Research Lab, Stanford University.
- Schmidt, S. R. (1991). Can we have a distinctive theory of memory? *Memory & Cognition*, 19(6), 523–542. <https://doi.org/10.3758/BF03197149>
- Schoemaker, P. J. (1991). The quest for optimality: A positive heuristic of science? *Behavioral and Brain Sciences*, 14(2), 205–215. <https://doi.org/10.1017/S0140525X00066140>
- Schulz, E., Konstantinidis, E., & Speekenbrink, M. (2018). Putting bandits into context: How function learning supports decision making. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 44(6), 927–943. <https://doi.org/10.1037/xlm0000463>
- Segal, U. (1987). The ellisberg paradox and risk aversion: An anticipated utility approach. *International Economic Review*, 28(1), 175. <https://doi.org/10.2307/2526866>
- Seta, J. J., Haire, A., & Seta, C. E. (2008). Averaging and summation: Positivity and choice as a function of the number and affective intensity of life events. *Journal of Experimental Social Psychology*, 44(2), 173–186. <https://doi.org/10.1016/j.jesp.2007.03.003>
- Shadlen, M. N., & Shohamy, D. (2016). Decision Making and Sequential Sampling from Memory. *Neuron*, 90(5), 927–939. <https://doi.org/10.1016/j.neuron.2016.04.036>
- Shafto, P., Goodman, N. D., & Griffiths, T. L. (2014). A rational account of pedagogical reasoning: Teaching by, and learning from, examples. *Cognitive Psychology*, 71, 55–89. <https://doi.org/10.1016/j.cogpsych.2013.12.004>
- Shapiro, L. A. (2017). Mechanism or bust? Explanation in psychology. *British Journal for the Philosophy of Science*, 68(4), 1037–1059. <https://doi.org/10.1093/bjps/axv062>
- Shapiro, L. A., & Spaulding, S. (2021). Embodied Cognition. In E. N. Zalta (Ed.), *The Stanford Encyclopedia of Philosophy* (Winter 2021). Metaphysics Research Lab, Stanford University.

REFERENCES

- Shepard, R. N. (1987). Toward a Universal Law of Generalization for Psychological Science. *Science*, 237(4820), 1317–1323. <https://doi.org/10.1126/science.3629243>
- Shiffrin, R. M., & Nosofsky, R. M. (1994). Seven plus or minus two: A commentary on capacity limitations. *Psychological Review*, 101(2), 357–361. <https://doi.org/10.1037//0033-295X.101.2.357>
- Siegel, S. (2021). The Contents of Perception. In E. N. Zalta (Ed.), *The Stanford Encyclopedia of Philosophy* (Fall 2021). Metaphysics Research Lab, Stanford University.
- Simon, H. A. (1955). A Behavioral Model of Rational Choice. *Quarterly Journal of Economics*, 69, 99–118. <https://doi.org/10.2307/1884852>
- Simon, H. A. (1982). *Models of Bounded Rationality: Empirically grounded economic reason*. MIT Press.
- Simonson, I. (1989). Choice based on reasons: The case of attraction and compromise effects. *Journal of Consumer Research*, 16(2), 158. <https://doi.org/10.1086/209205>
- Simpson, D. (2018). Justify my love.
- Skinner, B. F. (1953). *Science and human behavior*. Macmillan.
- Smith, R. L. (1984). Efficient Monte Carlo Procedures for Generating Points Uniformly Distributed Over Bounded Regions. *Operations Research*, 32(6), 1296–1308.
- Snijders, T. A. B., & Bosker, R. J. (2011). *Multilevel Analysis: An Introduction to Basic and Advanced Multilevel Modeling*. SAGE.
- Spektor, M. S., Gluth, S., Fontanesi, L., & Rieskamp, J. (2018). How similarity between choice options affects decisions from experience: The accentuation of differences model. *Psychological Review*, 126(1), 52–88. <https://doi.org/10.1037/rev0000122>
- Spence, M. (1973). Job Market Signaling. *The Quarterly Journal of Economics*, 87(3), 355–374. <https://doi.org/10.2307/1882010>
- Stan Development Team. (2020). Prior Choice Recommendations.
- Stanovich, K. E., & West, R. F. (2003). Individual differences in reasoning: Implications for the rationality debate? *Behavioral and Brain Sciences*, 26(4), 527. <https://doi.org/10.1017/S0140525X03210116>

