

Optimal decision agents?  
Biases during voluntary  
information sampling using  
eye-tracking



Hildward Vandormael

Department of Experimental Psychology

St. John's College

University of Oxford

A thesis submitted for the degree of

*Doctor of Philosophy*

Trinity Term 2017

## Table of Contents

|                                                                                              |            |
|----------------------------------------------------------------------------------------------|------------|
| <b>THESIS ABSTRACT.....</b>                                                                  | <b>5</b>   |
| <b>ACKNOWLEDGEMENTS.....</b>                                                                 | <b>6</b>   |
| <b>CHAPTER 1: PERCEPTUAL DECISION MAKING AND INFORMATION SAMPLING .....</b>                  | <b>7</b>   |
| 1.1. THEORETICAL MODELS OF PERCEPTUAL DECISION-MAKING.....                                   | 7          |
| 1.1.1 <i>Signal detection theory</i> .....                                                   | 8          |
| 1.1.2 <i>Sequential Probability Ratio Test</i> .....                                         | 10         |
| 1.1.3 <i>Drift Diffusion Model</i> .....                                                     | 11         |
| 1.3.4 <i>Limitations of current models</i> .....                                             | 15         |
| 1.2. SELECTIVE WEIGHTING OF INFORMATION DURING PERCEPTUAL DECISIONS .....                    | 18         |
| 1.2.1 <i>Selective weighting of information over time in human perceptual decisions</i> .... | 20         |
| 1.2.2 <i>Selective weighting of visual features during spatial averaging</i> .....           | 25         |
| 1.3. THE DETERMINANTS OF SACCADIC EYE MOVEMENTS .....                                        | 28         |
| 1.3.1. <i>Why do we look where we do?</i> .....                                              | 28         |
| 1.3.2. <i>Bottom-up influences on saccadic gaze control</i> .....                            | 29         |
| 1.2.3. <i>Top-down influences on saccadic gaze control</i> .....                             | 33         |
| 1.4. EYE MOVEMENTS AND PERCEPTUAL DECISION-MAKING.....                                       | 34         |
| 1.4.1 <i>Interactions between gaze direction and choice</i> .....                            | 35         |
| 1.4.2 <i>Eye movements may be allocated to reduce uncertainty</i> .....                      | 40         |
| <b>CHAPTER 2: MEASURING THE INFORMATION SAMPLING POLICY. ....</b>                            | <b>44</b>  |
| 2.1 TRADITIONAL APPROACHES TO MEASURE THE SAMPLING POLICY .....                              | 44         |
| 2.2 EYE-TRACKING AS A TOOL TO MEASURE THE SAMPLING POLICY .....                              | 46         |
| 2.2.1 <i>Width of the information processing spotlight around fixation</i> .....             | 46         |
| 2.2.2 <i>Retinal sensitivity mapping</i> .....                                               | 48         |
| 2.2.3 <i>The landscape approach</i> .....                                                    | 54         |
| 2.2.4 <i>The gaze-contingent approach</i> .....                                              | 58         |
| <b>CHAPTER 3: BENEFITS OF VOLUNTARY INFORMATION SAMPLING AND</b>                             |            |
| <b>DOMINANCE OF PERCEPTUAL FACTORS AS DRIVERS OF THE SAMPLING POLICY .60</b>                 |            |
| 3.1. INTRODUCTION.....                                                                       | 60         |
| 3.2. STUDY: PILOT 1.....                                                                     | 63         |
| 3.2.1 <i>Methods</i> .....                                                                   | 63         |
| 3.2.2. <i>Results</i> .....                                                                  | 67         |
| 3.2.3. <i>Conclusion for Study pilot 1</i> .....                                             | 71         |
| 3.3. STUDY: PILOT 2.....                                                                     | 72         |
| 3.3.1. <i>Methods</i> .....                                                                  | 72         |
| 3.3.2. <i>Results</i> .....                                                                  | 76         |
| 3.3.3. <i>Conclusion for Study Pilot 2</i> .....                                             | 88         |
| 3.4. STUDY: PILOT 3.....                                                                     | 89         |
| 3.4.1. <i>Methods</i> .....                                                                  | 89         |
| 3.4.2. <i>Results</i> .....                                                                  | 91         |
| 3.5 CONCLUSION.....                                                                          | 101        |
| <b>CHAPTER 4: ROBUST SAMPLING OF DECISION INFORMATION DURING</b>                             |            |
| <b>PERCEPTUAL DECISION MAKING.....</b>                                                       | <b>103</b> |
| 4.1 INTRODUCTION.....                                                                        | 103        |
| 4.2 METHODS .....                                                                            | 108        |

|                                                                                                           |            |
|-----------------------------------------------------------------------------------------------------------|------------|
| 4.3 RESULTS .....                                                                                         | 112        |
| 4.3.1 Effects of evidence strength $\mu$ and evidence reliability ( $\sigma$ ).....                       | 112        |
| 4.3.2 Gaze Trajectories. ....                                                                             | 114        |
| 4.3.3 Landscape Approach Allows Estimation of Processed Information.....                                  | 115        |
| 4.3.4 Landscape Analysis Soft vs. Hard approach .....                                                     | 118        |
| 4.3.5 Estimating Both Quantity and Quality of Processed Information.....                                  | 120        |
| 4.3.6 Decisions Depend on Gaze Location. ....                                                             | 123        |
| 4.3.7 Gaze Is Directed Preferentially to Inliers Over Outliers. ....                                      | 125        |
| 4.3.8 Humans show Robust Averaging of Numbers.....                                                        | 130        |
| 4.3.9 Linking Sampling Policy and Human Decisions.....                                                    | 132        |
| 4.3.10 Computational Simulations.....                                                                     | 137        |
| 4.3.11 Determinants of Human and Model Performance.....                                                   | 143        |
| 4.3.12 Mechanisms underlying the effect of the sampling policy on accuracy.....                           | 146        |
| 4.3.13 Confirmation bias.....                                                                             | 155        |
| 4.4 DISCUSSION .....                                                                                      | 158        |
| 4.4.1 Future studies .....                                                                                | 161        |
| 4.4.2 Limitations.....                                                                                    | 163        |
| <b>CHAPTER 5: VOLUNTARY INFORMATION SAMPLING WHEN DECIDING BETWEEN</b>                                    |            |
| <b>TWO CATEGORIES (EXPERIMENT ONE).....</b>                                                               | <b>169</b> |
| 5.1 INTRODUCTION.....                                                                                     | 169        |
| 5.1.1 Decisions from description vs. decisions from experience.....                                       | 170        |
| 5.1.2 Mechanisms underlying the description-experience gap: the importance of the<br>sampling policy..... | 174        |
| 5.1.3 The gaze cascade.....                                                                               | 177        |
| 5.1.4 Selective encoding.....                                                                             | 182        |
| 5.1.5 The attentional drift diffusion model.....                                                          | 183        |
| 5.2. CURRENT STUDY.....                                                                                   | 185        |
| 5.2.1 Experiment 1 Methods.....                                                                           | 187        |
| 5.3 RESULTS .....                                                                                         | 192        |
| 5.3.1 Effects of evidence strength $\mu$ and evidence reliability ( $\sigma$ ).....                       | 193        |
| 5.3.2 Framing effect on the sampling of information.....                                                  | 196        |
| 5.3.3 Q1. Sampling for a prolonged period of time and the effect on decision accuracy<br>.....            | 198        |
| 5.3.4 Sampling policy and the weighting of information: recency.....                                      | 201        |
| 5.3.5 Approximating the decision evidence: LPR.....                                                       | 204        |
| 5.3.6 Q2. Frequency bias: preference for the overly sampled category.....                                 | 206        |
| 5.3.7 Q3. Consistency effect: switching between categories.....                                           | 209        |
| 5.3.8 Computational modeling: recency, frequency, & consistency bias.....                                 | 215        |
| 5.3.9 Q4. Drivers of the sampling policy: information selection.....                                      | 215        |
| 5.3.9.1 Sampling policy: asymmetrical sampling and the equity policy.....                                 | 215        |
| 5.3.10 Q5. Category uncertainty as a driver of the sampling policy.....                                   | 220        |
| 5.3.11 Sampling policy: dwell times?.....                                                                 | 225        |
| 5.3.12 Sampling policy: when to stop sampling?.....                                                       | 227        |
| 5.4 DISCUSSION .....                                                                                      | 229        |
| <b>CHAPTER 6: VOLUNTARY INFORMATION SAMPLING WHEN DECIDING BETWEEN</b>                                    |            |
| <b>TWO CATEGORIES (EXPERIMENT TWO).....</b>                                                               | <b>233</b> |
| 6.1 INTRODUCTION.....                                                                                     | 233        |
| 6.1 CURRENT STUDY.....                                                                                    | 234        |
| 6.3 RESULTS .....                                                                                         | 240        |
| 6.3.1 Effects of evidence strength $\mu$ and evidence reliability ( $\sigma$ ).....                       | 240        |

|                                                                                                                      |            |
|----------------------------------------------------------------------------------------------------------------------|------------|
| 6.3.2 Q1. Sampling for a prolonged period of time and the effect on decision accuracy .....                          | 243        |
| 6.3.3 Sampling policy and the weighting of information: recency.....                                                 | 244        |
| 6.3.4 Q2. Frequency bias: preference for the overly sampled category.....                                            | 247        |
| 6.3.5 Q3. Consistency effect: switching between categories.....                                                      | 249        |
| 6.3.6 Drivers of the sampling policy: information selection.....                                                     | 257        |
| 6.3.6.1 Sampling policy: asymmetrical sampling and the equity policy.....                                            | 257        |
| 6.3.7 Q5. Category uncertainty as a driver of the sampling policy.....                                               | 261        |
| 6.3.8 Sampling policy: when to stop sampling?.....                                                                   | 264        |
| 6.4 DISCUSSION .....                                                                                                 | 266        |
| <b>CHAPTER 7: CONCLUSION.....</b>                                                                                    | <b>270</b> |
| 7.1 IDENTIFIED DRIVERS OF THE SAMPLING POLICY.....                                                                   | 270        |
| 7.1.1 Confirmation bias at the perceptual feature level.....                                                         | 270        |
| 7.1.2 Confirmation bias at the decision information level .....                                                      | 271        |
| 7.1.3 Robust sampling .....                                                                                          | 271        |
| 7.1.4 Summary sampling .....                                                                                         | 272        |
| 7.2.5 Equity sampling .....                                                                                          | 272        |
| 7.2 EFFECTS OF THE SAMPLING POLICY ON DECISION-MAKING .....                                                          | 273        |
| 7.3 DEVELOPMENT OF THE LANDSCAPE APPROACH TO ANALYSE EYE-MOVEMENT DATA .....                                         | 274        |
| 7.4 THE IMPORTANCE OF ACKNOWLEDGING THE SAMPLING POLICY AS AN INTEGRAL STAGE OF THE<br>DECISION MAKING PROCESS. .... | 275        |
| <b>APPENDIX.....</b>                                                                                                 | <b>277</b> |
| APPENDIX 1.1 CAPTURING THE REGENCY, FREQUENCY, AND CONSISTENCY BIAS .....                                            | 277        |
| Appendix 1.1 A Bayesian model implementing different implementations of a leaky<br>process.....                      | 277        |
| Appendix 1.2 An adaptive gain-field model .....                                                                      | 281        |
| APPENDIX 2.1 LIST OF THE EXPLORATORY ANALYSES IN THIS THESIS. ....                                                   | 285        |
| <b>REFERENCES.....</b>                                                                                               | <b>286</b> |

## **Thesis abstract**

This thesis aims to understand the policy by which humans sample information, and how this policy influences perceptual decision-making. We will address this question using a variety of methods, including behavioural measures, computational modelling, and eye tracking. First, we will highlight the importance of considering an explicit information sampling policy as an integral component of the decision-making process. Secondly, we will focus on the determinants of this information sampling policy. Finally, we will address potential methodological concerns and provide a novel analysis approach to combat these concerns.

In this first chapter, we provide a brief overview of theoretical models of perceptual decision-making, followed by a discussion about optimality and biases. Finally, we will discuss the necessity of treating information sampling as an integral component of decision-making via an overview of recent work on this topic.

## Acknowledgements

I would like to thank the ESRC and St. John's College for funding most of my D.Phil. St. John's has provided me with an incredibly stable and supportive environment.

My supervisor Chris and the Summerfield lab were instrumental to the completion of my thesis. I have been able to grow as a person and learn from an amazing group of people during my time in the lab. Thank you Sam Cheadle and Jan del Ojo Balaguer for helping me on my journey through the wonderful world of coding.

I would like to thank Santiago Herce-Castañón and Ornela de Gasperin for providing me with the necessary victories in sports when the academic results were not always as expected. I could not have asked for better friends. My family: Mama, Papa, Karel, and Kiri, who have always motivated me to reach for the stars.

Lastly, I would like to thank Vickie. I have never met a more loving, caring, and supportive individual; all tightly packed in such a small frame! Words cannot express how important you have been to me throughout this entire endeavour.

Thank you.

# Chapter 1: Perceptual Decision Making and Information Sampling

## 1.1. Theoretical models of perceptual decision-making

Perceptual decision-making is the process by which we detect, discriminate, and categorize information that is sampled from the environment via our senses (Hanks & Summerfield, 2017). Perceptual decision-making is classically divided into three distinct processes: perceptual encoding, information integration, and motor preparation. The first process, perception, is concerned with the mapping of external information signals onto internal sensory representations. Thus, perception can be seen as an initial interpretation of the external world by the organism. However, sensory signals can originate from multiple distinct sources. For example, they can occur at different spatial locations, or at different times. The second process - integration – refers to the computational mechanism by which information is combined across space and time, into a single quantity or “decision variable”. The third process, motor preparation, describes how the organism transforms this decision variable into a discrete motor act.

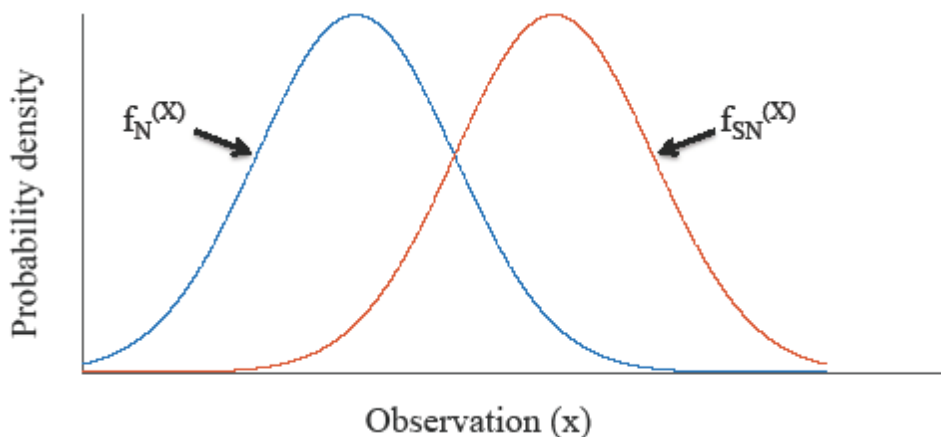
Let me illustrate with a simple example: you are at the supermarket wishing to buy fruit. You sample the environment by successively fixating portions of the environment, e.g. the different shapes and colours of the produce on the shelves, transforming the external information signals into internal representations (~perceptual encoding, process 1). You combine the information signals that originate from multiple sources of information in the visual field such as the

bananas, strawberries, and plums (~ information integration, process 2). Finally, the integrated information signal is transformed into a decision about which of the available fruit has the highest quality, and thus which to purchase (~ motor execution, process 3).

### **1.1.1 Signal detection theory**

Over the past decades, a vast array of distinct but interrelated quantitative models have been used to describe the computational processes underlying perceptual decision-making. The modern approach to understanding perceptual decisions can be traced back to Signal Detection Theory (SDT) (Green & Swets, 1974; Tanner Jr & Swets, 1954), which was developed in the 1940s to quantify the perceptual of noisy sensory signals, for example on a telephone line. The key intuition that distinguishes signal detection theory from earlier accounts is that we make perceptual decisions under uncertainty. Signal detection theory draws upon the foundational notion that true states of the world (e.g. H1: raspberries are better; H2: strawberries are better) are only partially observable, and must be estimated via noisy sensory signals. Thus when evaluating the quality (size, or colour) of fruit in the supermarket, our estimates will be imperfect, because they are corrupted by noise. This noise can arise both in the environment itself (e.g. the fruit may be only partially visible) and in early stages of the internal processing system (e.g. neural signals are intrinsically variable, giving rise to variability in estimates of fruit colour or size) (Barlow, 1961; Faisal, Selen, &

Wolpert, 2008). Internal noise has been estimated via the equivalent noise method (Pelli, 1981). However, note that recent work has argued that effects that were often attributed to internal noise, were in fact due to changes in sampling of the input or changes in gain (A. Baldwin, Baker, & Hess, 2016). Thus, in order to make decisions about which fruit to buy in the supermarket, we need to be able to calculate the likelihood of two possible states of the world  $p(H1)$  and  $p(H2)$ , conditioned on this noise. SDT accomplishes this by calculating the (log) likelihood ratio  $\log \left( \frac{p(x|H1)}{p(x|H2)} \right)$  that embodies the relative probability of either state of the world (**Fig. 1.1**).



**Fig. 1.1 Signal Detection Theory (SDT).** Figure adapted from Swets, Tanner, & Birdsall, 1961, with permission from APA. By comparing the likelihood that a particular observation ( $x$ ) is driven by noise ( $f_N(x)$ ) or by true signal ( $f_{SN}(x)$ ).

In SDT decisions are made when the likelihood ratio is compared to a criteria value. For an unbiased observer making choices about equally value and equiprobable stimuli, this criterion lies at 0 in log space (e.g. likelihood ratio of 1, or  $p(x|H1)=p(x|H2)$ ). However, SDT allows the flexible placement of this criterion

boundary so that the discrimination sensitivity between signal and noise can be maximised for a particular range of input values.

### **1.1.2 Sequential Probability Ratio Test**

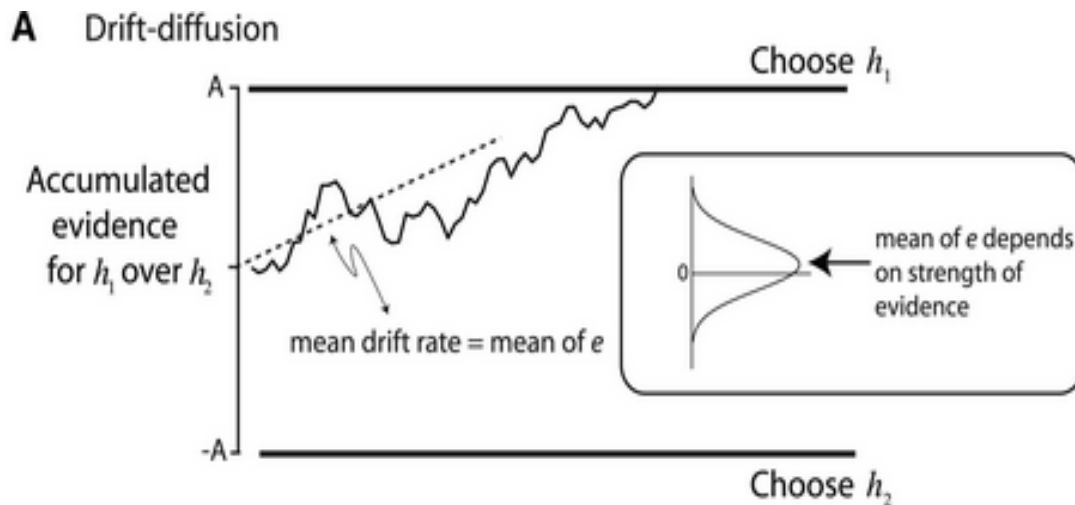
One limitation of SDT is that it is a static model, e.g. it assumes that all the available decision information arrives at once. It thus skips step ~2 in our decision model, whereby information needs to be combined over time and space. Critically, thus, it does not allow an estimation of *when* a decision would be made, i.e. a prediction about the response time that will be produced by an observer making a categorisation judgment.

The Sequential Probability Ratio test (SPRT) was the crucial development that extended SDT in a way that allowed decision latencies to be estimated (Wald & Wolfowitz, 1948). The SPRT implements a sequential Bayesian process that calculates the posterior likelihood of a particular hypothesis (e.g. H1) on the basis of multiple sequential sources of information (e.g. multiple encounters with a particular fruit, or multiple fixations across the supermarket aisle). Within this framework, the posterior likelihood of a particular hypothesis is proportional to the cumulative likelihood ratio of the observed cumulative evidence. Critically, the SPRT allows for a bound to be placed on the requisite evidence, corresponding to a desired level of accuracy. Information can thus be sampled and integrated

until this bound is reached, at which point a decision is made. Therefore, SPRT will find the most likely hypothesis (H1 vs. H2) in the shortest amount of time. Many subsequent models known as *sequential sampling models* explicitly built on this framework to model perceptual decision-making, drawing a link between current models and the original framework of SDT (Bogacz, Brown, Moehlis, Holmes, & Cohen, 2006).

### 1.1.3 Drift Diffusion Model

Among the most popular of this new class of sequential sampling model is the drift-diffusion model or DDM (Bogacz et al., 2006; Ditterich, 2006; Gold & Shadlen, 2007; Purcell, Schall, Logan, & Palmeri, 2012; Ratcliff, 1978). This model is based on the SPRT, in that it assumes that information is accumulated linearly (i.e. via summation) to a threshold. However, rather than assuming that the average quality of each category (e.g. the mean evidence produced under H1) and the level of noise in the world (e.g. the variability of the evidence produced under H1) are known, it fits them as free parameters to the data. The basic assumption is that the brain extracts information from the stimulus at each time point (drift) while this information is disturbed by noise. This information accumulates over time via summation (diffusion) towards a threshold, at which point a decision is made (**Fig. 1.2**). Thus, according to the DDM, decisions are optimised by sequential sampling, at the cost of delaying responding (Bogacz et al., 2006).



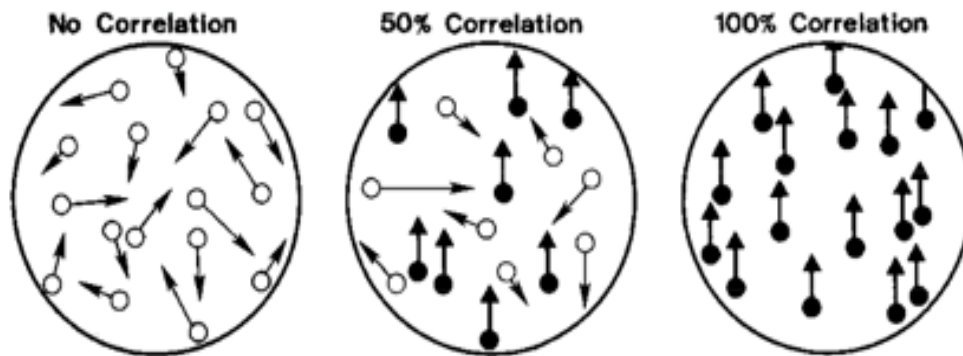
**Fig. 1.2 The Drift Diffusion Model (DDM).** Figure copied from (Shadlen & Roskies, 2012) under CC BY-NC 3.0 copyright. Within the DDM, the decision process is modelled as the integration of evidence from a continuous stream of information (curved line), starting from zero evidence (starting point), moving towards one of two potential decision boundaries (i.e. choose  $h_1$  or  $h_2$ ). The mean drift rate corresponds to the rate of evidence accumulation and depends on the strength of the evidence. Thus, a noisier stimulus will lead to a lower mean drift rate, and therefore more samples will be necessary to reach a decision boundary.

We can illustrate this type of model to the example of deciding whether you should buy raspberries ( $H_1$ ) or strawberries ( $H_2$ ). Let us first assume that raspberries and strawberries are, *a priori*, on average of equal quality. Therefore, the starting point for the accumulation of evidence would be equidistant from the two decision thresholds (buy strawberry vs. buy raspberry). As you sample the

environment with your eyes, the accumulation process can drift towards the H1 boundary (i.e. buy raspberries) or the H2 boundary (buy strawberries). The magnitude of the drift towards either boundary will depend on the strength of the evidence, which is given by the relative average quality of strawberries and raspberries on the shelves. The noise is given by variability in the environment, i.e. the extent to which the true state of the world (H1 vs. H2) is observable from the sensory information, and/or to which processing is corrupted by internal noise. The threshold can be set at a desired level that will vary the trade-off between decision speed and accuracy. It is worth noting that the starting point of the diffusion process need not be unbiased. For example, if you visit the supermarket regularly and the raspberries of higher quality, you might start with a prior belief that favours H1. This type of bias can be implemented in the DDM model by using a non-zero starting point for the accumulation of evidence: in this case the starting point would be closer to the 'choose H1' threshold.

The DDM has been used to model a wide range of tasks in which data about binary decisions and their latencies are available, including numeracy judgments, priming, social decisions, clinical research, memory, and many more (Klauer, Voss, Schmitz, & Teige-Mocigemba, 2007; McKoon & Ratcliff, 2012; Ratcliff, Thompson, & McKoon, 2015); see Voss et al., 2013 and Ratcliff, Smith, Brown, McKoon (2016) for a comprehensive overview (Ratcliff, Smith, Brown, & McKoon, 2016; Voss, Nagler, & Lerche, 2013). However, perhaps the most popular vehicle is a psychophysical task that was previously used for investigating motion

perception: the Random Dot Motion task (RDM) (Britten, Shadlen, Newsome, & Movshon, 1992) . The RDM task requires an observer to discriminate the overall direction of motion from a cloud of moving dots, e.g. up vs. down (**Fig 1.3**).



**Fig 1.3 The Random Dot Motion task (RDM).** Figure copied from Britten et al. (1992) with permission through RightsLink. Observers are required to discriminate the overall direction of motion from a cloud of moving dots. These dots can move in differing ‘motion coherence’, e.g. varying degrees of correlation between their movements. The motion coherence corresponds to the strength of the evidence.

In many tasks, including the RDM task, the DDM offers a good description of the decisions and decision latencies produced by humans and monkeys. Like other sequential sampling models, it correctly predicts fundamental features of decision latencies, such as the ways in which humans trade off speed and accuracy, or the rightward skew typically observed in RT distributions. With minor extensions,

it can also predict the relative speed of correct and error trials under different manipulations, such as variation in the prior probability of H1 and H2 (Ratcliff & Rouder, 1998).

#### **1.3.4 Limitations of current models**

In the SPRT and DDM, evidence accumulation is now portrayed as a sequential process over time. However, they also assume that evidence accumulation is a wholly passive process i.e. a forward passage of information from cumulative to cumulative quantities (Gottlieb, Hayhoe, Hikosaka, & Rangel, 2014). In most reported psychophysical studies, the observer typically sits in front of a screen and passively samples from a single, continuous, centrally presented stream of information about motion direction, and decides whether the average dot motion tends in one direction or another. Although this task offers a limited purview on the types of decision humans make in the real world, these psychophysical studies have often been argued to offer a window on the general mechanisms by which humans make perceptual decisions (Shadlen & Kiani, 2013). This contention is among those that I wish to question in the work described in this thesis. I will argue that when making decisions we *selectively sample* a subset of the available information. One key mechanism by which this occurs is the use of successive eye movements (fixations), which allow us to focus on one portion of egocentric space at the expense of another.

In contrast with psychophysical studies in which information integration is seen as a passive process, decisions about natural scenes are made by actively saccading (or covertly attending) to successive image locations and selective sampling the subset of information that is available there. Furthermore, eye movements in natural behaviour are tied into subsequent actions and serve to accumulate sensory evidence relevant to those actions (Tatler, Hayhoe, Land, & Ballard, 2011). For example, a foraging monkey might move its eyes freely to several locations in the forest canopy, and average the available fruit yield in order to decide where to move. Real-world category decisions, thus, are affected not only by the quality of the sensory information and the dynamics of integration, but also by the policy that dictates how information is sampled and weighted when forming a decision. However, the nature of this policy remains largely unknown.

Furthermore, an increase in our understanding about the nature of this sampling policy will afford us richer descriptions of the underlying mechanisms of decision-making. Indeed, researchers in the economics literature have called for more investigations in the sampling policy, in order to open up the black box of decision makers (Rubinstein, 2003). The information sampling stage is a crucial part of the decision-making process and the advantages of incorporating this stage in data collection are well described in recent work (Schulte-Mecklenbeck, Kühberger, & Ranyard, 2011). In their work, the authors argue that data from the information sampling stage allows improved differentiation between theoretical

accounts, as these accounts often yield different quantitative and testable predictions at the information sampling stage (e.g. see our discussion of the pursuit of non-instrumental information). Furthermore, investigation of the sampling policy allows researchers to further probe the mechanisms underlying individual differences (Schulte-Mecklenbeck et al., 2011).

We will argue that selective information sampling is particularly important because human processing capacity is limited (Broadbent, 1958; Burgess & Colborne, 1988; Dijksterhuis, Aarts, & Smith, 2005; Ebbinghaus, 1913; Nørretranders, 1998; Spiegel & Green, 1981; Wilson, 2004). For example in the classic model proposed by Broadbent (1958), only a subset of the information is passed through a selective processing filter, while the majority of available information is filtered out and ignored. In this view, attention serves to allocate our limited resources efficiently in order to avoid information overload. I will argue that the selective information sampling stage allows the decision maker to attend to only a subset of the information that is relevant to the task, leading to increased efficiency in decision making. To illustrate, when scanning the fruit section in the supermarket, it is impossible to encode all of the sensory information available on the shelves in a single instant, to accurately compute the ground truth value for the cumulative evidence. Instead, one must selectively sample a subset of the information, for example by focussing on particular locations or features of the fruit (size or colour) at the expense of others. In order to make the best choice, one might sample information about the attributes of a particular

item: are these fruits? are they red (ripe) or green (unripe)? Are they large (nutritious) or small (less nutritious)? Each attribute that is sampled contributes to the overall value of the candidate and has a corresponding influence your decision.

How is information selectively sampled during human perceptual decision-making? A great deal of research has been devoted to understanding the computational mechanisms by which information is integrated over time, and which neural computations may give rise to this process. However, as argued above, the selective information sampling stage during perceptual choices remains relatively unexplored, leaving an important gap in the literature. The goal of this thesis is to help fill this gap.

## **1.2. Selective weighting of information during perceptual decisions**

In this thesis, we will argue that this filtering of information encompasses two processes. One is an overt selection process, for example when observers move their eyes to distinct locations in space (Gottlieb et al., 2014). Overt attention, thus, is a key mechanism by which information is selectively sampled (see below). We will begin, however, by discussing evidence that information is internally filtered or reweighted during perceptual decisions, in a way not

currently accounted for by the DDM and other popular sequential sampling models that are based on the SPRT.

According to the DDM and SPRT, noisy cumulative decision information is accumulated linearly to a bound. This means that each sample of information carries a weight in the decision process that is proportional to its degree of sensory reliability. One useful analogy is to consider a task in which the cumulative samples of information are numbers, and the task is to determine which of two streams is drawn from a distribution with the higher average number. For a perfect integrator (e.g. a human using a calculator), each number will have an impact that is proportional to its value. Thus, the optimal policy is to add up the numbers with equivalent weight (e.g. 1), rather than, say, ignoring those numbers that are less than 3 or greater than 7 and integrating the others. Agents that deviate from this policy are discarding information, and will, under idealised conditions at least, perform more poorly on the task.

Recently, studies of perceptual decision-making in humans have asked whether humans give equal weight to discrete samples of information. To address this question, they have moved away from the psychophysical paradigm in which samples (e.g. frames of the RDM stimulus) are numerous and occur only fleetingly, making their weight difficult to assess empirically. Instead, they have used tasks in which observers view a sequence or array composed of discrete

items, each associated with a given feature value (e.g. gratings defined by tilt), and are asked to categorise the mean feature relative to a category boundary (e.g. the vertical axis of orientation). In this task, as in the numbers task, the best policy is to average the feature values with equivalent weight. These studies have revealed evidence for selective weighting of information during human perceptual decision making.

### **1.2.1 Selective weighting of information over time in human perceptual decisions**

Some of this evidence comes from tasks in which the samples of information occurred separately in time, and had to be averaged over a temporal sequence. For example, Wyart and colleagues found evidence for selective weighting over time in neural signals. The influence of a given sample of evidence on decisions was modulated by its temporal occurrence relative to a slow parietal signal, measurable with scalp EEG, that oscillated at approximately 2Hz (Wyart, De Gardelle, Scholl, & Summerfield, 2012). The authors suggested that information is weighted in this way to avoid interference between adjacent samples that might compete for processing resources, due to the limited capacity of human decision processes (Broadbent, 1958). In other words, information integration stage (stage ~2) is capacity limited, and involves a computational step that selective filters or downweights information in time. Interestingly, the temporal profile of this filtering mechanism matches that of an oscillatory rhythm observed

in nonhuman primates performing a multimodal perceptual task, which controlled the firing rates associated with a sensory input (e.g. rhythmic gain control) (Lakatos, Karmos, Mehta, Ulbert, & Schroeder, 2008). This finding in the visual domain was subsequently replicated for auditory and somatosensory decisions (Spitzer, Blankenburg, & Summerfield, 2016), and subsequent investigations addressed how this rhythmic weighting was modulated under divided and focussed attention (Wyart, Myers, & Summerfield, 2015).

Selective weighting might also depend on the statistics of the simulation sequence. Using an equivalent task, Cheadle and colleagues asked how the weight given to a sample of information depended on the state of the decision variable (Cheadle et al., 2014). Under the SPRT and DDM, each sample of information should be treated independently, so that the ordering of the samples in the stimulation sequence should not affect the judgment made by the observer. They found, by contrary, that samples of information that were consistent with the current state of the decision variable carried more weight, as if participants were selectively filtering information that contradicted a current hypothesis. This led to a suboptimal, confirmatory bias in decision-making. Neural correlates of this downweighting of inconsistent samples were also observed in three classes of physiological signal: pupil diameter, scalp EEG, and fMRI signals.

These studies suggest that during perceptual decision-making, information is selectively weighted according to its position in a stimulation sequence, and according to the statistics of the information that has occurred thus far. These effects of temporal context on the weighting of information in perceptual decisions are not predicted by the SPRT or DDM in their classic form. Rather, they imply that humans deploy active selection mechanisms when making perceptual decisions. The authors describe a new sequential sampling model variant, called the adaptive gain control model, that can account for this “suboptimal” weighting of information in response to capacity limits.

Limitations in processing capacity have been argued to be one reason why humans show biases in decision-making. In the following section I will briefly describe two further experimental findings that highlight the importance of emphasising the information sampling stage during decision-making. In particular, I will focus on the fact that prolonged sequential sampling does not always lead to better decisions, as one would predict under the sequential sampling framework (Bastardi & Shafir, 1998; Dijksterhuis, 2004; Drugowitsch, Moreno-Bote, Churchland, Shadlen, & Pouget, 2012; Kahneman, 2011).

Using a more descriptive task, Bastardi and Shafir showed how sampling extra information could lead to worse decisions (Bastardi & Shafir, 1998). In this study, participants were presented with either a scenario containing (i) all the

information about the scenario at once or (ii) a subset of the information and participants could choose to pursue additional information. In this second scenario, the information was split up into two sections. Participants could sample the information, but had to wait some time before they could sample some extra non-instrumental information. The authors found that pursuing the additional non-instrumental (e.g. not relevant) information could lead to inappropriate (re-) weighting of information and thus lead to decisions incompatible with participants' true preference.

Similarly, research on hypothesis testing has focused on knowledge-driven models of information selection (Lindley, 1956). This literature has also focused on uncertainty reduction (or information gain): a potential sample is valued to the extent that it is likely to reduce uncertainty about the true hypothesis (Nelson, 2005; Skov & Sherman, 1986; Slowiaczek, Klayman, Sherman, & Skov, 1992). However, it has been shown that in certain environments, heuristics predict the same decisions as a sampling policy driven by information gain (Klayman & Ha, 1987, 1989; Navarro & Perfors, 2011). One such heuristic, the positive test strategy, is an active sampling policy by which we seek out information that confirms our current hypothesis (Klayman & Ha, 1989; Nickerson, 1998). This confirmatory bias is similar to that observed in the study by Cheadle et al (2014).

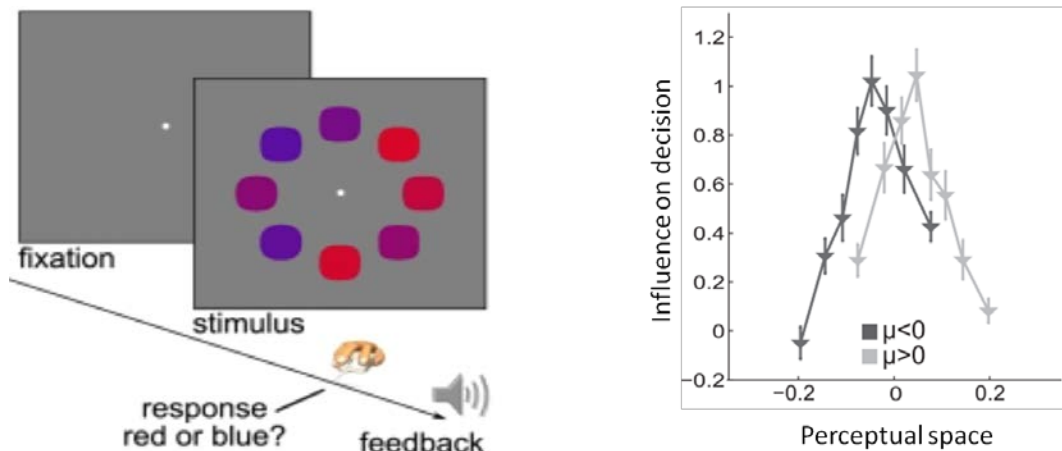
Similarly, research on preference reversals has clearly shown that the weights assigned during preference construction are (i) not consistent or reliable and (ii) context dependent (Huber, Payne, & Puto, 1982; Lichtenstein & Slovic, 1973; Pettibone & Wedell, 2007; Simonson, 1989; Slovic & Lichtenstein, 1983; Tsetsos, Usher, & Chater, 2010). Preference reversals occur when the same decision is systematically responded to in opposing ways in two otherwise irrelevant contexts. Imagine a scenario in which you are deciding whether you would buy strawberries or raspberries, and prefer strawberries. Subsequently, you notice another fruit (plums) nearby. Although you prefer both strawberries and raspberries to plums, the mere presence of this third option in the choice set causes you to change your mind and purchase raspberries. What happened? This preference reversal is a violation of one of the principles of axiomatic rationality, that decisions should be independent from irrelevant alternatives (IIA) (Luce, 1977; Von Neumann & Morgenstern, 1947). Violations of IIA are typically attributed to distortions in evaluative processes, such as loss aversion (Kahneman & Tversky, 1979), to heuristic decision policies (Tversky & Simonson, 1993) or to internal neurophysiological mechanisms, such as value normalisation (Louie, Khaw, & Glimcher, 2013). However, a new sequential sampling model, the selective integration model, has been developed which explains preference reversals in terms of selective weighting of information during evidence integration (Tsetsos, Chater, & Usher, 2012; Tsetsos et al., 2016). The selective integration model proposes that information competes for limited resources, as if more salient choice attributes are preferentially sampled, and are thus more likely to impact decisions. Therefore, resources have to be allocated asymmetrically to

samples of information about each option. A critical recent development is the realisation that in the presence of noise that arises “late” in the decision process, i.e. that is not due to irreducible variability in the external world, or to noise in sensory encoding – but instead to uncertainty that arises during or beyond information integration – selective integration can paradoxically maximise reward. In other words, the policy that maximises reward for an ideal (“perfect”) integrator is not necessarily the policy that maximises outcomes for humans. This is a theme that we will return to later in the thesis.

### **1.2.2 Selective weighting of visual features during spatial averaging**

In natural scenes, information can arrive from multiple, spatially distinct locations, and must be integrated to select a course of action. Is information selectively weighted during averaging of spatial information? Using a different task, de Gardelle and Summerfield (2011) presented participants with a multi-element array, composed of 8 coloured shapes (“squircles”) arranged in a ring around fixation (De Gardelle & Summerfield, 2011). Observers were asked categorise the array item according to their average colour (more red vs. more blue) or average shape (more circle vs. more square). The authors reported that participants selectively downweighted extreme samples of perceptual evidence, i.e. those that were more outlying with respect to the category boundary. For example, when classifying red vs. blue, they gave less weight to those items which were very red or very blue (**Fig 1.4**; (De Gardelle & Summerfield, 2011)).

Control analyses verified that this was not simply due to nonlinearities in colour space.



**Fig 1.4 Experimental design of the multi-element averaging task.** Figure copied from de Gardelle & Summerfield (2011) under CC BY-NC 3.0 copyright. On each trial, subjects had to average the colour or shape from an eight element array, in order to classify the display as either predominantly red or blue (left). The responses of subjects were less influenced by more extreme colours (e.g. further away from 0 in perceptual space) in comparison to less extreme colours (e.g. closer to 0), providing evidence of the selective weighting of decision information (right).

Similar findings have also been reported for other classes of stimuli, even when the features in question are high-dimensional, such as facial expression (Haberman & Whitney, 2010). These studies, and subsequent replications of this

finding (Li, Castañón, Solomon, Vandormael, & Summerfield, 2017) suggest that humans engage in a particular form of selective weighting: they downweight outlying information, that which lies at the extremes of feature space. The authors term this policy ('robust averaging'), as the discounting of outliers is a major method by which statistical inference is made robust to the presence of aberrant data points. Importantly, finding cannot be accounted for by the SPRT or DDM without further modification, although it is consistent with the adaptive gain framework proposed by the authors (Cheadle et al., 2014; Summerfield & Tsetsos, 2015).

However, this work leaves open a key question. In the spatial averaging task, observers were encouraged to maintain their gaze at a central fixation spot, and performance was not dependent on the presentation duration of the array. However, the authors did not measure or control eye movements during task, so they were unable to make strong claims about whether the robust averaging observed was due to discarding of information at an internal weighting stage, or whether participants selectively sampled inliers over outliers with their gaze. More generally, the role of eye movements in the selection process that underlies perceptual decisions remains a relatively unexplored topic (see below). A major goal of this thesis is thus to elucidate how selective integration sampling depends on the allocation of gaze across a spatial array.

### **1.3. The determinants of saccadic eye movements**

The abovementioned studies suggest that humans engage in selective weighting of information in spatial arrays that attempt to mimic the complexity of natural scenes. However, as mentioned above, another key mechanisms by which information is selectively sampled is through saccadic eye movements. A growing number of studies have begun to address this question, and these In this section, I briefly summarise an extensive literature that has sought to understand the determinants of saccadic eye movements, to understand how humans choose to place their gaze when viewing a sensory stimulus.

#### **1.3.1. Why do we look where we do?**

In a real-world setting, information arises in crowded natural scenes, replete with multiple features and objects that compete to drive which information is sampled, and ultimately drives decisions. The determinants of saccadic eye movements (why do we look where we do?) have been a pertinent question in the literature for decades (Buswell, 1935; Alfred L Yarbus, 1967). This question has been extensively investigated in a variety of paradigms including visual search (see (Eckstein, 2011) for a review), reading (see (Starr & Rayner, 2001) for a review), smooth pursuit (see (Schütz, Braun, & Gegenfurtner, 2011) for a review), and natural scene perception (Henderson & Hollingworth, 1999; Torralba, Oliva, Castelhana, & Henderson, 2006). However, more than 60 years later, the determinants of saccadic eye movements are still not fully understood and no

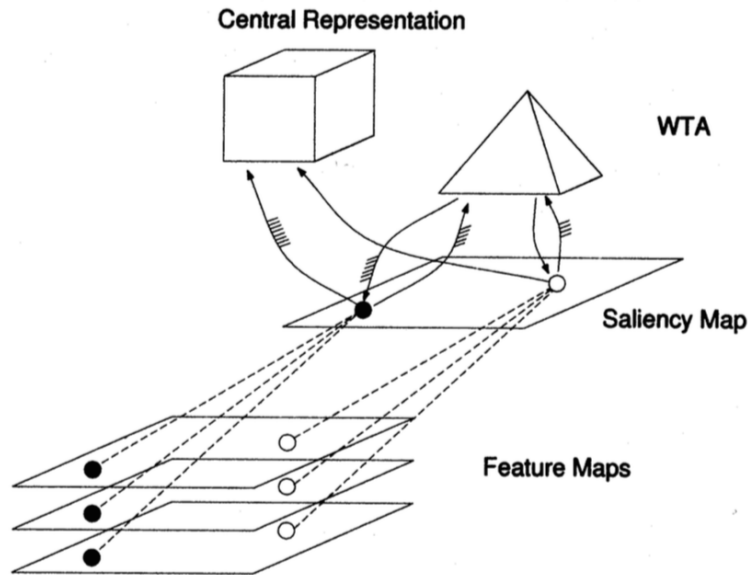
satisfactory consensus has been reached on the determinants of gaze placement (Tatler, 2009). In this line of work, a distinction is usually made between bottom-up (i.e. driven by properties of a stimulus) and top-down (i.e. driven by knowledge of the observer) influences on the information sampling stage (i.e. the sampling policy). We will begin by discussing bottom-up factors.

### **1.3.2. Bottom-up influences on saccadic gaze control**

Evidence for bottom-up influences is provided by studies of saccadic control which revealed that attention is automatically captured by perceptually salient information, such as a target defined by a unique colour in a visual search experiment (A. L. Yarbus, 1967). In a typical visual search task, participants need to identify a target, which is surrounded by several distractors. The target and distractors can be differentiated on the basis of simple visual dimensions (e.g. “find the red vertical line amongst blue vertical lines and red horizontal lines”). As one would expect, increasing the number of distractors leads to longer search times. However, when the target can be detected on the basis of a single feature, there is little to no cost of increasing the number of distractors (Treisman & Gelade, 1980). This “pop-out” effect is a classic illustration of a bottom-up influence on the sampling policy.

Furthermore, it has been argued that attention can be automatically captured by semantically deviant information, such as an object in an unusual context, like a sofa on the beach (Võ & Henderson, 2010), even when this item does not “pop out” perceptually. However, this contention remains controversial.

The mechanisms underlying these saliency-driven effects on information sampling have been extensively investigated and an influential computational model was developed (Itti & Koch, 2000; Itti, Koch, & Niebur, 1998; Niebur & Koch, 1996). This saliency map model is derived from feature integration theory (Treisman & Gelade, 1980): individual feature maps combine to form a salience map and a winner-takes-all network then guides visual attention (Itti & Koch, 2000; Koch & Ullman, 1987; Niebur & Koch, 1996); **Fig. 1.5**).



**Fig 1.5. Illustration of a simple saliency map model.** Figure copied from Itti & Koch (2000) with permission under RightsLink. Retinal input drives parallel computation of early visual features such as colour or luminance in a set of pre-attentive feature maps. Thus, each feature map has activity at feature specific locations on the map, and this activity from all of these feature maps is then combined in a saliency map. The activity in the topographic saliency map inputs to a winner-take-all (WTA) network that detects the most salient location of the map and directs attention towards it. Thus, only features from the 'winning' location on the saliency map reach a central representation.

A large part of the literature on information sampling has focused on this saliency map model, investigating which features should be part of the map, and how these features are then combined at a later stage (Einhäuser & König, 2003; Jansen, Onat, & König, 2009; Nothdurft, 2000; Zhao & Koch, 2011). Recently,

researchers have started to use machine-learning techniques to investigate the determinants of saccadic eye movements. The main advantage of this approach is that it is no longer necessary to specify which features should be part of the map or how these features are combined (Kienzle, Franz, Schölkopf, & Wichmann, 2009).

The saliency map model has been very popular because of its neurological plausibility (Carandini et al., 2005; Kusunoki, Gottlieb, & Goldberg, 2000; Thompson & Bichot, 2005), while achieving remarkable accuracy in predicting fixation locations, for such a simple model that only works on individual features (57-68% accuracy, (Betz, Kietzmann, Wilming, & Koenig, 2010). It is important to note, that even though such a simple model can predict fixation well, it still fails to explain 32% to 43% of the remaining fixations. Thus, the saliency map model does not explain all features of the sampling policy. In line with this observation, researchers have expanded these models of saccadic control with object processing modules, oculomotor biases like the center bias, and top-down influences to greatly improve performance (Cerf, Frady, & Koch, 2009; Tatler & Vincent, 2009).

### **1.2.3. Top-down influences on saccadic gaze control**

Eye movements in natural behaviour are allocated in a way that allows us to achieve our current goals. In natural behaviour, eye-movements are tightly linked to subsequent actions and serve to accumulate sensory evidence relevant to those actions (Tatler et al., 2011). Task demands or plans are one of the most prominent top-down influences on information sampling (see (Hayhoe & Ballard, 2005; M. F. Land, 2006) for reviews). For example, if the target in a visual search task is known to be a certain colour, attention is guided preferentially towards that colour, resulting in a more efficient search (Wolfe, Cave, & Franzel, 1989). The role of task demands in saccadic control has been demonstrated in a variety of paradigms, ranging from laboratory tasks on the one hand to basic everyday tasks, reading and sports activities on the other hand (Ballard, Hayhoe, & Pelz, 1995; Hayhoe, Mennie, Sullivan, & Gorgos, 2005; M. Land, Mennie, & Rusted, 1999; Rayner, 1998).

Secondly, stimulus value plays an important role in saccadic control. Reward associations modify the ability of stimuli to automatically bias attention in monkeys (Peck, Jangraw, Suzuki, Efem, & Gottlieb, 2009) and humans (Anderson, Laurent, & Yantis, 2011; Hickey, Chelazzi, & Theeuwes, 2010). Furthermore, saccade speed and precision are influenced by reward (Manohar et al., 2015; Milstein & Dorris, 2007, 2011; Takikawa, Kawagoe, Itoh, Nakahara, & Hikosaka, 2002). Research utilising visual search paradigms has shown that the

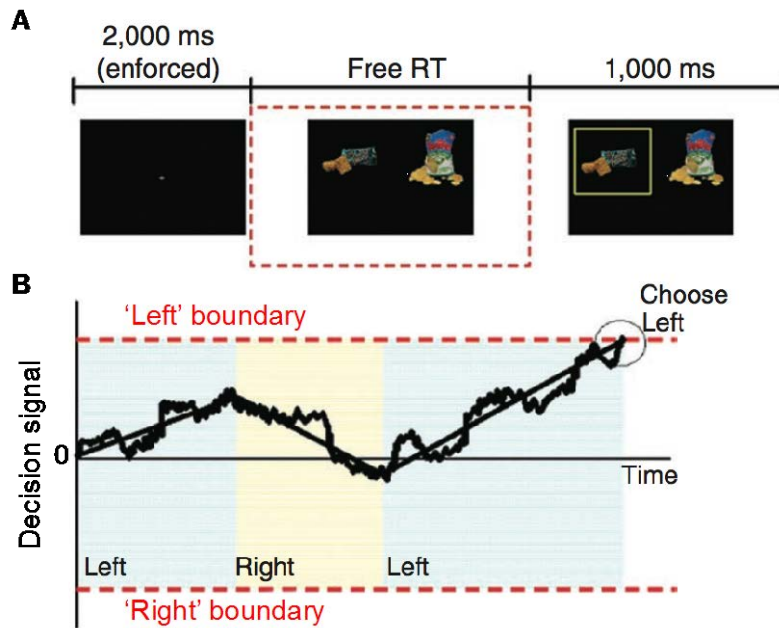
sampling policy is modified by reward associations and that this is not due to simply an increase in arousal (Chelazzi, Perlato, Santandrea, & Della Libera, 2013; Yantis, Anderson, Wampler, & Laurent, 2012). This line of research suggests that our brain keeps track of the value of potential saccade destinations and uses this representation to guide our gaze. However, there is an important caveat to this work. In laboratory experiments involving nonhuman primates, value is often operationalized as a liquid reward that ensues when a correct eye movement is made. However, in natural environments, eye movements alone do not automatically induce rewards. Rather, eye movements in the natural world serve to sample information from the environment, in the service of improving our knowledge of the world or furthering long-term goals.

#### **1.4. Eye movements and perceptual decision-making**

Recent work has begun to attempt to understand how eye movements might modulate decisions about visual stimuli, using both psychophysical tasks such as the RDM, and paradigms that involve choosing among stimuli defined by economic value, such as consumer products. In the next section, I will summarise recent developments in the field that have begun to help us understand how eye movements allow selective sampling of information during human decision-making.

### 1.4.1 Interactions between gaze direction and choice

One key development has been to attempt to understand the ubiquitous finding that humans tend to look preferentially at an item that they subsequently choose, or conversely, that humans may prefer to choose an item that they have looked at more often. In one important study, subjects were shown a display of two food items and were allowed to move their eyes freely among the two options the items for as long as they wanted before indicating their preferred choice (**Fig 1.6a**), rather than having to maintain central fixation like in the RDM task. In this study, items that were fixated more frequently were chosen more often. The authors extended the DDM framework by incorporating a process level explanation of the information stage, and in particular the selective weighting of information (Krajbich, Armel, & Rangel, 2010; Krajbich & Rangel, 2011). The results were accurately captured with a variant of the DDM that incorporated one crucial difference: the slope at which the decision signal evolves over time (i.e. the drift rate parameter) depends on the value of the currently fixated location (i.e. which item is sampled, **Fig 1.6b**). On a neural level, the authors argue that fixations affect this slope by amplifying the relative value signal for the fixated food item in the medial orbital frontal cortex (Lim, O'Doherty, & Rangel, 2011).



**Fig 1.6 Experimental design of a two alternative forced choice task.** Figure copied from Krajbich, Armel, & Rangel (2010) with permission through RightsLink. **(a)** On each trial, subjects had to fixate a central fixation point for 2 seconds. Subsequently, two food items were displayed on the screen and subjects could sample these items for an unlimited time. A yellow box highlighted the chosen option for 1 second after subjects made their choice. **(b)** The decision signal evolves over time until it reaches a boundary (left in this case). Crucially, the slope of the decision signal (up or down) is dictated by which stimulus (left or right) is sampled.

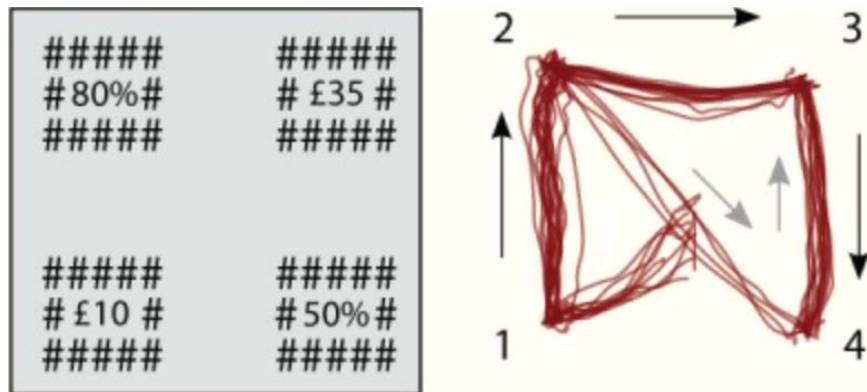
This study demonstrates that eye movements may be an important modulator of the “gain” of the decision process, i.e. the rate at which information is accumulated towards a particular choice. However, the authors argue that the direction of causality in this process is from eye movements to choices: in other

words, we “choose what we look at” rather than “looking at what we [will] choose”. The evidence they describe for this direction of causality is weak. Further, one fundamental assumption in their work is that saccade-based selective sampling is intrinsically stochastic, rather than guided by a fixed policy that privileges one sort of information over another. However, the information selection process does not have to be intrinsically stochastic – indeed, as noted above, the constraints of limited capacity suggest that the best policy will be one which selectively samples the most diagnostic aspects of the available information, in the service of enhancing the efficiency of decisions. For example, which food item is fixated when and for how long, can itself be influenced by its subjective value or current internal state of the agent. Similarly, when you are browsing through the available fruit items in the supermarket, you do not sample information randomly. You might sample information about various fruit locations in turn, spend more time on red fruit, or shy away from information about fruit in a particular location or from a particular colour. Indeed, in an earlier study involving choosing among faces on the basis of their perceived attractiveness, decisions and information gathering via saccades mutually drive one another in a reciprocal ‘gaze cascade’ (Shimojo, Simion, Shimojo, & Scheier, 2003). We will return to these issues in later chapters, and consider them in more detail.

Thus, based on our overview of the information sampling policy, one would expect bottom up but also top down-factors such as attention to actively guide our information sampling (i.e. a sampling policy) and thus influence the decision

making process. For example during medical diagnosis, clinicians choose which medical factors to attend to based on their hypotheses in order to make clinical decisions (Groopman & Prichard, 2007). In the machine learning literature, research on information sampling focuses on ‘active sensing’. Active sensing is defined as the process of efficiently choosing what views and samples to acquire additionally to improve the overall decision (Ahmad & Yu, 2013; Markant, Settles, & Gureckis, 2016). Similarly, a patient does not undergo all possible tests, but these tests are selected based on the evidence collected up to a particular point. Furthermore, research on hypothesis testing has shown that both stimulus and knowledge appear to play a role in determining how people collect information (Slowiaczek et al., 1992). In conclusion, a growing body of decision making research is starting to acknowledge information sampling as a critical, and active, stage in the decision making process. In what follows, we discuss the studies that have suggested that selective sampling with the eyes may be an “active” process.

For example, Manohar & Hussain (2013) sought to disentangle two classes of eye movement policy that have been confounded in the previous literature: sampling for information from sampling for value (Manohar & Husain, 2013). In this task subjects had to choose between two gambles presented on the left and right side of the screen (**Fig 1.7**). Each gamble is characterised by a probability of winning and a monetary value, and subjects could sample the available information without a time limit while their gaze was tracked.



**Fig 1.7 Example stimulus display and scan path** (figure copied from **Manohar & Hussain, 2013 under CC-BY**). Left panel: example trial display. Each gamble (left vs. right) is characterized by both a probability of winning and a monetary value. Right panel: example scan path of a participant on a particular trial. The numbers represent the sampling order (i.e. the stimulus on the bottom left was fixated first).

Manohar & Hussain showed that in this task, sampling was first guided towards information (i.e. to reduce uncertainty), and that this policy shifted to expected value (reward) of the stimuli later on. Furthermore, replicating previous studies, participants tended to fixate the stimulus they eventually chose near the end of the trial.

### **1.4.2 Eye movements may be allocated to reduce uncertainty**

The notion that the eyes may be moved to locations that specifically reduce uncertainty has been addressed in a number of previous studies. One classic report investigated the optimal allocation of fixations (i.e. optimal sampling) in a noisy perceptual display (Najemnik & Geisler, 2005; Yang, Lengyel, & Wolpert, 2016). Selecting a certain gaze position allows us to process the information available at that position, contributing to a belief update about the stimulus (in this case, the presence or absence of a target grating in the noise). Using a Bayesian ideal observer model, this information gain can be precisely quantified and compared to the information gained by an ideal target selector (in this case a Bayesian model). Optimal gaze allocation is defined by having the ideal target selector consider all of the possible fixation locations, given its current knowledge of the landscape (i.e. posterior probabilities, visibility maps, previous fixation locations...). The ideal searcher then computes the gaze trajectory based on a particular quantity, such as maximising the information gain of each fixation, or maximising the probability of correctly fixating a target stimulus. A related model was described by Itti and Baldi (Itti & Baldi, 2009). Here, the authors measured human scan trajectories in natural scenes. They compared the findings to those of a Bayesian model that allocated the eyes to locations that maximised the posterior update, i.e. the greatest divergence between prior and posterior belief about the presence or absence of a target item ("Bayesian surprise"). The authors found that humans matched the model well. This line of research has

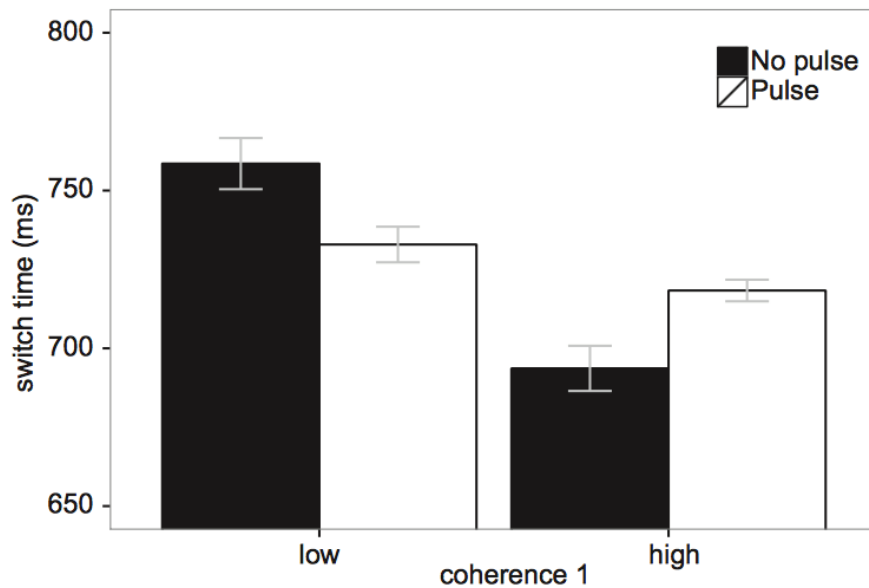
provided more evidence that scanning paths closely resemble active sensing strategies (i.e. acquired information influences future gaze position). Furthermore, these approaches allow researchers to compare allocation of fixations from subjects with the allocation of fixations from an ideal observer model that maximises the information gain at each fixation. Yang, Lengyel, & Wolpert (2016) have shown that perceptual noise, prior biases, and inaccurate eye movements are the main sources of sub-optimality in fixation allocation (Yang et al., 2016). This line of research opens up exciting research avenues on how we can improve fixation allocation (i.e. the sampling policy) in order to ultimately improve decision-making.

Another line of research has considered the sampling policy as an active process and investigated the influence of uncertainty on the sampling policy (Cassey, Evens, Bogacz, Marshall, & Ludwig, 2013; Ludwig & Evens, 2016). Ludwig and Evens utilised a paradigm in which information is available from two sources, each corresponding to a RDM stimulus, but that only became visible once the gaze was oriented towards it (i.e. was gaze-contingent) (**Fig. 1.8**). Participants were able to switch their eyes freely between the two sources in order to decide for which stimulus the motion energy was translated in a more clockwise direction.



**Fig 1.8 Example stimulus display (figure copied from Ludwig & Evens, 2017 under CC-AL permission).** Two sources of information were available on each trial (top and bottom). Crucially information was sampled in a gaze-contingent way, i.e. when a particular RDM-pattern was fixated (purple ellipse), the dots started moving and information could be accumulated.

Moreover, the authors embedded an early coherence pulse on half of the trials in the stimulus that was sampled first, in order to directly manipulate the uncertainty about that stimulus. More specifically, if the coherence of the first sampled RDM stimulus was low, the pulse would increase coherence briefly while if the coherence was high it would be decreased briefly. Importantly, when the coherence of the first sampled pattern was decreased the switch point was delayed, while if the coherence increased the switch point occurred earlier (**Fig. 1.9**).



**Fig 1.9** The influence of introduced (un-) certainty on switch times (figure copied from Ludwig & Evens, 2017 under CC-AL). When the coherence of the first sampled pattern was decreased due to the pulse, the switch point was delayed (right), while the switch point occurred earlier when the coherence increased (left).

Thus, these authors showed that the sampling strategy (e.g. the point at which observers switched between the two sources of information) in their task was under direct influence of on going cognitive processing, in the sense that subjects' sampling policy was directly guided by their uncertainty about a particular source of information.

## **Chapter 2: Measuring the information sampling policy.**

The previous section provides us with evidence from the past literature that information sampling is an active process. I thus argue that incorporating an information sampling policy stage in models of the decision making process is an important next step for the field.

However, how do we measure which information is sampled? When investigating the information sampling policy, one must always strike a balance between simplicity and complexity: ideally, a design would be both ecologically valid, but still allow the information sampling policy to be carefully measured. On the one hand you want to have a precise measure of what information is sampled at each point in time, but on the other hand you want the information sampling to be as similar as possible to that which occurs in natural environments.

### **2.1 Traditional approaches to measure the sampling policy**

Over the years, the information sampling policy has been measured in a variety of ways. This tradition started out with descriptive methods in which subjects were required to 'think aloud' or to physically retrieve information in a laboratory environment (Ericsson & Simon, 1980; Payne, 1976). Subsequent approaches constrained the sampling of items via mouse-clicks or touch screens, e.g. instructing participants to 'click on the stimulus to reveal information' (Payne & Braunstein, 1978; Rieskamp & Hoffrage, 2008; Schkade & Johnson, 1989;

Zellman, Kaye-Blake, & Abell, 2010). Therefore, any information that is used by a decision maker must have been revealed beforehand, in a response-contingent fashion (**Fig. 2.1**). However, this approach is rather cumbersome, may be subject to various demand characteristics, and is not suited to investigating the sampling mechanisms that occur with the eyes during natural scene perception.

|                  | Company 1                                | Company 2                                                                           | Company 3                                | Company 4                                |
|------------------|------------------------------------------|-------------------------------------------------------------------------------------|------------------------------------------|------------------------------------------|
| Share Price      | <input type="text"/>                     | <input type="text"/>                                                                | <input type="text"/>                     | <input type="text"/>                     |
| Recognition Rate | <input type="text"/>                     | <input type="text"/>                                                                | <input type="text"/>                     | <input type="text"/>                     |
| Dividend         | <input type="text"/>                     | 4  | <input type="text"/>                     | <input type="text"/>                     |
| Share Capital    | <input type="text"/>                     | <input type="text"/>                                                                | <input type="text"/>                     | <input type="text"/>                     |
| Investments      | <input type="text"/>                     | <input type="text"/>                                                                | <input type="text"/>                     | <input type="text"/>                     |
| Employees        | <input type="text"/>                     | <input type="text"/>                                                                | <input type="text"/>                     | <input type="text"/>                     |
| Choose One       | <input type="button" value="Company 1"/> | <input type="button" value="Company 2"/>                                            | <input type="button" value="Company 3"/> | <input type="button" value="Company 4"/> |

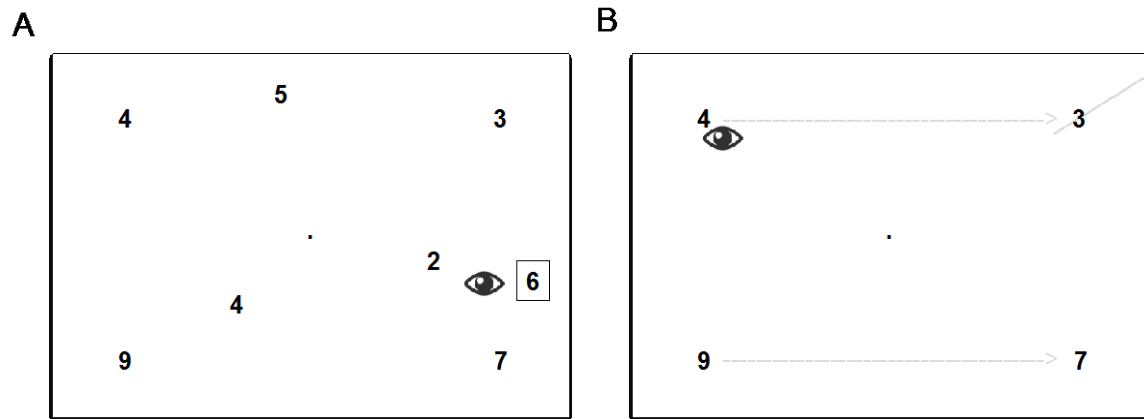
**Fig 2.1** An example stimulus display of a ‘mouse-click’ study (adapted from Rieskamp & Hoffrage, 1999). Subjects were shown an information board including 4 companies and their six (hidden) pieces of information about a particular attribute (left column). Subjects could sample the information about the different companies and had to indicate their choice (by clicking on one of the companies on the bottom row) when they were ready. In this example, information about Company 2’s dividends is revealed via a mouse-click.

## 2.2 Eye-tracking as a tool to measure the sampling policy

Technological advancements in the development of eye-tracker systems subsequently allowed researchers to accurately track the placement of participants' gaze, while making eye tracking less intrusive. This has the potential to allow measurement of real world information sampling with eye movements.

### 2.2.1 Width of the information processing spotlight around fixation

However, even if we know the exact location of where subjects in our task are fixating, how do we know which information they actually have access to? For example, imagine that a subject can view a computer display filled with numbers, and he or she has to decide whether the average of these numbers is higher or lower than a reference value, e.g. 5 (**Fig. 2.2.A**). Even if an accurate measure of gaze position is available (e.g. the eye on **Fig. 2.2**), the researcher still has to decide which information the subject has access to: can the subject process both the number 2 and 6? What about number 7? How accurate is the information that is extracted from each number and does this depend on its distance to fixation? Does the type of stimulus affect this relationship of information quality and distance to fixation (e.g. identity versus colour of the number)? These are among the challenges that will be considered in the current chapter.



**Fig. 2.2 Example stimulus displays and estimating the width of information processing.** *Panel A.* Subjects can scan ('sample') the computer display that is filled with numbers and have to decide whether the average of the numbers is higher or lower than 5. The eye represents gaze fixation at a particular point in time. Using the nearest neighbour approach, number 6 is the number closest to the fixation location and is thus classified as the processed information for this particular fixation *Panel B.* In this design the stimuli are sufficiently spread out to improve estimation accuracy of the processable information for a particular fixation location. Spatial gaze biases such as sampling from left to right and top to bottom are depicted by the grey arrows.

In some cases, this problem has been addressed by developing paradigms in which stimuli are sufficiently spread out in space in order to make it impossible to view multiple items at once; for example a stimulus in each corner of the computer screen (Manohar & Husain, 2013). One potential problem with this approach is that it can potentially exacerbate the influence of arbitrary spatial

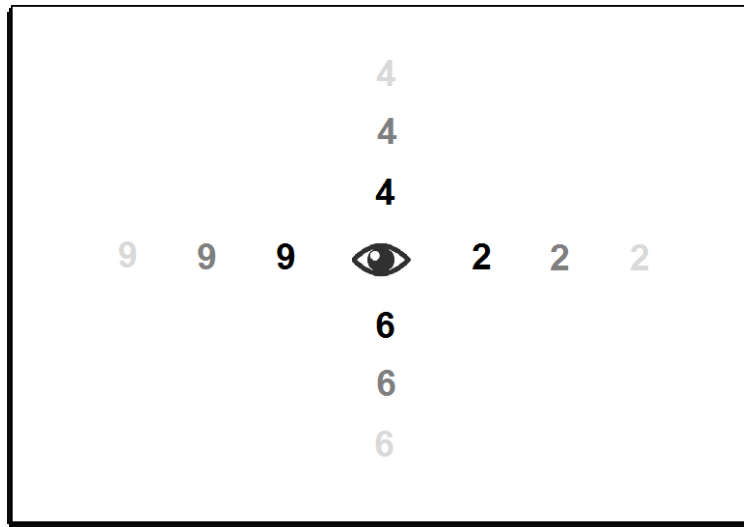
biases (**Fig. 2.2.B**; (Durgin, Doyle, & Egan, 2008; Glaholt, Wu, & Reingold, 2010).

Other approaches have tried to use more ecologically valid paradigms in which subjects can use peripheral vision to guide their information sampling. The difficulty then becomes estimating the amount of information available to subjects at each point in time as they scan through the environment. One approach is to assume that only one stimulus is processed at a time, and that only the stimulus closest to fixation is processed ('Nearest Neighbour', **Fig. 2.2.A**). However, this approach may fail to provide an accurate measurement of the processed information when multiple stimuli are close to the fixation location (e.g. number 2 is almost as close to the fixation as number 6, but in this approach no information about number 2 is processed). Moreover, this assumption seems rather ad-hoc, and recent methodological advances have allowed researchers to more accurately estimate the available information at each time point.

### **2.2.2 Retinal sensitivity mapping**

One such methodological advancement has made it possible for researchers to map the retinal sensitivity. Retinal sensitivity describes how sensitive the retina is to information in the periphery, and can thus be used as a measure of peripheral information processing. In this approach, the experimental task of interest is

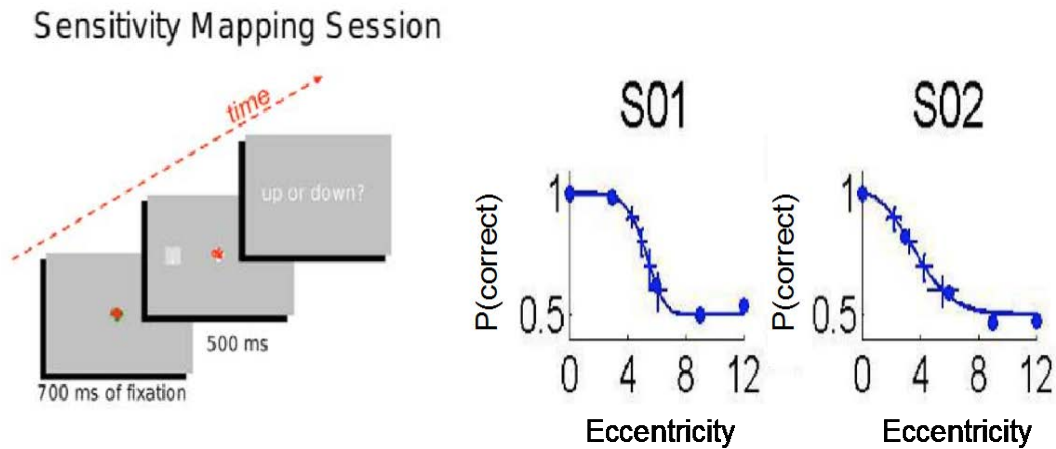
preceded by a retinal sensitivity mapping task. In these tasks, participants are instructed to maintain central fixation, while a target is presented in the periphery. The retinal eccentricity of that target is varied, and the observer's ability to correctly judge the target is recorded (**Fig. 2.3**)



**Fig 2.3 An example stimulus display showcasing the ‘Retinal Sensitivity Mapping’ approach.** Stimuli are presented at varying retinal eccentricities (e.g. lighter shades of grey indicate a higher eccentricity) and subjects need to identify the stimulus while maintaining central fixation. Multiple numbers are presented in different shades of grey for illustration purposes only. A real retinal sensitivity mapping session would only show one number at a time, in the same colour, and at differing eccentricities.

For example, Morvan & Maloney used this approach to quantify the information subjects could take into account around their fixation (Morvan & Maloney, 2012).

Retinal sensitivity functions were estimated for each subject (**Fig. 2.4**; right panel). The sensitivity can be low (i.e. only foveated information is processed) or high (i.e. information in the far periphery is additionally processed).

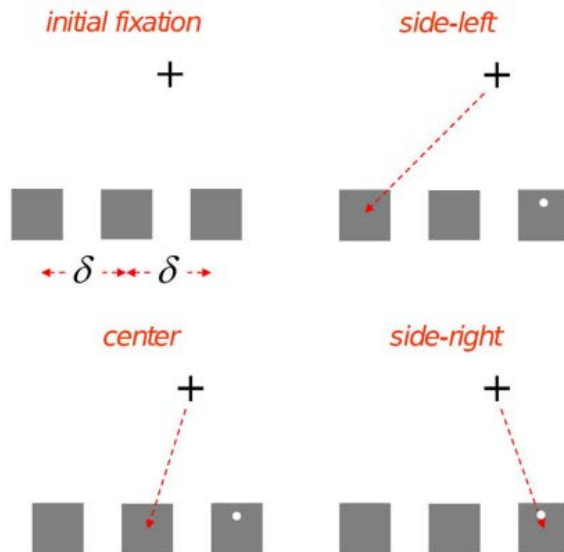


**Fig 2.4 Retinal sensitivity mapping session, figure adapted from Morvan & Maloney (2012) under CC-AL license.** Fixation is maintained at a central fixation cross for 700ms (red dot), after which a target is presented in the periphery for 500ms. Subjects still need to maintain central fixation, while they judge this target (in this case whether the white target dot is up or down). The target stimulus is presented at different levels of eccentricity (i.e. distance from the central fixation cross). The ability of an observer to correctly evaluate the target stimulus for different levels of eccentricity allows the estimation of retinal sensitivity functions specific for each subject (right panel, retinal sensitivity functions for two subjects). When retinal sensitivity is high, information from the periphery can be processed, and as sensitivity decreases, the amount of processable information from the periphery decreases.

It is important to note here that it may be important to use a retinal sensitivity mapping task that closely resembles the subsequent target experimental task, in order to obtain generalisable retinal sensitivity estimates. For example, the retinal sensitivity for making a particular classification about a stimulus (e.g. arrow up or down) might be different when that same stimulus needs to be classified on another dimension (e.g. colour or size of the arrow). Similarly, the retinal sensitivity for processing the number 8 when the task is to decide that a number is an 8 or not, might be different when the task is to average that number with previously sampled numbers.

Interestingly, retinal sensitivity mapping does not always lead to an accurate measure of how much information a subject will process around their fixation in the experimental task, even when authors take measures to use an appropriate retinal sensitivity mapping session (Morvan & Maloney, 2012). Morvan & Maloney showed that observers do not always employ their extrafoveal retinal sensitivity to optimize their visual search. To do this, the authors first estimated the retinal sensitivity of each subject (see above; **Fig. 2.4**). Next, observers had to perform a decision task (**Fig. 2.5**). In this task, observers are shown 3 grey squares that are presented at a trial-by-trial varying eccentricity ( $\delta$ , top left panel). Observers start the trial by saccading to the left square ('side-left'), the 'center', or to the right square ('side-right'). Following this initial saccade, the stimulus (e.g. a dot pointing up or down) briefly appears in one of the side locations. Importantly,

the probability of being correct thus depends on the location observers saccade to and the location of the stimulus.



**Fig 2.5 Decision task and possible saccadic strategies (figure copied from Morvan & Maloney, 2012 under CC-AL license).** In this task, observers are shown 3 grey squares that are presented at a trial-by-trial varying eccentricity ( $\delta$ , top left panel). Observers start the trial by saccading to the left square ('side-left'), the 'center', or to the right square ('side-right'). Following this initial saccade, the stimulus (e.g. a dot pointing up or down) briefly appears in one of the side locations.

Crucially, since the authors estimated the retinal sensitivity of each subject and varied the eccentricity of the 3 grey boxes, it is possible to calculate the 'optimal' saccadic strategy for a given subject and a given eccentricity of the boxes. Thus, when both the 'side-left' and 'side-right' boxes fall within an observer's retinal

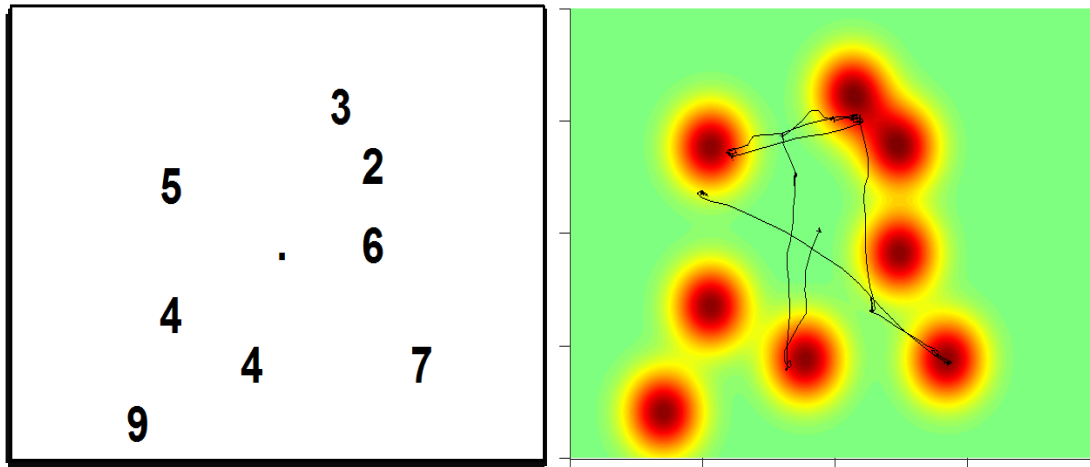
sensitivity, the optimal strategy would be to saccade towards the middle box. However, as eccentricity of the boxes increases, the optimal strategy will shift towards side selecting saccades. The degree to which this shift should occur depends on a subjects' retinal sensitivity. Strikingly, the authors showed that observers did not adapt their strategy as a function of eccentricity and thus did not respond in a way that was consistent with their optimal retinal sensitivity function. Importantly, if observers were forced to make the 'optimal' saccade, performance markedly increased to the optimal level.

Given that observers may fail to adjust their fixation locations to optimally utilise their retinal sensitivity, inferring the processed information from retinal sensitivity will not always be accurate. For this reason, we developed an alternative approach to independently estimate the information that is sampled by a given fixation. The spatial processing window is estimated from an independent aspect of the data collected during the main task. This 'landscape analysis' allows the researcher to quantify how much information a subject is able to process around their fixation in real time.

To illustrate, imagine again a task in which multiple numbers are spatially distributed on a computer screen and observers are tasked to indicate whether the average of these numbers are above or below 5 (**Fig. 2.6**, left panel).

### 2.2.3 The landscape approach

Instead of estimating the retinal sensitivity, one can treat the stimulus display as a 'landscape'. This landscape analysis converts the stimulus array on each trial to a topography of to-be-sampled values spanning the pixel space constituting the screen. The landscapes are created by convolving each stimulus location (i.e., the x and y coordinates of each number in pixel space) with a bivariate Gaussian basis function. Each Gaussian has a standard deviation  $s$  and height  $g$ , with the latter reflecting some information conveyed by the stimulus. In the simplest case,  $g = 1$  and  $s$  is a constant (e.g., 50 pixels). An example of such a landscape is shown in the right panel of **Fig. 2.6** (for the corresponding stimulus array shown in the left panel). The estimated quality trajectory at time  $t$ , or  $QT_t$ , represents the total information available to the participant at any time point  $t$ , can thus be quantified as the landscape height at the current eye position  $(x, y)$  at that time point (black line in **Fig. 2.6** right panel). Across each trial, this provides a vector of values that indicate whether the participant is looking directly at a stimulus ( $QT_t, \sim 1$ ) or away from a stimulus ( $QT_t, \sim 0$ ) at each point in time following array onset.