- Steffens, J., & Guastavino, C. (2015). Trend effects in momentary and retrospective soundscape judgments. *Acta Acustica united with Acustica*, 101(4), 713–722. <https://doi.org/10.3813/AAA.918867>
- Stewart, N. (2009). Decision by sampling: The role of the decision environment in risky choice. *Quarterly Journal of Experimental Psychology*, 62(6), 1041–1062. <https://doi.org/10.1080/17470210902747112>
- Stewart, N., Chater, N., & Brown, G. D. (2006). Decision by sampling. *Cognitive Psychology*, 53(1), 1–26. <https://doi.org/10.1016/j.cogpsych.2005.10.003>
- Stewart, N., Chater, N., Stott, H. P., & Reimers, S. (2003). Prospect relativity: How choice options influence decision under risk. *Journal of Experimental Psychology: General*, 132(1), 23–46. <https://doi.org/10.1037/0096-3445.132.1.23>
- Stewart, N., Reimers, S., & Harris, A. J. L. (2015). On the origin of utility, weighting, and discounting functions: How they get their shapes and how to change their shapes. *Management Science*, 61(3), 687–705. <https://doi.org/10.1287/mnsc.2013.1853>
- Stiglitz, J. E. (1975). The theory of screening, education, and the distribution of income. *American Economic Review*, 65(3), 283–300.
- Stone, A. A., Broderick, J. E., Kaell, A. T., DelesPaul, P. A., & Porter, L. E. (2000). Does the peak-end phenomenon observed in laboratory pain studies apply to real-world pain in rheumatoid arthritics? *Journal of Pain*, 1(3), 212–217. <https://doi.org/10.1054/jpai.2000.7568>
- Strevens, M. (2003). Making Things Happen: A Theory of Causal Explanation. *Philosophy and Phenomenological Research*. <https://doi.org/10.1111/j.1933-1592.2007.00012.x>
- Strijbosch, W., Mitas, O., van Gisbergen, M., Doicaru, M., Gelissen, J., & Bastiaansen, M. (2019). From experience to memory: On the robustness of the peak-and-end-rule for complex, heterogeneous experiences. *Frontiers in Psychology*, 10(JULY), 1–12. <https://doi.org/10.3389/fpsyg.2019.01705>
- Summerfield, C., & Tsetsos, K. (2015). Do humans make good decisions? *Trends in Cognitive Sciences*, 19(1), 27–34. <https://doi.org/10.1016/j.tics.2014.11.005>

REFERENCES

- Sunstein, C. R., & Vermeule, A. (2008). Conspiracy Theories. <https://doi.org/10.2139/ssrn.1084585>
- Suppes, P. (1970). *A Probabilistic Theory of Causality*. North-Holland Publishing Company.
- Szollosi, A., Donkin, C., & Newell, B. R. (2022). Toward nonprobabilistic explanations of learning and decision-making. *Psychological Review*, No Pagination Specified–No Pagination Specified. <https://doi.org/10.1037/rev0000355>
- Taleb, N. N. (2008). *The black swan: The impact of the highly improbable*. Penguin Books Limited.
- Tannenbaum, D., Fox, C. R., & Ülkümen, G. (2017). Judgment extremity and accuracy under epistemic vs. aleatory uncertainty. *Management Science*, 63(2), 497–518. <https://doi.org/10.1287/mnsc.2015.2344>
- Tanner, M. A., & Wong, W. H. (1987). The Calculation of Posterior Distributions by Data Augmentation. *Journal of the American Statistical Association*, 82(398), 528–540. <https://doi.org/10.1080/01621459.1987.10478458>
- Taylor, K. A. (1995). Testing credit and blame attributions as explanation for choices under ambiguity. *Organizational Behavior and Human Decision Processes*, 64(2), 128–137. <https://doi.org/10.1006/obhd.1995.1095>
- Taylor, M. J., Vlaev, I., Maltby, J., Brown, G. D., & Wood, A. M. (2015). Improving social norms interventions: Rank-framing increases excessive alcohol drinkers' information-seeking. *Health Psychology*, 34(12), 1200–1203. <https://doi.org/10.1037/hea0000237>
- Teigen, K. H. (1988). The language of uncertainty. *Acta Psychologica*, 68(1-3), 27–38. [https://doi.org/10.1016/0001-6918\(88\)90043-1](https://doi.org/10.1016/0001-6918(88)90043-1)
- Templin, M. C. (1957). *Certain language skills in children* (Vol. 26). University of Minnesota Press.
- Tenenbaum, J. B., & Griffiths, T. L. (2001). Generalization, similarity, and Bayesian inference. *Behavioral and Brain Sciences*, 24(4), 629–640. <https://doi.org/10.1017/S0140525X01000061>

- Thagard, P. (2020). Cognitive Science. In E. N. Zalta (Ed.), *The Stanford Encyclopedia of Philosophy* (Winter 2020). Metaphysics Research Lab, Stanford University.
- Thomas, D. L., & Diener, E. (1990). Memory accuracy in the recall of emotions. *Journal of Personality and Social Psychology*, 59(2), 291–297. <https://doi.org/10.1037/0022-3514.59.2.291>
- Thomas, D., Olsen, D., & Murray, K. (2018). Evaluations of a sequence of affective events presented simultaneously. *European Journal of Marketing*, 52(3/4), 866–881. <https://doi.org/10.1108/EJM-09-2016-0526>
- Thunnissen, D. P. (2003). Uncertainty classification for the design and development of complex systems. *3rd Annual Predictive Methods Conference*, 16.
- Todd, P. M., & Gigerenzer, G. (2000). Précis of simple heuristics that make us smart. *Behavioral and Brain Sciences*, 23(5), 727–780. <https://doi.org/10.1017/S0140525X00003447>
- Tokdar, S. T., & Kass, R. E. (2010). Importance sampling: A review. *Wiley Interdisciplinary Reviews: Computational Statistics*, 2(1), 54–60. <https://doi.org/10.1002/wics.56>
- Tooby, J., & Cosmides, L. (1998). Evolutionizing the Cognitive Sciences: A Reply to Shapiro and Epstein. *Mind & Language*, 13(2), 195–204. <https://doi.org/10.1111/1468-0017.00073>
- Tremblay, L., & Schultz, W. (1999). Relative reward preference in primate orbitofrontal cortex. *Nature*, 398(6729), 704–708. <https://doi.org/10.1038/19525>
- Tsetsos, K., Chater, N., & Usher, M. (2012). Salience driven value integration explains decision biases and preference reversal. *Proceedings of the National Academy of Sciences*, 109(24), 9659–9664. <https://doi.org/10.1073/pnas.1119569109>
- Tsetsos, K., Moran, R., Moreland, J., Chater, N., Usher, M., & Summerfield, C. (2016). Economic irrationality is optimal during noisy decision making. *Proceedings of the National Academy of Sciences*, 113(11), 3102–3107. <https://doi.org/10.1073/pnas.1519157113>
- Tversky, A. (1972). Elimination by aspects: A theory of choice. *Psychological Review*, 79(4), 281–299. <https://doi.org/10.1037/h0032955>