**Fig. 2.6 Landscape analysis illustration.** The landscape analysis converts the stimulus array (left panel) to a topography of to-be-sampled values spanning the pixel space constituting the screen (see text). The estimated quality trajectory, or  $QT_t$ , which represents the total information available to the participant at any time point  $t$ , can thus be quantified as the landscape value at the current eye position  $(x, y)$  at that time point (black line). This thus provides a vector of values that indicate whether the participant is looking directly at a stimulus ( $QT_t \sim 1$ , red/yellow area) or away from a stimulus ( $QT_t \sim 0$ , green area) at each point in time following array onset.

The parameter  $s$  of these Gaussians can be seen as a proxy of how much information participants can harvest around their fixation. For convenience of display, this analysis is conducted by reading out a topography value in the stimulus space (**Fig. 2.6**). However, this is equivalent to a Gaussian filter spotlight of width  $s$ , centered at the  $x, y$  position of gaze, moving across the landscape and integrating the sum of all values under the spotlight.

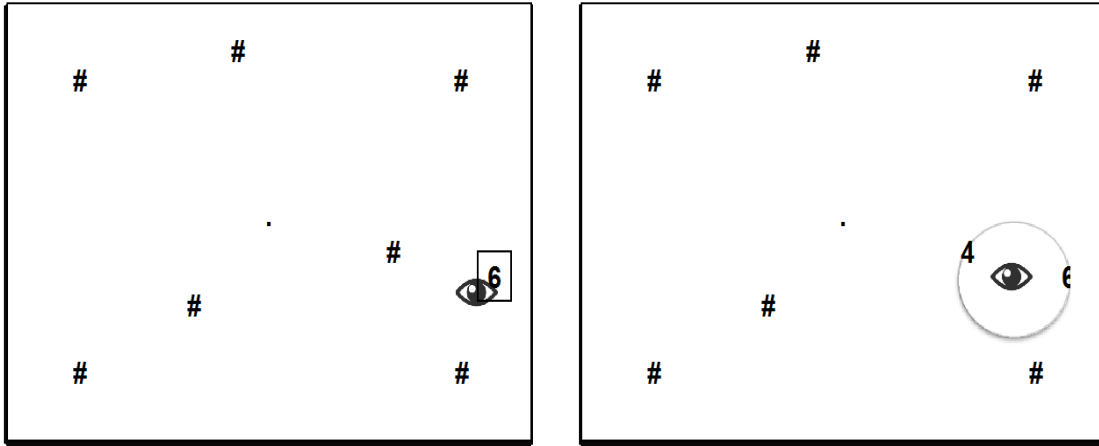
In this approach, the values under the spotlight are scaled by their distance to the center of the spotlight. Critically, the width of the spotlight  $s$  need not be assumed but can be estimated by model fitting. The trajectory quality  $QT$  is exhaustively computed across a space of possible parameter values of  $s$  (range 0–100) to identify the value of  $s$  that maximizes  $QT$ . In other words, this approach assumes that participants move their eyes to locations that maximize available information ( $QT$ ), instead of empty locations on the screen. Note that an excessively narrow spotlight would lead to items even close to the point of gaze being neglected, whereas an excessively broad spotlight would lead to large portions of uninformative space being sampled, reducing the overall level of information ( $QT$ ) obtained. The aperture half-width  $s$  that maximizes  $QT$  is the one that maximizes this trade-off. This approach thus defines the width of the information-processing spotlight within the same task but in a way that is independent of the decision values (numbers relative to the reference in this case).

This computation of the width of the information-harvesting spotlight was developed for four reasons. First, the work by Morvan and Maloney (2012) suggests that the use of retinal sensitivity mapping for a design like **Fig 2.5** might lead to an overestimation of the degree to which information is processable around each fixation (i.e. the width of the processing spotlight (Morvan & Maloney, 2012)). Second, this landscape analysis does not require the experimenter to run an additional retinal sensitivity mapping session for every

experiment. The landscape approach, analysing the Quality trajectory, shows that researchers do not necessarily need to estimate this retinal sensitivity for each experiment but could instead use a proxy of 'utilised' retinal sensitivity. Third, this approach assumes that information falling further from the point of gaze has a weaker influence on choices than that falling directly under a fixation point rather than assuming that all information within a fixed window contributes equally to choices. Fourth, our approach has the major benefit that it allows us to characterize, in real time on each trial, information that is available to the participant and thus to explore how stimulation and gaze interact in determining choices. It is also immune to otherwise arbitrary choices that might have to be made during a more standard fixation analysis, such as the size and shape of the window that surrounds each stimulus or the criteria for determining dwell times. Finally, the main benefit of the landscape methodology is that it allows the researcher to not only quantify the degree to which an element is sampled (QT) but also to construct trajectories that indicate the decision value (or any other construct the researcher might be interested in) in real time. Thus, after first estimating the width of the information processing spotlight (i.e. the parameter  $s$  of the Gaussian), one can set the height of the Gaussian  $g$  to any construct the researcher would like to associate with each stimulus (e.g. the decision value, contrast, luminance etc.). Chapter 3 includes the first application of this landscape methodology in a pilot study. Chapter 4 includes 4 eye-tracking studies that used this landscape analysis approach in more detail, and highlights additional benefits of the landscape approach.

### 2.2.4 The gaze-contingent approach

Finally, in the gaze-contingent approach the available information is directly tied (i.e. contingent) to gaze location on the screen. Thus, the display of the computer screen changes as a function of where the subject is looking (**Fig. 2.11**). This approach is very adaptable as it allows researchers to manipulate the amount of information available around fixation. On the one hand this approach has been used as a more ecologically valid variant of the mouse-clicking studies, i.e. only stimuli that are fixated are revealed on the screen (**Fig. 2.11**, left panel). On the other hand, it is possible to combine gaze-contingent approaches with retinal sensitivity mapping, in order to create a sliding 'visibility' window with a subject-specific size (**Fig. 2.11**, right panel). This technique thus allows researchers to directly manipulate the information that is available from the periphery in a parametric way. Chapter 5 includes two eye-tracking studies that utilised the gaze-contingent window approach.



**Fig 2.11 An example stimulus display showcasing the ‘Gaze-Contingent’ approach.** Stimuli are masked (#) when not fixated. The left panel shows the ‘window’ approach: a stimulus is revealed as soon as eye-gaze is fixated in a window around the stimulus. The right panel shows the ‘sliding visibility window’ approach: a sliding window follows eye-gaze around and information that falls under this window is revealed. This sliding window can be determined by the researcher beforehand or estimated via for example retinal sensitivity mapping or the landscape approach.

## **Chapter 3: Benefits of voluntary information sampling and dominance of perceptual factors as drivers of the sampling policy**

### **3.1. Introduction**

When making decisions, humans and other primates move their eyes freely to gather information about their environment. In chapter 1, we described a large literature that explored factors that determine which information is sampled. In particular, research that focuses on determining where our eyes fall during natural scene perception and visual search has concluded that deviant or surprising perceptual information attracts attention and gaze. Various studies have investigated the sampling policy in differing domains, such as category learning (Rehder & Hoffman, 2005) and classification tasks (Meier & Blair, 2013).

Recently, work in the perceptual decision making domain has started to focus on understanding how the sampling policy might modulate decisions about visual stimuli and which factors drive this sampling policy. However, little is known about the policy by which participants sample information during choice tasks that involve integrating decision information from multiple locations across the visual field. Moreover, the importance of investigating evidence accumulation as an active sampling process has been gaining traction in recent years (Gottlieb et al., 2014).

In this first experimental chapter we will discuss three different pilot experiments. In these pilot studies, we adapted a formerly developed paradigm that investigated the influence of multiple sources of perceptual uncertainty on perceptual summary judgements (Michael, de Gardelle, Nevado-Holgado, & Summerfield, 2013) into a controlled paradigm that allows sampling to be measured, including with the landscape analysis (see Pilot 3 below). These experiments, although preliminary, allowed us to investigate the influence of the sampling policy on decision making with reference to the results from previous established work. The main goal of these pilot experiments was to investigate how sampling policies in terms of i) duration, and ii) order of sampling, influence evidence integration from multiple locations across visual field. The other goal of these pilot experiments is to allow us to develop a suitable, well-controlled perceptual decision making paradigm that will allow us to investigate the sampling policy in the subsequent studies that form the bulk of the work in this thesis. This work is rather preliminary, but we present it here to set the context for the more elaborate studies that make up the subsequent chapters.

More specifically, other studies that have shown detrimental or no beneficial effects of prolonged sampling, have little to no control over the sampling duration, which information is sampled when, and the objective decision value of each piece of information (Bastardi & Shafir, 1998; Dijksterhuis, 2004; Kahneman,

2011). We aimed to develop a perceptual decision making design in which each piece of evidence is quantified by an objective decision quality (e.g. each piece of evidence is associated with a distinct decision value). Having such a design allowed us to more accurately measure the sampling policy *per se* (which information is processed at which time during the trial), and also allows us to infer how sampling policies influence decision process (e.g. how performance varies as a function of which information is sampled and how long the information is sampled for).

The stimuli we adopted in this task, ‘squircles’, were based on the work from de Gardelle and colleagues (De Gardelle & Summerfield, 2011). These stimuli contained two important feature dimensions (e.g. colour and shape), which can be simultaneously and orthogonally manipulated. This allowed us to develop a task in which the decision value of a stimulus (i.e. combined feature values from colour and shape) can be independent of its individual feature values (see methods for details); this allowed us to separate the influence of perceptual features and decision values on sampling policies. We were interested in exploring whether we could find evidence for influences on the sampling policy from both the perceptual features of the stimuli as well as from its decision information.

Furthermore, we knew that there are two sources of perceptual uncertainty that

contribute to the difficulty of a perceptual decision (Michael et al., 2013). These two sources of uncertainty are i) the strength of the evidence (i.e. the proximity of decision information to the category boundary) ii) the reliability of evidence (i.e. the variability of the feature values present on a particular trial). Here, we are interested in knowing how the decision environment, quantified by these two sources of uncertainty, influences the sampling policy and performance.

## **3.2. Study: Pilot 1**

### **3.2.1 Methods**

*Sample.* The experiment took place in the Department of Experimental Psychology at the University of Oxford. Twenty-six subjects volunteered to participate in experiment 1. Subjects reported normal or corrected vision. The experiment took approximately 50 minutes. Participants provided written consent before the experiment and received 10£ for their participation. The experiment was approved by the local ethics committee (approval no. MSD-IDREC-C1-2009-1).

*Stimuli.* The stimuli used in the experiment (squircles) were generated using the PsychToolBox for MATLAB and presented on a 17" CRT screen (resolution: 1024x768). The stimuli varied across 2 feature dimensions, each with their own continuum: colour (red to blue) and shape (circle to square). Thus, each squiracle

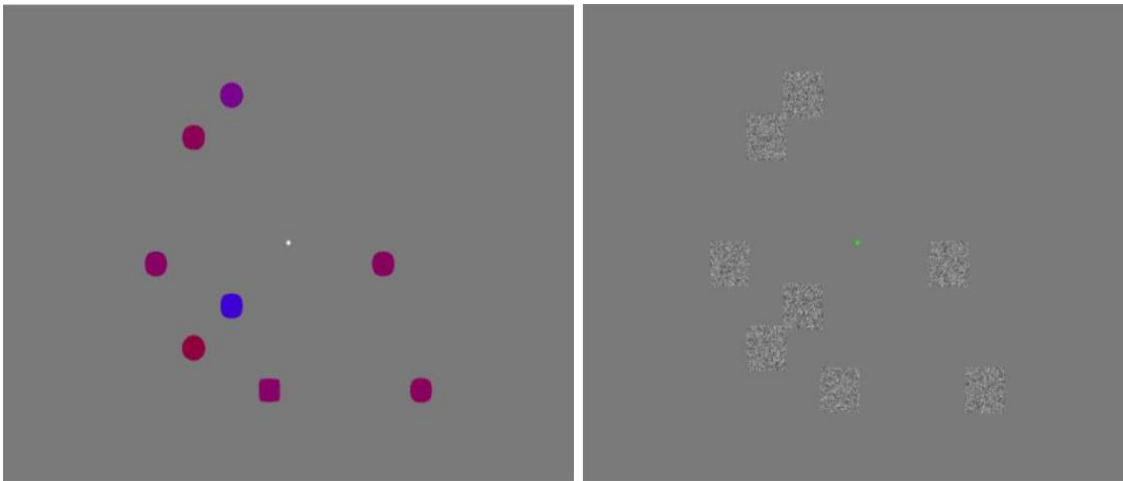
was defined by two independent parameters that controls the feature of two feature dimensions: a (momentary) colour value and a shape value. The colour value defined the colour of the squircle in Red-Green-Blue (RGB) values between red ([1,0,0]) and blue ([0,0,1]) by following a linear transition in RGB space ([C, 0, 1-C]). Gamma correction was applied to these values. Furthermore, the squircles were hyper-elliptic shapes that varied parametrically between the shape of a square and a circle. These varying shapes were constructed while ensuring that each shape occupied the same area on the screen (i.e. the size of the stimuli was the same).

In this study, unlike in previous work (de Gardelle and Summerfield, 2011), momentary decision values were defined by a combination of both colour and shape values. Colour and shape values were both parameterised in the theoretical range [-1,1] with a mean of zero. We used an XOR rule to define the momentary decision value, so that decision values were computed as the product of the (zero-centred) feature values. Thus, for a given participant, positive decision values might correspond to there being “more red squares/blue circles” whereas negative values might correspond “more blue squares/red circles”. This manipulation allowed us to orthogonalise the influence of individual feature values (colour and shape) and momentary decision values (shape/colour combinations) on choices (see below).

Colour and shape values were generated by generating random samples from

normal distributions with pre-specified mean ( $\mu$ ) feature values ([-0.1, -0.05, 0.05, 1]) and variance ( $\sigma$ ) values ([0.15, 0.30]). The DV corresponds to the evidence strength (e.g. further away from the category boundary at 0 corresponds to higher evidence strength), and the variance of the generated feature values on a particular trial corresponds to the evidence reliability. These values for evidence strength and reliability were determined in a brief pilot where we used a staircase procedure to adjust the difficulty (via increasing / decreasing the evidence strength) for each individual participant until an accuracy of 70% was achieved.

Experiment one consisted of 10 blocks each containing 48 trials. On each trial, eight stimuli (e.g. squircles) were presented at semi-randomised locations within an invisible 8x8 grid (**Fig. 3.1**). The squircles were masked after a variable presentation duration (see below) and the fixation point in the middle turned green, indicating to participants that they had to make a response. Participants then had to press the left or right mouse-key to make a perceptual decision based on the combined shape and colour features of each stimulus (i.e. the decision rule was based on an XOR problem; see above). Participants received feedback after each response: a low tone for incorrect, and a high tone for correct responses. Additionally, after each block, participants were reminded about the correct response mapping and saw their achieved accuracy rate from the previous block.



**Fig. 3.1 Example trial of pilot 1.** On each trial, eight 'squircles' were presented for a variable duration (1000, 2000, or 3000ms) in random (semi-restricted) locations. After this presentation, the fixation point in the middle turns green and the items are masked, indicating to participants that they have to make a response. Participants then have to press the left or right mouse-key to make a perceptual decision based on the combined shape and colour features of the stimulus array. For example, say the response-mapping for this participant is pressing left if on average the display represented more red circles and blue squares, or right for red squares and blue circles. Thus, the correct answer on this trial is to press left, which is determined by the sign of the average decision value (red circle/blue square). Participants received feedback after each response: a low tone for incorrect-, and a high tone for correct- responses.

In addition to determining suitable parameter values for our planned follow-up experiments, we wanted to investigate the effect of presentation duration on performance. As one would expect, longer presentation duration allows participants to sample more information and thus integrate more evidence, which

means performance should rise with longer durations. However, we originally conjectured that performance might plateau once the duration is long enough for participants to sample information to their capacities, or even decrease due to detrimental effects of oversampling the available information. We also manipulated the spatial location of the stimuli (i.e. the maximum distance from the middle of the screen). We expected that a larger spread between the stimuli would decrease performance because more information would fall out of the attentional span due to the limited presentation duration. Furthermore, by having the stimuli that are more spread out, we force people to apply more sequential information sampling strategies, unlike in the original studies where the squircle stimuli were simultaneously presented in a circular array.

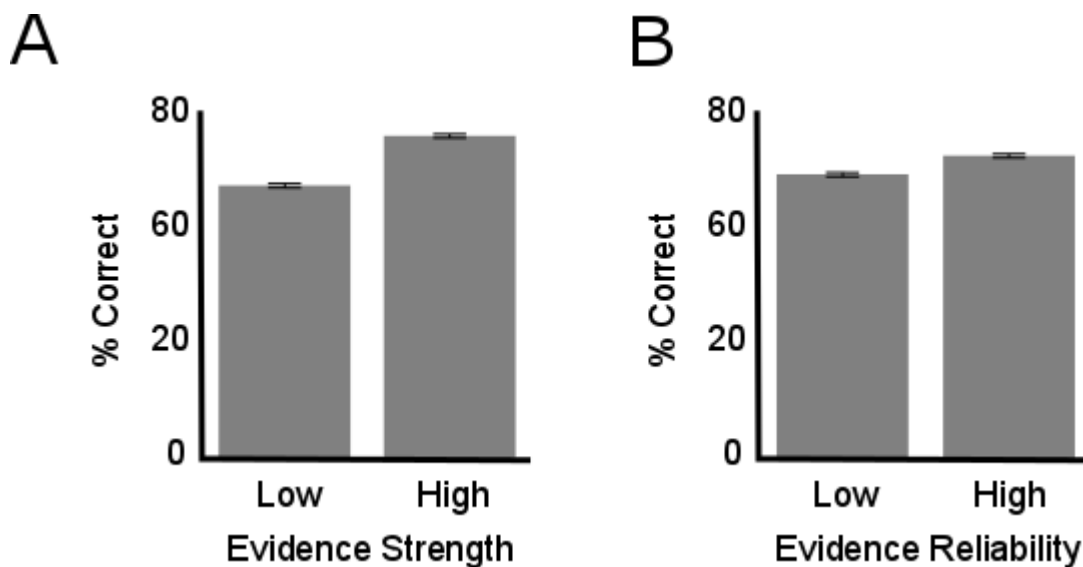
### **3.2.2. Results**

Since we employed a fixed response interval we did not analyse the reaction time (RT) data. For the present analysis we excluded 6 out of 26 participants with mean accuracies lower than 60%, resulting in an average accuracy of 69.4% (n=20).

#### *3.2.2.1 Independent effects of evidence strength and reliability.*

**Fig. 3.2** shows the effects of evidence strength (absolute decision value) and evidence reliability (variance) on performance. As expected, performance

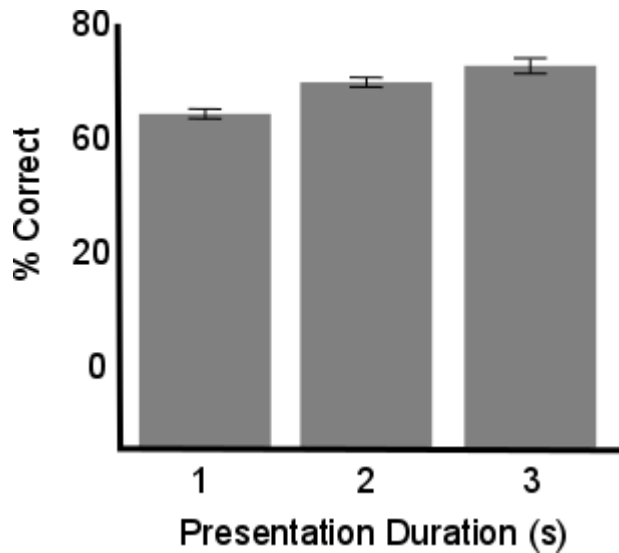
increased on high mean trials compared to low mean trials [ $F_{(1,19)}=22$ ;  $P<0.001$ ]. Performance decreased on high variance (low reliability) trials compared to low variance trials [ $F_{(1,19)}=20$ ;  $p<0.001$ ]. The interaction between variance and mean was not significant [ $F_{(1,19)}=0.37$ ;  $p>0.05$ ]. These results replicate the original work in a different setting, highlighting these two distinct sources of uncertainty (Michael et al., 2013).



**Fig. 3.2** Left: Performance increases when evidence strength is higher (high mean) compared to when evidence strength is lower (low mean). Right: Performance increases when the decision relevant evidence (the squircles) is more reliable (low variance) compared to when evidence is less reliable (higher variance). Error bars represent the standard error of the mean (SEM).

### 3.2.2.2. Presentation duration and eccentricity of the display

Performance increased with longer durations (**Fig.3.3**,  $F_{(2,50)} = 55$ ;  $P < 0.001$ ).

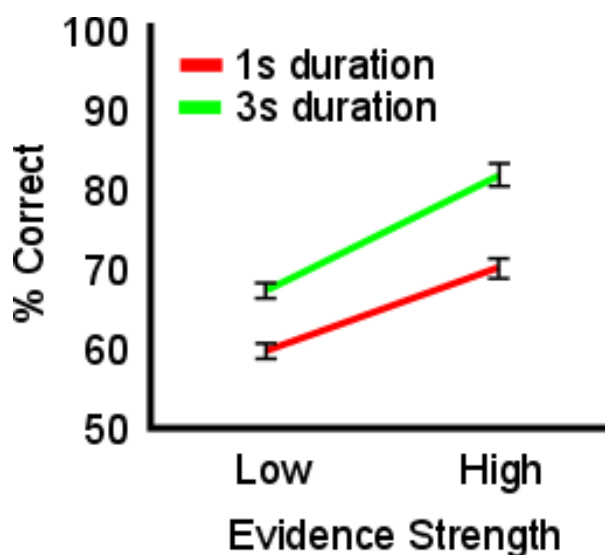


**Fig. 3.3 Main effect of presentation duration.** Performance increased for longer presentation durations. Error bars represent the standard error of the mean (SEM).

The effect of the eccentricity of the display (i.e. how spread out the stimuli were from the central fixation point) was not significant [ $F_{(1,25)}=3$ ;  $P > 0.05$ ]. The interaction between presentation duration and eccentricity was not significant ( $P > 0.05$ ). Thus, participants sampled the stimuli in an even fashion across space.

The interaction between presentation duration and mean decision value (evidence strength) was not significant ( $P > 0.05$ ). However, when comparing

only short (1s) with long (3s) duration trials, the interaction between duration and evidence strength was significant ( $P < 0.01$ ). This could either signify a benefit of a 3 second presentation for high mean trials or a cost for low mean trials, relative to a 1 second presentation (**Fig. 3.4**). While it is not standard to report subsidiary contrasts when the overall interaction is not significant, we chose to report it here because a similar pattern of results was present in subsequent pilot studies.



**Fig. 3.4** Interaction between evidence strength and presentation duration. This significant interaction effect can either signify a benefit of a 3 second presentation for high mean trials or a cost for low mean trials, relative to a 1 second presentation. Error bars represent the standard error of the mean (SEM).

No interaction was observed between presentation duration and evidence reliability ( $P > 0.05$ ). Thus, there was no influence of evidence reliability under different presentation durations, unlike for evidence strength.

### **3.2.3. Conclusion for Study pilot 1**

The goal of the first experiment was to test the general paradigm: making complex two-dimensional multi-attribute judgments in a perceptual decision task. Furthermore we determined suitable parameter values (eccentricity, evidence strength, and reliability) for future experiments.

Previous work emphasized the importance of taking both evidence strength and evidence reliability into account when designing decision making studies (De Gardelle & Summerfield, 2011). Like optimal (Bayesian) agents base their judgments on the strength and reliability of decision-relevant evidence, we find independent effects of the evidence strength and reliability in this first pilot study, replicating previous work that performance increased when evidence strength and reliability increased. This proof-of-concept study gave us confidence that we could use this paradigm for the more elaborate studies described below.

### 3.3. Study: Pilot 2

Even though the first pilot provides us with a controlled experimental paradigm with short duration trials and an objective decision value and choice criterion, we still do not have a measure of which information is processed at which time (i.e. which items participants are sampling). Thus, we developed a second paradigm that incorporates voluntary resampling (e.g. the ability for observers to resample the information) in order to provide an accurate measure of the sampling policy.

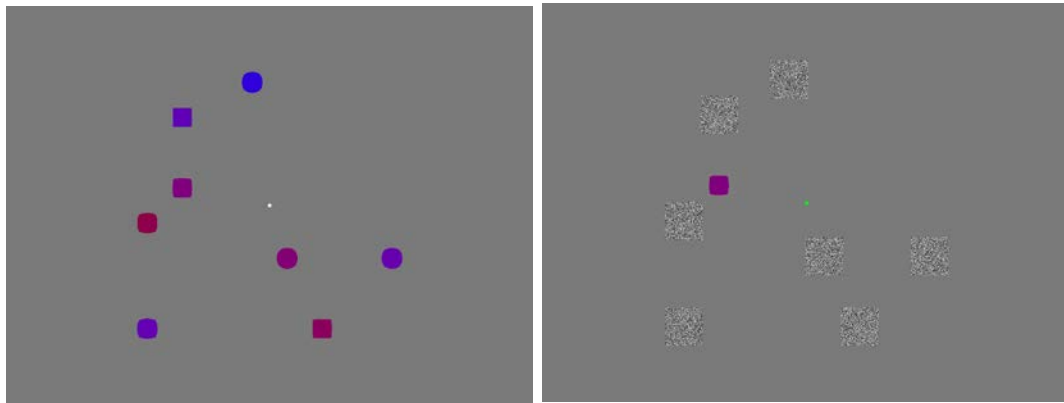
#### 3.3.1. Methods

*Sample.* The experiment took place in the Department of Experimental Psychology at the University of Oxford. Twenty-three subjects took part in experiment 2. Participants could only participate if they participated in experiment 1, as we wanted to use observers who are familiar with the stimulus distributions. Subjects reported normal or corrected vision. The experiment took approximately 50 minutes. Participants provided written consent before the experiment and received 10£ for their participation. The local ethics committee approved the experiment (approval no. MSD-IDREC-C1-2009-1).

*Stimuli.* We used the same stimuli and the same restricted random presentation as in the first experiment. Experiment two consists of 8 blocks each containing 36 trials (288 trials per subject). The experimental paradigm for the second pilot is an adapted version of the first paradigm (**Fig. 3.5**). In

order to investigate the sampling policy during the experiment, we presented the stimuli for a shorter duration (600 ms) and then masked each stimulus with a random noise mask. Next, a fixation dot in the middle of the screen turned green indicating to observers that they have entered a re-sampling stage. There are three conditions at this stage, masked squircles reappeared in two out of those three conditions. During a fixed resampling interval (2000, 4000 or 6000ms) participants could either i) voluntarily resample stimuli themselves (i.e. free resampling), ii) could not resample any stimuli (no resampling), or iii) stimuli were resampled randomly for them by the computer (i.e. fixed resampling). When participants engage in the free resampling stage, they could click with the mouse on a specific mask in order to reveal the underlying stimulus. On other trials, the computer randomly revealed stimuli for the observers. To avoid any bias due to the unequal number of resampled stimulus, the number of stimuli and duration that the stimulus stayed on screen in this stage mimicked trials where observers were able to resample the display themselves. The pool of possibilities started out with the sampling characteristics of the observers' free resampling trials during the practice block. The sampling characteristics of the subsequent free resampling trials were gradually added (via replacement) to this pool as the experiment continued. On the last type of trials, observers were unable to interact with the masks (e.g. clicking on the mask had no effect). All observers were fully informed about the task design and performed a block of practice trials. After this resampling interval of variable duration, the dot in the middle turned red indicating participants should make a decision by pressing 'd' or 'f' on the keyboard. Unlike in the first pilot, participants had to make a response

with the keyboard instead of the mouse, because we did not want participants making a decision by accident when they tried to resample an element exactly at the time the red dot appeared.



**Fig. 3.5 Example of a trial in pilot 2.** After a brief presentation of 600ms the 8 squircles were masked and the green dot indicates that participants could try to reveal stimuli during a variable resampling interval (2000, 4000, 6000ms). Participants could resample stimuli by clicking with the mouse on the mask covering it (in this case the participant clicked on the mask revealing the purple square). Participants could resample as many stimuli as they preferred during the interval. When participants decided to sample another stimulus, the previous stimulus was masked again. After the variable search interval, the dot in the middle of the screen turns red, prompting a categorical decision from the participant. Again, the response mapping for this participant is to press left if on average the display represented more red circles and blue squares, or right for red squares and blue circles. Thus, the correct answer on this trial is to press right (red square / blue circle).

This modified paradigm allowed us to measure the sampled information on each trial (during the search interval). For this particular pilot, we aimed to investigate whether the possibility to reveal stimuli is beneficial or detrimental for performance (in contrast to not being able to resample or not being able to choose which elements are resampled). We predicted that the ability to resample freely is beneficial for performance in comparison to not being able to sample and fixed sampling. However, we wanted to explore whether under certain conditions the ability to sample could also have detrimental effects.

More importantly, this design allowed us to quantify which information is sampled, allowing the investigation of potential drivers of the sampling policy. In particular, we aimed to investigate whether both the perceptual characteristics as well as the decision values of the stimuli influenced the sampling policy during an experimental design in which observers are sampling multiple stimuli in a very short time-frame. Previous literature has displayed evidence for a positive test strategy in humans – i.e. the tendency to seek out information that confirms your current hypothesis (Klayman & Ha, 1989; Nickerson, 1998). The current task design allowed us to make conclusions whether information sampling is more driven by the perceptual features or the decision values of the previously sampled information.

### 3.3.2. Results

We did not analyse our RT data due to the fixed response interval. For the present analysis we excluded 2 out of 23 participants with mean accuracies lower than 60%, resulting in an average accuracy of 69% ( $n=21$ ).

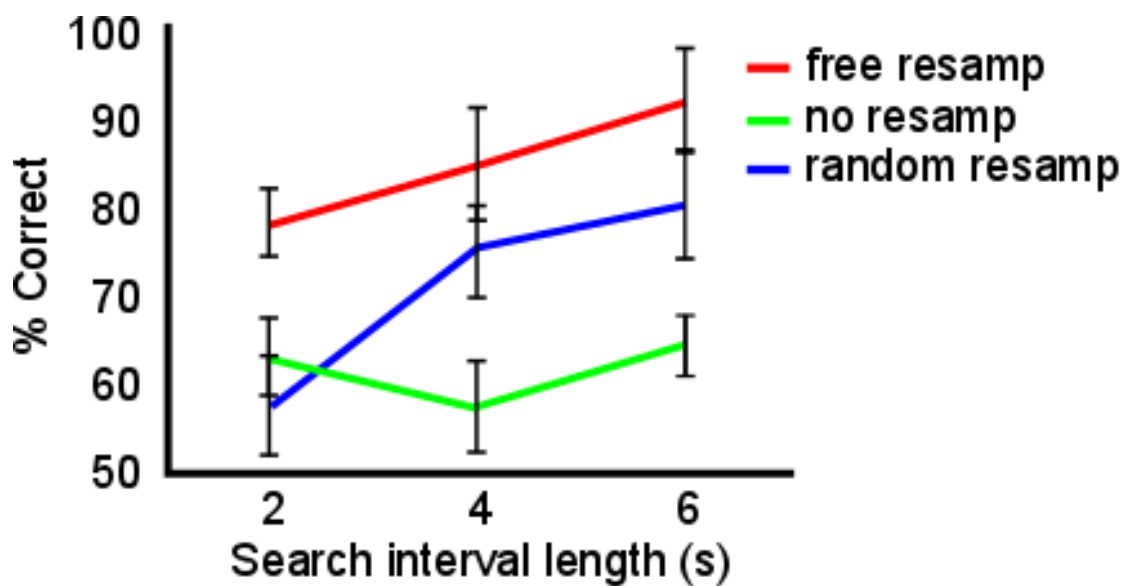
#### 3.3.2.1. *Effect of stimulus distribution*

Even with such a short presentation duration and the inclusion of a re-sampling step, similar effects of mean decision value and variance in the stimulus on performance were observed in experiment 2. Performance increased on high mean (evidence strength) trials compared to low mean trials ( $F_{(1,20)}=12$ ;  $P<0.01$ ). Performance decreased on high variance (low reliability) trials compared to low variance trials ( $F_{(1,20)}=9$ ;  $P<0.01$ ). The interaction between variance and mean was not significant ( $P>0.05$ ). These effects replicate earlier findings of pilot 1 as well as earlier studies.

#### 3.3.2.2. *Resampling type and search interval length*

Simply analysing the influence of resampling type (free, fixed or no resampling) on performance yielded a significant main effect ( $F_{(1,20)}=20$ ;  $P<0.001$ ). Post-hoc tests showed that the difference in performance between random resampling and no resampling was significant ( $F_{(1,20)}=5.9$ ;  $P<0.05$ ), with performance being higher for random resampling. This means that if participants were allowed to integrate evidence for longer (due to resampling),

their performance was better. Performance increased if subjects could resample freely compared to no resampling ( $F_{(1,20)}=63$ ;  $P<0.01$ ). Finally, performance increased when observers could resample freely compared to when stimuli were randomly resampled for them ( $F_{(1,20)}=14$ ;  $P<0.01$ ), indicated by **Fig. 3.6**.



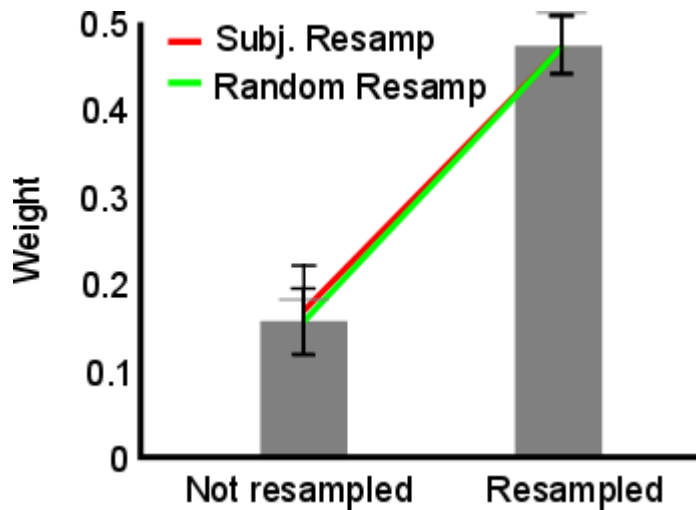
**Fig. 3.6** The interaction between resampling type and the length of the search interval. The green line indicates participants could not resample. The blue line indicates that items were randomly revealed to the participant during the search interval. Finally, the red line indicates participants could resample voluntarily (i.e. free resampling). Error bars represent the standard error of the mean (SEM).

These results are a first indication that participants are resampling the stimuli in an efficient way as they improved their performance. Moreover, performance increased for longer resampling intervals ( $F_{(2,40)}=6.2$ ;  $P<0.01$ ). Being able to search voluntarily (free resampling) is beneficial (**Fig 3.6**). However, this beneficial effect levels off: there is no difference between 4 and 6 second intervals ( $F_{(1,20)}=2$ ;  $P>0.05$ ). As expected, search interval has no effect when no resampling is possible ( $P>0.05$ ) leading to an interaction between sampling condition and interval ( $P<0.05$ ). In summary, these results indicate that resampling stimuli is beneficial for observers and that observers resampled evidence in an efficient way (although this beneficial effect is capped).

Next, we investigated whether resampled stimuli were more predictive of choice than stimuli that were not resampled. We used probit regression to test whether observers' choices (0 = category 1, 1 = category 2) were better predicted by resampled stimuli compared to stimuli that were not resampled

$$p(category2) = \phi \left[ b + \sum_{k=1}^{8-n} w1_k \cdot (DV_{Not\_R}) + \sum_{n=1}^n w2_n \cdot (DV_R) \right]$$

where  $DV_{Not\_R}$  is the decision value of not resampled elements and  $DV_R$  is the decision value of  $n$  resampled elements. **Fig. 3.7** shows us that resampled stimuli predict the final choice more strongly than stimuli which were not resampled, indicated by the higher beta weights from the probit regressions ( $F_{(2,40)}=5.2$ ;  $P<0.01$ ).



**Fig. 3.7 Weight given to resampled and not resampled items split up according to free (red) and random (green) resampling.** More weight is given to resampled items (grey bars). However, this effect does not change when people could choose which items to resample compared to random resampling (i.e. no significant difference between the coloured lines). Error bars represent the standard error of the mean (SEM).

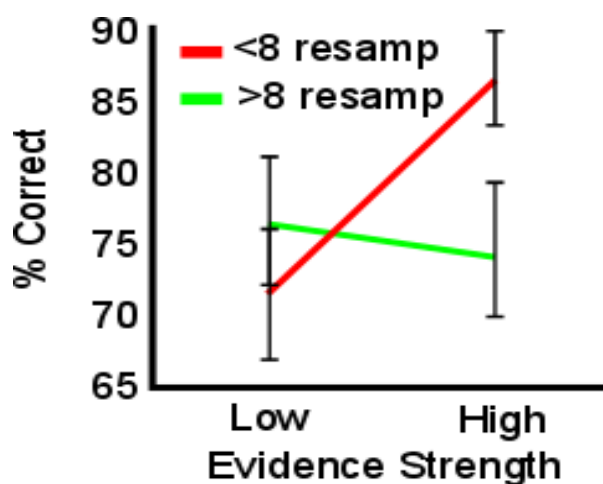
However, this is independent of whether the resampled items were voluntarily chosen (i.e. people can resample themselves) or resampled randomly by the computer. The weight given to resampled stimuli is the same for free vs. random resampling ( $P > 0.05$ ). Thus, observers performed better when they were able to choose which elements to resample but the weight given to resampled stimuli was the same for free and random resampling. One possibility is that when participants re-sample items, they were able to resample more 'valuable' items (e.g. items with higher decision values, or items for which they had a worse representation after the initial 600ms presentation).

Next, we focused on understanding the effects of resampling under different evidence strength and evidence reliability conditions.

### *3.3.2.3. Resampling and evidence strength*

We wondered whether the benefits of information sampling differed as a function of the quality of the information on offer, i.e. with the mean and variance of the momentary decision values. We classified trials as 'high' or 'low' number of items resampled based on a median split per subject. Resampling more items is beneficial compared to resampling fewer items irrespective of the evidence strength (median split,  $F_{(1,20)}=6.9$ ;  $P < 0.05$ ). Surprisingly, the interaction between amount sampled and evidence strength was not significant ( $P > 0.05$ ), as one could imagine that when the evidence is ambiguous, one might resample more information to become more certain of their choice. Next, we split up this interaction for each length of the search interval. The main effect of resampling many items disappeared: resampling many items is not beneficial when you only have 2 seconds and 4 seconds to resample stimuli ( $P > 0.05$  for both). Moreover, the main effect of evidence strength disappeared under 2 second search intervals: strikingly, participants did not perform better on trials when the strength of evidence was stronger ( $P > 0.05$ ). Snap judgments (i.e. judgments under time pressure) seem to be made without considerations of evidence strength. Indeed, we found the main effect of evidence strength under 4 seconds search intervals ( $F_{(1,20)}=5.7$ ;  $P < 0.05$ ). The interaction between amount resampled and evidence strength was not significant ( $P > 0.05$  for both 2s and 4s search intervals).

There was no main effect of amount resampled ( $P > 0.05$ ) under 6 second search intervals. Replicating the 4s search interval, with sufficient search time, there was a significant influence of evidence strength ( $F_{(1,20)}=5.1$ ;  $P < 0.05$ ). Interestingly, for the six second search interval, the interaction between evidence strength and amount sampled was significant ( $F_{(1,20)}=6.1$ ;  $p < 0.05$ ). No difference is observed between sampling many items or just a few when the strength of the evidence is low ('many' and 'few' were defined via a medium split). However, when the strength of the evidence is high, sampling many items becomes detrimental compared to sampling only a few items (Fig. 3.8).



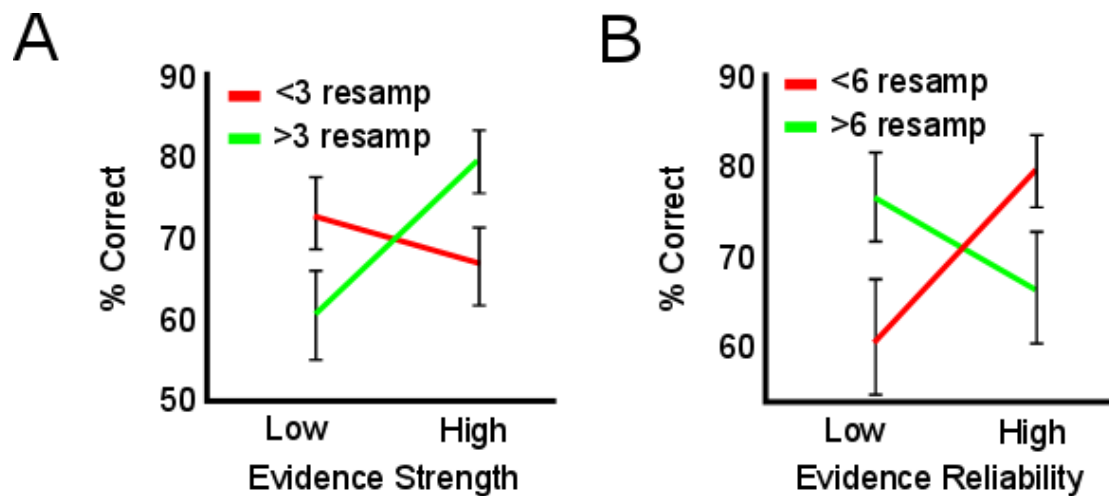
**Fig. 3.8** Interaction between evidence strength and the amount of resampled stimuli on 6 second search interval trials. No difference is observed between sampling many items or just a few when the strength of the evidence is low. However, when the strength of the evidence is high, sampling many items becomes detrimental compared to sampling only a few items. Error bars represent the standard error of the mean (SEM).

#### 3.3.2.4. Resampling and evidence reliability

Resampling many items is beneficial compared to resampling few items irrespective of evidence reliability ( $F_{(1,20)}=6.9$ ;  $P < 0.05$ ). No main effect of evidence reliability was observed ( $F_{(1,20)}=2.6$ ;  $P > 0.05$ ). The interaction between amount sampled and evidence reliability was not significant ( $F_{(1,20)}=0.2$ ;  $P > 0.05$ ).

We found no main effect of amount sampled in all search intervals ( $p > 0.05$  for all conditions), nor effect of evidence reliability under 2 second and 4 second search intervals ( $p > 0.05$ ). However, there was a significant interaction between evidence reliability and amount resampled in during the two second search interval: sampling many items was beneficial for performance when the evidence is reliable but detrimental when the evidence is less reliable (2 second: [ $F_{(1,20)}=9.7$ ;  $P<0.01$ ]; **Fig. 3.9.** left panel).

Interestingly, this interaction reversed during the 4 second search interval: sampling many items was detrimental for performance when the evidence is more reliable, but beneficial when the evidence is less reliable [ $F_{(1,20)}=6.8$ ;  $P<0.05$ ]. This interaction was in the same direction for the six second search interval condition but was not significant ( $P>0.05$ ). Thus, whether prolonged sampling is beneficial or detrimental is not only affected by the evidence reliability, but also by the length of the search interval.



**Fig. 3.9. Interaction between evidence reliability and amount resampled.**

Left: Two second search intervals. Sampling many items is beneficial for performance when the evidence is reliable but detrimental when the evidence is less reliable. Right: Four second search intervals. The effect might reverse: sampling many items is detrimental for performance. Error bars represent the standard error of the mean (SEM).

### 3.3.2.5. Confirmation bias

Next, we investigated if participants employed confirmation biases during the task. We distinguished three types of confirmation biases. First, a resampling bias based on the current estimation of the average colour feature value (i.e. are you more likely to sample a red element than a blue element when the average colour of the previously sampled items is red?). Second, a resampling bias based on the current estimation of the average shape feature value (i.e. are you more likely to sample a circle than a square when the

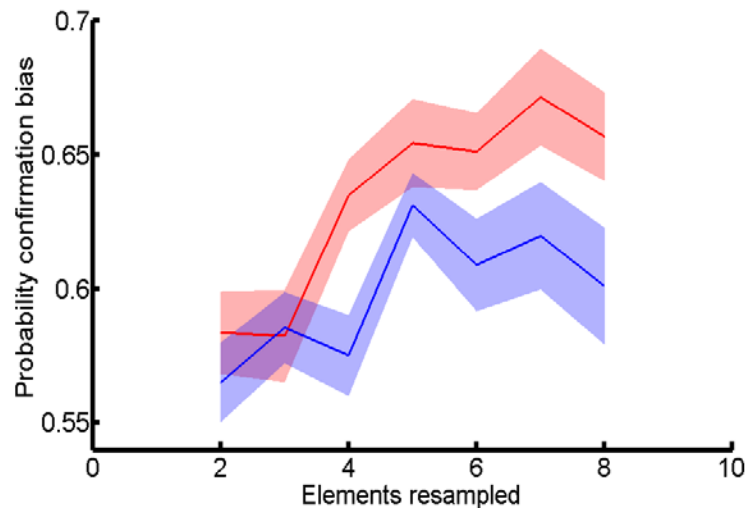
average shape of the previously sampled items resembled a circle more than a square?). Thus, these first two types of confirmation biases are driven by perceptual characteristics of the stimuli (e.g. not the decision value per se). Finally, we differentiated a confirmation bias based on the current estimation of the decision value (i.e. if you have just sampled an element with a positive decision value, are you more likely to sample another element with a positive decision value?).

In order to quantify these biases we first adjusted the decision/shape/colour values of each item on a specific trial so that they are relative to the mean of these values on that specific trial. We did this to factor out the overall bias in evidence towards one or the other category, which would otherwise confound the confirmation bias analysis. Since, without the unadjusted decision values a confirmation bias may result from simply resampling at random, just because there is more evidence for a particular category. Next, we ordered these adjusted values of the resampled items for each trial based on when they were sampled. We then compared each value to the cumulative mean value on that specific trial. Thus, we could see for each resampled item if it confirmed participants' current assessment of the stimulus array or not. Resampled items which confirmed participants' current assessment were given a value of 1, while resampled items which did not confirm the current assessment received a value of 0 (i.e. a binary measure of confirmation bias).

Next, the mean of these values across trials was calculated to measure the confirmation bias. For example: imagine a trial on which an observer resampled six items by clicking on them and we want to calculate the decision value confirmation bias. First we order the decision values of the resampled items from the first to the last resampled element. Second, we compare the decision value of sample  $n$  to the cumulative mean decision value (i.e. the average decision value of samples 1 to  $n-1$ ). If the decision value of the resampled element confirms this cumulative mean decision value it's assigned a value of 1. Otherwise it is assigned a value of 0. We repeat this process for each resampled element and end up with a matrix of ones (resampled items which confirmed the current assessment) and zeros (resampled items which did not confirm the current assessment) for each trial. By taking the mean of these values across trials we have a measure of the confirmation bias for each subject. For example a value of 0.7 for the second resampled element would indicate that there is a 70% chance for the second resampled element to confirm a participants' current assessment of the stimulus array on a given trial. The colour and shape confirmation biases are calculated using an identical approach.

We found evidence for a strong colour and shape confirmation bias, but no confirmation bias based on decision value was observed (**Fig. 3.10**). Thus, the sampling policy in this task was driven by the independent perceptual characteristics of the stimuli, but not the decision value (e.g. the combination of these perceptual characteristics). We reasoned that perhaps participants

cannot process the decision values fast enough to guide sampling processes or perhaps any possible decision value bias is dominated by the sampling biases associated with the independent perceptual characteristics.



**Figure 3.10. Confirmation bias for colour (red) and shape (blue).** The y-axis shows the probability that the resampled element confirms the current assessment of the display (i.e. either the current assessment of the colour or the shape). The x-axis shows which element is resampled (i.e. what is the probability that the 4<sup>th</sup> resampled element confirms your current assessment of the display?). Shaded areas around the lines represent the standard error of the mean (SEM).

Next we analysed the influence of low/high confirmation trials on performance for each search interval length. Low/high confirmation trials are defined as trials for which the confirmation bias is lower/higher than the median confirmation bias. Unexpectedly, neither the colour nor the shape confirmation bias showed any influence on performance (all  $p > 0.05$ ).

Since sampling many items on trials with long search intervals and high evidence strength is detrimental to performance we also investigated the confirmation biases for these specific trials. However, no significant differences in these biases were observed for long search interval trials with high versus low evidence strength. Thus, the shape and colour confirmation biases could not explain the detrimental effect of a longer search interval.

Similarly we investigated the confirmation biases for two second search interval trials where resampling many items is beneficial when the evidence is reliable, but detrimental when the evidence is less reliable. However, again no significant differences in these biases were observed for short search interval trials with high versus low evidence reliability. We also investigated the confirmation biases on four second search interval trials where resampling many items is detrimental when the evidence is reliable and beneficial when the evidence is less reliable. We found that the colour confirmation bias was stronger for high evidence reliability trials. Suggesting that the detrimental effects of sampling many items might be due to the colour confirmation bias. However, subsequent analyses suggest that this is not the case: performance on these trials did not differ for low and high confirmation bias trials ( $P > 0.05$ ). We also found that the shape confirmation bias was stronger for high evidence reliability trials. However, performance on these trials did again not differ between low and high bias trials ( $P > 0.05$ ).

### **3.3.3. Conclusion for Study Pilot 2**

Even though performance increased in the free resampling condition compared to the random or no resampling conditions, resampling many items was detrimental to performance under certain conditions (potentially due to oversampling). Interestingly, sampling many items becomes detrimental on trials where the resampling interval is long (presentation duration/search interval) and the evidence strength is strong.

When participants can sample the display for only two seconds, sampling many items is beneficial for performance when the evidence is reliable, but detrimental when the evidence is less reliable. Interestingly, this effect reverses when participants can sample the display for four seconds: resampling many items is detrimental when the evidence is reliable, and beneficial when the evidence is less reliable.

Furthermore, different biases determining the sampling policy during perceptual decision making were identified: a shape confirmation bias and a colour confirmation bias. However, these biases did not seem to influence performance nor could they explain the aforementioned detrimental effects of oversampling.

### 3.4. Study: Pilot 3

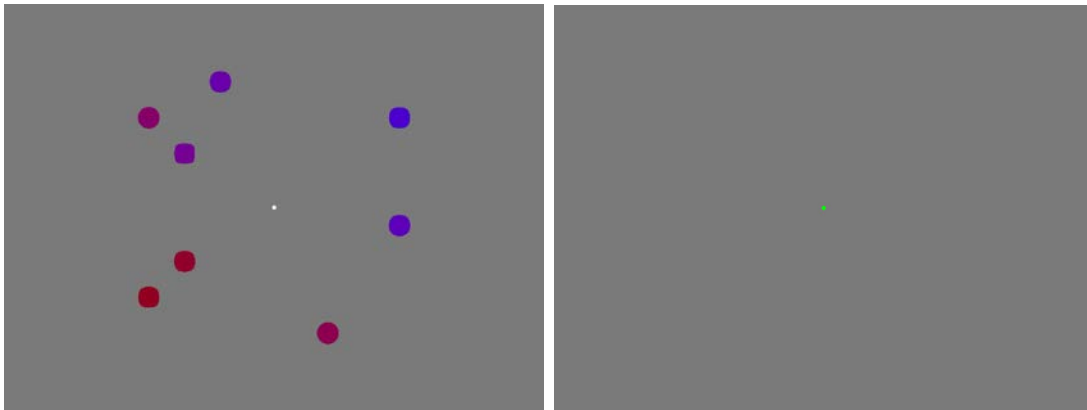
The previous experiment has a number of limitations. Performance in the search experiment might be influenced by having to use more resources (planning and executing motor movements), reducing the benefit of this condition. Secondly, resampling in pilot 2 was quite artificial, and perhaps a more natural way is how we move our eyes and how this influences our choices; Thus, we next conducted an eye-tracking experiment in order to investigate the sampling processes participants engage in during a perceptual decision making task.

#### 3.4.1. Methods

*Sample.* The experiment took place in the Department of Experimental Psychology at the University of Oxford. Seventeen subjects took part in experiment 3. Subjects reported normal or corrected vision. The experiment took approximately 50 minutes. Participants provided written consent before the experiment and received 10£ for their participation. The local ethics committee approved the experiment (approval no. MSD-IDREC-C1-2009-1).

The paradigm used in pilot 3 is again a close adaptation of the first pilot. The third pilot consisted of 10 blocks each containing 48 trials. The stimuli (squirrels) were the same as the ones used in the previous pilots. Again, 8 squirrels were presented on semi-random locations. These locations were

restricted in minimum distance from each other. The squircles are presented for a variable duration (see **Fig. 3.11** for details). Unlike pilot 1, we did not manipulate the eccentricity of the display (maximum distance from the point in the middle). We chose to employ the largest spread of the display used in pilot 1 in order to elicit more eye-movements and sequential processing of the stimuli. Once again, the task was to integrate the shape and colour values of each stimulus and make a perceptual decision based on their combined features. A calibration session followed every two blocks in order to achieve accurate eye-tracking measures.



**Figure 3.11. Example trial experiment 3 (Eye-tracker).** Eight squircles are presented for a variable duration (1,2 or 3s) in random restricted locations. After this variable duration the dot turns green, indicating that participants have to make a response. We chose not to mask the stimuli in this experiment. Participants then had to press the left or right mouse key to make a perceptual decision based on the combined shape and colour features of each stimulus.

### 3.4.2. Results

Since we employed a fixed response interval we did not analyse our RT data. For the present analysis we excluded 3 out of 17. One participant did not move his eyes from the middle of the screen, one participant wore glasses that resulted in unusable eye-tracking data, and one participant had a mean accuracy lower than 60% ( $n=14$ ).

#### 3.4.2.1. *The influence of stimulus distribution on performance*

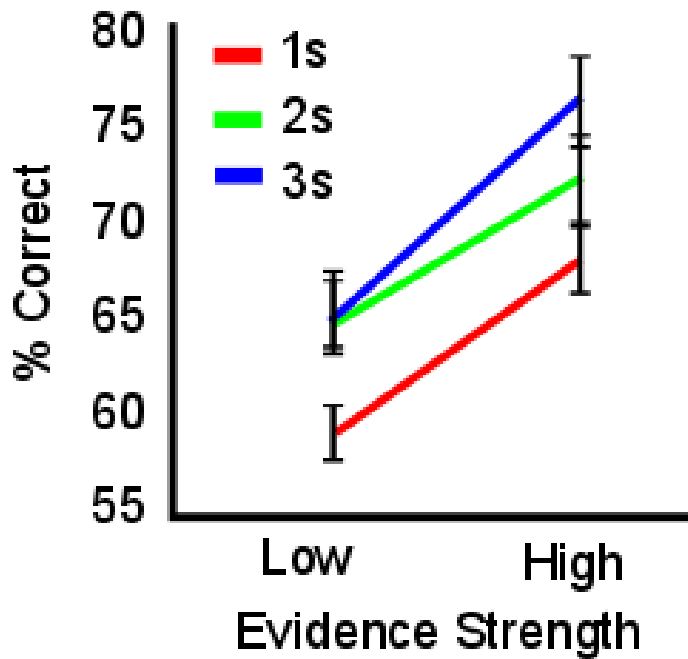
Similar effects of evidence strength (mean decision value) and evidence reliability (variance) on performance were observed in the third pilot. Performance increased on high mean trials compared to low mean trials ( $F_{(1,13)}=64$ ;  $P<0.001$ ). Performance decreased on high variance trials compared to low variance trials ( $F_{(1,13)}=17$ ;  $P<0.001$ ). The interaction between variance and mean was not significant ( $P>0.05$ ).

#### 3.4.2.2. *The influence of presentation time on performance*

As expected, performance again increased with longer durations, just like in earlier pilot data ( $F_{(2,26)}=8.5$ ;  $P<0.05$ ). No interaction was observed between presentation duration and variance ( $P>0.05$ ). Closer inspection of these results inform us that, contrary to experiment 1, performance did not vary according to variance for low duration trials ( $P>0.05$ ). Performance did decrease on high variance trials for medium duration trials ( $F_{(1,13)}=15$ ;

$p < 0.05$ ). Performance also decreased for higher variance trials on high duration trials ( $F_{(1,13)} = 11$ ;  $P < 0.05$ ). This casts doubt on the small increase in performance on high variance / high duration trials observed in pilot 1. However, this could also result from an increased familiarity with the task since all subjects previously completed experiment 1 and 2.

The interaction between presentation duration and mean decision value was significant: when the evidence strength is low, performance decreased more strongly for long duration trials (**Fig. 3.12**,  $F_{(2,26)} = 3.8$ ;  $P < 0.05$ ). Performance increased for longer durations (3s) when the strength of the evidence was high ( $F_{(1,13)} = 7.7$ ;  $P < 0.05$ ), but this benefit disappeared when the strength of the evidence was low ( $F_{(1,13)} = 0.019$ ;  $P > 0.05$ ). Similar to the same effect in pilot 1, this could suggest a cost of sampling more information on low evidence strength trials.

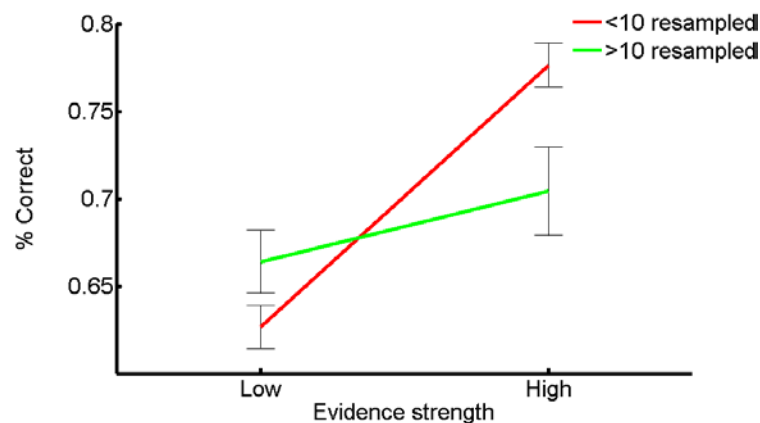


**Figure 3.12 Interaction between evidence strength and presentation duration.** Replicated effect of pilot 1: this significant interaction effect can either signify a benefit of a 3 second presentation for high mean trials or a cost for low mean trials, relative to a 1 second presentation. Error bars represent the standard error of the mean (SEM).

#### 3.4.2.3. Resampling and evidence strength

Overall, resampling more items is beneficial compared to resampling a few items irrespective of the evidence strength (median split,  $F_{(1,13)}=24$ ;  $p<0.001$ ). The interaction between amount sampled and evidence strength was not significant ( $p>0.05$ ). Next, we again split up this interaction by presentation duration.

The amount of resampled items showed no significant influence on performance for all presentation durations ( $p > 0.05$ ). Like in pilot 2, there was no main effect of evidence strength under short presentation duration. We only found significant effects of evidence strength on performance under 2s ( $F_{(1,13)}=11$ ;  $p<0.05$ ) and 3s presentation durations ( $F_{(1,13)}=12$ ;  $P<0.01$ ). Lastly, we found a significant interaction between amount sampled and evidence strength only for the longest presentation duration trials, similar to experiment 2 (**Fig. 3.13**;  $F_{(1,13)}=3.8$ ;  $P<0.01$ ). Just as in the previous experiment, sampling many items seems to be detrimental (or less beneficial than sampling only a few) on trials with long presentation durations and high evidence strength.

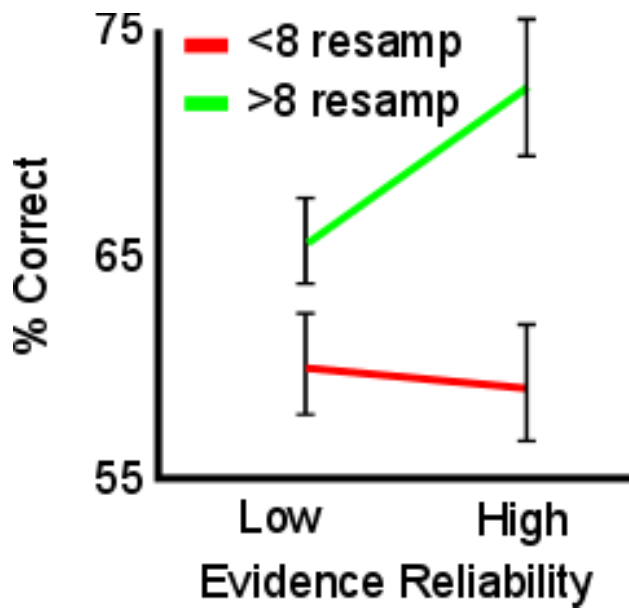


**Fig. 3.13** Interaction between evidence strength and the amount of resampled stimuli on 3 second resampling interval trials. Replication of the effect from pilot 2. No difference is observed between sampling many items or just a few when the strength of the evidence is low. However, when the strength of the evidence is high, sampling many items becomes detrimental compared to sampling only a few items. Error bars represent the standard error of the mean (SEM).

#### 3.4.2.4. Resampling and evidence reliability

Resampling many items is beneficial compared to resampling few items irrespective of evidence reliability ( $F_{(1,13)}=27$ ;  $P<0.001$ ). A main effect of variance was observed ( $F_{(1,13)}=13$ ;  $P<0.01$ ). The interaction between amount sampled and evidence reliability was also significant: the negative effect of low reliability on performance is stronger when more items are resampled ( $F_{(1,13)}=4.7$ ;  $P<0.05$ ).

The interaction between evidence reliability and amount of resampled stimuli was only significant in the 1 second presentation duration condition: sampling many items is beneficial for performance when the evidence is reliable and this benefit decreased when the evidence becomes less reliable ( $F_{(1,13)}=5.1$ ;  $P<0.05$ ; **Fig. 3.14**). There was a similar trend for the 2 second presentation duration trials ( $F_{(1,13)}=4.5$ ;  $P_s = 0.053$ ). These effects are in line with a similar pattern that was observed in the second pilot.



**Fig 3.14 Interaction between evidence reliability and amount resampled for 1 second presentation duration trials.** Sampling many items is beneficial for performance when the evidence is reliable and this benefit decreases when the evidence becomes less reliable. Error bars represent the standard error of the mean (SEM).

#### 3.4.2.5. Confirmation biases

Next, we investigated if we could find evidence for the perceptually driven confirmation biases (e.g. colour & shape) from the second pilot during this task. Moreover, we were interested to investigate whether changes to the design allowed the discovery of a confirmation bias based on the decision values. Similar methodology as in pilot 2 was used to calculate the confirmation biases. However, due to changes in the paradigm we could not capture the sampled items as accurately when observers fixate a certain

location on the screen (How much information can observers process in proximity to their fixation? How well is information processed further away from the fixation?).

As a first step, we characterised discrete fixations larger than 100 ms using the Eyelink pre-set criteria (velocity threshold of 30, acceleration threshold of 8000, and a saccade motion threshold of 0.1). Next, the eye-tracking data were analysed by treating the stimulus array on each trial as a scene or 'landscape' defined over the (1024 x 768) pixel space. This landscape approach has been described in detail in chapter 2.

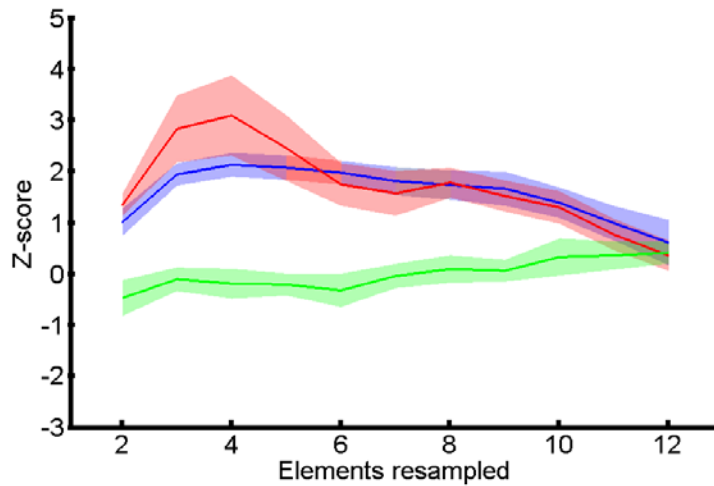
Landscapes were created by summing bivariate Gaussian functions centered on the coordinate locations of each stimulus in pixel space. First, the width of the information processing spotlight  $s$  was estimated independently for each subject by calculating the width the extent to which items were viewed compared to uninformative space ( $s \sim 50$ , see chapter 2 for details). Second, the height  $g$  of each Gaussian function was set to the decision value of the stimulus at that location. Therefore, on each fixation each stimulus was assigned an adjusted decision value based on the position of this fixation on the landscape. By comparing the sign of the decision information value of each fixation, we could check whether each fixation confirmed (i.e. same sign of the information value;  $=1$ ) or disconfirmed (i.e., opposite sign;  $= 0$ ) observers' current hypothesis (i.e. confirmation, defined as the running mean of accumulated evidence so far). After calculating the average of these values

per trial, this yields a quantitative measure of whether participants adapted their sampling policy as a function of the sampled evidence on each trial.

Next, we constructed a null-distribution specific to our pilot, in order to factor out confirmation biases that could result just from chance or due to the method itself. For instance when there is a small range saccade from X to Y in a random direction, the adjusted decision values of X and Y will be highly correlated. Similarly, the overall values on the landscape are biased towards one side of the category boundary (e.g. due to the manipulation of evidence strength). These factors could lead to the discovery of a confirmation bias, but it is a false result. The null-distribution was constructed by randomly permuting the numbers of each trial and repeating the analysis 100 times, as an extra careful measure to control for any possible correlations in the values. Thus for each random permutation, each fixation on a given trial was associated with the decision value from another randomly selected stimulus location on that trial, while the information sampling characteristics remained the same. By comparing the confirmation bias statistics (i.e. those generated in our pilot) to the statistics under the null distribution (i.e. those generated with the randomly permuted values), we could calculate z-scores and we can assess the reliability of fixation-to-fixation dependencies due to only cognitive factors.

We again distinguished three types of confirmation biases based on shape, colour or decision value. We replicated the effects of the second pilot: a colour

and shape confirmation bias, but no confirmation bias based on decision value was observed (**Fig. 3.15**).



**Fig. 3.15. Confirmation bias for colour (red), shape (blue) and decision value (green).** The y-axis shows the z-score difference between the confirmation bias statistics (i.e. those generated in our pilot with non-permuted values) and the null distribution. The x-axis indicates which element is resampled. Shaded areas around the lines represent the standard error of the mean (SEM).

Next we analysed the influence of low/high confirmation trials on performance for each presentation duration. Both the colour and shape confirmation biases showed no significant influence on performance (all  $p > 0.05$ ).

Since sampling many items on trials with long presentation durations (3s) and high evidence strength was again detrimental to performance (~experiment 2)

we also investigated the confirmation biases for these specific trials. However, none of the discovered biases showed a significant influence on performance nor were they stronger/weaker for these specific trial types ( $P > 0.05$ ).

Since sampling many items on trials with long presentation durations (3s) and high evidence strength is detrimental to performance (~pilot 2) we also investigated the confirmation biases for these specific trials. However, again, none of the discovered biases have a significant influence on performance nor do they seem to be stronger/weaker for these specific trials ( $p > 0.05$ ).

Similarly we investigated the confirmation biases for one second duration trials where sampling many items was beneficial for high evidence reliability. The decision value confirmation bias has no significant influence on performance nor was it stronger for these trials (all  $p > 0.05$ ).

Thus, these perceptually driven confirmation biases did not seem to influence performance nor could they explain the detrimental effects of resampling in certain conditions.

### 3.5 CONCLUSION

First, our results further emphasize the importance of distinguishing evidence strength and evidence reliability in decision making paradigms, as previously shown (De Gardelle & Summerfield, 2011).

More importantly, these three preliminary experiments provide clear evidence for beneficial effects of increased resampling intervals and the opportunity to resample stimuli voluntarily. However, under certain conditions resampling many items becomes less beneficial and sometimes even detrimental. Our results suggest that resampling many items becomes less beneficial as evidence reliability decreases. Furthermore, resampling many items is no longer beneficial to performance for reliable evidence when the resampling interval becomes too extensive. Additionally, sampling many items becomes detrimental when the resampling interval is long (presentation duration/search interval) and the evidence strength is high.

Finally, we identified two biases driving the sampling policy: a shape and a colour confirmation bias. However, we failed to relate these perceptual biases to the observed detrimental effects of resampling more information

We were hoping to demonstrate effects of the accumulated decision values on the sampling policy in these pilots. We chose to investigate the confirmation bias, as it is one of the most well established biases in the

literature (albeit in paradigms that are quite different from the ones employed in this chapter). However, the accumulated evidence did not influence the sampling policy. Why did we only find evidence for perceptual characteristics driving the sampling policy but not the decision values? A first possibility is that the perceptual characteristics are stronger drivers of the sampling policy and therefore completely diminish any form of guidance of the sampling policy through decision information. Second, it is possible that combining the perceptual characteristics into a decision value is more difficult (or more time consuming) in the periphery. Since observers sampled evidence with their eyes, the intervals between subsequent samples of evidence were extremely short. Thus, these short intervals might have contributed to the null effect of decision information on the sampling policy. However, note that this is exactly what we were hoping to test. Is it possible to find evidence for an influence of decision information on the sampling policy, even when the interval between subsequent samples is so short (e.g. like in typical decision making tasks)?

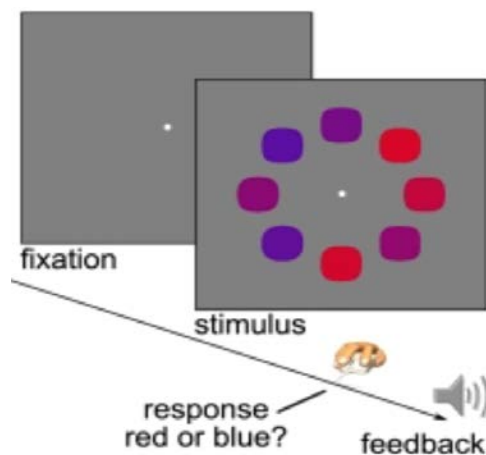
## **Chapter 4: Robust sampling of decision information during perceptual decision making**

### **4.1 Introduction**

In this chapter, we aimed to characterise the sampling policy that dictates how information is sampled and weighted when forming a decision. We again investigated the sampling process by measuring eye movements while human participants performed a categorization task that involved averaging decision information in a spatial array. Remember that the perceptual features of the stimuli in experiment 3 from the previous chapter (colour and shape) dominated the sampling policy, and that the combined decision value did not affect the sampling policy. Because of this, we opted to use fully visible symbolic numbers as stimuli in this chapter instead of the 'squircle' stimuli (see methods). It is important to note in this regard, that a large body of evidence has shown that numbers can be treated as noisy estimates of magnitude, which can be averaged just like continuous sensory signals during decision making (Brezis, Bronfman, & Usher, 2015; Feigenson, Dehaene, & Spelke, 2004; Hyde & Spelke, 2011; Pica, Lemer, Izard, & Dehaene, 2004). Moreover, the short duration for which numerical information was available and the low average dwell times observed for each number (precluding overt arithmetic approaches) in the experiments described in this chapter, further reinforce this assumption.

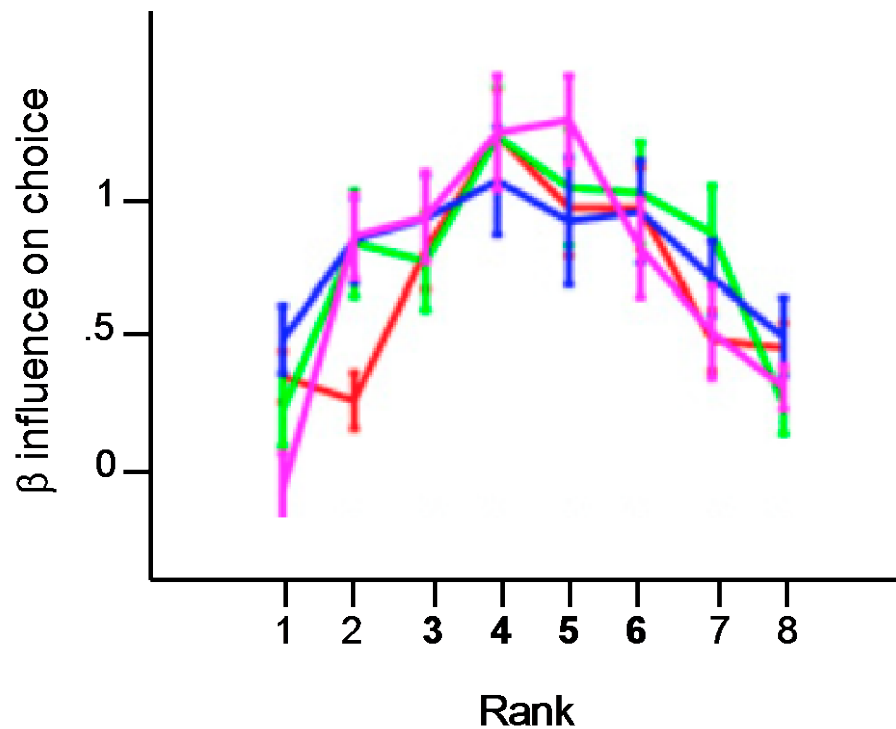
The use of visible symbolic numbers as stimuli, allowed us to assess how gaze placement was determined by the statistics of decision information, rather than the perceptual arrangement of the stimulus array. In this chapter, we will focus on two key determinants of the sampling policy: the influence of the value of the decision information sampled thus far (i.e. confirmation bias) and the influence of extreme values (i.e. robust sampling).

When averaging perceptual features across multiple locations under central fixation, participants tend to downweight or ignore extreme values, such as a deviant expression in a crowd of faces (Haberman & Whitney, 2010) or outlying color values in an array of shapes (De Gardelle & Summerfield, 2011). De Gardelle & Summerfield presented observers with a multi-element averaging task very similar to our tasks described in the previous chapter (**Fig. 4.1**).



**Fig 4.1 Experimental design of the multi-element averaging task.** Figure copied from de Gardelle & Summerfield (2011) under CC BY-NC 3.0 copyright. On each trial, subjects had to average the colour or shape from an eight element array, in order to classify the display as either predominantly red or blue. The responses of subjects were less influenced by more extreme colours in comparison to less extreme colours, providing evidence of the selective weighting of decision information.

Observers had to discriminate the average value on a feature dimension (colour / shape) for eight stimuli that were simultaneously presented on a computer screen (e.g. is the average colour of the stimuli more red or blue? Is the average shape closer to a square or a circle?). The authors used logistic regression in order to estimate the weight (e.g. the strength of influence) of each element on choice (**Fig. 4.2**).



**Fig 4.2 Robust averaging of decision information.** Figure adapted from **de Gardelle & Summerfield (2011)**. Inlying information (rank 3-6, bold) exhibit a stronger influence on choice than outlying information (rank 1-2-7-8, non-bold), demonstrated by an inverted U-shape.

Inlying and outlying elements were defined based on their sorted rank: the four centrally ranked elements were classified as inliers (e.g. close to the category boundary), while the outer four elements were classified as outliers (e.g. far away from the category boundary). De Gardelle & Summerfield found that inliers (ranks 3-6) contributed more to choices than outliers (ranks 1-2-7-8), i.e. a downweighting of outlying values (inverted U shape).

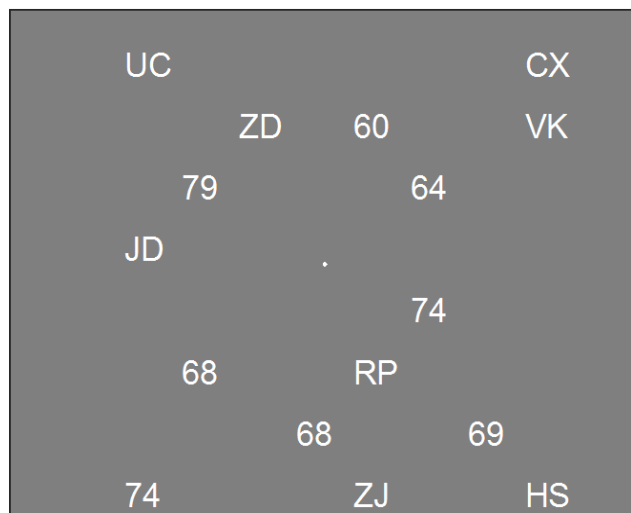
This has led to the proposal that “inlying” items (those that are closer to the category boundary) are processed with higher gain than outliers, i.e. ‘robust averaging’ (De Gardelle & Summerfield, 2011; Michael et al., 2013). However, previous studies have left unresolved the question of whether robust averaging occurs because inlying items are sampled (i.e., overtly attended) by preference or whether the influence of outliers is attenuated at a later decision stage. For example, one model proposes that observers simply compute a posterior likelihood ratio conditional on the past distribution of features and their associated responses. This is equivalent to using a nonlinear decoder that “squashes” outlying information, muting its influence in the average and prolonging response latencies when decision information is heterogeneous (De Gardelle & Summerfield, 2011)

To test this theory, we developed the landscape approach discussed in chapter 2 to more accurately capture the sampled information in real time. We outline this methodological approach and its applications in more detail in this chapter. Using this methodology, we replicated the robust averaging effect and additionally found that humans prefer to fixate the inlying items in a stimulus array during a perceptual decision-making task—a policy that contradicts the well-established preference for fixating deviant items in visual search tasks (Towal, Mormann, & Koch, 2013; Võ & Henderson, 2010). Moreover, humans prefer to fixate items that fit with their current assessment of the stimulus (e.g. a confirmation bias).

## 4.2 Methods

*Sample.* Sixty healthy human participants with normal or corrected to normal vision were recruited in two experiments (experiment 1,  $n = 30$ ; experiment 2,  $n = 30$ ). Participants provided written consent beforehand, and the study was approved by the Oxford University Medical Sciences Division Ethics Committee (approval no. MSD-IDREC-C1-2009-1). There were only minor differences between the two experiments (highlighted below), and so we collapse across their results in this chapter. Six participants whose performance failed to differ from chance were excluded from all further analyses, leaving  $n = 54$ .

*Stimuli.* On each trial, participants viewed an array composed of eight two digit numbers. Stimuli were generated using PsychToolbox ([psychtoolbox.org](http://psychtoolbox.org)) for MATLAB (Mathworks) and presented on a 17-inch screen (resolution, 1024 x 768) viewed from a distance of 60 cm. Numbers ( $\sim 1.15^\circ$  horizontal and  $\sim 0.95^\circ$  vertical visual angle) were presented at randomly selected locations within an invisible 8 x 8 grid. Eight additional randomly selected grid locations contained letter pairs, which were distracters. These task-irrelevant items (e.g. letters) introduced a gentle crowding effect and encouraged shifts of gaze during decision formation. No letter or number pairs were presented within the central four grid locations. An example stimulus array can be seen in **Fig. 4.3**.



**Fig 4.3 An example stimulus.** Each stimulus display consists of eight numbers and 8 distractor letter pairs distributed within an invisible 8 x 8 grid. Participants could sample the numbers freely by moving their eyes and had to indicate whether the average of all of the numbers on the screen was higher or lower than the reference value. The display could be sampled for up to 3,000 ms (experiment 1) or up to 5,000 ms (experiment 2).

At the start of each block of 96 trials, participants were informed of the reference value, against which they compared the average number for all trials across that block. Subsequently, each trial began with the onset of a central fixation cross for 500 ms. This was followed by the stimulus array, which remained on the screen for either up to 3,000 ms (experiment 1) or up to 5,000 ms (experiment 2). During array presentation, participants freely moved their eyes across the screen and chose whether the reference number or the average of the number array was higher (lower). The decision framing

(decide whether numbers were higher vs. lower than the reference) was counterbalanced across participants within each experiment. The decision framing manipulation was introduced because we did not want any of our results to be confounded with a logarithmic (Weber-like) compression of the number line or a bias towards a particular part of the number space. Participants could respond at any point during the array presentation by pressing one of two keys ("F" or "J"), and they received auditory feedback in the form of a high or low tone lasting 100 ms (800/400 Hz, correct/incorrect). In experiment 2 only, participants were additionally paid in proportion to their accuracy. In the high frame, participants received a bonus corresponding to their choice (reference value or the average of the numbers on that trial, in pence) on 20 randomly selected trials; Thus, participants would maximise their reward by correctly choosing the average of the numbers when this average was higher than the reference (and vice versa when the reference was higher than the average). In the low frame, the reward of participants started at £20, from which their choice was deducted on 20 randomly selected trials. Participants earned ~£10 on average for their participation. Participants performed eight blocks of 96 trials (768 trials in total) in an experiment lasting approximately 50 minutes.

*Design.* The reference number  $R$  on each block was drawn randomly from a uniform distribution between 25 and 75. Numbers were drawn from a Gaussian distribution with mean  $\mu$  and SD  $\sigma$ , and participants were asked to compare their average to a reference value  $R$  that varied across blocks ( $\mu = R$

$\pm 2$  or  $R \pm 4$  in low and high  $\mu$  trials, respectively, and  $\sigma = 6$  or  $12$  in low and high  $\sigma$  trials, respectively). Sample values were resampled to ensure that the sample mean and variance fell within the desired value for that trial with a tolerance of 0.7 and 0.1, respectively. Each of the eight trial types occurred equally often, and their presentation order within a block was random.