REFERENCES

- Tversky, A., & Edwards, W. (1966). Information versus reward in binary choices. *Journal of Experimental Psychology*, 71(5), 680–683. <https://doi.org/10.1037/h0023123>
- Tversky, A., & Kahneman, D. (1974). Judgment under uncertainty: Heuristics and biases. *Science*, 185(4157), 1124–1131. <https://doi.org/10.1126/science.185.4157.1124>
- Tversky, A., & Kahneman, D. (1992). Advances in prospect theory: Cumulative representation of uncertainty. *Journal of Risk and Uncertainty*, 5(4), 297–323. <https://doi.org/10.1007/BF00122574>
- Ulam, S. M. (1976). *Adventures of a mathematician*. Scribner.
- Ülkümen, G., Fox, C. R., & Malle, B. F. (2016). Two dimensions of subjective uncertainty: Clues from natural language. *Journal of Experimental Psychology: General*, 145(10), 1280–1297. <https://doi.org/10.1037/xge0000202>
- Ungemach, C., Stewart, N., & Reimers, S. (2011). How incidental values from the environment affect decisions about money, risk, and delay. *Psychological Science*, 22(2), 253–260. <https://doi.org/10.1177/0956797610396225>
- Vanunu, Y., Hotaling, J. M., & Newell, B. R. (2020). Elucidating the differential impact of extreme-outcomes in perceptual and preferential choice. *Cognitive Psychology*, 119(May 2019). <https://doi.org/10.1016/j.cogpsych.2020.101274>
- Vehtari, A., Gelman, A., & Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*, 27(5), 1413–1432. <https://doi.org/10.1007/s11222-016-9696-4>
- Vehtari, A., Gelman, A., Simpson, D., Carpenter, B., & Bürkner, P.-C. (2020). Rank-normalization, folding, and localization: An improved $\hat{R}$ for assessing convergence of MCMC. *Bayesian Analysis*, -1(-1), 1–27. <https://doi.org/10.1214/20-ba1221>
- Volz, K. G., Schubotz, R. I., & Von Cramon, D. Y. (2004). Why am I unsure? Internal and external attributions of uncertainty dissociated by fMRI. *NeuroImage*, 21(3), 848–857. <https://doi.org/10.1016/j.neuroimage.2003.10.028>
- Volz, K. G., Schubotz, R. I., & Von Cramon, D. Y. (2005). Variants of uncertainty in decision-making and their neural correlates. *Brain Research Bulletin*, 67(5), 403–412. <https://doi.org/10.1016/j.brainresbull.2005.06.011>

- von Neumann, J. (1951). Various techniques used in connection with random digits. In A. S. Householder, G. E. Forsythe, & H. H. Germond (Eds.), *Monte carlo method* (pp. 36–38). US Government Printing Office.
- von Neumann, J., Morgenstern, O., & Rubinstein, A. (1944). *Theory of games and economic behavior*. Princeton University Press.
- von Restorff, H. (1933). Über die wirkung von bereichsbildungen im spurenfeld. *Psychologische Forschung*, 18(1), 299–342. <https://doi.org/10.1007/BF02409636>
- Vul, E., Goodman, N., Griffiths, T. L., & Tenenbaum, J. B. (2014). One and done? Optimal decisions from very few samples. *Cognitive Science*, 38(4), 599–637. <https://doi.org/10.1111/cogs.12101>
- Wagner, D. A. (1975). The effects of verbal labeling on short-term and incidental memory: A cross-cultural and developmental study. *Memory & Cognition*, 3(6), 595–598. <https://doi.org/10.3758/BF03198223>
- Wakker, P. P. (2001). Testing and characterizing properties of nonadditive measures through violations of the sure-thing principle. *Econometrica*, 69(4), 1039–1059. <https://doi.org/10.1111/1468-0262.00229>
- Walasek, L., & Stewart, N. (2015). How to make loss aversion disappear and reverse: Tests of the decision by sampling origin of loss aversion. *Journal of Experimental Psychology: General*, 144(1), 7–11. <https://doi.org/10.1037/xge0000039>
- Walasek, L., & Stewart, N. (2018). Context-dependent sensitivity to losses: Range and skew manipulations. *Journal of Experimental Psychology: Learning, Memory, and Cognition*. <https://doi.org/10.1037/xlm0000629>
- Walker, A. R., Navarro, D. J., Newell, B. R., & Beesley, T. (2021). Protection from uncertainty in the exploration/exploitation trade-off. *Journal of Experimental Psychology: Learning, Memory, and Cognition*. <https://doi.org/10.1037/xlm0000883>
- Walker, W., Harremoës, P., Rotmans, J., van der Sluijs, J., van Asselt, M., Janssen, P., & Krayen von Krauss, M. (2003). Defining uncertainty: A conceptual basis for uncertainty management in model-based decision support. *Integrated Assessment*, 4(1), 5–17. <https://doi.org/10.1076/iaij.4.1.5.16466>

REFERENCES

- Wallace, W. P. (1965). Review of the historical, empirical, and theoretical status of the von Restorff phenomenon. *Psychological Bulletin*, 63(6), 410–424. <https://doi.org/10.1037/h0022001>
- Walters, D. J., Ulkumen, G., Tannenbaum, D., Erner, C., & Fox, C. R. (2022). Investment behaviors under epistemic and aleatory uncertainty. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.3695316>
- Warren, P. A., Gostoli, U., Farmer, G. D., El-Deredy, W., & Hahn, U. (2018). A re-examination of “bias” in human randomness perception. *Journal of Experimental Psychology: Human Perception and Performance*, 44(5), 663–680. <https://doi.org/10.1037/xhp0000462>
- Watson, J. B. (1925). *Behaviorism*. Read Books Ltd.
- Watts, A. W. (1957). *The Way of Zen*. Random House.
- Waxman, S. R., & Markow, D. B. (1995). Words as invitations to form categories: Evidence from 12- to 13-month-old infants. *Cognitive Psychology*, 29(3), 257–302. <https://doi.org/10.1006/cogp.1995.1016>
- Weber, B. J., & Chapman, G. B. (2005). Playing for peanuts: Why is risk seeking more common for low-stakes gambles? *Organizational Behavior and Human Decision Processes*, 97(1), 31–46. <https://doi.org/10.1016/j.obhdp.2005.03.001>
- Weber, D. L., Green, D. M., & Luce, R. D. (1977). Effects of practice and distribution of auditory signals on absolute identification. *Perception & Psychophysics*, 22(3), 223–231. <https://doi.org/10.3758/BF03199683>
- Weber, E., & Kirsner, B. (1997). Reasons for rank-dependent utility evaluation. *Journal of Risk and Uncertainty*, 14(1), 41–61. <https://doi.org/10.1023/A:1007769703493>
- Weiskopf, D. A. (2011). Models and mechanisms in psychological explanation. *Synthese*, 183(3), 313–338. <https://doi.org/10.1007/s11229-011-9958-9>
- Wetzel, L. (2018). Types and tokens. In E. N. Zalta (Ed.), *The Stanford Encyclopedia of Philosophy*. Metaphysics Research Lab, Stanford University.
- Wheeler, M. (2005). *Reconstructing the Cognitive World: The Next Step*. MIT Press.
- Whitehouse, K. (2007). One ‘Quant’ Sees Shakeout For the Ages – ‘10,000 Years’. *Wall Street Journal*.

- Wilson, T. D., & Gilbert, D. T. (2003). Affective forecasting. *Advances in Experimental Social Psychology*, 35(3), 345–411. [https://doi.org/10.1016/S0065-2601\(03\)01006-2](https://doi.org/10.1016/S0065-2601(03)01006-2)
- Wilson, T. D., & Nisbett, R. E. (1978). The Accuracy of Verbal Reports About the Effects of Stimuli on Evaluations and Behavior. *Social Psychology*, 41(2), 118–131. <https://doi.org/10.2307/3033572>
- Wimsatt, W. C., & Wimsatt, W. K. (2007). *Re-Engineering Philosophy for Limited Beings: Piecewise Approximations to Reality*. Harvard University Press.
- Winston, J. S., Vlaev, I., Seymour, B., Chater, N., & Dolan, R. J. (2014). Relative valuation of pain in human orbitofrontal cortex. *Journal of Neuroscience*, 34(44), 14526–14535. <https://doi.org/10.1523/JNEUROSCI.1706-14.2014>
- Wirtz, D., Kruger, J., Scollon, C. N., & Diener, E. (2003). What to do on spring break? The role of predicted, on-line, and remembered experience in future choice. *Psychological Science*, 14(5), 520–524. <https://doi.org/10.1111/1467-9280.03455>
- Wisniewski, N. J., Truong, G., Handy, T. C., & Chapman, C. S. (2017). Reaching reveals that best-versus-rest processing contributes to biased decision making. *Acta Psychologica*, 176(September 2016), 32–38. <https://doi.org/10.1016/j.actpsy.2017.03.006>
- Wood, A. M., Brown, G. D., & Maltby, J. (2012). Social norm influences on evaluations of the risks associated with alcohol consumption: Applying the rank-based decision by sampling model to health judgments. *Alcohol and Alcoholism*, 47(1), 57–62. <https://doi.org/10.1093/alcalc/agr146>
- Woodward, J. (2000). Explanation and Invariance in the Special Sciences. *The British Journal for the Philosophy of Science*, 51(2), 197–254. <https://doi.org/10.1093/bjps/51.2.197>
- Woodward, J., & Woodward, J. F. (2003). *Making Things Happen: A Theory of Causal Explanation*. Oxford University Press.
- Wright, A. A., Santiago, H. C., Sands, S. F., Kendrick, D. F., & Cook, R. G. (1985). Memory Processing of Serial Lists by Pigeons, Monkeys, and People. *Science*, 229(4710), 287–289. <https://doi.org/10.1126/science.9304205>