*Eye-Tracking.* All eye movement data were recorded monocularly using an SR Research EyeLink 1000 eye-tracking system at a sampling rate of 1,000 Hz. Participants placed their heads on a chin rest that was mounted on the table at a fixed distance of 68 cm from the screen. Before the start of each new block, a standard calibration/validation procedure was executed, in which participants tracked a dot that appeared at 13 evenly spaced locations across the screen. Data were converted to vectors of x, y position in the frame of reference of the screen exported to MATLAB using the EyeLink toolbox and analyzed using in-house scripts. Missing data (due to blinks or eye movements away from the screen) were ignored (i.e., set to NaN) in all analyses. Eye movement traces were cut off at reaction time for all analyses described below. All of our analyses assume that the eye-tracker faithfully estimates the pixel position of the gaze on the screen.

We analyzed the eye-tracking data using both a combination of conventional methods (e.g., identifying the location and dwell time of individual fixations) and a technique that treated the stimulus array as a topography of decision information that was sampled continuously with the gaze (e.g. landscape analysis). The landscape analysis is described in detail in chapter 2 but will be

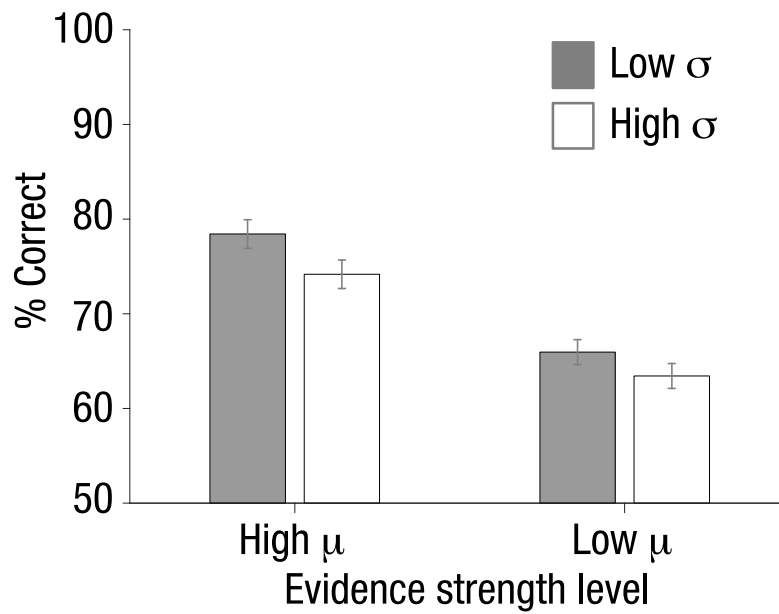
applied and extended upon in the results section.

## 4.3 Results

We report data collapsed over two experiments that differed in terms of the response deadline and monetary incentive. Experiment 1 utilised a 3,000 ms deadline with a fixed £10 payment, while Experiment 2 utilised a 5,000 ms deadline with a performance dependent payment.

### 4.3.1 Effects of evidence strength $\mu$ and evidence reliability ( $\sigma$ )

Overall mean accuracy was 70.5% ( $n = 54$ ). Performance depended on both the mean ( $\mu$ ) and variance of the distribution from which the numbers were drawn (**Fig. 4.4**). Accuracy increased on high  $\mu$  trials compared with low  $\mu$  trials [ $F(1, 53) = 76, P < 0.001$ ] and decreased on high  $\sigma$  (more variable evidence) trials compared with low  $\sigma$  trials [ $F(1, 53) = 47, P < 0.001$ ], as previously reported (de Gardelle & Summerfield, 2011; Michael, de Gardelle, Nevado-Holgado, & Summerfield, 2015).



**Fig 4.4 Influence of evidence strength ( $\mu$ ) and reliability ( $\sigma$ ) on accuracy.**

Performance increased for higher evidence strength and higher evidence reliability. Bars are shown with 95% confidence intervals.

The influence of evidence strength ( $\mu$ ) and evidence reliability ( $\sigma$ ) on reaction times (RT) is depicted in **Table 4.5**, split by experiment due to the different response deadline.

**Table 4.5 Influence of evidence strength ( $\mu$ ) and reliability ( $\sigma$ ) on reaction times.**

| Exp.            | $\sigma$      | High $\mu$       | Low $\mu$        | Effect on RT                                | Significance                       |
|-----------------|---------------|------------------|------------------|---------------------------------------------|------------------------------------|
| 1<br>(3,000 ms) | High $\sigma$ | 1.57 $\pm$ 0.27s | 1.61 $\pm$ 0.27s | $\mu \uparrow \rightarrow RT \downarrow$    | $F_{(1,25)} = 31$ ,<br>$P < 0.001$ |
|                 | Low $\sigma$  | 1.62 $\pm$ 0.27s | 1.67 $\pm$ 0.28s | $\sigma \uparrow \rightarrow RT \downarrow$ | $F_{(1,25)} = 33$ ,<br>$P < 0.001$ |
| 2<br>(5,000 ms) | High $\sigma$ | 2.05 $\pm$ 0.55s | 2.1 $\pm$ 0.60s  | $\mu \uparrow \rightarrow RT \downarrow$    | $F_{(1,27)} = 19$ ,<br>$P < 0.001$ |
|                 | Low $\sigma$  | 2.11 $\pm$ 0.58s | 2.18 $\pm$ 0.62s | $\sigma \uparrow \rightarrow RT \downarrow$ | $F_{(1,27)} = 25$ ,<br>$P < 0.001$ |

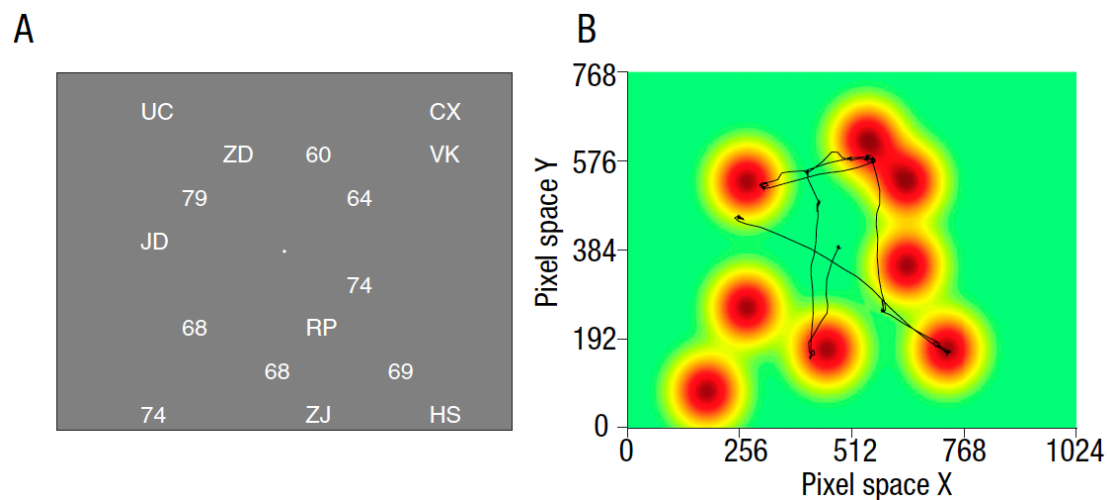
Mean reaction times (RTs) and SD are shown per condition. RTs decreased on high  $\mu$  trials compared to low  $\mu$  trials. RTs decreased on high  $\sigma$  trials compared with low  $\sigma$  trials, which is in contrast with previous studies of feature-averaging using colour and shape (Michael, de Gardelle, Nevado-Holgado, & Summerfield, 2015). No significant interaction effects were found for either accuracy or RT.

#### 4.3.2 Gaze Trajectories.

By tracking gaze during task performance, we were able to ask how the information sampled with the eyes related to decisions. More specifically, we were interested in (i) whether participants sampled by preference those numbers that were inlying or outlying (e.g. 'robust sampling') and (ii) whether participants sampled by preference those numbers that confirmed their current hypothesis about the stimulus (e.g. a confirmation bias).

### 4.3.3 Landscape Approach Allows Estimation of Processed Information

As described above, we developed an approach to analyzing eye tracking data (e.g. the landscape analysis, chapter 2) that estimates the information that is available to the participant at each point in time (the “Quality of the trajectory,” or  $QT_t$ ) across each trial  $t$  (see **Fig. 4.6**). Remember that the landscape analysis converts the stimulus array (**Fig. 4.6** panel A) on each trial to a topography of to-be-sampled decision values spanning the  $1024 \times 768$  pixel space constituting the screen (**Fig. 4.6** panel B).



**Fig 4.6 Example of the landscape analysis and the Quality Trajectory (QT) for a particular trial.** The landscape analysis converts the stimulus array (A) on each trial to a topography of to be sampled decision values spanning the  $1024 \times 768$  pixel space constituting the screen (B).

The landscapes for each trial are created by convolving each stimulus location (i.e., the x and y coordinates of each number in pixel space) with a bivariate Gaussian basis function. Each Gaussian has an SD  $s$  and height  $g$ , with the latter reflecting some information conveyed by the stimulus. In the simplest case (panel B),  $g = 1$  and  $s$  is a constant (e.g., 50 pixels). An example of such a landscape is shown in B (for the corresponding stimulus array shown in panel A). The estimated quality trajectory, or  $QT_t$ , which represents the total information available to the participant at any time point  $t$ , can thus be quantified as the landscape value at the current eye position  $(x, y)$  at that time point (black line in B). Across each trial, this provides a vector of values that indicate whether the participant is looking directly at a stimulus ( $QT_t \sim 1$ , red shaded area) or away from a stimulus ( $QT_t \sim 0$ , green shaded area) at each point in time following the array onset. The parameter  $s$  of the Gaussian can be seen as a proxy of how much information participants can harvest around their fixation. Thus, an increasing value of  $s$  results in an increase in size of the red 'information harvesting' blobs on panel B; while decreasing the value of  $s$  results in a decrease of the size of these blobs. As gaze shifts away from the center of the information-harvesting spotlight the amount of information decreases steadily from 1 (strongest shade of red) to 0 (green).

This approach is equivalent to conceiving of the gaze as a “soft” (or smooth) information-processing spotlight moving across the screen. Importantly, the width of the spotlight was estimated independently of further analyses under

the assumption that participants' gaze maximized the acquisition of total sensory evidence. The width of the information-processing spotlight is determined by searching for values that maximized "quality trajectory" (QT)—that is, the extent to which items were viewed irrespective of their task-relevant values. To do this, we exhaustively computed the QT across a space of possible parameter values (range 0-100) for  $s$  and identified the value of  $s$  that maximised the QT (**Fig. 4.7**). In other words, we made the assumption that participants move their eyes to locations that maximise the available information (QT), irrespective of the decision values themselves. Note that an excessively narrow spotlight would lead to items even close to the point of gaze being neglected, whereas an excessively broad spotlight would lead to large portions of uninformative space being sampled, reducing the overall level of information (QT) obtained. The aperture size  $s$  that maximises QT is the one that maximises this trade-off. The resulting function was found to be concave with a peak at 50 pixels (**Fig. 4.7**). The QT-maximising SD of the Gaussian for the entire cohort (50 pixels) was used for subsequent analyses. Crucially, this approach thus defines the width of the information-processing spotlight within the same task but in a way that is independent of all key analyses described in this chapter.

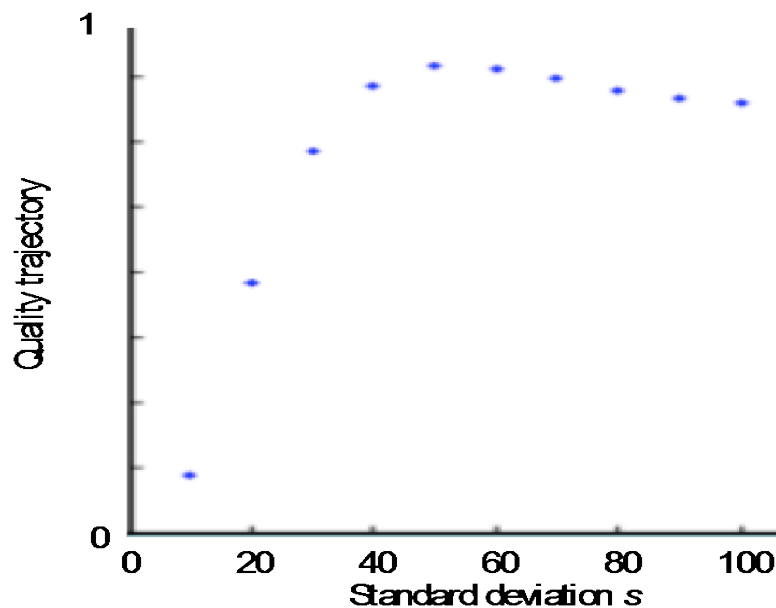


Fig 4.7 Quality trajectory as a function of the width of the information processing spotlight ( $s$ ). The QT-maximising SD of the Gaussian for the entire cohort was found to be at a value of  $s = 50$ .

#### 4.3.4 Landscape Analysis Soft vs. Hard approach

Importantly, this landscape analysis approach allows the researcher to build in a variety of a priori assumptions. For example, in the methodology described above we are assuming that the quality or the amount of information that is available to a participant directly depends on the distance between the fixation and the stimulus location. Thus, a fixation that is slightly distant from the stimulus location (e.g. a more yellow shade in **Fig. 4.6**) will yield a lower quality of information than a fixation that is closer to the center of the stimulus (e.g. a more red shade). Interestingly, this assumption has received independent support from neural recordings in the medial orbitofrontal cortex of the macaque, where neurons code for stimulus value (e.g. 3 vs. 1 drop of

juice), but do so in a fashion that is multiplicatively scaled by the distance between the item and the currently-fixated location (McGinty, Rangel, & Newsome, 2016).

However, it is possible to define a "hard" spotlight approach - that is, an aperture of fixed width with no spatial smoothing. Thus, any fixation that would fall within an information processing blob yields equal information, while any fixation falling outside of this boundary yields no information at all (e.g. a hard filter approach would result in no colour shading in **Fig. 4.6**). The distinction between a soft and hard spotlight approach is particularly important in relation to the earlier discussion about retinal sensitivity (chapter 2): sampling policies are not always well predicted by mapping retinal sensitivity based on independent data from a task with different constraints (Morvan & Maloney, 2012). Moreover, since the way in which the retinal sensitivity is utilised is stimulus and task dependent, the ability to test two different a priori assumptions with this approach is highly beneficial. For example, the information that is extracted from the periphery can be different for stimulus modality (e.g. colour or shape), as well as other dimensions such as numbers versus squircles (J. Baldwin, Burleigh, Pepperell, & Ruta, 2016; Hansen, Pracejus, & Gegenfurtner, 2009). Thus, when sampling a display of numbers one can imagine that the information is characterised by an 'all or nothing' approach: you either know the identity of the number or you have no information at all. While if you are sampling a display of squircles, you might have more detailed information about the shape of nearby stimuli but not

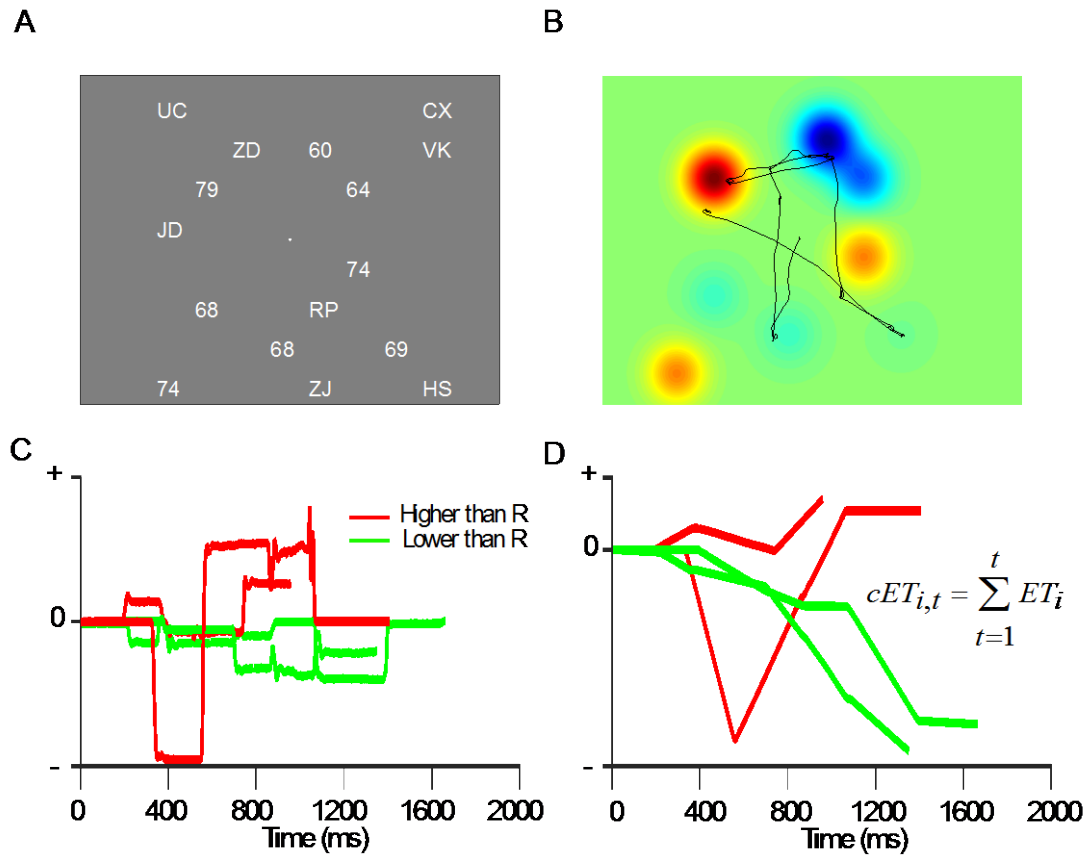
about their colour (J. Baldwin et al., 2016). The distinction between a hard and soft spotlight approach thus allows the researcher to parametrically vary the quality and quantity of the available information for each fixation. Furthermore, the estimation of the quality and quantity of the available information can be estimated in a stimulus and task specific way. Crucially, using the hard spotlight approach for the analyses in this chapter yielded qualitatively identical results to the soft spotlight approach.

#### **4.3.5 Estimating Both Quantity and Quality of Processed Information**

So far we have only demonstrated the use of the landscape approach for estimating the width of the information processing spotlight. However, the landscape approach is not solely a tool for estimating how much information people can harvest around their fixations. Remember that the bivariate Gaussians situated on each stimulus location are not only characterised by their width  $s$  (e.g. the width of the information processing spotlight) but also by their height  $g$ . When estimating the QT, the height was kept constant at 1 and was the same for each stimulus. However, the height can be set to reflect any quantity of the stimulus that the researcher is interested in. As researchers we are not only interested in whether information is processed at a particular location, but we are also interested in the quality and quantity of the processed information.

Thus, in a subsequent step, we used a variant of this method to obtain real-

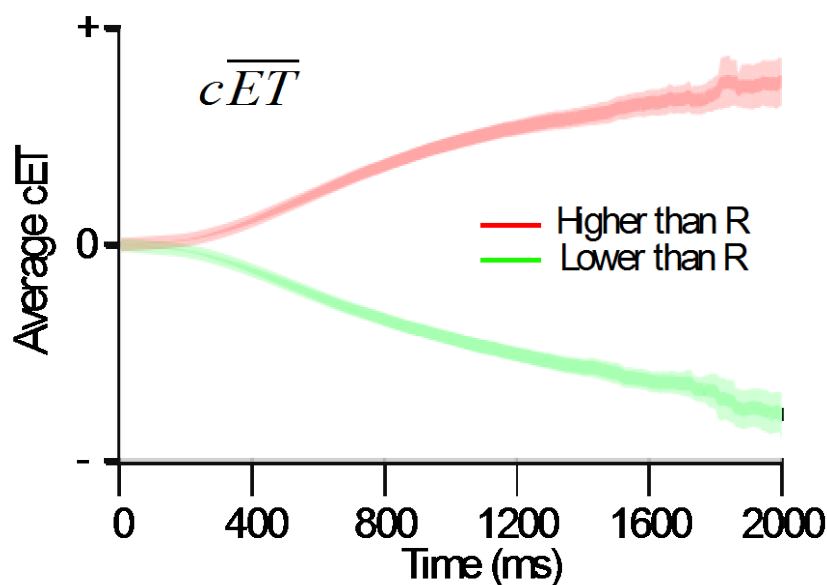
time measures of the decision information available at each point in the trial, a quantity we call ET. In order to quantify the cumulative decision evidence, the height of the Gaussian  $g_k$  associated with each item  $k$  is set to  $V_k - R$ , where  $V_k$  is the value (number) of each item and  $R$  is the reference value for that block. This results in a smooth topography of decision values on each trial, with peaks (red) located at numbers higher than the reference value and troughs (blue) located at numbers lower than the reference value (**Fig. 4.8** panel B). Positive red peaks thus favour the decision “higher than reference,” negative blue troughs favour the decision “lower than reference” (**Fig. 4.8** panel B). The amplitude of peaks (red) and troughs (blue) depended linearly on the difference between each number and the reference value. Importantly, these values were always inverted for participants in the lower framing group, such that positive deflections in each landscape were always in the frame of reference of the choice, not the absolute number. Time point by time point measures of decision information are then obtained by defining gaze trajectory as the temporal sequence of landscape values falling under the observer’s gaze at each time point  $t$ . Thus, the evidence trajectory (ET) constituted the time-series of values obtained as the eyes moved across the stimulus array on each trial (**Fig. 4.8** panel C). The cumulative evidence trajectory (cET) reflects the total amount of evidence collected over time (**Fig. 4.8** panel D).



**Fig 4.8 Example of the landscape analysis and the Evidence Trajectory (ET) for a particular trial.** The landscape analysis converts the stimulus array (A) on each trial to a topography of to be sampled decision values (B). Note that unlike on Fig. 4.4., the height of the Gaussian  $g$  is no longer the same for each number location. Instead, the height  $g$  for each number is determined by the signed difference between each number  $V_k$  and the reference for that block  $R$  ( $V_k - R$ ). This results in the landscape on panel B, with positive red peaks indicating evidence in favour of the decision “higher than reference” and negative blue troughs in favour of the decision “lower than reference”. The integral of gaze (black line) under these landscapes is then used to calculate the ET on panel C. (C) ET over time for four example trials (lines). (D) The cumulative ET over time ( $t$ ) for these four example trials ( $i$ ).

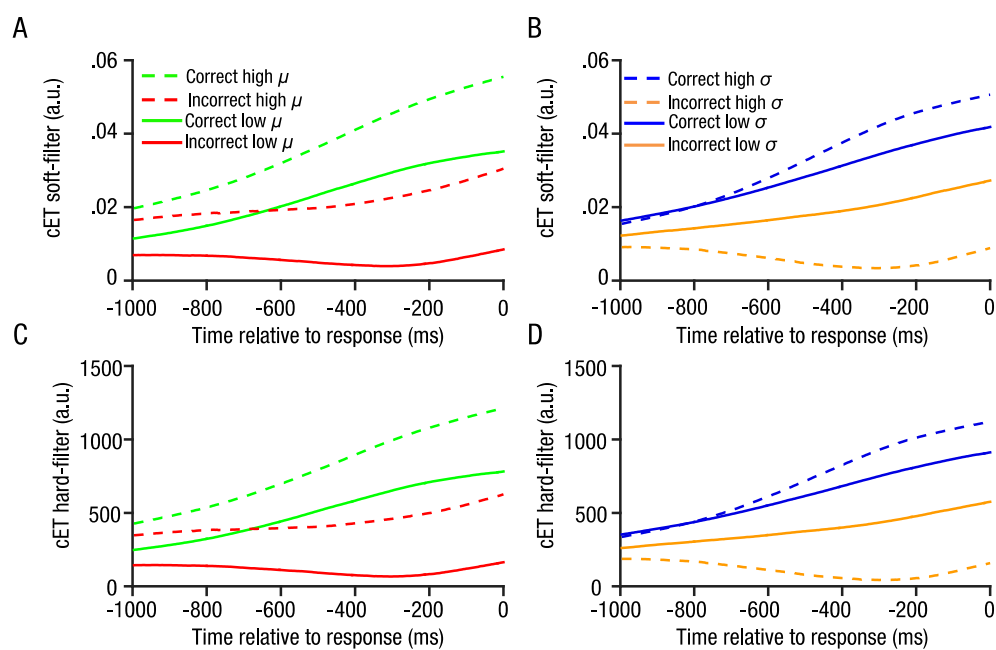
#### 4.3.6 Decisions Depend on Gaze Location.

Using this approach, we first compared cumulative ET on trials where participants responded “higher” and those where they responded lower, correcting for differences in available information on “high” and “low” trials. As expected, choices depended strongly on the decision information sampled, with average cumulative ET building up steadily until the response: participants responded lower when the sampled evidence was lower than the reference and higher when the sampled evidence was higher than the reference (**Fig. 4.9**). In other words, consistent with previous findings, choices were strongly affected by where participants looked (Krajbich et al., 2010; Krajbich & Rangel, 2011; Manohar & Husain, 2013).



**Fig 4.9 Cumulative ET time series averaged across trials.** Cumulative ET is on average positive when participants responded “higher than reference” (red), and negative when participants responded “lower than reference” (green).

**Fig. 4.10** shows the cumulative ET locked to the response as a function of evidence strength ( $\mu$ , panel A) and evidence reliability ( $\sigma$ , panel B). cET is higher on correct trials compared to incorrect trials. Interestingly, cET does not converge to a similar decision threshold for different values of evidence strength and reliability. This is addressed later on in the chapter.



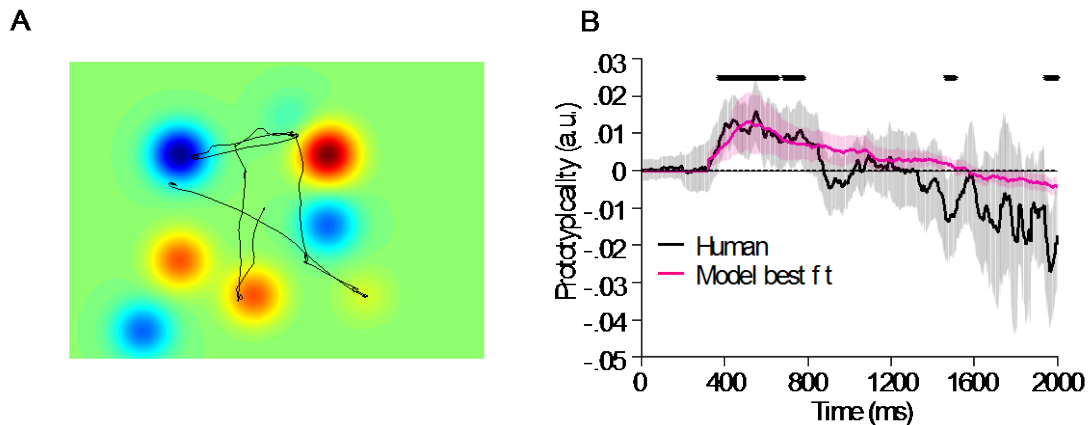
**Fig 4.10 Response locked cumulative ET time series as a function of evidence strength and evidence reliability.** Cumulative ET is on average positive when participants responded “higher than reference” (red), and negative when participants responded “lower than reference” (green). The plots are again split by accuracy, mean ( $\mu$ , A – C) and variance ( $\sigma$ , B – D) of the numbers on each trial. Qualitatively identical ETs were obtained under soft (A and B) and hard (C and D) information harvesting spotlights. A hard information spotlight assumes that information within that spotlight that is

proximal and distal to the point of gaze contributes equally to decisions.

Importantly, using a “hard” information processing spotlight - that is, an aperture of fixed width with no spatial smoothing obtained qualitatively identical results (**Fig. 4.10** panel C & D). All of the analyses presented below were replicated with a hard information processing spotlight.

#### **4.3.7 Gaze Is Directed Preferentially to Inliers Over Outliers.**

Next, we used a similar landscape approach to determine the extent to which participants' gaze favoured items that were inlying or outlying (numerically closer or further from the reference). we generated a new set of landscapes in which the topography on each trial was determined by the rank of the absolute value of each element (number) within the stimulus array, so that more inlying elements had higher values and more outlying elements had lower values (**Fig. 4.11.A.**; similar results were obtained when this analysis was carried out using the absolute values instead of the ranks). The integral of the eye movement trajectory under these landscapes, which we call the “prototypicality trajectory” (PT), thus encoded participants' time-varying preference to fixate those elements that were numerically similar (inliers; denoted by positive values) or dissimilar (outliers; negative values) to the reference.



**Fig 4.11 Prototypicality trajectories** . (A) Prototypicality landscape for the example trial from Fig. 4.1. This landscape is created by setting the height of the Gaussian ( $g_k$ ) associated with each number location equal to the prototypicality value  $P_k = 4.5 - \text{tiedrank}(|V_k - R|)$ . (B) PT over time. Inlying items were viewed by preference from ~400 – 800 ms post stimulus presentation, indicated by a significantly positive PT value (black). Black bars mark significant clusters ( $P < 0.05$ , corrected for multiple comparisons). The magenta line shows the predicted values for PT over time for the best fitting sampling strategy, as determined per subject, of the model (see *Computational simulations* below).

We found that inlying items were sampled preferentially, from 400 to 800 ms post stimulus presentation, indicated by a significantly positive PT value (**Fig. 4.11.B.**, black line and shading). This result survived correction for multiple comparisons across time points using a well-validated cluster-based approach (Maris & Oostenveld, 2007).

Next, we turned to a more conventional analysis technique and we evaluated this preference for inliers using a different approach. First, we characterised

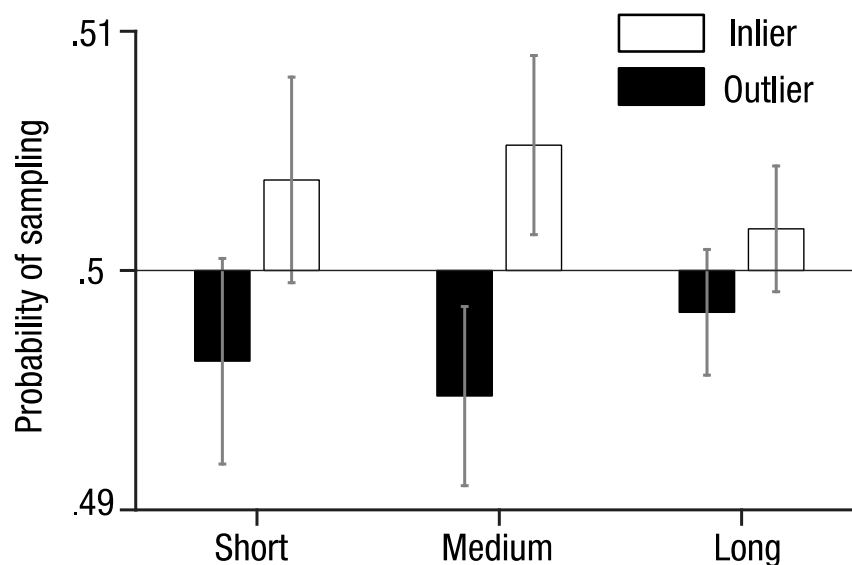
discrete fixations larger than 100 ms using Eyelink pre-set criteria (velocity threshold of 30, acceleration threshold of 8000, and a saccade motion threshold of 0.1). On average, participants made 7.9 (SD= 3.7; exp.1) and 5.9 (SD = 2.2; exp.2) fixations with a mean dwell time of 191.7ms (SD = 56.2ms; exp.1) and 197.2ms (SD = 30ms; exp.2).

Next, on each fixation, we defined the item being sampled as the item that was closest to the current gaze location. Numbers on each trial were sorted in ascending order yielding a 1-8 rank ordering. This allowed the calculation of the probability of gaze moving from an item of one rank (i.e., ordinal position in the array) to another. For example, the probability of gaze moving from rank 3 items to rank 4 items could be calculated. Remember that the four centrally ranked elements (e.g. ranks 3-6) were classified as inliers, while the other four elements (e.g. ranks 1-2 and 7-8) were classified as outliers.

To avoid any biases that may arise from non independence of nearby items (e.g. if by chance inliers would be grouped closer together than outliers over the entire experiment, in pixel space), we constructed a null distribution specific to our experiment by randomly permuting item location assignments on each trial and repeating the analysis 1,000 times. Next, the difference between gaze transition probabilities (e.g. the probability of making a saccade transition from a rank X to rank Y number) in our dataset and under the null distribution is computed. The main effect of preference to sample inliers (sampling policy) was significant [ $F(1, 53) = 8.58, P < 0.01$ ]. Thus, subjects were more likely to sample elements of ranks 3-6 compared to elements of ranks 1-2-7-8. Demonstrating this preference to sample inliers via a fixation based approach shows that the PT effect is not due to for example increased

perceptual difficulty in estimating whether inlying items are above or below the reference.

However, one might wonder how participants manage to know the identity of items in the periphery and whether this is possible at different eccentricities. In other words, does the distance between the currently and subsequently fixated items modulate this effect? To test this, we calculated the probability for a single fixation to go to an inlier/outlier, split for short, medium, and long Euclidean distance in pixels on the screen between the currently and subsequently fixated items (**Fig. 4.12**).



**Fig 4.12. Preference to fixate inlying items: the probability to sample inliers/outliers on each fixation in the real human data in comparison with the null distribution.** Probabilities higher than 50% show an increase of sampling relative to the null distribution (and vice versa). The main effect of preference to sample inliers was significant [ $F_{(1,53)} = 8.58, P < 0.01$ ]. The medium Euclidean distance fixations largely drove this effect [ $t_{(53)} = 2.74, P < 0.01$ ], whereas the effect was not significant for short and long distances ( $P > 0.05$ ). The main effect of distance [ $F_{(1.9, 100.67)} = 1.71, P > 0.05$ ] and its interaction with sampling policy [ $F_{(1.89, 100.38)} = 1.168, P > 0.05$ ] were not significant. Bars are shown with 95% confidence intervals across the cohort.

As expected, this effect was largely driven by fixations of medium distance [ $t_{(53)} = 2.74, P < 0.01$ ] in comparison with short and long distances (both  $P > 0.05$ ). The reason why this effect is not present for fixations of a short distance is because these fixations consist almost exclusively of re-fixations of

the same item. The main effect of distance [ $F(1.9, 100.67) = 1.71, P > 0.05$ ] and its interaction with sampling policy [ $F(1.89, 100.38) = 1.168, P > 0.05$ ] were not significant.

These findings obtained using discrete fixations rather than the landscape approach demonstrate that observers exhibited a tendency to shift their gaze toward inliers rather than merely dwelling for longer while viewing these items. Unsurprisingly, these effects diminished for long Euclidean distances between previous and current fixations, as the decision information associated with distant items cannot be evaluated through parafoveal processing. Thus, the data suggest that participants' parafoveal processing is not unbounded but occurs within a fixed range, and the sampling bias (moving to inlying information) is only possible for items falling within this range.

#### **4.3.8 Humans show Robust Averaging of Numbers**

Next, we investigated whether participants' decisions show evidence of robust averaging—that is, a decision policy in which outliers carry less weight than inliers.

Probit regression was used to predict human choices (higher or lower than the reference) as a function of three sets of predictors: (i) the decision values of each number ranked from lowest to highest, (ii) QT for each number  $QT_k$  (e.g.

the time spent sampling element  $k$ ; once again, with  $k$  ranked from lowest to highest), and (iii) the interaction between these two quantities, all entered as competitive predictors in the predictor matrix (**Fig. 4.13A-C**). The QT for each number  $QT_k$  was computed as the sum of the integral under an element-specific QT landscape—that is, a landscape constructed using each element in isolation (with the height of the Gaussian set to 1)—up to the response time.  $QT_k$  is then proportional to the time spent sampling element  $k$ . The regression model we used was as follows, where in each case:

$$p(high) = \phi \left[ b + \sum_{k=1}^8 w1_k \cdot (V_k - R) + \sum_{k=1}^8 w2_{k,t} \cdot TQ_{k,t} + \sum_{k=1}^8 w3_{k,t} \cdot (TQ_{k,t} \cdot (V_k - R)) \right]$$

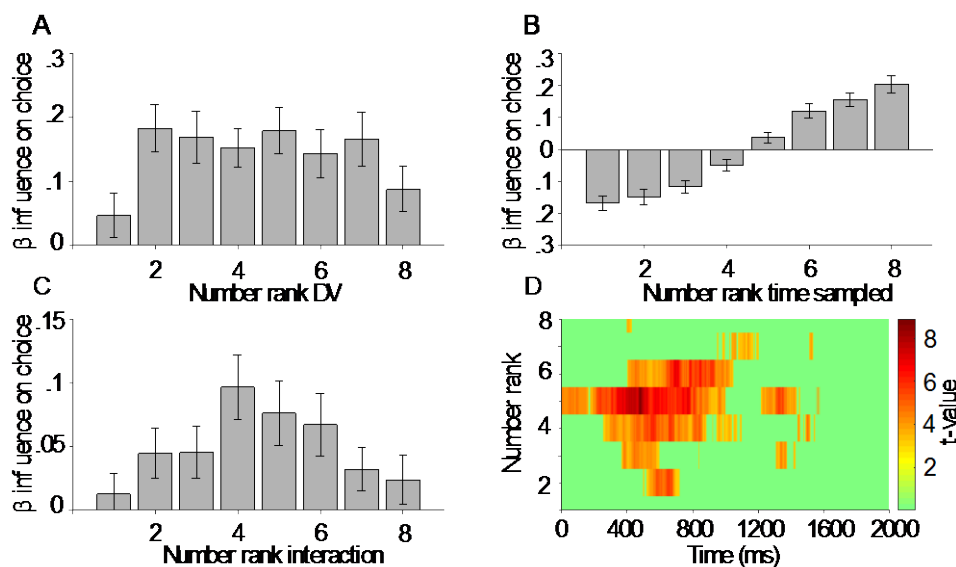
where  $\phi$  denotes the cumulative normal distribution,  $k$  the sample number, and  $t$  the time point of the QT vector. For the data shown in Fig. 14 A–C, we dropped the indexing by  $t$  and used the sum of QT for all time points up to the response.

The weighting function plotted in **Fig. 4.13.A** shows the main effect of stimulus—that is, the weight given to each number as a function of its rank within the numerical array. The inverted-u shape indicates that outliers are downweighted. This observation is qualified by a statistical comparison between the weight given to inlying ranks [3, 4, 5, 6] and outlying ranks [1, 2,

7, 8],  $t(53) = 3.72$ ,  $P < 0.001$ . This finding of robust averaging with numbers complements earlier reports with colour, shape, and facial expression, confirming that outliers carry less influence during perceptual averaging across domains and task settings (De Gardelle & Summerfield, 2011; Michael et al., 2013).

#### **4.3.9 Linking Sampling Policy and Human Decisions**

Next, we return to the landscape analysis and show how the main effect of sampled information influenced choices, finding that the resulting estimates grew with rank: sampling low numbers was associated with a propensity to respond low, and sampling high numbers was associated with responding high (**Fig. 4.13.B**).



**Fig 4.13 Robust averaging of decision information.** (A) Weights associated with each item rank, calculated by sorting the numbers associated with each item, and using probit regression to predict participants' choices. The inverted U shape over ranks shows evidence for robust averaging (larger weights for inliers). (B) Weights associated with the time spent sampling each number rank. More time spent sampling low-ranked numbers led to a tendency to respond lower than higher and vice versa. These weights are negative for the first four number ranks because these number ranks contain on average mostly numbers lower than the reference. Thus, as you spend more time sampling the number ranks that are more likely to contain numbers lower than the reference, you are more likely to respond 'low' (e.g. a negative coefficient). Similarly, these weights are positive for the final four number ranks because these number ranks contain on average mostly numbers higher than the reference. (C) Weights associated with the interaction between number rank and the time spent sampling these numbers. The inverted U shape over ranks shows evidence that robust averaging is mediated by gaze duration. (D) Robust averaging due to the sampling policy as a function of

time. Ranks are shown on the  $y$  axis and  $t$  values are shown by the green – red colour scale. Robust averaging due to preferential sampling of inliers was maximal during the 400 – 800 ms time frame after stimulus onset (larger weights for inliers). Bars on A – C are shown with 95% confidence intervals.

This is expected from the ET analyses above (**Fig. 4.9**), which show that the information sampled was strongly predictive of the response (high vs. low). Critically, the third predictor (interaction term) in the regression allowed us to assess how the weight given to each stimulus rank was modulated by gaze duration, over and above the influence of stimulus value. Ranks where this term differs positively from zero are those where directing gaze to the relevant numbers amplifies their weight multiplicatively. For example, if gaze were merely a filter that acted equally for all numbers irrespective of their rank, we would see weights that are equal for all ranks. However, this is not what is observed (**Fig. 4.13.C**). The time spent sampling inlying items enhances the effect that they will have on choice (compared with outlying items) above and beyond what would be predicted simply by their values alone [inlying vs. outlying:  $t(53) = 5.10$ ,  $P < 0.001$ ].

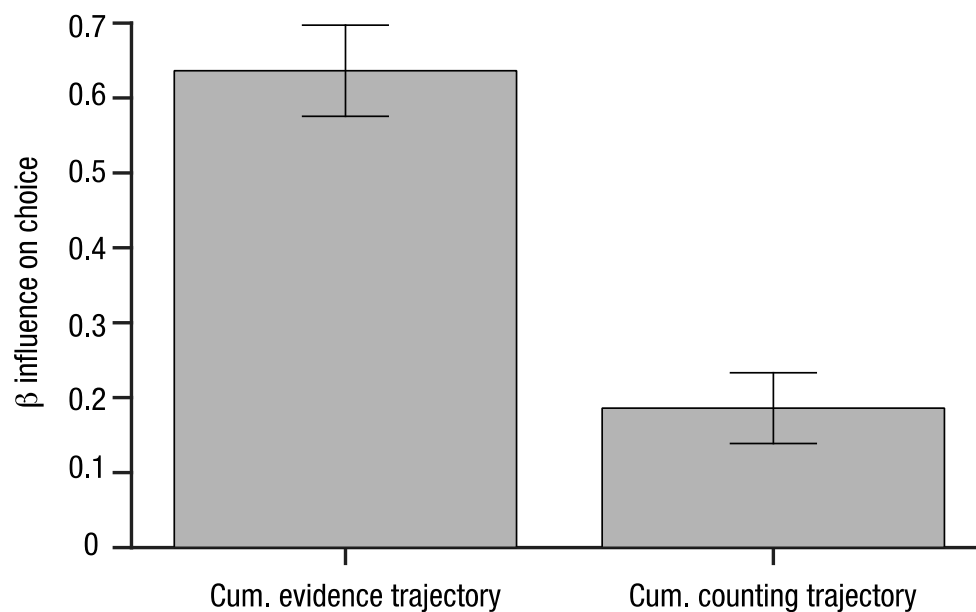
Note that this finding is not simply due to a logarithmic (Weber-like) compression of the number line, because it occurred both for those participants that were asked whether the average was higher than the reference [ $t(28) = 4.99$ ,  $P < 0.001$ ] and those that reported whether the average was lower than the reference [ $t(24) = 2.18$ ,  $P = 0.039$ ], a framing

variable that we counterbalanced over participants. This result also persisted when we used the average (rather than total) ET calculated over the trial, suggesting that it is not simply due to different sampling policies being deployed on trials with short and long RTs.

Given that the prototypicality analysis (above) showed that inlying items were preferentially viewed 400–800 ms poststimulus, we examined whether this downweighting, due to preferential sampling, is also most apparent in this time period. To do this, we conducted the previous regression for each sampling time point individually and tested which of these resulting coefficients (for the Gaze  $\times$  Stimulus interaction) were significantly different from zero. In **Fig. 4.13.D**, these are plotted ( $t$  values  $> 3$ ) on a heat map, showing that this downweighting of outliers effect is also maximally depended on gaze during this 400–800 ms timeframe. These results suggest that robust averaging is due at least in part to the sampling policy with which participants actively sample inlying rather than “outlying” information. Together, **Fig. 4.13 A and C** show that (i) inlying stimuli (i.e., those close to the reference) have a stronger influence on choice than outlying stimuli (i.e., those that are much larger or much smaller than the reference) and (ii) that this effect is multiplicatively enhanced by gaze. Thus, any tendency that inlying stimuli might have had to drive choices more strongly is exaggerated when they are fixated, but this is not true for outlying stimuli. In other words, human observers are not only “robust averagers” but also “robust samplers”—selectively viewing inlying category information and thereby ensuring that this

information is a more potent driver of choice.

In further control analyses, we ruled out the possibility that this effect was due to a “counting” strategy, in which all items above or below the reference contributed equally to decisions (**Fig. 4.14**). In past work, our laboratory has shown that down-weighting of outliers can be explained by a nonlinear (sigmoidal) transformation of the decision information before averaging, over and above any counting (de Gardelle & Summerfield, 2011). To test this alternative explanation in the current study, we generated new cumulative counting landscapes in which the decision value is the same positive value (+1) for all values above the reference and the same negative value (-1) for all values below. We then used both the cumulative ET values and the cumulative counting ET values from each trial (averaged up to the response) to predict human choices, in a competitive regression. Predictors were standardized (z-scored) to allow comparison. The resulting coefficients (betas) for these two predictors are shown in **Fig. 4.14**. Both the cumulative ET and the cumulative counting ET predicted choice [ $t_{(53)} = 20.53$ ,  $P < 0.001$  and  $t_{(53)} = 7.72$ ,  $P < 0.001$ , respectively]. This is both consistent with previous work (in which counting accounted for some variance in the data) and also with the results described in **Fig. 4.13.A**, which show that there is a residual “squashing” that occurs beyond the sampling stage (i.e. over and above the Gaze x Decision Information weights, we still see robust averaging). Importantly, the cumulative ET was a significantly better predictor of choice than the cumulative counting ET [ $t_{(53)} = 9.08$ ,  $P < 0.001$ ].



**Fig 4.14 Counting strategy as an alternative explanation.** Both the cumulative ET and the cumulative counting ET predicted choice. Importantly, the cumulative ET was a significantly better predictor of choice than the cumulative counting ET. Bars are shown with 95% confidence intervals. Thus, robust sampling cannot be explained by an alternative counting strategy, in which participants only keep track of the number of elements higher/lower than the reference

#### 4.3.10 Computational Simulations

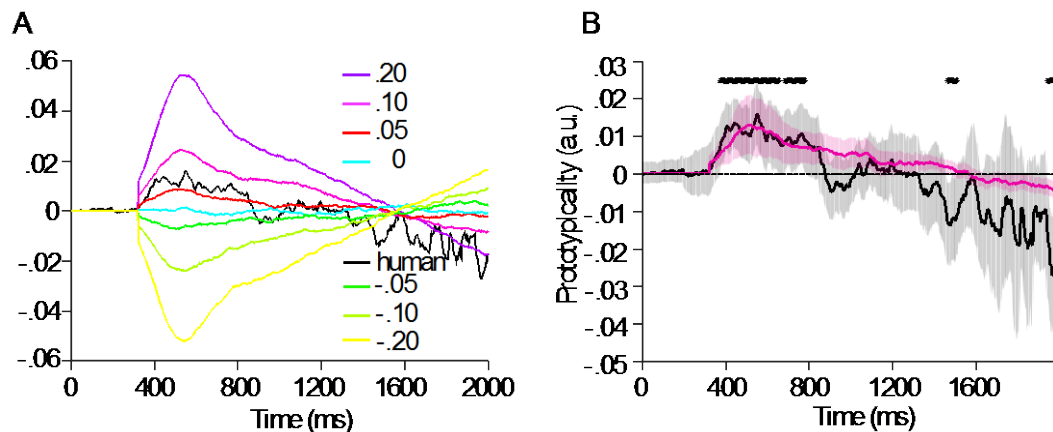
Next, to validate the approach, I turned to computational simulations. The reasoning for the use of this computational approach is two-fold. First, the aim is to simulate different degrees of robust sampling (i.e. sampling policy) to estimate how the ET and PT should vary under differing sampling policies. Moreover, these simulated sampling policies will allow us to quantify its

effects on performance for an ideal observer model. Secondly, a computational approach will result in an additional piece of evidence for ET as an adequate measurement of evidence accumulation.

To do this, distributions of decision information (i.e., numbers) identical to those viewed by human participants (i.e., same  $\mu$  and  $\sigma$ ) were used. Next, simulated ET and PT vectors consisting of alternating periods of fixation and eye movements are constructed. These simulated ET and PT vectors are prefixed by a delay period resulting in a simulated sampling via gaze process similar to the one exhibited by the human participants. The duration of all of these periods (initial delay, fixation, eye movement) was sampled from the distributions exhibited by human participants. During simulated fixations, simulated ET was set to the decision values  $V_k - R$  (where  $V_k$  is the numerical value of item  $k$ ), and during delay and eye movements, ET values were set to zero. Simulated PT was constructed in the same manner by setting PT during a fixation of stimulus  $k$  to its value ranked within the array, as for the human analyses. Model choices occurred when the cumulative simulated ET reached a fixed threshold, as in standard integration-to-bound models (Ratcliff & McKoon, 2008). The sign of the cumulative ET at this fixed threshold determined model choice and thus accuracy. To ensure that human PT and model PT values were on the same scale, we multiplied the model PT with the average QT for the corresponding participant.

This framework allows a manipulation of the model sampling policy—that is,

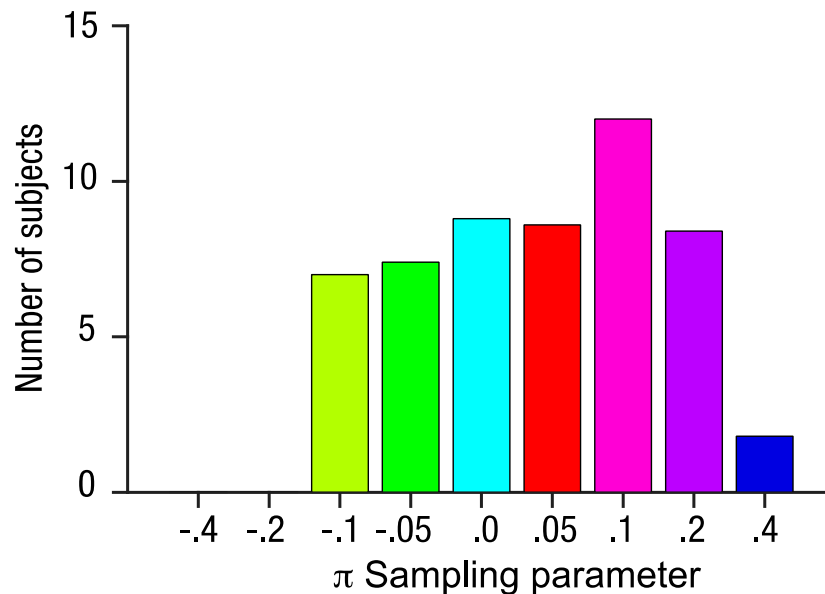
to simulate the order in which the array items were sampled. In this case, we were interested in simulating the degree of robust sampling. We implemented this sampling policy via a single parameter  $\pi$  that dictated whether the observer preferred to view inliers before outliers ( $\pi > 0$ ), outliers before inliers ( $\pi < 0$ ), or sampled the elements in random order ( $\pi \sim 0$ ). The parameter  $\pi$  encodes the (positive or negative) correlation between the observed sampling order and the most “prototypical” ordering—that is, from most inlying to most outlying. The parameter  $\pi$  can thus range from -1 to +1. In **Fig. 4.15.A**, we plot simulated PT for different values of  $\pi$ . The best fitting parameterization was determined per subject by calculating the mean squared error (MSE) between human and simulated PT for each value of  $\pi$ . The model included two extra parameters that were fixed and not fit per subject: a threshold parameter indicating the decision threshold and a parameter for the standard deviation of zero mean Gaussian noise that detailed the amount of noise corrupting the DV.



**Fig 4.15 PTs for the human and model data.** (A) Simulated PTs for different sampling strategies ( $\pi$ ) in the model. More positive values indicate a preference to sample inlying items first; more negative values indicate a preference to sample outlying items first. The black line corresponds to the average PT over time for the human data, indicating that humans exhibit a sampling policy that is best predicted by a positive  $\pi$  value. (B) PT over time. Inlying items were viewed by preference from ~400 – 800 ms post stimulus presentation, indicated by a significantly positive PT value (black). Significant clusters ( $P < 0.05$ , corrected for multiple comparisons) are marked by black bars. Shading around the lines signify 95% confidence intervals. The magenta line shows the predicted values for PT over time for the best fitting sampling strategy (determined per subject) of the model.

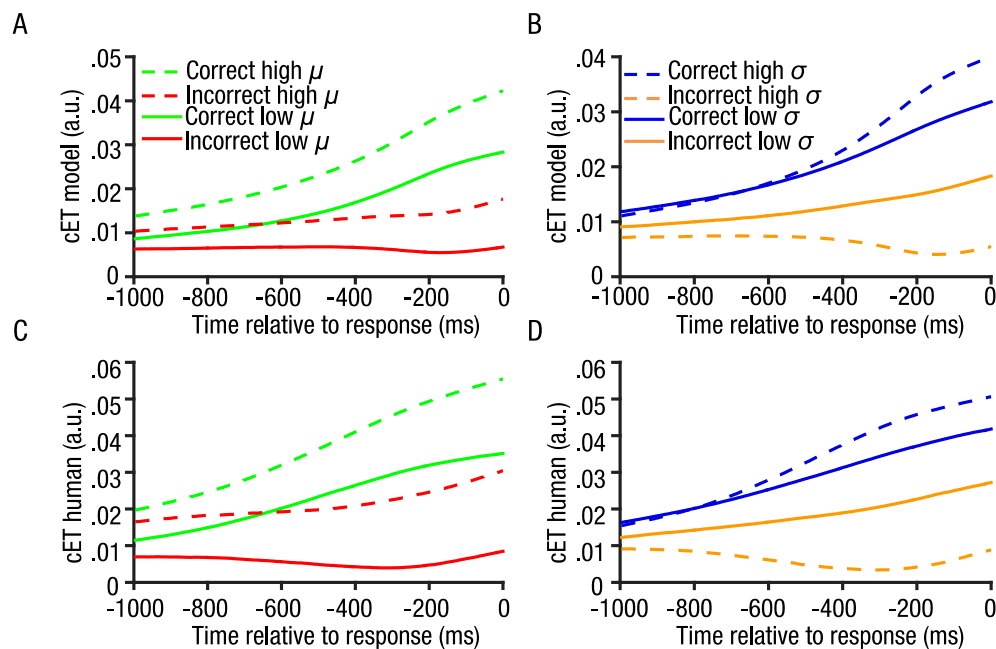
**Fig. 4.15.B** shows the simulated values of PT for the best fitting average  $\pi$  value over subjects (magenta line) plotted against human data. Human PT curves were best accounted for by positive  $\pi$  values, consistent with a policy that involved sampling inliers before outliers and thus with the analyses

described above. **Fig. 4.16** shows the distribution of best fitting values for  $\pi$  over subjects.



**Fig 4.16. Histogram of the best fitting parameter values  $\pi$  from the model, over the human cohort.** Because of random noise in the model, results varied slightly each time the model was fit. Therefore, the results shown were obtained by fitting the model 10 times and averaging the number of subjects. Note that the x axis is an ordinal, not interval scale. The best fitting sampling parameter was significantly more positive than 0 ( $P < 0.05$ ; on average, ~57% of participants had a positive best fitting  $\pi$  – i.e., values of either 0.05, 0.1, 0.2, or 0.4 – whereas 27% had a negative best-fitting  $\pi$  - i.e. values of either -0.05 or -0.1), indicating a tendency to preferentially sample inlying items.

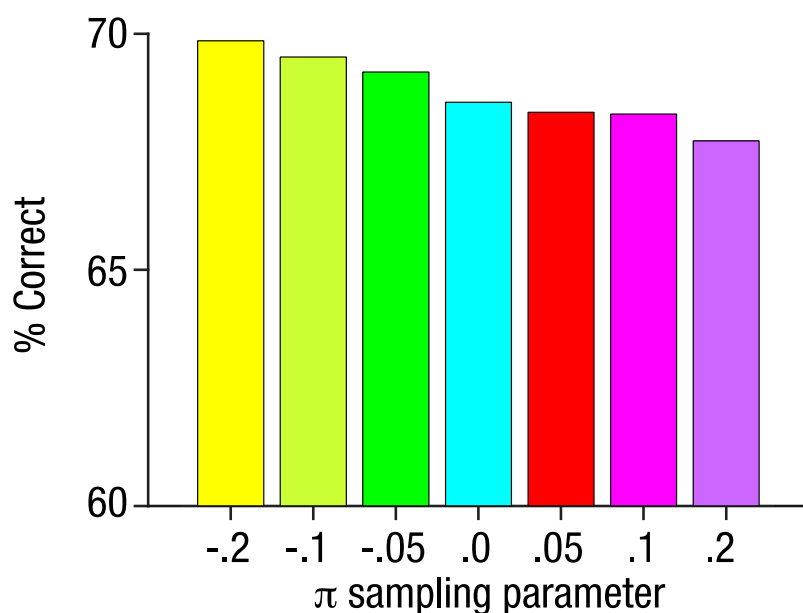
Next, using the best fitting  $\pi$  values for each subject, we returned to the analysis of ET and compared the human and simulated ET values for the different conditions—that is, for different means  $\mu$  and SDs  $\sigma$  of the decision values with respect to the reference. We observed a striking convergence between the human and model ET values (**Fig. 4.17**), providing an independent validation of our best fitting model.



**Fig 4.17 Cumulative ETs from the landscape analysis for the simulated model (A and B) and human data (C and D).** The plots are split by accuracy, mean ( $\mu$ , A – C) and variance ( $\sigma$ , B – D) of the numbers on each trial. The simulated and human cumulative ET show very similar qualitative patterns.

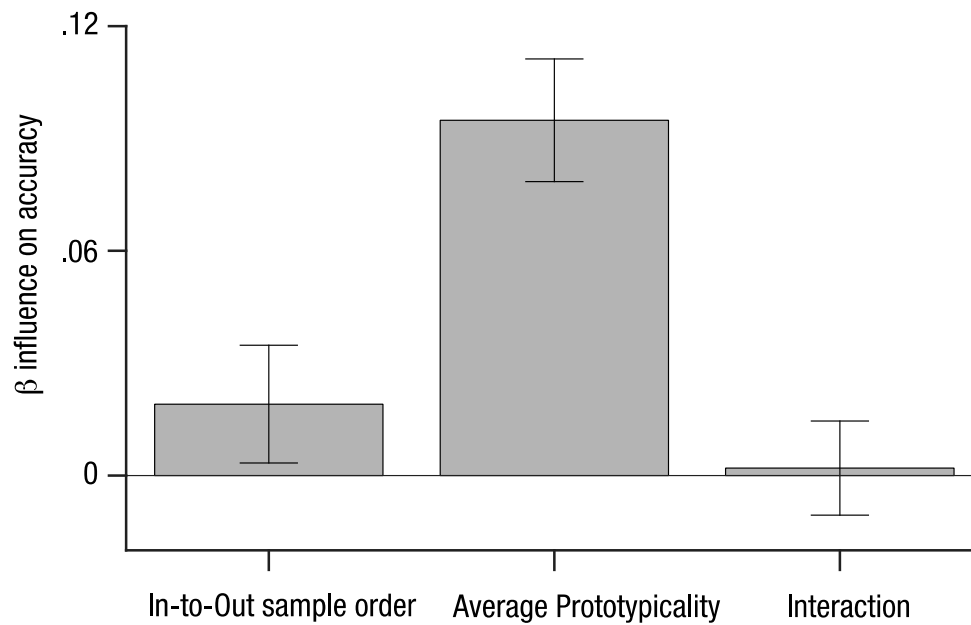
#### 4.3.11 Determinants of Human and Model Performance.

Finally, we tested how model performance (accuracy) varied with  $\pi$ . Interestingly, performance for the model linearly decreased as the sampling policy  $\pi$  changed from a policy to sample outlying information first to a policy to sample inlying information (**Fig. 4.18**). Thus, for an “ideal” model—in which decisions are only limited by the quality of information that is sampled from the screen—a policy that preferentially samples inlying information incurs a model performance cost relative to random sampling.



**Fig 4.18 Accuracy of the computational model under different sampling policies ( $\pi$ ).** Performance is lower for a sampling policy preferentially sampling inliers first (positive  $\pi$ ) compared with a sampling policy preferentially sampling outliers first (negative  $\pi$ ).

To test whether this was the case for human observers, we conducted an additional analysis that evaluated how estimated human values of  $\pi$  predicted accuracy at the single-trial level. We computed single-trial estimates of  $\pi$  by correlating the order in which items were actually sampled with their ranked equivalent and then used logistic regression to estimate how these values predicted choice accuracy (**Fig. 4.19**). Both the sampling policy  $\pi$  [ $t(53) = 2.37$ ,  $P < 0.05$ ] and the average prototypicality [ $t(53) = 11.3$ ,  $P < 0.001$ ] were positive predictors of accuracy, indicating that unlike the model, participants performed better as they exhibited more robust sampling.

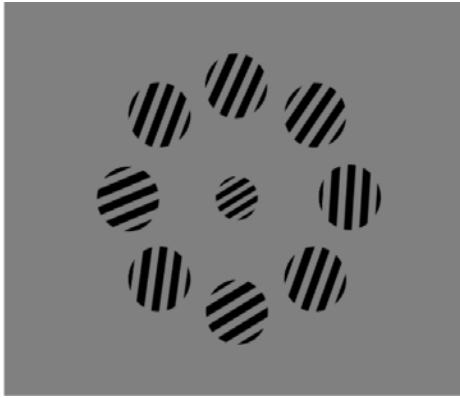


**Fig 4.19 Predicting human accuracy based on the sampling policy (preference for in- vs. outlying elements).** Single trial estimates of  $\pi$  are computed by correlating the order in which items were actually sampled with their ranked equivalent. To quantify the effect of sampling policy on accuracy in our dataset, we regressed this single-trial value of  $\pi$  (given the eye-movement trajectory on each trial) and the average prototypicality against accuracy – that is, we used a probit regression to determine the extent to which these variables predicted whether participants would be correct (1) or not (0). Both the sampling policy [ $t_{(53)} = 2.37$ ,  $P < 0.05$ ] and the average prototypicality [ $t_{(53)} = 11.3$ ,  $P < 0.001$ ] show positive parameter estimates, indicating that participants were more likely to be correct when they pursued a policy of sampling inliers over outliers (positive values of  $\pi$ ).

#### **4.3.12 Mechanisms underlying the effect of the sampling policy on accuracy**

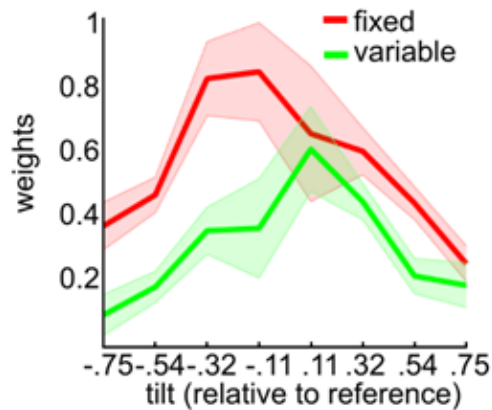
What is the reason for this apparent divide between what is optimal for the model and the participants in this study? Why would humans adopt an integration policy that discards important decision information and why do they perform better doing so?

Recent work by our lab has recently addressed this question (Li et al., 2017). In this paper, we show that robust averaging is indeed suboptimal for a perfect integrator. However, under certain circumstances, robust averaging paradoxically increases performance. In order to investigate the mechanisms underlying the effect of nonlinear weighting of information on accuracy, Li et al. (2017) asked 24 participants to take part in two psychophysical testing sessions. Participants were asked to average the orientation of a circular array of gratings and compare this estimated average to a central reference grating (**Fig. 4.20**).



**Fig 4.20. Demonstration of the circular stimulus array.** Participants had to estimate the average orientation in the array and state whether this average orientation was clockwise or counter clockwise of the central reference grating.

As expected, participants showed evidence of robust averaging: more weight was given to the inlying items in both experimental sessions (a “fixed reference sessions” where the reference orientation remained constant within a block and a “variable reference session” where the reference orientation varied on a trial-by-trial basis; **Fig. 4.21**, inverted-U shape).

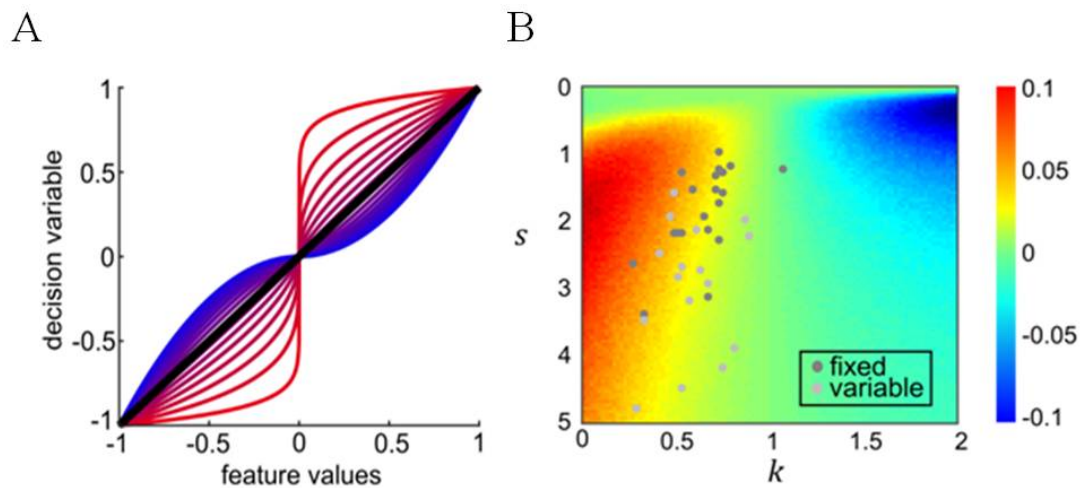


**Fig 4.21 Robust averaging of a circular array of gratings (figure copied from Li et al., 2017 under CC-AL license).** The average regression coefficients associated with the orientations relative to the reference (reference = 0 in radians) are plotted for the session in which the reference was fixed (red) or variable (green). Robust averaging, exhibited by an inverted-U shape, was found both when the reference grating was fixed for an entire block of trials as well as when the reference grating varied from trial to trial.

In order to investigate the relationship between robust averaging and accuracy, a simple mathematical model that involves using a power law transducer to nonlinearly transform the decision value of an item was used. It is important to note that, unlike in the model described previously, robust averaging is not modelled by the order in which elements are sampled. Instead, the model computes a decision value by transforming the feature values ( $X$ ) with a nonlinear function parameterised by an exponent  $k$ .

$$DV = \sum_{i=1}^8 \text{sign}(X_i) \cdot |X_i|^k$$

Using this formula, one can simulate different implementations of robust averaging as a function of  $k$ : from overweighting inliers ( $k < 1$ , blue), to overweighting outliers ( $k > 1$ , red; **Fig. 4.22.A.**). Choice probabilities were generated by passing this DV through a sigmoidal choice function with an inverse-slope  $s$ . This parameter  $s$  can be seen as an estimation of zero-mean Gaussian noise that corrupts the DV at a post averaging stage, termed as "late" noise



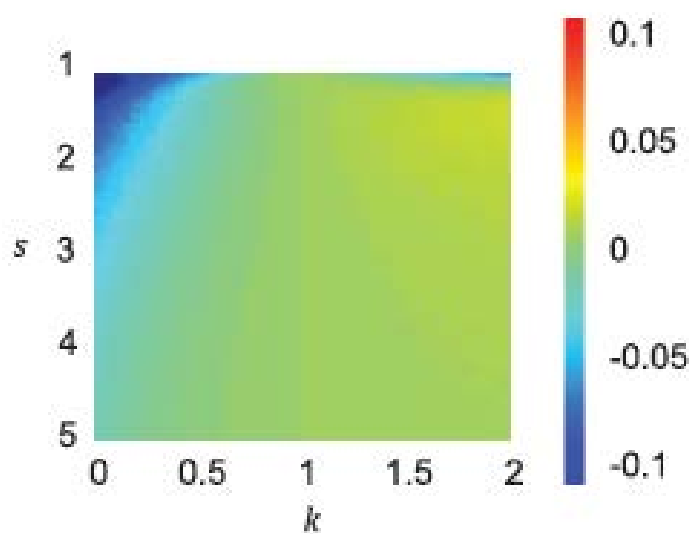
**Fig 4.222** Different implementations of robust averaging (left) and behavioural fits (right). Figure copied from Li et al., 2017 under CC-AL license. (A) The figure shows the transformation of feature values (x axis) to decision variables (y axis), using different parameterisations of  $k$  (coloured lines). The black diagonal line represents when  $k = 1$ , where feature values are linearly transformed into decision variables. Red tinted lines correspond to cases where  $k < 1$ , where inlying feature values are being transformed into greater decision variables (i.e. overweighting of inliers). Blue tinted lines correspond to cases where  $k > 1$ , where inlying feature values are transformed into lower decision variables (i.e. downweighting of inliers). (B) The surface plot shows the simulated model accuracy difference between the model where items are being nonlinearly transformed using the power law transducer and a perfect linear integrator model. The grey dots on the plot refer to the best fitting parameterisation of  $k$  and  $s$  in human subjects.

Model fitting of human trial-by-trial choices resulted in best-fitting values for  $k$  that were lower than 1. Using the best fitting parameterisation of the model,

the model was able to accurately reconstruct the weighting profile and accuracy levels exhibited by the participants, demonstrating that this model is able to capture the single subject level robust averaging effect. Li et al. (2017) describe multiple variations of models but in this chapter we will focus on the simple "Late noise only power model". This model has only 2 parameters:  $k$  (the exponent) and  $s$  (the level of late noise). Like Li et al., (2017), I do not address noise in the encoding of each stimulus because the grating stimuli were presented in full contrast just like the number stimuli described earlier in the chapter and thus these number stimuli should be relatively easy to encode.

This model implementation allows the researcher to explicitly simulate how model performance varies under different levels of late noise ( $s$ ) and different levels of robust averaging ( $k$ ). Remember, in the computational model described earlier, model performance was only calculated at different levels of robust sampling and not for different levels of late noise. **Fig. 4.22.B.** shows the difference in accuracy between the power model and a model that gives equal weight to all of the elements ( $k = 1$ ). The warm coloured areas on the left of the plot indicate a performance benefit of robust averaging ( $k < 1$ ). In other words, robust averaging can be beneficial for protecting against late noise. The reason for this protective effect is that robust averaging essentially pushes the decision variable of a given inlying feature value away from the boundary. Thus, that feature is less likely to cross the boundary (e.g. less likely to contribute to the wrong decision) in the presence of late noise.

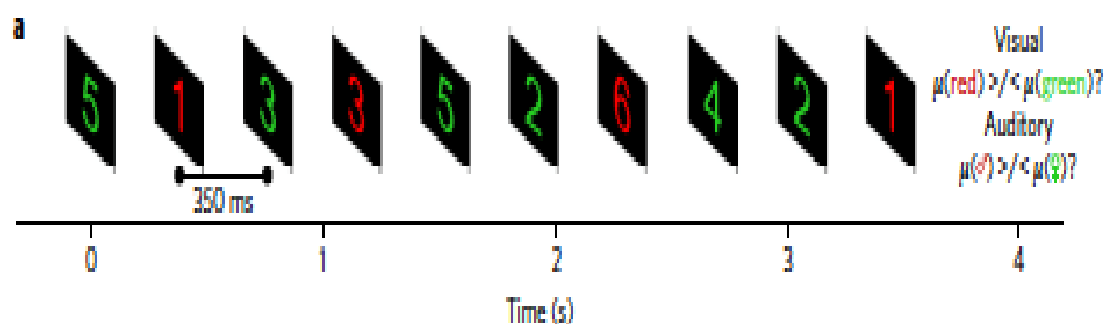
However, given limited resources, one has to allocate resources and weight or sample information with discretion. In fact, another simulation has shown that these performance benefits only occur when the overall distribution of features is close to Gaussian (e.g. when the inlying items are the most frequent items but not when the distribution of features is uniform in shape; **Fig. 4.23**).



**Fig 4.23. Simulated model accuracy difference when features are from an uniform distribution (figure copied from Li et al., 2017 under CC-AL license).** The same figure as Fig. 4.23B but now the model accuracy difference is computed by using uniformly distributed features. Note the difference between Fig. 4.23B and this figure. This figures shows that when features are distributed uniformly, there is no performance benefit with robust averaging ( $k < 1$ ).

Recent work from the lab has tested this explicitly using a sequential symbolic

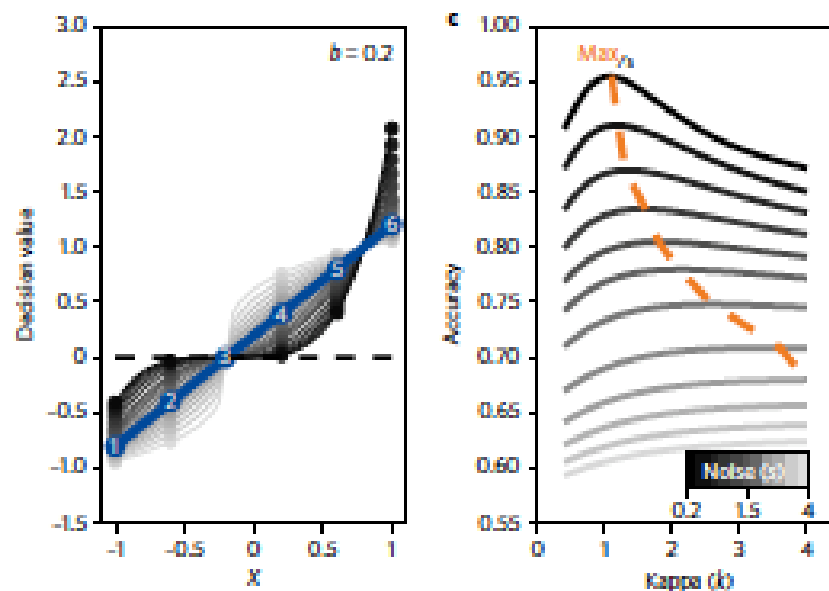
number integration task (Spitzer, Waschke, & Summerfield, 2017). In that task, subjects were shown 10 random symbolic numbers ranging from 1 to 6 that were drawn from two uniform distributions on each trial. Their task was to discriminate whether the red numbers that were drawn from one distribution were higher or lower than the green numbers that were drawn from the other distribution (**Fig. 4.24**). Spitzer et al. modelled the data with the same power model as described earlier, but added an extra “bias” parameter that is able to capture an asymmetric weighting bias (**Fig.4.25A**).



**Fig 4.24. Task schematics** (figure copied from Spitzer et al., 2017 with permission through Rightslink). Participants were presented a sequence of 10 numbers in red and green coloured fonts. Their task was to report whether the average of red or green numbers was higher.

Just like in Li et al. (2017), these authors simulated how accuracy values vary as a function of different values of  $k$  and  $s$  (**Fig. 4.25B**). They found that under a zero noise situation, the optimal strategy is to weight all numbers equally (i.e. just like a perfect integrator would). However, as the level of late noise increased, the optimal  $k$  also increased. Note that this effect is in the opposite direction from the effect reported earlier (Li et al., 2017). This work thus implies that when stimuli are drawn from uniform distributions, one should

allocate more resources to the most diagnostic items, i.e. the outliers in the case of a discrimination task. This is further supported by the fitting results in which the model with a  $k > 1$  yielded a better fit than a perfect integrator model (Spitzer et al., 2017)



**Fig 4.25 Power-law transducer with the bias parameter and simulated model accuracy (figure from Spitzer et al., 2017 with permission through Rightslink).** (A) Feature values  $X$  were mapped onto Decision values by a power-law transducer characterised by an exponent  $k$ . However, the difference of these transfer functions compared to the ones in Fig.4.23A is that the inflection point of these transfer functions is shifted left by 0.2 (i.e. a bias term). This bias term introduces the asymmetric weighting bias on numeric digits in individual subjects. (B) Simulated model accuracy under different level of  $k$  (x-axis) and late noise –  $s$  (lines). Lighter coloured lines correspond to higher late noise. As indicated by the orange dotted line, the optimal  $k$ , which is determined by the maximum performance one could achieve at a given level of late noise, increases as the level of late noise

*increases. This shows that when features are drawn uniformly, one show overweigh outliers especially in the presence of late noise.*

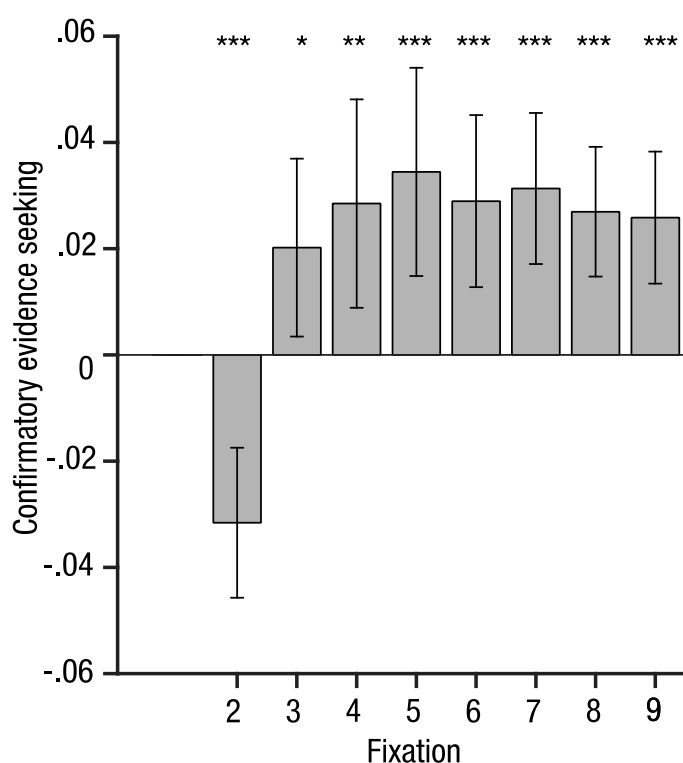
In conclusion, these two studies suggest that nonlinear weighting can be a systematic mechanism for countering decision noise during the integration process. These studies combined with the robust sampling results indicate that the weighting and / or sampling strategies that humans use are highly adaptive and flexible enough to meet the demands of the environmental statistics. The overall distribution of features was Gaussian in both the robust sampling study and the Li et al. study. Therefore, given limited resources (e.g. time constraint of the experiment / sampling, metabolic cost of making a saccade, working memory load...) and according to the predictions of the power model, humans should exhibit a robust sampling policy. Indeed, it was an optimal strategy (e.g. increase in performance) under these aforementioned constraints. Such a flexible gain (resources) allocation process that can result in these types of transfer functions has been described in chapter 1 (Cheadle et al., 2014).

#### **4.3.13 Confirmation bias**

As mentioned in chapter 3, previous studies have found that fixations are non-Markov, i.e. that the information sampled on fixation  $k$ , whilst dependent on fixation  $k-1$ , is independent of earlier samples of information. However, in chapter 3 we only found a confirmatory bias based on colour and shape. Moreover, we could not find evidence for a confirmatory bias based on the

running hypothesis. To test for a confirmation bias in our data set, we again turned to the identified discrete fixations. We quantified the information value of each fixation based on their coordinates on the ET landscape. By comparing the sign of the information value of each fixation, we could check whether each fixation confirmed (i.e. same sign of the information value; =1) or disconfirmed (i.e., opposite sign; =-1) their current hypothesis (i.e. confirmation, defined as the running mean of the fixations so far). After calculating the average of these values per trial, this yields a quantitative measure of whether participants adapted their sampling policy as a function of the sampled evidence on each trial. Next, we constructed a null-distribution specific to our experiment by randomly permuting the numbers of each trial and repeating the analysis 1000 times, as an extra careful measure to control for any possible correlations in the values. This is important because since we manipulated the strength of the evidence, the overall values on the landscape are biased towards one side of the category boundary (e.g. on a trial with an evidence strength of  $\mu = R+2$ , the landscape is biased towards peaks; while on a trial with an evidence strength of  $\mu = R-2$ , the landscape is biased towards troughs). Thus on each permutation, each fixation had a different ET, while the information sampling characteristics remained the same. We then compared the average confirmation of each fixation (sorted by position in which they occurred) in our experiment to the one found under this null-distribution by calculating the average difference score between them for each fixation (positive scores indicating a confirmation bias, negative scores indicating a disconfirmation bias, and zero indicating that the sampling was independent of the previously sampled information). **Fig. 4.26** shows these

average difference scores and highlights evidence for a confirmation bias, i.e. a tendency to sample information that confirms the running hypothesis for fixations 3-9 (all  $P < 0.05$ ;  $d=0.43$ ). In other words, participants viewed by preference information that confirms their running hypothesis over all sampled viewed thus far.



**Fig. 4.26 Confirmation bias with respect to the running hypothesis.** The second fixation significantly disconfirmed the running hypothesis (equivalent to the previous fixation in this case), while all the fixations thereafter significantly confirmed the running hypothesis. \*  $P = 0.021$ ; \*\*  $P = 0.006$ ; \*\*\*  $P < 0.001$ . Bars signify 95% confidence intervals.

## 4.4 Discussion

In real-world settings, information arises in crowded natural scenes, replete with multiple features and objects that compete to drive decisions. Studies of saccadic control have revealed that attention is automatically captured by perceptually salient information, such as a target defined by a unique colour in a visual search experiment (A. L. Yarbus, 1967), and may also be captured by semantically deviant information, such as an object in an unusual context like a sofa on the beach (Võ & Henderson, 2010).