REFERENCES

- Wu, C. M., Schulz, E., Speekenbrink, M., Nelson, J. D., & Meder, B. (2018). Generalization guides human exploration in vast decision spaces. *Nature Human Behaviour*, 2(12), 915–924. <https://doi.org/10.1038/s41562-018-0467-4>
- Wulff, D. U., Mergenthaler-Canseco, M., & Hertwig, R. (2018). A meta-analytic review of two modes of learning and the description-experience gap. *Psychological Bulletin*, 144(2), 140–176. <https://doi.org/10.1037/bul0000115>
- Yechiam, E., & Hochman, G. (2013). Loss-aversion or loss-attention: The impact of losses on cognitive performance. *Cognitive Psychology*, 66(2), 212–231. <https://doi.org/10.1016/j.cogpsych.2012.12.001>
- Zednik, C. (2017). Mechanisms in cognitive science. *The Routledge Handbook of Mechanisms and Mechanical Philosophy*. Routledge.
- Zeigenfuse, M. D., Pleskac, T. J., & Liu, T. (2014). Rapid decisions from experience. *Cognition*, 131(2), 181–194. <https://doi.org/10.1016/j.cognition.2013.12.012>

Appendix A

Additional details for utility-weighted sampling simulations

The experiences of participants in decisions from experience tasks depend on their choices and this introduces an additional element of complexity to our simulation. One approach would be to simulate the outcomes that were experienced by the actual participants at certain points throughout the experiment. Using this approach, the behaviour of the model would be influenced by the actual behaviour of the participants and this would leave us with two choices. Either describe this behaviour in detail or allow our analysis to be influenced by numerous implicit attributes.

Neither of these alternatives is adequate and given that our goal in this chapter was to examine the behaviour of utility-weighted sampling rather than any specific experiment, we simulated the outcomes presented by Ludvig et al. (2014) based on the assumptions that 1) each simulated estimate and choice occurred when the participant had completed 80% of the experiment and 2) one of the available options was selected randomly on each of the preceding trials. Although these assumptions are somewhat arbitrary, they ensure

APPENDIX A. ADDITIONAL DETAILS FOR UTILITY-WEIGHTED SAMPLING SIMULATIONS

that there was sufficient experience with each outcome and that the experienced average is roughly equal to the average of the possible outcomes.

How does this compare with participants' actual experience? Once the average experienced outcome stabilises the specific trial number has minimal impact. The risk preferences on previous trials would also have a minimal impact because both options had the same expected value. Neglecting the catch trials used by Ludvig et al. (2014) has a larger effect because the participants would be more likely to select the better option on these trials. The average would be above zero, and therefore, would be further away from the negative outcomes than the positive outcomes—this would result in an asymmetric distribution of estimates.

The full simulation code is available at <https://github.com/joelholwerda>. This includes additional simulations for the other experiments conducted by Ludvig et al. (2014), a simulation in which participants always selected the better option rather than randomly selecting an option, and simulations that examine the robustness of our conclusions regarding the third attribute of utility-weighted sampling.

Appendix B

Details on the epistemic and aleatory rating scale (EARS)

We used the 10-item version of the EARS (Ülkümen et al., 2016) to measure participants' interpretation of uncertainty. In Experiment 7a, these items were phrased in past tense with reference to either the red or blue option encountered in the choice task. Before completing the EARS, participants were presented with the following instructions:

You have finished the first section of the experiment. The second section will involve questions about the **outcomes (numbers of points)** that you experienced when you chose options during the game. Each question will refer either to when you selected a red or blue option. Please answer each question carefully. Click Next when you are ready to begin.

The following instructions were displayed on the screen above the list of EARS items:

Thinking now of outcomes (numbers of points) that you received when you chose a [red, blue] option, please answer the following questions:

APPENDIX B. DETAILS ON THE EPISTEMIC AND ALEATORY RATING SCALE (EARS)

After observing participants' responses to EARS items for the safe option, we were concerned that the responses were based on their uncertainty regarding the visual appearance of options rather than their outcomes. In the subsequent experiments, we attempted to mitigate this potential issue by presenting participants with a specific image associated with each option (in the unique images condition, this image was one they had not encountered in the task). Prior to completing the EARS, they were presented with the following instructions:

On the next screen, you will be asked to imagine that you are going to select the option presented on the [left, right]-hand side of the screen. You will be asked a number of questions about the **outcome (number of points)** that would occur if you were to select that option.

The following instructions were displayed on the screen above the list of EARS items:

Imagine that you are going to select the option presented here. Please answer the following questions regarding the **outcome (number of points)** that would result from your choice:

The EARS items were presented in a randomised order and were rated on a seven-point scale with the endpoints labelled as “Not at all” and “Very much”. The items in experiment 7b and 8 were as follows (these items were phrased in the past tense in Experiment 7a):

1. The outcome is something that has an element of randomness.
2. The outcome is unpredictable.
3. The outcome is determined by chance factors.
4. The outcome could play out in different ways on similar occasions
5. The outcome is in principle knowable in advance
6. The outcome is something that has been determined in advance.