However, the task at hand is a key determinant of saccadic control, and the eyes are drawn to those scene elements that are task-relevant, or valuable, or to those locations where the most information can be gleaned (Eckstein, Schoonveld, Zhang, Mack, & Akbas, 2015; Itti & Baldi, 2009; Manohar & Husain, 2013; Najemnik & Geisler, 2008; Navalpakkam, Koch, Rangel, & Perona, 2010; Towal et al., 2013). Nevertheless, most decision-making studies either use a single, centrally presented source of information (such as a dot motion stimulus) precluding the detailed investigation of the policy by which participants sample decision information or assume that gaze allocation is a stochastic process that drives a preference for fixated alternatives, rather than vice versa (Krajbich et al., 2010; Krajbich & Rangel, 2011). Recently, researchers have highlighted the importance of investigating evidence accumulation as an active sampling process (Gottlieb et al., 2014). In this chapter, we reveal two aspects of the sampling policy that determines where gaze is allocated during decision-making.

First, when faced with multiple spatially distinct sources of decision information, observers tend to ignore those sources that carry the most extreme information and focus instead on the inlying decision information—that is, that falling closest to the boundary that segregates one choice from another (or close to the mode of the overall information presented on one trial; our approach was not designed to distinguish these possibilities).

This is surprising given that participants exhibit a well-characterized policy to saccade toward perceptually deviant information during natural scene viewing, or visual search (Towal et al., 2013; Võ & Henderson, 2010). The focus on inlying elements influenced human choices, with gaze driving a multiplicative effect to weight inlying elements more heavily, which in turn drove “robust averaging” in perceptual choices made during the experiment. Robust averaging of perceptual information has been described before, for faces (Haberman & Whitney, 2010) and for colours and shapes (De Gardelle & Summerfield, 2011; Michael et al., 2013).

Here, we add to these findings, revealing that humans engage in robust averaging of symbolic number. Previous studies have left unresolved the question of whether robust averaging occurs because inlying items are sampled (i.e., overtly attended) by preference or whether the influence of outliers is attenuated at a later decision stage. For example, one model

proposes that observers simply compute a posterior likelihood ratio conditional on the past distribution of features and their associated responses. This is equivalent to using a nonlinear decoder that “squashes” outlying information, muting its influence in the average and prolonging response latencies when decision information is heterogeneous (De Gardelle & Summerfield, 2011). However, here we show that the robust averaging is driven at least in part by the policy with which gaze is directed during information integration. In other words, humans are not just robust averagers but robust samplers of decision information.

Why do humans sample and average robustly? We describe a computational simulation of the sampling policy adopted during our spatial averaging task that was able to faithfully recreate the relationship between the evidence sampled and the choices made by humans. Exploring model performance under a range of policies (spanning a tendency to sample inliers or outliers by preference) revealed that for an ideal model limited only by the information available on the screen, sampling outliers rather than inliers by preference boosted performance. This is consistent with the observation that when assigning samples to one of two Gaussian-distributed classes, eccentric or “off-channel” features are most diagnostic of the correct response (Navalpakkam & Itti, 2007; Scolari & Serences, 2009).

However, in the human data, a tendency to sample inliers was a positive predictor of accuracy at the single-trial level. This implies that a robust

sampling policy has evolved to counter additional sources of loss that arise during the decision process, such as “late” noise that corrupts the information of integration, perhaps due to capacity limits on attention and working memory processes. For example, robust policies may lead to more stable estimates of decision information, just as nonparametric statistical inference is robust to the inclusion of aberrant data points. Other recent work has suggested that apparent “suboptimalities” in human judgment may in fact maximize performance when this late noise is taken into account (Howes, Warren, Farmer, El-Deredy, & Lewis, 2016; Tsetsos et al., 2016).

In summary, the current results highlight that perceptual judgments are not only constrained by the quality of the sensory information and the dynamics of integration but also by the sampling policy, which dictates how the information is selected for entry into the decision process.

#### **4.4.1 Future studies**

However, other questions remain open. In particular, the information processing steps by which fixations are guided to subsequent inliers remain unclear in our work, although this presumably relies on parafoveal processing of nonfixated locations. One possibility is that covert attention is first shifted to candidate items and the eyes are then “pulled” to this location (Henderson & Ferreira, 1990; Inhoff & Rayner, 1986). Moreover, recent work has shown that

foveal analysis and peripheral selection operate in parallel (Ludwig, Davies, & Eckstein, 2014).

Nor does our study directly address how saccade placement varies according to the spatial arrangement of items on the screen. For example, participants may prefer to saccade toward the center of mass of clusters of objects (He & Kowler, 1989). Thus, gaze can indeed be driven by the spatial arrangement of items on the screen, but gaze is additionally driven by the decision values of the available evidence, exhibited in this work via a preference to sample inlying samples. Future work could focus on investigating the effect of manipulating the spatial arrangement of items on the sampling policy.

Additionally, the effect of the interplay between the decision environment (i.e., task characteristics) and cognitive constraints on the relative (sub-) optimality of various sampling policies is an interesting avenue for future research. For example, it would be interesting to see whether these sampling policies are amplified, as more information is available in the periphery. Likewise, one could passively feed subjects sequences of stimuli generated by different sampling strategies to quantify their relative (sub-) optimality in different decision environments. This would also open an interesting avenue of research investigating whether the benefit of certain sampling policies is linked to task design (i.e. voluntary vs. involuntary sampling).

#### 4.4.2 Limitations

First, the approach we use here has some drawbacks. Notably, we assume that all information falling under the gaze trajectory contributes to the decision, in direct proportion to its viewing time, an assumption that is probably incorrect; it is known that participants are blind to information arriving at the retina during saccades (39). However, using a standard analysis approach, in which we defined fixations using Eyelink pre-set criteria (velocity threshold of 30, acceleration threshold of 8,000, and a saccade motion threshold of 0.1), saccadic trajectories accounted for only 18.75% of time points, and as these were largely made across the space separating stimuli, they contributed only 0.00004% to the final decision value, so any bias that arises by accounting for inputs from between fixations is likely to be negligible. However, our approach has the major benefit that it allows us to characterize, in real time on each trial, information that is available to the participant and thus to explore how stimulation and gaze interact in determining choices. It is also immune to otherwise arbitrary choices that might have to be made during a more standard fixation analysis, such as the size and shape of the window that surrounds each stimulus or the criteria for determining dwell times. Moreover, we provide multiple validation analyses that show that our approach offers sensible information about how people move their eyes and how this determines choices. Nevertheless, to verify that our findings were not in some way specific to our use of this “landscape” approach, we repeated the analysis using a more conventional method of defining an aperture with fixed width. In this “hard” filter approach, the landscape is composed of a series of circular “cylinders” placed at the  $x, y$  position of each item  $k$ , with height  $c_k$

corresponding to the quantity  $V_k - R$  and width  $s$ , where  $s$  is a free parameter fit to the QT as for standard landscape analysis. The QT values are computed in the same way as for the “soft” filter approach—that is, as the integral under the eye movement trace. The two methods thus differ in that here, any time the gaze trajectory falls within the space defined by a cylinder, the ET value is incremented by  $c_k$ , rather than depending on its distance to the  $x, y$  position. All of the results were replicated using a hard information- harvesting spotlight (i.e., a spotlight of fixed width with no spatial smoothing), and the ETs under both approaches were qualitatively identical (see **Fig 4.10**).

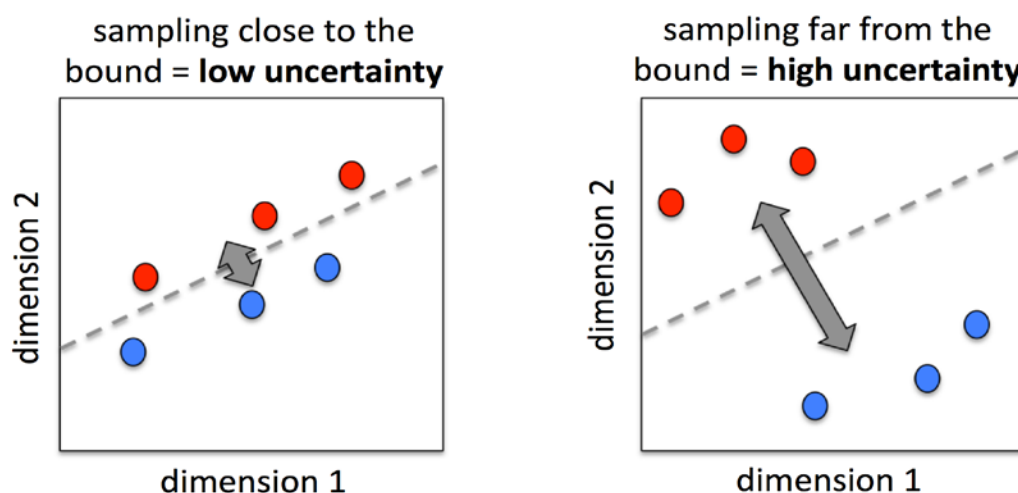
Second, our theory and results point to two similar but non-identical interpretations of the data: (i) humans downweight outliers relative to the reference provided on each trial, and (ii) humans downweight outliers relative to the mean of the information viewed thus far on each trial. Our study was not designed to arbitrate among these claims (Note however, that participants were not sampling exclusively at the reference; see the section entitled “Confirmation Bias”). Furthermore, the adaptive gain theory formulated by the Summerfield lab, developed over several recent papers (Cheadle et al., 2014; De Gardelle & Summerfield, 2011; Summerfield & Tsetsos, 2015) argues that during decision-making, neural resources (or “gain”) is allocated to those portions of the decision space in which the information is most likely to occur, a form of “efficient coding”. In the current task, the expectation (i.e. mean) of the distribution is initially unknown to the participant but is signalled approximately by the reference, which is fully known to participants because it

remains constant over the block. Deviations of the mean information from the reference (e.g. by  $\pm 2$  or  $\pm 4$ ) become apparent on each trial during the saccadic sampling process, over which the participant estimates the mean of the numbers and judges whether it is higher or lower than the reference. In fact, either interpretation would be consistent with the overall adaptive gain framework. Previous work involving sequential, experimenter-controlled presentation of discrete samples, has shown evidence that the gain control process adjusts dynamically as information is revealed (Cheadle et al., 2014). This manifests as a “consistency” bias, whereby higher gain is allocated to those samples that are more similar to the running average of all features viewed thus far in the sequence (Cheadle et al., 2014). Thus it is likely that participants begin by allocating gain to the reference feature, and gradually adjust their gain over time in our experiments in this chapter. Note that the confirmation bias finding does indicate that participants are indeed adapting their sampling policy as new information arrives with each successive fixation. This is consistent with an adjustment of the sampling policy towards the mean of the trial, rather than to the reference alone.

It is important to note that robust sampling is very different from the optimal policy that emerges as participants gradually learn a category boundary through repeated sampling (Markant et al., 2016). Gureckis et al. have shown that during category learning, participants sample preferentially from close to the category boundary (e.g. the reference in our task), as is optimal to do so in their task. The key difference is that in the Gureckis papers, the bound

(reference) is unknown to the participant, and is *learned* by taking successive samples; whereas in our task, participants were fully aware of the bound (i.e. reference value on each trial), and no further learning about the bound occurs during sampling. When the boundary is unknown (e.g. in the Gureckis case), it is optimal to sample close to the boundary, to refine estimates of the precise form of the boundary.

In **Fig. 4.27**, we illustrate the problem described by Gureckis et al: categorising stimuli from 2 categories (red and blue dots) that vary parametrically along 2 dimensions (x and y axes) where the true bound (dashed grey line) is unknown to the observer. An optimal policy involves assessing the category membership of samples that are close to the bound (left panel).



**Fig. 4.27** Sampling distance to the bound under low / high uncertainty.

Critically, in the Gureckis case, this policy is optimal because it will reduce uncertainty about the position of the bound (e.g. short grey arrow). However, a policy that samples items further from the bound (right panel) will fail to decrease uncertainty about the position of the bound, as the true bound can take a number of possible positions in between the sampled red and blue dots, as indicated by the longer grey arrow. This adaptive learning policy is well known in the machine learning literature, where it is employed to solve categorisation problems where the axis of classification is unknown, but must be learned gradually via supervised training (e.g. in a support vector machine).

However, this is very different from the situation in our task, where the reference (i.e. bound) does not have to be learned, but is fully known to the participants on each trial. In the Gureckis paper, as in much of the category learning literature, the task is to distinguish whether each sample belongs to categories A or B (i.e. red or blue), relative to an unknown bound. By contrast, in our task, there is only one category (e.g. A) and the task is to say whether the numbers fall one side or another of the bound (i.e. reference). Thus, in our task, participants cannot actually perform the task unless they are aware of the bound.

Moreover, in the Gureckis task, participants receive feedback after each sample; this allows them to gradually adjust their estimate of the position of the bound, an adjustment that occurs more rapidly under an optimal sampling policy. In our task, participants receive no feedback until the end of the trial, i.e. until all numbers have been sampled; thus, there is *no learning about the bound* that occurs during the trial, because there is no further information provided that allows that learning to proceed. We thus think that explanations that appeal to the bound-learning policy described by Gureckis and colleagues are extremely unlikely to account for our prototypicality effect.

This is why, as our computational simulations show (at least for an observer who has no “late” noise, i.e. under the same assumptions made by Gureckis et al.), the policy that should maximise accuracy in our task (with known bound) is one in which outliers, rather than inliers, are sampled by preference.

## **Chapter 5: Voluntary information sampling when deciding between two categories (experiment one).**

### **5.1 Introduction**

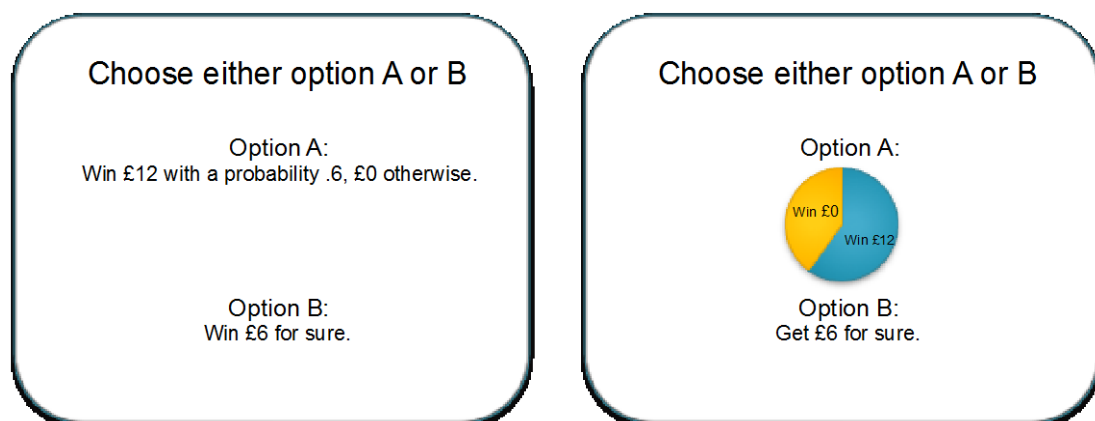
In natural environments, animals are often confronted with decisions between multiple options with differing value. Making the best decision can often be difficult, frustrating and time consuming. Say for example that you are thinking about buying a new laptop computer. This decision is everything but easy, as the laptops will come in various price ranges, sizes, qualities, and the 'value' of each laptop is not easily quantified.

So far, this thesis has only considered one class of decision: choosing whether a stimulus should be classified as belonging to category X or Y (e.g. is the average of the numbers displayed on the screen higher or lower than the reference category? Is the average colour of the stimulus display red or blue?). This is equivalent to evaluating the properties of a laptop and deciding whether to buy it or not.

However, many ecological decision problems involve the sampling of information about multiple distinct alternatives and choosing among them. Therefore, a second class of decision involves a comparison between two or more options. Confronted with two laptop computers, do you buy the Mac or the PC? In this chapter, we delve deeper into this particular class of decisions.

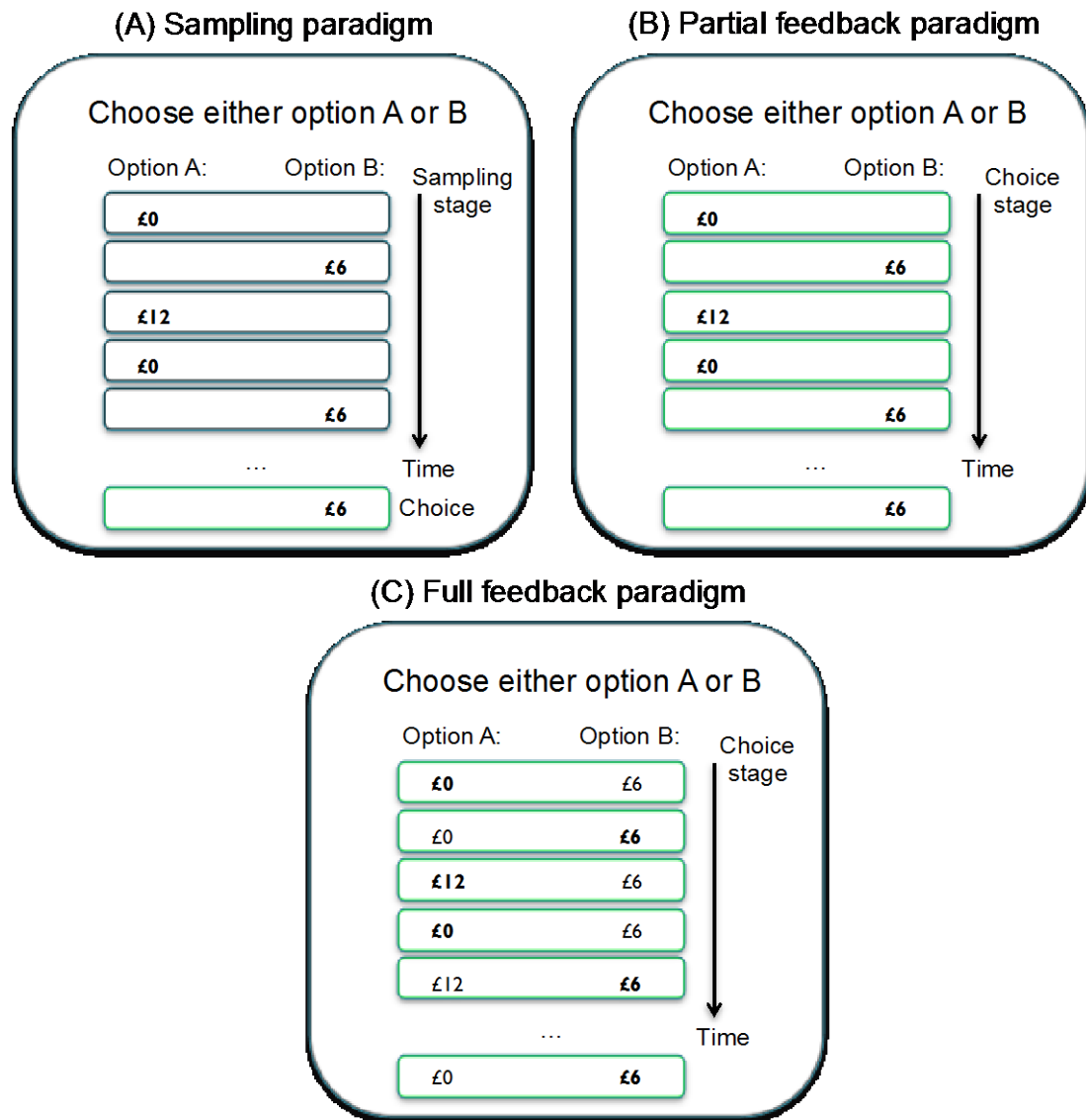
### 5.1.1 Decisions from description vs. decisions from experience

One particular paradigm, the risky gambles paradigm, has often been used to investigate the decision making processes in play when humans decide between two options (Allais, 1953; Barkan, Zohar, & Erev, 1998; Barron & Erev, 2003; Camerer, 2000; Erev, 1998; Kahneman & Tversky, 1979, 1984; Weber, Shafir, & Blais, 2004). Originally, this literature was mostly focussed on “decisions from description”. Typically, respondents are provided a summary description of two options (**Fig 5.1**).



**Fig 5.1. Decisions from description examples.** The outcomes and probabilities of each option are provided and the information is often conveyed either numerically (left) or visually (right). These decisions from description are usually one-shot problems, e.g. one choice per problem and often no feedback is given (Hertwig & Erev, 2009).

However, when making decisions outside of the laboratory we are often not presented with clear outcomes associated with each decision and their probabilities of occurring. Instead, we often learn about potential outcomes and their associated probabilities by making choices and receiving trial-and-error feedback. In short, we need to rely on our own experiences in order to infer the state of the environment. Recently, researchers have started to investigate these “decisions from experience” in humans (Barron & Erev, 2003; Erev & Barron, 2005; Hertwig, Barron, Weber, & Erev, 2006); see (Hertwig & Erev, 2009; Newell & Rakow, 2007) for reviews). These decisions from experience paradigms often employ the so-called ‘bandit’ tasks, adapted from a literature in animal learning. In this paradigm, the outcomes and probabilities of each option are initially unknown and respondents need to gather information about each option by repeatedly choosing among the available options and receiving feedback. The paradigm consists of an initial sampling stage, in which the respondent explores option A and B without cost via button presses. Each button press results in a random draw from a variable distribution associated with the chosen option. The outcome for their particular choice is then shown on screen. Three variations of this paradigm have been used in the literature: the sampling paradigm, the partial-feedback paradigm and the full-feedback paradigm (**Fig. 5.2**; (Hertwig & Erev, 2009).



**Fig 5.2 Decisions from experience examples.** In this paradigm, the outcomes and probabilities of each option are unavailable and respondents need to gather information about each option by repeatedly sampling the available options. The paradigm consists of an initial sampling stage, in which the respondent explores option A and B without cost via button presses. Each button press results in a random draw from the chosen distribution. The outcome for their particular choice is then shown on screen (bolded in the figure). (A) The sampling paradigm. In this paradigm only the eventual choice has an effect on the payoff (green box). (B) The partial feedback paradigm. In

this paradigm, every draw has an effect on the payoff (green box). (C) The full feedback paradigm. In this paradigm, every draw has an effect on the payoff (green box), and non-chosen outcomes are presented on each draw as well (non-bolded).

In the sampling paradigm, draws during the sampling stage have no effect on the respondents' payoffs **Fig. 5.2.A**; (Hertwig et al., 2006; Weber et al., 2004). In the partial feedback and full feedback paradigms, each draw contributes to the respondents' payoffs. In the full feedback paradigm, each draw will result in information from both the chosen and unchosen distributions (**Fig. 5.2.C**; (Yechiam & Busemeyer, 2006), while in the partial feedback paradigm only information about the chosen distribution is displayed (**Fig. 5.2.B** (Barron & Erev, 2003; Erev & Barron, 2005)).

Importantly, the work on risky gambles has demonstrated a robust description – experience gap. When making decisions from description, rare events have a stronger impact than they deserve according to their objective probabilities. However, when making decisions from experience, rare events have a weaker impact (See (Hertwig & Erev, 2009) for a review).

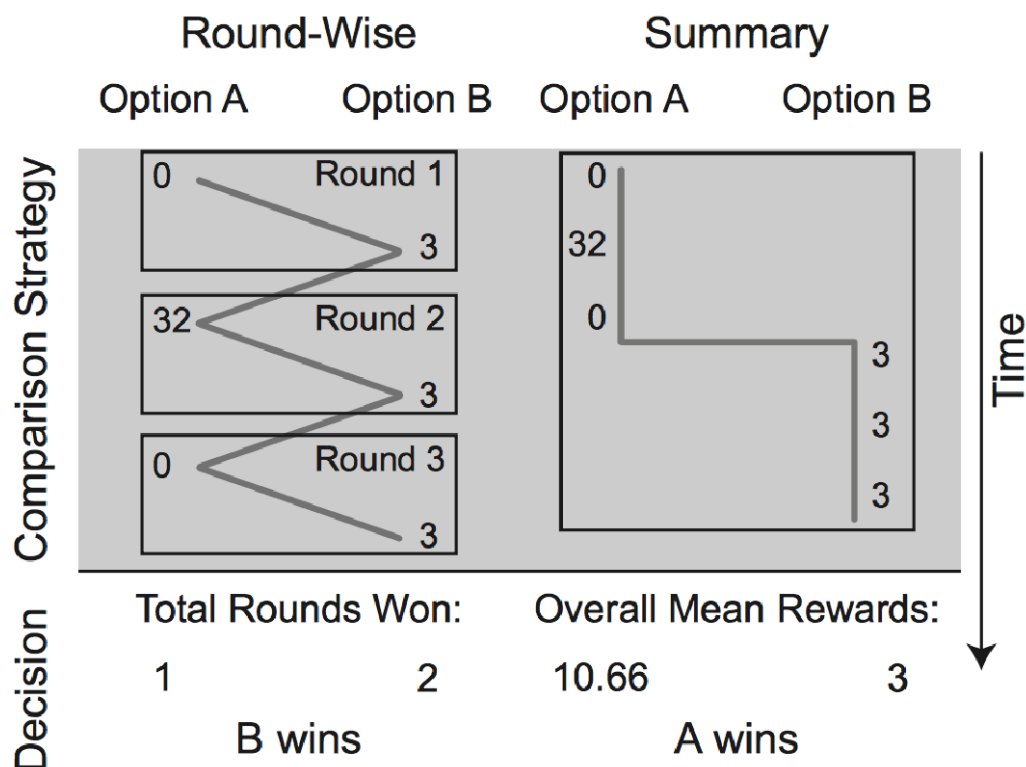
### 5.1.2 Mechanisms underlying the description-experience gap: the importance of the sampling policy

Hertwig and Erev (2009) described possible mechanisms underlying the description – experience gap (Hertwig & Erev, 2009). Potential causes included sample size (limited number of samples), recency (overweighting of recent samples), estimation error (underestimation of the frequency of rare samples), and format-dependent cognitive processes.

Recently another underlying mechanism of the description – experience gap was explored that is particularly important for this thesis. Hills & Hertwig (2010) proposed that the description – experience gap could be related to people's sampling policy (Hills & Hertwig, 2010).

The authors distinguished between two sampling policies in their decisions from experience paradigm: round-wise and summary sampling (**Fig. 5.3**). In the round-wise sampling strategy, choices alternate (e.g. switch) between the two available options (**Fig. 5.3.A**), while summary sampling focuses on sampling an option extensively before switching to the other option (**Fig. 5.3.B**). Moreover, the authors link these sampling strategies with the way individuals would make their final decisions. When sampling round-wise, individuals are hypothesised to use a counting strategy (i.e. counting the number of rounds (comparisons) option A and B had won, and choosing the

overall winner). Under summary sampling, individuals are hypothesised to calculate the average for each option and choosing the highest average.



**Fig 5.2 Round-wise (left) and summary sampling (right) comparison strategies** (figure copied from Hills & Hertwig 2010 with permission through RightsLink). The round-wise sampling strategy is characterised by sampling that alternates between the two available options (left), while summary sampling focuses on sampling an option extensively before switching to the other option (right). When sampling round-wise, individuals are hypothesised to use a counting strategy (i.e. counting the number of rounds or comparisons option A and B had won and choosing the overall winner). Under summary sampling, individuals are hypothesised to calculate the average for each option and choosing the highest average.

Hertwig & Erev (2009) re-analysed data from four published studies and confirmed their hypothesis that the decisions made by respondents depended on their specific sampling strategy (e.g. respondents who used a round-wise sampling strategy were more likely to use the counting strategy and vice versa for the summary sampling strategy). Moreover, the sampling strategy (round-wise or summary sampling) additionally influenced the impact of rare events, over and above what would be expected from the limited samples explanation (Hertwig & Erev, 2009). More specifically, a round-wise sampling policy led to more underweighting of rare events compared to a summary policy. As a side note, we recall that in the previous chapter it was important to demonstrate that the counting strategy was not an alternative explanation for our robust sampling findings. Thus, the sampling policy can have widespread consequences on the decision making process itself when deciding between two options or categories.

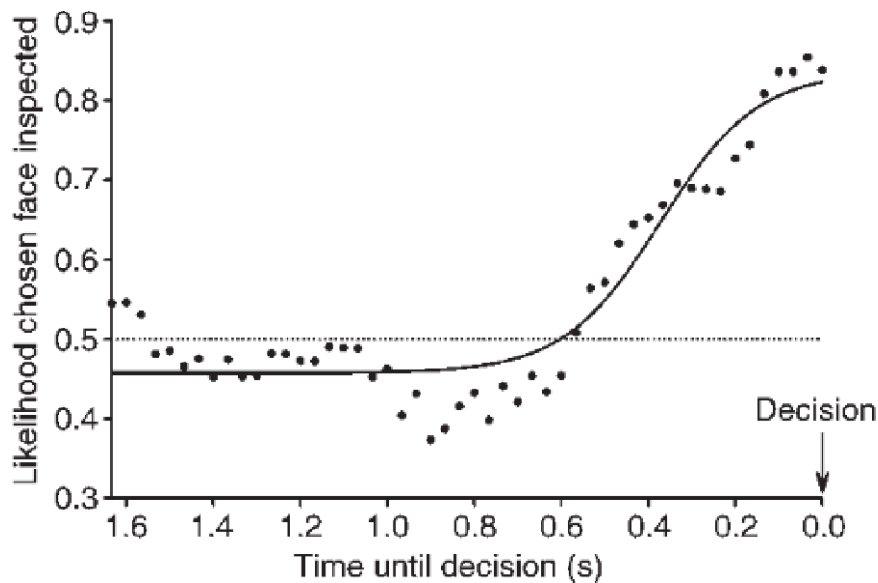
This study demonstrated that observers utilised different sampling policies (i.e. round-wise / summary sampling) and that these sampling policies are associated with different decision strategies (i.e. counting / averaging). However, it is unclear why and under which circumstances observers choose to employ a particular sampling policy. For example, does more consistent sampling (e.g. summary sampling minimising the number of switches between categories) lead to higher decision accuracy? Similarly, when observers do not utilise round-wise sampling, do they switch at particular points in the evidence sampling sequence?

Note that we are assuming an equivalence here between sampling information with a response and sampling with the eyes. Furthermore, it might be argued that in the partial-feedback and full feedback paradigms the sampling occurs between trials instead of within a trial. However in our opinion, these paradigms can be seen as the equivalent of one-trial paradigms: you sample evidence for each option, and these samples can either have an effect on your payoff (feedback paradigms), or not (sampling paradigm).

In the next section, we briefly describe some studies on preference formation and value based decision making that illustrate the influence of the sampling policy on the decision making process when deciding between two options or categories.

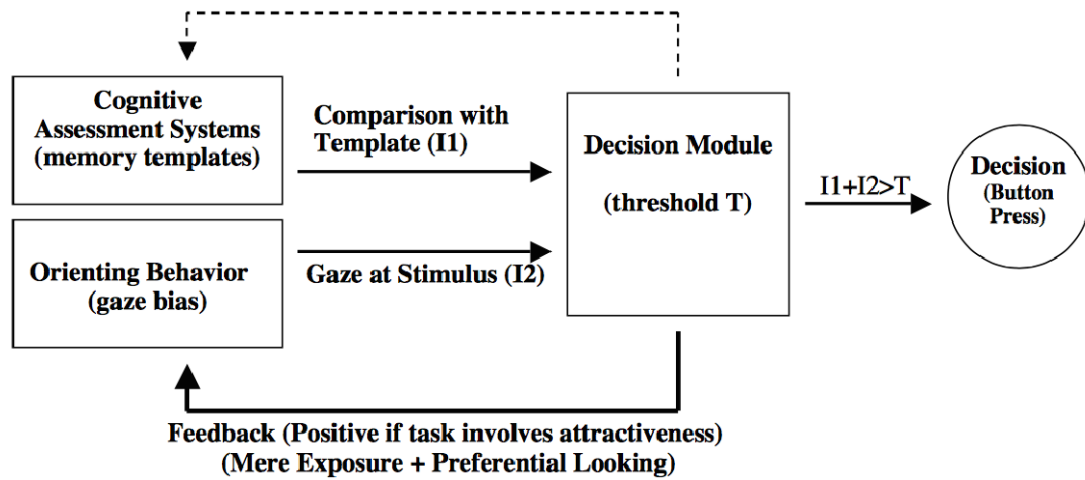
### **5.1.3 The gaze cascade**

Shimojo, Simion, Shimojo, & Scheier (2003) investigated the role of the sampling policy in preference formation (Shimojo et al., 2003). Shimojo et al. employed a task in which participants had to make a two alternative forced choice about two face stimuli (e.g. choose the most attractive face). In experiment 1, observers were free to sample the two faces for as long as they wanted and had to indicate their preference with one of two buttons. Interestingly, the results showed a gaze-cascade effect: a progressive bias in the observers' gaze toward the stimulus they ultimately chose (**Fig. 5.3**)



**Fig 5.3 Example gaze likelihood curve (figure copied from Shimojo et al., 2003 with permission through RightsLink).** The likelihood of inspecting the ultimately chosen face is shown as a function of time left until the decision was made. The likelihood of inspecting the ultimately chosen face gradually increased up until the decision was made (up to about an 83% likelihood; e.g. the gaze cascade).

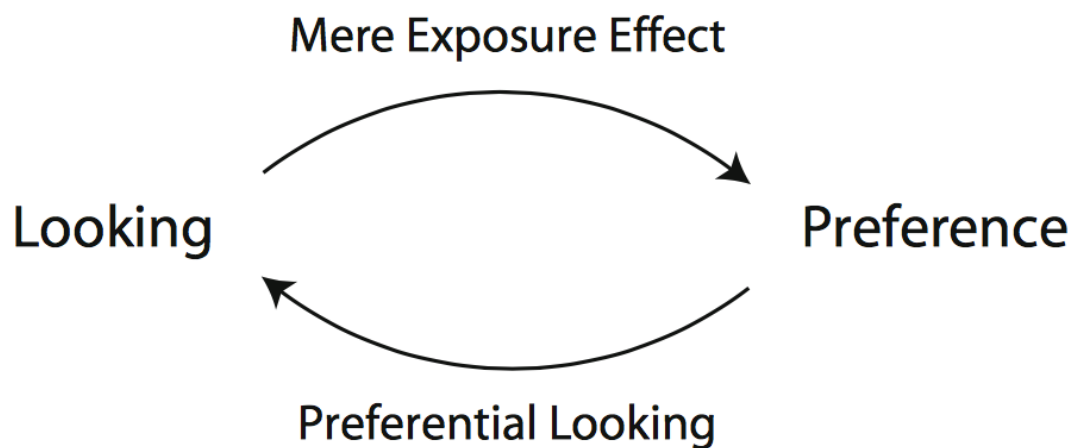
The authors proposed a dual contribution model of preferential decision making with two inputs of parallel information processing, one from the cognitive assessment systems and the other from the orienting behaviour system (i.e. the sampling policy; **Fig. 5.4**). The cognitive assessment system encompasses task relevant judgments (i.e. attractiveness of the faces), while the orienting behaviours encompass the sampling policy. The decision module itself is influenced by both systems.



**Fig 5.4** The dual contribution model of preferential decision making (figure copied from Shimojo et al., 2003, supplementary information with permission through RightsLink).

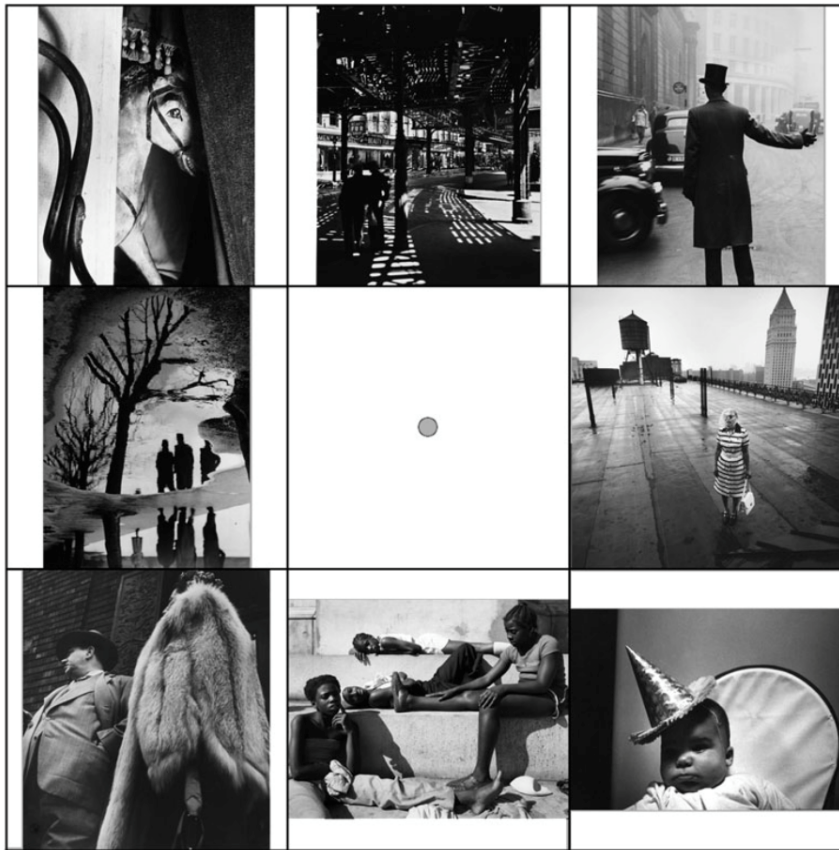
Shimojo et al. proposed two crucial processes related to the orienting behaviours and their influences on the decision module: the mere exposure effect and preferential looking (Shimojo et al., 2003). The mere exposure effect refers to the finding that merely looking at a stimulus increases your preference for that stimulus (Kunst-Wilson & Zajonc, 1980; Moreland & Zajonc, 1977, 1982; Zajonc, 1968). Preferential looking refers to the finding that observers tend to sample the stimulus they like for longer (Birch, Shimojo, & Held, 1985; Fantz, 1964). Thus, preferential looking and the mere-exposure effect instantiate a positive feedback loop (**Fig. 5.5**; (Simion & Shimojo, 2007)). Further evidence for this model was provided in their second experiment. In this experiment, only one face was available on the screen at any time while

the presentation duration itself was manipulated (e.g. the sampling policy was forced in the sense that one face was viewed preferentially). Shimojo et al. found that biasing observers' gaze via manipulation of the presentation duration influenced their choices towards the overly sampled stimuli.



**Fig 5.5** Positive feedback loop between preferential looking and the mere exposure effect (Simion & Shimojo, 2007). Figure copied from Glaholt & Reingold, 2009 with permission under RightsLink.

Glaholt and Reingold took this line of research further and tested whether pre-exposing some of the stimuli would lead to the pre-exposed stimuli to be chosen more often, as would be predicted by the dual contribution model **Fig 5.6** (Mackenzie G Glaholt & Eyal M Reingold, 2009; Mackenzie G. Glaholt & Eyal M. Reingold, 2009).



**Fig 5.6 Example stimulus display in the 8-AFC task** (figure from Glaholt & Reingold, 2009 with permission under RightsLink). Observers had to indicate which image they preferred (or which image was most unusual in a control condition). Crucially, four of the eight images were shown to the observers before they were presented with all eight images (pre-exposure).

Interestingly, this was not what the authors observed. Even though they did replicate the traditional gaze cascade effect, they additionally found that pre-exposure led to shorter dwell times and that pre-exposure had no effect on the likelihood of a stimulus being chosen (e.g. contrary to what would be expected under the dual contribution model). Moreover, Glaholt & Reingold found that the gaze cascade effect was present in both preference and non-

preference control conditions (e.g. 'choose the most unusual face'). Additionally, other authors have recently provided more evidence supporting that the gaze cascade effect is not solely present in preferential judgment tasks (Saito, Otani, & Kinjo, 2015; Schotter, Berry, McKenzie, & Rayner, 2010).