-
7. The outcome is knowable in advance, given enough information.
 8. The outcome is something that well-informed people would agree on.
 9. The outcome is something that could be better predicted by consulting an expert.
 10. The outcome is something that becomes more predictable with additional knowledge or skills.

After observing outcomes associated with the risky option in each round in Experiment 9, participants made a choice between this risky option and a number of points presented on the screen. After making this decision—and before they were told the outcome—we presented them with an abridged 4-item version of the EARS that included items 1, 3, 7, and 8 (Tannenbaum et al., 2017).

Reliability estimates

For each experiment, we examined the reliability of the epistemic and aleatory subscales of the EARS. We used the MBESS package (K. Kelley, 2007) to calculate McDonald's omega using the Green and Yang (2009) method for ordered categorical variables and 95% confidence intervals using bias-corrected and accelerated (BCA) bootstrapping (K. Kelley & Pornprasertmanit, 2016). The interpretation of omega is similar to the more commonly reported Cronbach's alpha. The main difference is that omega does not make the often erroneous tau equivalence assumption that factor loadings are equal for all items (FLORA).

APPENDIX B. DETAILS ON THE EPISTEMIC AND ALEATORY RATING SCALE (EARS)

Table B.1: Reliability estimates for each subscale

Experiment	Option	Factor	Omega	95% CI	
				Lower	Upper
1	Risky	Aleatory	0.83	0.78	0.87
1	Risky	Epistemic	0.87	0.83	0.89
1	Risky	Aleatory	0.86	0.81	0.89
1	Risky	Epistemic	0.78	0.71	0.83
2	Safe	Aleatory	0.83	0.76	0.88
2	Safe	Epistemic	0.80	0.72	0.85
2	Safe	Aleatory	0.84	0.77	0.89
2	Safe	Epistemic	0.86	0.80	0.90
3	Risky	Aleatory	0.86	0.83	0.89
3	Risky	Epistemic	0.82	0.79	0.84

Note:

Omega is the categorical McDonald's omega. 95% confidence intervals were estimated using BCA bootstrapping.

Appendix C

Factor analysis for the epistemic and aleatory rating scale (EARS)

We used numerous methods for dimensionality reduction on participants' responses to better understand the internal structure of the EARS items. In this section, we report the results of the confirmatory factor analyses that were mentioned in the discussion section. The code for item correlations, cluster analysis, and exploratory factor analysis can be accessed at <https://github.com/joelholwerda>.

We used the lavaan package (Rosseel, 2012) to assess the factor loadings of EARS items 1-4 onto an aleatory uncertainty factor and EARS items 5-10 onto an epistemic uncertainty factor. We conducted this analysis separately for the safe and risky options in Experiment 7 and 8 to gauge whether participants were responding to the items differently across options or experiments. Goodness of fit statistics for each analysis are displayed in Table C.1. Factor loadings for each item are displayed in Table C.2 for Experiment 7 and Table C.3 for Experiment 8.

Some general patterns can be observed when examining these factor loadings. The first three items of the aleatory subscale loaded strongly onto their factor whereas the fourth item was not as closely related. This suggest that “play out in different ways in

APPENDIX C. FACTOR ANALYSIS FOR THE EPISTEMIC AND ALEATORY RATING SCALE (EARS)

Table C.1: Goodness of fit indicators for confirmatory factor analysis

Experiment	Option	χ^2	df	CFI	RMSEA	SRMR
1	Risky	75.25	34	0.99	0.07	0.05
1	Safe	74.10	34	0.99	0.07	0.05
2	Risky	64.03	34	0.99	0.08	0.07
2	Safe	57.11	34	1.00	0.07	0.06

Note:

CFI is the comparative fit index, RMSEA is the root mean-square error of approximation, and SRMR is the standardized root mean square

similar conditions” is capturing something different from the other aleatory items. This might allow the fourth item to capture different aspects of aleatory uncertainty but might also result from ambiguity in the intended meaning of “similar occasions.” A similar conclusion could be drawn for items 6 (“determined in advance”) and 9 (“consult an expert”) that were not loaded as strongly onto the epistemic uncertainty factor as the other four items on the subscale.

The factor loadings for the safe options were roughly similar for the safe and risky options, especially for the aleatory items. There was more variability amongst the epistemic items but there were no clear patterns that were consistent across experiments. The one exception to this was that item 9 (“consult an expert”) was more weakly related to the other epistemic items for the safe option. This item also had the lowest loading for the risky options and we suspect that this was because participants were unsure who this supposed expert in computer-based decision-making tasks could possibly be.

Table C.2: Confirmatory factor analysis loadings for Experiment 7

Option	Factor	Item	Unstandardised	Standardised	Residual
Risky	Aleatory	1	1.00 (—)	0.83	0.31
Risky	Aleatory	2	1.08 (0.04)	0.90	0.19
Risky	Aleatory	3	0.93 (0.04)	0.78	0.39
Risky	Aleatory	4	0.62 (0.05)	0.52	0.73
Risky	Epistemic	5	1.00 (—)	0.78	0.38
Risky	Epistemic	6	0.77 (0.06)	0.61	0.63
Risky	Epistemic	7	1.17 (0.05)	0.92	0.16
Risky	Epistemic	8	0.91 (0.05)	0.71	0.49
Risky	Epistemic	9	0.67 (0.06)	0.52	0.73
Risky	Epistemic	10	1.06 (0.04)	0.83	0.31
Safe	Aleatory	1	1.00 (—)	0.84	0.30
Safe	Aleatory	2	1.01 (0.03)	0.84	0.29
Safe	Aleatory	3	0.98 (0.03)	0.82	0.32
Safe	Aleatory	4	0.85 (0.04)	0.71	0.50
Safe	Epistemic	5	1.00 (—)	0.86	0.27
Safe	Epistemic	6	0.84 (0.05)	0.72	0.48
Safe	Epistemic	7	0.90 (0.03)	0.77	0.41
Safe	Epistemic	8	0.91 (0.04)	0.78	0.40
Safe	Epistemic	9	0.15 (0.06)	0.13	0.98
Safe	Epistemic	10	0.54 (0.05)	0.46	0.79

Note:

Dashes (—) indicate the standard error was not estimated. The correlation between the Epistemic and Aleatory factors was 0.78 for the risky option and 0.80 for the safe option.

APPENDIX C. FACTOR ANALYSIS FOR THE EPISTEMIC AND ALEATORY
RATING SCALE (EARS)

Table C.3: Confirmatory factor analysis loadings for Experiment 8

Option	Factor	Item	Unstandardised	Standardised	Residual
Risky	Aleatory	1	1.00 (—)	0.85	0.28
Risky	Aleatory	2	1.04 (0.05)	0.88	0.22
Risky	Aleatory	3	0.88 (0.05)	0.74	0.45
Risky	Aleatory	4	0.67 (0.07)	0.57	0.68
Risky	Epistemic	5	1.00 (—)	0.64	0.59
Risky	Epistemic	6	0.54 (0.11)	0.34	0.88
Risky	Epistemic	7	1.06 (0.10)	0.68	0.54
Risky	Epistemic	8	1.34 (0.11)	0.85	0.27
Risky	Epistemic	9	1.04 (0.12)	0.66	0.56
Risky	Epistemic	10	1.06 (0.11)	0.67	0.55
Safe	Aleatory	1	1.00 (—)	0.87	0.25
Safe	Aleatory	2	0.96 (0.04)	0.84	0.30
Safe	Aleatory	3	0.90 (0.05)	0.78	0.39
Safe	Aleatory	4	0.66 (0.07)	0.57	0.68
Safe	Epistemic	5	1.00 (—)	0.74	0.46
Safe	Epistemic	6	0.88 (0.06)	0.65	0.58
Safe	Epistemic	7	1.17 (0.05)	0.86	0.26
Safe	Epistemic	8	1.17 (0.06)	0.86	0.25
Safe	Epistemic	9	0.42 (0.09)	0.31	0.90
Safe	Epistemic	10	1.10 (0.06)	0.81	0.35

Note:

Dashes (—) indicate the standard error was not estimated. The correlation between the Epistemic and Aleatory factors was 0.64 for the risky option and 0.87 for the safe option.

Appendix D

Details of the dynamic Ellsberg urn task

In Chapter 6, we examined whether observable variability increases the amount of epistemic uncertainty associated with risky options and whether this gives rise to the avoidance of risky options in these contexts. In addition to this primary research question, we were also curious whether preferences regarding epistemic uncertainty generalise across other contexts such as Ellsberg’s (1961) urn task. Fox et al. (2021) recently demonstrated that aversion to ambiguity—unknown outcome probabilities—in this task reflects preferences regarding epistemic uncertainty. Therefore, we predicted that participants who avoid ambiguous options in the urn task would also avoid options in our experiment that they associate with epistemic uncertainty.

To investigate this, we presented participants in Experiment 7b and 8 with a dynamic version of the task proposed by Ellsberg (1961). This task required them to choose between an option where the probabilities associated with each outcome were *known* and another option where they were *unknown*. Specifically, we presented participants with two virtual boxes—each containing 100 balls that were either orange or purple—and proposed the following choice: “The box on the left-hand side contains 50 orange balls and 50 purple

APPENDIX D. DETAILS OF THE DYNAMIC ELLSBERG URN TASK

balls. The box on the right-hand side contains an unknown proportion of orange and purple balls. One ball will be randomly selected from the box you choose. If the ball is orange you will win one dollar but if it is purple you will win nothing.”

Participants were subsequently presented with a sequence of three similar choices between pairs of options where the proportion of orange and purple balls depended on their previous choices. If the box with known outcome probabilities was selected on the previous round, this option was made *less* attractive by reducing the number of winning (orange) balls. Contrastingly, if the box with unknown outcome probabilities was selected, the option with known outcome probabilities was made *more* attractive by increasing the number of winning balls.

The number of winning balls in the box with known outcome probabilities was increased or decreased by 20 balls in the second round (e.g., 30 winning and 70 losing). This was further increased or decreased by 10 winning balls in the third round (e.g., 40 winning and 60 losing) and 5 winning balls in the final round (e.g., 45 winning and 55 losing). The choice that participants made between the first pair of boxes is analogous to the standard urn task but the dynamic procedure also allowed us to determine the indifference point between these options with an accuracy of ± 5 balls.

We used this indifference point as a measure of participants’ epistemic uncertainty *preferences* and their responses to the epistemic EARS items for the risky option in the main decision task as a measure of their *interpretation* of uncertainty. If preferences regarding epistemic uncertainty generalise across tasks, we would expect that participants who avoided epistemic uncertainty in the Ellsberg task *and* perceived higher epistemic uncertainty associated with the risky option in the main decision task would be more likely to avoid the risky option. In other words, we predicted that the proportion of choices for the risky option could be predicted using the interaction between the responses to the Ellsberg task and epistemic EARS items.

We examined this hypothesis using Bayesian hierarchical logistic regression predicting the choice participants’ made on each trial. The indifference point on the Ellsberg,

average responses to the epistemic EARS items, and their interaction were included as fixed predictor variables. The intercept was allowed to vary for each participant. Based on this model, there was very little evidence that the interaction between responses to the Ellsberg task and epistemic EARS items influenced participants' choices between the safe and risky option. The probability that this interaction was in the predicted direction was 66.32% in Experiment 7b (median = -0.03, 95% CI = [-0.18, 0.11]) and 62.73% in Experiment 8 (median = -0.01, 95% CI = [-0.08, 0.06]). This estimate is compatible with sampling error and does not provide strong evidence that responses in the Ellsberg task are able to predict participants' responses to epistemic uncertainty.

Appendix E

Variance reduction using outcome sequences

One of the most common approaches to studying decision-making is to simply *describe* options to participants but it has become increasingly clear over the past two decades that people respond differently when they instead learn about options through *experience* (Wulff et al., 2018). Most of our decisions are based on experience rather than explicit descriptions but studying experience-based decisions has its own methodological difficulties. In contrast with choices where the probabilities are described (e.g., 50% chance of \$10, otherwise nothing), the same underlying probability distribution could result in a large number of possible experienced sequences. In the simplest possible case where people are observing the outcome of a single coin toss, there are only two possible outcomes (H or T). For two sequential coin tosses there are four possible sequences (HH or TT or HT or TH), for three coin tosses there are eight, and so on. A combinatorial explosion!

Given that the participants in our first experiment made a sequence of 110 choices and there were 80 possible outcomes associated with the risky option, no two participants would ever receive the same sequence if outcomes were generated randomly. This already daunting issue is exacerbated in partial-feedback experiments because participants' choices

influence the outcomes they observe. Each participant is wandering through a garden of forking paths. One participant might experience the worst possible outcome the first time they choose the risky option and avoid it for the remainder of the task whilst another experiences a sequence of favourable outcomes. It seems plausible that we could decrease the within-condition variability in participants' responses by decreasing the variability of their experience during the task.

One simple approach would be to assign each option a predetermined sequence of outcomes. If we generated a predetermined sequence of 110 outcomes for each option, we could ensure that every participant that chooses the risky option on a given trial receives the same outcome. We employed a similar approach but attempted to address two potential issues. The first issue is that roughly half the participants choose the safe option on the first trial and the other half choose the risky option. Suppose that the first two outcomes in the sequence for the risky option were 20 points and 75 points. For the participants that selected the risky option as their first choice, their first impression would differ considerably from those who waited until their second choice. To address this issue, we presented participants with outcomes based on the number of times they had selected each option. The first time they selected an option, they experienced the first outcome in the sequence for that option (e.g., 20 points) regardless of whether they chose it on the first trial or the 50th. The next time they selected that option, they experienced the second outcome in the sequence (e.g., 75 points) and so on.

The second issue with the simple approach is that using a single distribution impacts the ability to generalise the findings. It is possible that an experimental manipulation causes a difference between conditions using one sequence (e.g., one where the first outcome is good) but not using another sequence (e.g., when the first outcome is bad). The experimenter might reach a different conclusion based entirely on the sequence. Therefore, we attempted to reduce the within-condition variability whilst remaining confident that our results were not subject to the quirks of a single sequence by randomly assigning each participant to one of ten possible outcome sequences. We employed this method in Experiment 7a and used model comparison to determine whether using predetermined

sequences successfully reduced the within-condition variability.

We examined three Bayesian hierarchical logistic regression models that each predicted the choice participants' made on each trial. The first (baseline) model included condition (unique images or identical images) as a fixed predictor variable and allowed the intercept to vary for each participant. The outcome sequence was not included as a predictor in this model. The second (intercept) model allowed the intercept to vary for each distribution and the third (slope) model allowed the intercept and the slope for condition to vary for each distribution. These models suggest that including the outcome sequence as a predictor has no appreciable influence on our ability to estimate the effect of our primary manipulation. Firstly, we predicted that including the outcome sequence in the model would explain some variability, and therefore, that we would observe a decrease in the variability in the varying intercepts for the participants. The standard deviation was 0.65 for the baseline model and 0.64 for the intercept and slope models.

Secondly, there was a difference of less than 0.01 in the standard deviation of the posterior distributions (estimation error) of the condition parameter across the three models. Finally, we conducted ten-fold cross validation using the loo package (Vehtari et al., 2017) to estimate the expected predictive accuracy of each model. The expected log predictive density (ELPD) was highest for the baseline model suggesting that it is expected to be slightly better in terms of predictive accuracy. The difference was -0.8 for the intercept model and -2.4 for the slope model with a standard error of 1.5. Therefore, based on these three pieces of evidence, we concluded that using predetermined sequences does not offer a promising method to reduce within-condition variability in participants choices.