#### **5.1.4 Selective encoding**

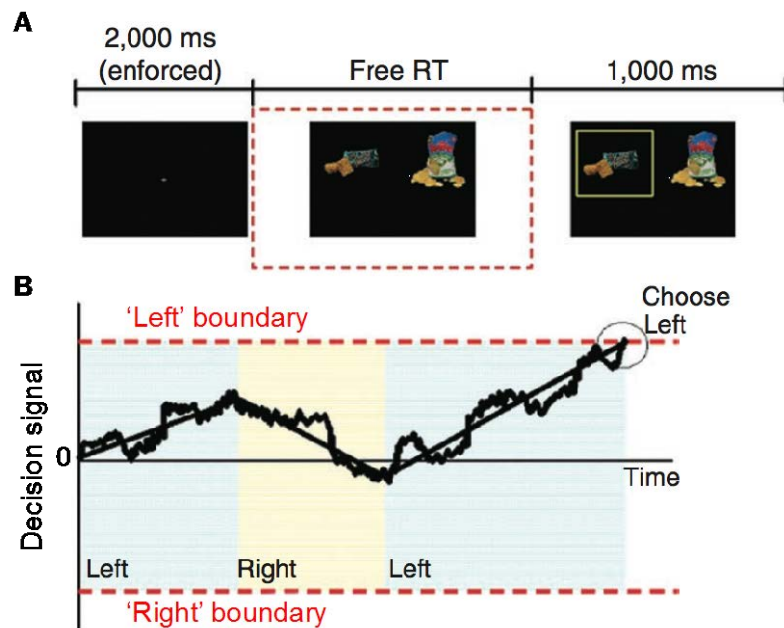
This led Glaholt & Reingold to conclude that the gaze cascade effect is not specific to preference judgements but instead that it is the selective encoding of stimuli that drives eye movements during perceptual decision making. In turn, the sampling policy (i.e. eye movements) influences the decision making itself. This selective encoding manifests itself by a greater allocation of resources to the features of the stimuli that are task relevant. Thus, when performing a preference judgment where the observers have to indicate which face they 'like' the most, selective encoding and a bias towards looking at the face the observer likes the most will be additive. However, when performing a 'dislike' preference judgement selective encoding would yield an opposing effect. This opposing effect in turn leads to the reduced gaze-cascade effect in the original Shimojo study (Schotter et al., 2010). In summary, the gaze cascade effect has been demonstrated in a variety of tasks and is domain general. Moreover, the main driver of the gaze cascade is selective encoding, resulting in a greater allocation of resources to features that are task relevant and in turn influencing the decision making process.

### 5.1.5 The attentional drift diffusion model

Previous chapters in this thesis have provided evidence that the sampling policy directly affects the decision process during (categorical) perceptual decision making. In this introduction, we have described evidence that the sampling policy directly affects the decision process during two alternative forced choice tasks (e.g. the sampling policy as an explanation for the decision-experience gap and the gaze cascade effect). Moreover, studies have shown that it is possible to directly bias choices by manipulating the sampling policy exogenously (experiment 2 in (Armel & Rangel, 2008; Shimojo et al., 2003)). Thus, the sampling policy plays a crucial role in the decision making process as argued in chapter 1. The classic DDM approach mentioned in chapter 1 ignored the role that visual attention plays during the choice process. However as discussed in chapter 1, Krajbich et al. acknowledged this crucial role and implemented visual attention in the classic DDM model (Krajbich et al., 2010; Krajbich & Rangel, 2011).

Remember that in this work, subjects were shown a display of two food items and were allowed to move their eyes freely between the two options the items for as long as they wanted before indicating their preferred choice (**Fig 5.7.a**). In this study, items that were fixated more frequently were chosen more often. The authors extended the DDM framework by incorporating a process level explanation of the information stage, and in particular the selective weighting of information (Krajbich et al., 2010; Krajbich & Rangel, 2011). The results were accurately captured with a variant of the DDM that incorporated one

crucial difference: the slope at which the decision signal evolves over time (i.e. the drift rate parameter) depends on the value of the currently fixated location (i.e. which item is sampled, **Fig 5.7b**). Additionally, it is easy to imagine how this attentional drift diffusion model framework can fit within the gaze cascade and selective encoding literature cited earlier in this chapter.



**Fig 5.7 Experimental design of a two alternative forced choice task.**

Figure from Krajbich, Armel, & Rangel (2010). **(a)** On each trial, subjects had to fixate a central fixation point for 2 seconds. Subsequently, two food items were displayed on the screen and subjects could sample these items for an unlimited time. A yellow box highlighted the chosen option for 1 second after subjects made their choice. **(b)** The decision signal evolves over time until it reaches a boundary (left in this case). Crucially, the slope of the decision signal (up or down) is dictated by which stimulus (left or right) is sampled.

However, this work has two important limitations. First, the model treats the sampling process as effectively random. So far, this thesis has provided ample evidence that the sampling process is not random. Furthermore, in this chapter we have provided evidence for non-random sampling in 2-AFC tasks specifically (e.g. round-wise versus summary sampling, preferential looking, selective encoding...).

Second, these value-based decision making studies have tried to control for value by having observers rate each stimulus (e.g. in this case the food items) at the beginning of the experiment. However, it is still very difficult to have an accurate measure of exactly what information is extracted during the sampling of a particular choice problem. For example, the value of a food item might be influenced by the value of the competing food item on that particular choice problem.

Thus, it still remains unclear whether the sampling policy itself can affect which food item (e.g. category) is chosen, independently from the value of each category.

## **5.2. Current study**

In this final experimental chapter, we aim to delve deeper into the characteristics and determinants of the sampling policy when deciding between two categories or options, and how the sampling policy affects these decisions themselves. To do this, we developed a perceptual decision making task in which participants could sample multiple sources of available

information from two distinct categories. Such a design allowed us to explicitly quantify the evidence associated with each sample of information. Unlike chapter 4, we opted to use a gaze-contingent stimulus display (see below). This design allowed participants to still guide their (category based) sampling on peripheral information, while explicitly limiting the information that they could sample concurrently to one category.

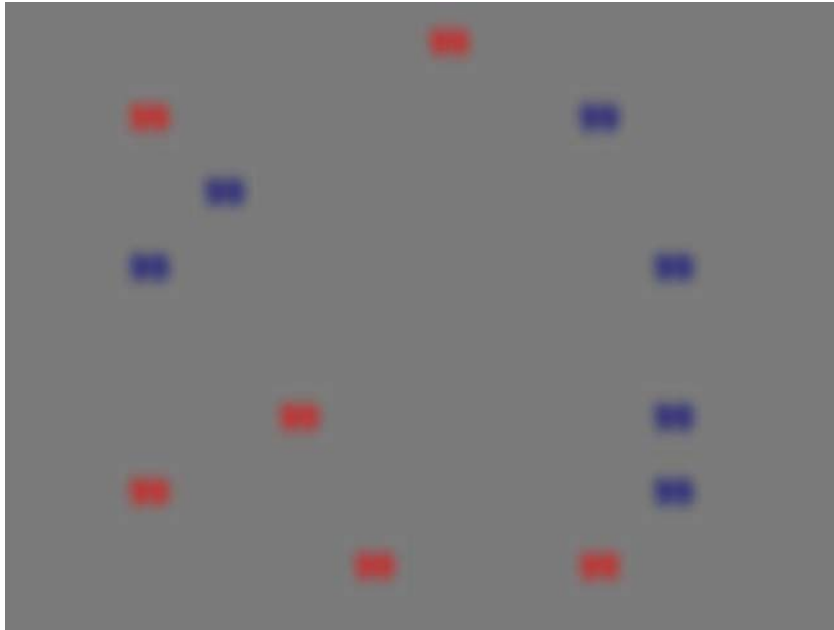
We will tackle the following questions in this chapter:

- Q1. Does sampling for a prolonged period of time increase decision accuracy when deciding between two categories?
- Q2. Does the sampling policy itself affect which category is chosen, *independently* from the value of each category?
- Q3. Does consistent sampling increase decision accuracy when deciding between two categories?
- Q4. During sampling, how does the accumulated evidence influence the categorical identity of the next sample (e.g. is the category of the next sample influenced by your current hypothesis)?
- Q5. During sampling, how does the uncertainty associated with one category determine the probability that a new sample will be drawn from that category (e.g. is uncertainty a driver of the sampling policy)?

### 5.2.1 Experiment 1 Methods

*Sample.* The experiment took place in the Department of Experimental Psychology at the University of Oxford. Thirty-six subjects volunteered (12 males and 24 females, age between 18 and 30 years old) to participate in experiment 1. Subjects reported normal or corrected to normal vision and had no history of any psychological disorders. Participants provided written consent before the experiment. The experiment was approved by the local ethics committee (approval no. MSD-IDREC-C1-2009-1).

*Stimuli.* On each trial, participants viewed an array composed of twelve two-digit numbers. Stimuli were generated using PsychToolbox ([psychtoolbox.org](http://psychtoolbox.org)) for MATLAB (Mathworks) and presented on a 17-inch screen (resolution, 1024 x 768) viewed from a distance of 60 cm. Twelve numbers ( $\sim 1.15^\circ$  horizontal and  $\sim 0.95^\circ$  vertical visual angle; Six numbers per category) were presented at randomly selected locations within an invisible 8 x 8 grid. Each category was associated with a specific colour (red or blue) and each number on the screen was initially masked but still conveyed its category colour. When gaze fell upon a masked number, the number was revealed while observers sampled the stimulus resulting in a gaze-contingent stimulus display. As soon as their gaze shifted to another number, the sampled number was masked once again. No number pairs were presented within the central four grid locations. An example stimulus array can be seen in **Fig. 5.8**.



**Fig 5.8 Example stimulus display.** Twelve masked numbers were present on each trial. The masks still conveyed category information but no decision information. When observers fixated a particular mask, the underlying number was revealed as long as observers maintained fixation.

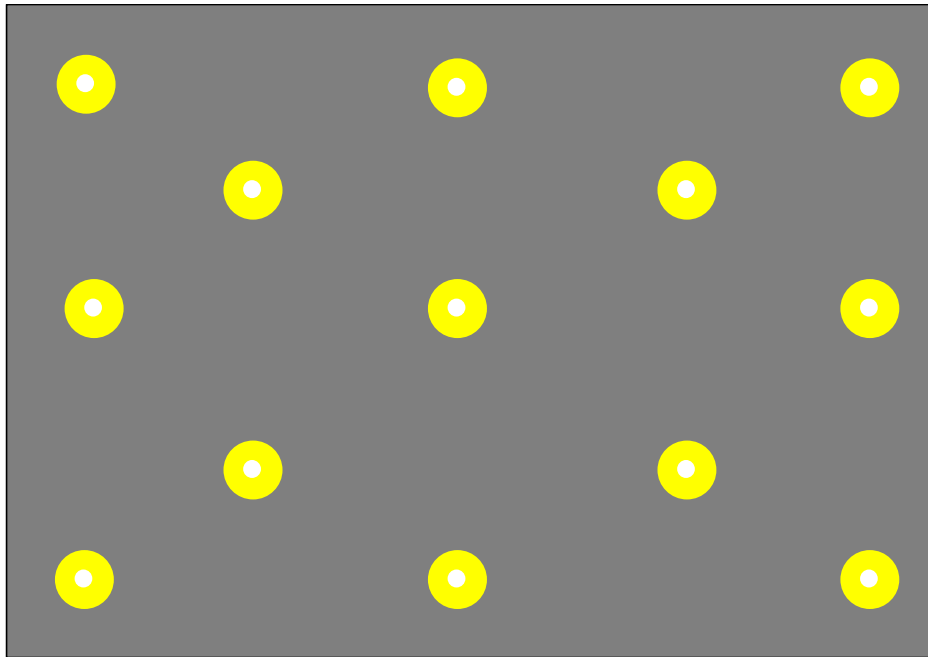
Experiment one consisted of 8 blocks each containing 50 trials, lasting approximately 50 minutes. Subjects were made aware that the first block of 50 trials was a practice block in order to ensure that subjects were familiar with the task and the gaze-contingent setup. During the seven non-practice blocks, performance on the task would affect the reward subjects would earn at the end of the experiment: rewards ranged from £5 (performance at chance level) to £15.

Each trial began with the onset of a central small grey fixation circle for 1,000 ms. After this 1000 ms delay, the (masked) numbers followed and remained on screen for up to 5,000 ms. During array presentation, participants freely moved their eyes across the screen and could choose at any time which

category (red or blue) exhibited the highest (lowest) average of the number array. The 5000 ms response deadline was chosen in order to increase the influence of the sampling policy on performance (i.e. subjects were not able to sample all elements on every trial).

We again used a framing manipulation in order to minimize potential biases towards a particular part of the number space. Depending on the framing condition, subjects had to decide which category (red or blue) had the highest or lowest mean. The response mapping, framing, and the colours (red and blue) were counterbalanced between subjects. A red text saying 'Too slow' would appear for 50 ms in the middle of the screen when subjects failed to respond within the 5000 ms interval. Auditory feedback (a 100 ms tone) was given on each trial: a low tone (400 Hz) indicated an incorrect response while a high tone (1200 Hz) indicated a correct response. The inter trial interval (ITI) consisted of a fixation cross that was presented in the middle of the screen for a total duration of 1000 ms.

*Eye tracking and gaze contingency.* We recorded all eye movement data using an SR research EyeLink 1000 eye-tracking system (SR Research, Ontario, CA, USA). The right eye was tracked monocularly at a sampling rate of 1000 Hz. Since gaze-contingent designs are very sensitive to accurate calibration, we chose to run a 13-point calibration and validation procedure at the start of each block (**Fig. 5.9**).



**Fig 5.9** Locations for the fixation targets for the thirteen point calibration / validation procedure, using the eyelink eye tracker system.

The calibration-validation procedure was repeated until a satisfactory precision was achieved. Additionally, a chin rest was placed at a fixed distance of 68 cm from the screen in order to minimize the influence of head movements on data quality. Missing data (due to blinks or eye movements away from the screen) were discarded (i.e., set to NaN) in all analyses. Eye movement traces were cut off at reaction time for all analyses described below. All of our analyses assume that the eye-tracker faithfully estimates the pixel position of the gaze on the screen.

Gaze contingency was implemented through a self-developed MATLAB script. A blurry Gaussian patch (of the same colour as the underlying number) with a radius of 25 x 25 pixels masked each two-digit number, in order to conceal its

identity while still conveying information about the category (red or blue; **Fig. 5.8**). A 50 x 50 pixels frame was defined around each two-digit number. When a fixation fell within this frame for more than 60 ms, the mask covering the stimulus was removed and the underlying stimulus was revealed. This minimum fixation time was implemented in order to prevent fast eye-movements (e.g. saccades) from revealing stimuli unintentionally. Whenever fixations no longer fell within the predefined frame, the mask once again concealed the stimulus.

*Design.* A reference number  $R$  was drawn randomly from a uniform distribution with the restriction that no one-digit or three-digit numbers would be present in the number array (e.g. 10-99). The numbers were drawn from two Gaussian distributions with mean  $\mu_1$  and  $\mu_2$  and SD  $\sigma$ . Participants were asked to compare the average of each category and the difficulty varied between trials:  $\mu_2 - \mu_1 = \pm 2$  or  $\mu_2 - \mu_1 = \pm 4$  in low and high  $\mu$  trials, respectively, and  $\sigma = 3$  or  $6$  in low and high  $\sigma$  trials, respectively. Each of the eight trial types occurred equally often, and their presentation order within a block was random. Sample values were resampled to ensure that their mean and variance fell within the desired value for that trial with a tolerance of 0.7 and 0.1, respectively.

### 5.3 Results

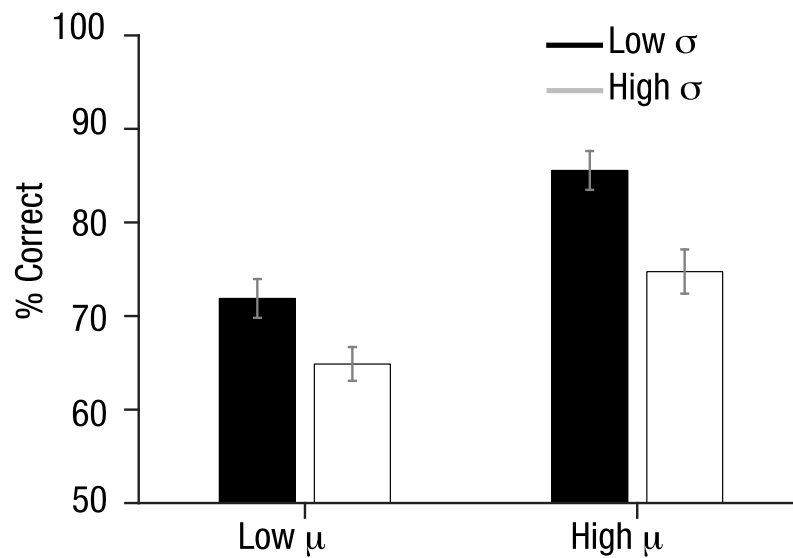
We excluded the data from the practice block in every analysis reported in this section. Moreover, we excluded trials where the deadline was exceeded (1.7%) and trials in which not a single stimulus was revealed (0.001%). Instead of using fixation durations, we opted to use the duration in milliseconds that a stimulus was revealed on the computer screen (i.e. automatically excluding any fixation that did not reveal a stimulus as well as more accurately capturing the time frame in which evidence was able to be integrated).

First, we will present descriptive results for accuracy, reaction times, sampling rates and dwell times in order to give the reader a general impression of how subjects perform the task. Next, we will focus on two major aspects of the sampling policy: selective weighting of information and information selection. Finally, we will briefly formulate and discuss theories and models to capture these two major aspects of our data. We will again collapse our results over the framing manipulation, as the pattern of results was similar in both frames. In order to improve readability, we will refer to the ‘winner’ and ‘loser’ category in the results section. The winner is the high category in the high frame and the low category in the low frame. The loser is the low category in the high frame and the high category in the low frame.

The overall mean accuracy was 74.3% ( $n = 36$ ) and the mean RT was 3380 ms ( $SD = 750$  ms). An average of 8 stimuli were revealed on each trial ( $SD = 2.98$ ) with a mean dwell time of 229.3 ms ( $SD = 106$  ms) per fixation. Moreover, of these 8 revealed stimuli, only 0.6 stimuli on average were resampled. Note that participants responded considerably before the response deadline (5000 ms) and that participants did not sample all of the available information (12 elements).

### **5.3.1 Effects of evidence strength $\mu$ and evidence reliability ( $\sigma$ )**

Performance depended on both the mean ( $\mu$ ) and variance of the distribution from which the numbers were drawn (**Fig. 5.10.**). Accuracy increased on high  $\mu$  trials compared with low  $\mu$  trials [ $F(1, 35) = 173$ ,  $P < 0.001$ ] and decreased on high  $\sigma$  (more variable evidence) trials compared with low  $\sigma$  trials [ $F(1, 35) = 92$ ,  $P < 0.001$ ], as previously reported for tasks involving averaging of a single category (De Gardelle & Summerfield, 2011; Michael et al., 2013; Vandormael, Castañón, Balaguer, Li, & Summerfield, 2017). However, unlike in these previous reports, the interaction between  $\mu$  and  $\sigma$  was significant [ $F(1, 35) = 6.1$ ,  $P < 0.05$ ].



**Fig 5.10 Influence of evidence strength ( $\mu$ ) and reliability ( $\sigma$ ) on accuracy.** Performance increased for higher evidence strength and higher evidence reliability. Bars are shown with 95% confidence intervals.

The influence of evidence strength ( $\mu$ ) and evidence reliability ( $\sigma$ ) on reaction times (RT) is depicted in **Table 5.11**. Reaction times decreased on high  $\mu$  trials compared to low  $\mu$  trials in both frames [ $F_{(1,35)} > 49.5$ ;  $p < 0.001$ ]. The RTs were not significantly influenced by  $\sigma$  ( $F_{(1,35)} = 3.3$ ;  $P > 0.05$ ).

**Table 5.11 Influence of evidence strength ( $\mu$ ) and reliability ( $\sigma$ ) on reaction times**

| $\sigma$      | High $\mu$       | Low $\mu$        | Effect on RT                             | Significance                         |
|---------------|------------------|------------------|------------------------------------------|--------------------------------------|
| High $\sigma$ | 3328 $\pm$ 74 ms | 3377 $\pm$ 77 ms | $\mu \uparrow \rightarrow RT \downarrow$ | $F_{(1,35)} = 49.5$ ,<br>$P < 0.001$ |
| Low $\sigma$  | 3313 $\pm$ 77 ms | 3442 $\pm$ 69 ms | n.s.                                     | $F_{(1,35)} = 3.3$ ,<br>$P > 0.05$   |

Mean reaction times (RTs) and SD are shown per condition. RTs decreased on high  $\mu$  trials compared to low  $\mu$  trials (both frames)

The influence of evidence strength ( $\mu$ ) and evidence reliability ( $\sigma$ ) on the sampling rate is depicted in **Table 5.12**. Neither  $\mu$  nor  $\sigma$  had a significant influence on the number of samples taken by participants ( $P > 0.05$ ).

**Table 5.12 Influence of evidence strength ( $\mu$ ) and reliability ( $\sigma$ ) on the number of samples**

| $\sigma$      | High $\mu$      | Low $\mu$       | Effect on # of samples | Significance                        |
|---------------|-----------------|-----------------|------------------------|-------------------------------------|
| High $\sigma$ | 8.11 $\pm$ 0.32 | 8.07 $\pm$ 0.31 | $\mu$ : n.s.           | $F_{(1,35)} = 2.86$ ,<br>$P > 0.05$ |
| Low $\sigma$  | 8.1 $\pm$ 0.32  | 7.99 $\pm$ 0.31 | $\sigma$ : n.s.        | $F_{(1,35)} = 0.92$ ,<br>$P > 0.05$ |

Mean number of samples and SD are shown per condition. Neither  $\mu$  nor  $\sigma$  had a significant effect on the number of samples taken by participants.

The influence of evidence strength ( $\mu$ ) and evidence reliability ( $\sigma$ ) on sample latency (dwell time) is depicted in **Table 5.13**. Evidence strength has a small but significant effect on sample latency: dwell times decreased on high  $\mu$  trials compared to low  $\mu$  trials [ $F_{(1,35)} = 5.5$ ;  $p < 0.05$ ]. Evidence reliability ( $\sigma$ ) had no significant influence on dwell time [ $F_{(1,35)} = 0.5$ ;  $P > 0.05$ ].

**Table 5.13** Influence of evidence strength ( $\mu$ ) and reliability ( $\sigma$ ) on sample latency (dwell time).

| $\sigma$      | High $\mu$         | Low $\mu$          | Effect on sample latency                           | Significance                       |
|---------------|--------------------|--------------------|----------------------------------------------------|------------------------------------|
| High $\sigma$ | 235.5 $\pm$ 3.8 ms | 237.6 $\pm$ 3.7 ms | $\mu \uparrow \rightarrow \text{dwell} \downarrow$ | $F_{(1,35)} = 5.5$ ,<br>$P < 0.05$ |
| Low $\sigma$  | 235.1 $\pm$ 3.7 ms | 236.8 $\pm$ 3.6 ms | $\sigma$ : n.s.                                    | $F_{(1,35)} = 0.5$ ,<br>$P > 0.05$ |

Mean sample latency and SD are shown per condition. Dwell times decreased on high  $\mu$  trials compared to low  $\mu$  trials;  $\sigma$  had no significant effect on the mean dwell time.

### 5.3.2 Framing effect on the sampling of information

Due to the framing manipulation, samples can be classified as frame consistent (e.g. a sample from the high category in the high frame condition) or frame inconsistent (e.g. a sample from the low category in the high frame

condition). The framing manipulation had a small but significant effect on the number of elements sampled from each category: more frame consistent elements were sampled than frame inconsistent (**Table 5.14**; [ $t_{(35)} = 4.8$ ;  $P < 0.001$ ]). Dwell times were not influenced by this manipulation [ $t_{(35)} = 1.5$ ;  $P > 0.05$ ].

**Table 5.14 Influence of framing on the number of samples taken from each category and dwell times**

| D.V.        | Frame consistent       | Frame inconsistent     | Effect                                                    | Significance                     |
|-------------|------------------------|------------------------|-----------------------------------------------------------|----------------------------------|
| # Samples   | $4.09 \pm 1.12$        | $3.96 \pm 1.14$        | Consistent $\uparrow \rightarrow$<br># samples $\uparrow$ | $t_{(35)} = 4.8$ ,<br>$P < 0.05$ |
| Dwell times | $234.1 \pm 24.8$<br>ms | $232.5 \pm 24.5$<br>ms | Consistent: n.s.                                          | $t_{(35)} = 1.5$ ,<br>$P > 0.05$ |

Mean number of samples / dwell time and their SD are shown per condition. More frame consistent samples were taken than frame inconsistent samples. Dwell times were not affected by frame consistency.

In summary, when the evidence strength was high (high  $\mu$ ), observers were more accurate, spent less time on each sample, and were faster to respond compared to when the evidence strength was low. The number of samples taken was not influenced by evidence strength. When the evidence reliability was high (low  $\sigma$ ) observers were more accurate. Evidence reliability

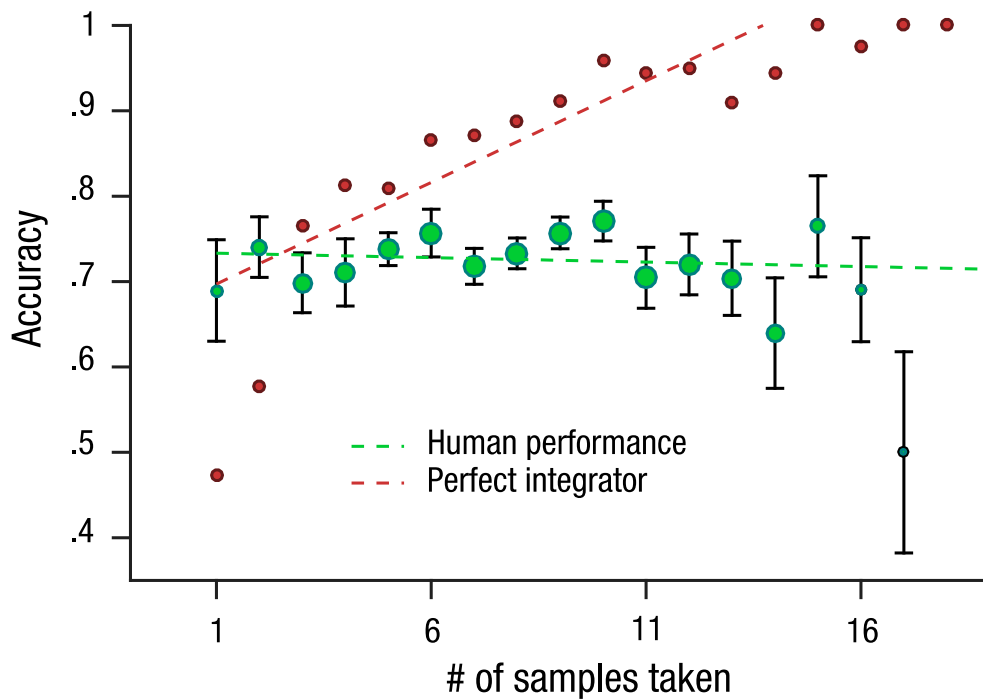
had no significant influence on sample latency, reaction times, or the number of samples taken.

### **5.3.3 Q1. Sampling for a prolonged period of time and the effect on decision accuracy**

Next, just like in chapter 3, we take a brief look at the effects of prolonged sampling on decision accuracy. Within the Drift-Diffusion Model (DDM) framework in perceptual decision making performance is optimised through repeated sampling (Bogacz et al., 2006). In the classic DDM model information accumulation proceeds without any decay of information. Therefore, if allowed to sample indefinitely, subject accuracy will always increase unless there is no difference between the stimuli. However as addressed in chapter 1, sampling and deliberating for a prolonged period of time does not always lead to improved decision making, partly due to limitations in our processing capacity (Broadbent, 1958; Burgess & Colborne, 1988; Dijksterhuis et al., 2005; Ebbinghaus, 1913; Nørretranders, 1998; Spiegel & Green, 1981; Wilson, 2004). Swensson (1972) were one of the first to explicitly test this ceiling effect on prolonged sampling (Swensson, 1972). Swenson showed an asymptotic effect of sampling duration and accuracy: sampling for a prolonged period of time increased accuracy but only up to a certain point (i.e. asymptotic performance). The asymptotic performance was acknowledged and led to the development of the diffusion with drift variance (DDV) and the leaky accumulator model (Ratcliff, 1978; Ratcliff & Rouder, 1998; Ratcliff, Van Zandt, & McKoon, 1999; Usher & McClelland, 2001).

Within the leaky accumulator model, the degree to which prolonged sampling increases performance reflects a balance between the benefit of sampling additional information and the loss of accumulated information. The amount of 'leak' depends on the task (e.g. the environment) and on individual differences (Glaze, Kable, & Gold, 2015; Ossmy et al., 2013; Usher & McClelland, 2001).

**Fig. 5.15** shows evidence for an asymptotic level of performance in our task: performance levels out in the 74% range and does not increase as more samples are taken (e.g. the number of samples taken is not a significant predictor of performance [ $t_{(35)} = -1.31$ ;  $P > 0.05$ ]).



**Fig 5.15 Asymptotic level of performance.** Human performance is shown in green, with the size of the bubbles reflecting the number of subjects per number of samples taken (i.e. most subjects take between 5-10 fixations). Performance of a perfect integrator without loss is shown in red. Note that the perfect integrator sees exactly the same decision information as the human observers, i.e. is endowed with an identical sampling policy. Weighted least-squares regression lines are shown via the dotted lines. Bars are shown with 95% confidence intervals.

Thus in the current dataset, sampling for a prolonged period of time does not have affect decision accuracy in a beneficial nor detrimental way.

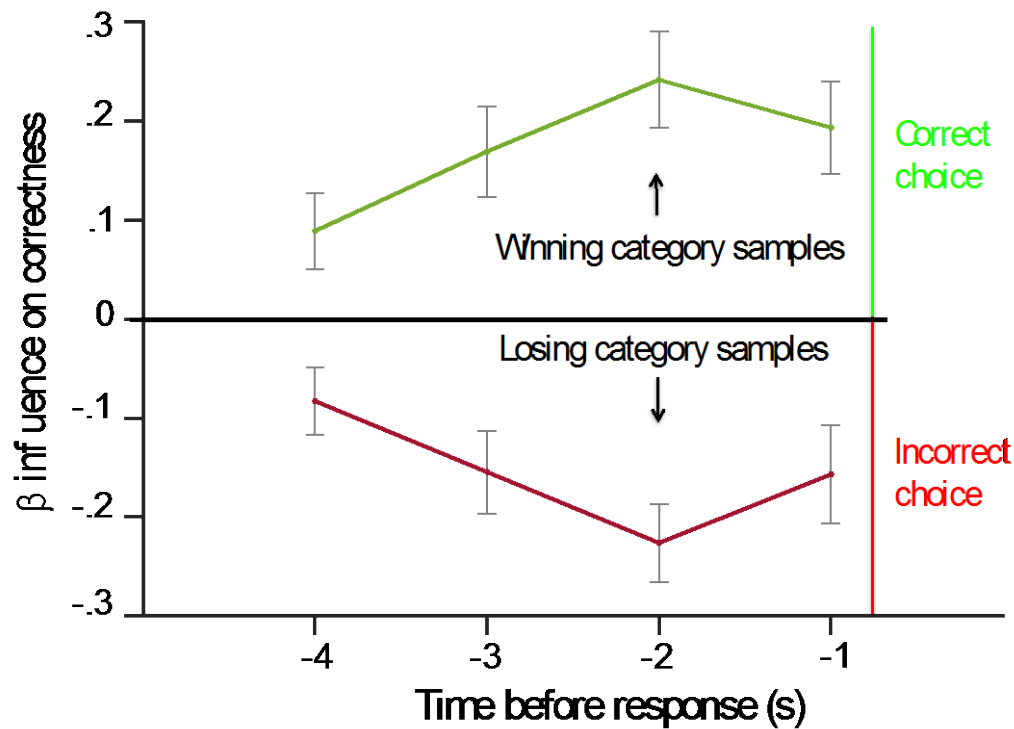
### 5.3.4 Sampling policy and the weighting of information: recency

This asymptotic level of performance and the difference in accuracy between human performance and a perfect integrator is often attributed to either primacy (early information has a strong impact on choice while later information is partially ignored) or recency (later information has a strong impact on choice while early information is forgotten). The leaky accumulator provides a mechanistic explanation for both the primacy and recency effects through mutual inhibition and leak (Dreher & Tremblay, 2016; Tsetsos et al., 2012; Usher & McClelland, 2001).

We turned to logistic regression in order to investigate the possible sources of the asymptotic level in performance. More specifically, we regressed accuracy (i.e. correctly chosen the winner = 1, incorrectly chosen the loser = 0) on the decision information (i.e. the sampled numbers) as a function of when they were sampled relative to when the decision was made (**Fig. 5.16**). The sampled decision information is normalized with respect to the trial mean in order to combine trials with different levels of evidence strength and reliability in all of our analyses. We controlled for individual variability in response times (due to the free-response paradigm) by aligning fixations relative to when the response occurred instead of when the stimulus appeared. The decision information was binned based on when a sample was revealed relative to when the response was made (i.e. 0-1000 ms, 1000-2000 ms, 2000-3000 ms or 3000-4000 ms before the response). Moreover, we created these regressors separately for the winning and losing category, resulting in 8

regressors. **Fig. 5.16** illustrates that samples closer to the response exhibited a larger influence on choice than fixations further away from the response. Furthermore, samples from the losing category made observers more likely to choose the losing category (i.e. negative  $\beta$  weights), while samples from the winning category increased the likelihood of choosing the winning category (i.e. positive  $\beta$  weights).

Thus, observers exhibited a recency effect: information sampled closer to the response was weighted more strongly than information sampled earlier on the trial.



**Fig 5. 16 Information sampled closer to the response has a stronger influence on choices (i.e. a recency effect).** The bottom four regressors correspond to losing category samples while the top four regressors correspond to winning category samples. Samples from the winning category are accompanied by positive  $\beta$  weights (e.g. higher likelihood of choosing the winning / correct category) while losing category samples are accompanied by negative  $\beta$  weights (e.g. higher likelihood of choosing the losing / incorrect category). Importantly, information sampled near the end of the trial was more influential (larger  $\beta$  weights) than information sampled at the beginning of the trial. Bars are shown with 95% confidence intervals.

### 5.3.5 Approximating the decision evidence: LPR

Earlier work from the Summerfield lab highlighted the importance of not only taking the evidence strength into account when studying perceptual decision making, but also the evidence reliability (De Gardelle & Summerfield, 2011). As described earlier, de Gardelle & Summerfield found independent effects of both the evidence strength and the reliability. More importantly, a probabilistic optimal model in which observers integrate the LPR (the logarithm of the posterior ratio) as a proxy for the integrated evidence captured all of their reported behavioural effects.

Furthermore, quantifying the integrated evidence as the LPR yielded a better behavioural fit than simply considering the average feature value ( $\mu$ ) or the ratio of the average feature value and its SD ( $\mu$  divided by  $\sigma$ ). The key idea is that in order to take into account both the evidence strength and reliability of a particular piece of information, we will need to transition from simply considering the stimulus value (i.e. number) in its native space to considering it in probability space. More specifically, each coloured number gives the observer an indication of the underlying state space. In our implementation, each coloured number gives the observer information about whether this colour is associated with the winning category or the losing category (e.g. dependent on the frame high or low).

Given a particular number, we can aim to estimate these probabilities ( $H_{\text{win}}$  or

$H_{lose}$ ). De Gardelle & Summerfield (2011) demonstrated how the LPR is the same as the LLR when the prior probabilities of each category are the same. They started with rewriting the posterior probabilities ( $p(H_{win} | x)$  and  $p(H_{lose} | x)$ ) via Bayes' rule:

$$\frac{p(H_{win} | x)}{p(H_{lose} | x)} = \frac{p(x | H_{win}) \times p(H_{win})/p(x)}{p(x | H_{lose}) \times p(H_{lose})/p(x)}$$

where  $x$  is a value in the decision space. When the prior probabilities of  $H_{win}$  and  $H_{lose}$  are the same (as is the case in their experiment and in ours), the posterior ratio and the likelihood ratio are identical:

$$\frac{p(H_{win} | x)}{p(H_{lose} | x)} = \frac{p(x | H_{win})}{p(x | H_{lose})}$$

In experiment one, observers have to infer which category is the winner in an environment with low or high evidence strength (mean) and low or high evidence reliability (variance). Therefore, the underlying state space will encompass the category (winner vs. loser), evidence strength (low vs. high), evidence reliability (low vs. high), and number (decision) space. Each trial starts with a flat prior across the state space, and the first sample updates this prior by the likelihoods of that first sample under the state space. The posterior is then normalized to ensure that the posterior sums to 1. The LPR<sub>s</sub> for each sample  $s$ , is then computed by marginalising over the evidence

strength, reliability and number space and calculating:

$$LPR_s = \log\left(\frac{H_{win}}{H_{lose}}\right)$$

This procedure is repeated for each subsequent sample with the prior of each sample being equal to the posterior of the previous sample.

Therefore, we will be using the LPR when referring to the cumulative decision evidence in subsequent analyses. Moreover when predicting subjects' choices in this experiment via logistic regression, the LPR was more predictive of choices than the average feature value ( $[t_{(35)} = 7.4; P < 0.001]$ ), and the ratio of the average feature value and the SD ( $[t_{(35)} = 6.5; P < 0.001]$ ).

### **5.3.6 Q2. Frequency bias: preference for the overly sampled category**

Next, we looked at whether items that were fixated more frequently were chosen more often (Krajbich et al., 2010; Krajbich & Rangel, 2011; Shimojo et al., 2003). Importantly, in our gaze-contingent design we are not only able to characterise which items were sampled and for how long, but are also able to more accurately capture the exact 'information value' observers have accrued on each trial while having an objectively correct answer.

In order to investigate whether observers were biased towards the overly sampled category (i.e. a 'frequency bias'), we again turned to a probit regression. We used probit regression to predict observers' accuracy (i.e.

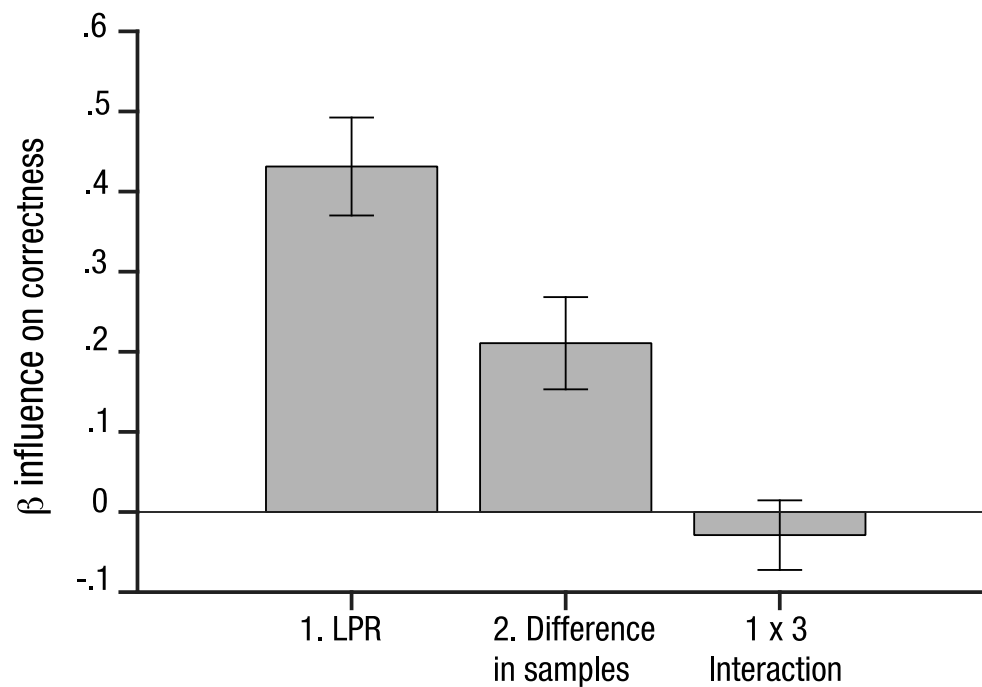
choosing the winner versus choosing the loser) as a function of three sets of predictors: (i) the cumulative evidence captured by the LPR after the last sample, (ii) the frequency bias measured by the difference between the number of items sampled from the winner minus the number of items sampled from the loser, and (iii) the interaction between these two quantities. The regression model we used was as follows, where in each case:

$$p(\text{correct}) = [b + w1 \cdot LPR_{last} + w2 \cdot (n_{winner} - n_{loser}) + w3 \cdot (LPR \cdot (n_{winner} - n_{loser}))]$$

where  $LPR_{last}$  denotes the LPR after the last sampled element before making a response, and  $n_{winner}$  denotes the number of elements sampled from the winner (and vice versa for  $n_{loser}$ )

Crucially, this regression model allows us to separate out the influence of the decision information (i.e. the LPR) and the frequency bias (i.e. the number of samples taken from the winner minus the number of samples taken from the loser). The LPR positively predicts the likelihood of correctly choosing the winner, providing extra support for the LPR as an approximation of the cumulative evidence [ $t_{(35)} = 13.8$ ;  $P < 0.001$ ; **Fig. 5.17**]. More importantly, the category that was sampled more frequently was also chosen more often (i.e. a frequency bias, [ $t_{(35)} = 7.2$ ;  $P < 0.001$ ]). Thus, over and above the cumulative evidence, the overly sampled category was chosen more often (e.g. an LPR independent frequency bias). Moreover, the interaction of the cumulative evidence (LPR) and the frequency bias was not significant [ $t_{(35)} =$

1.3;  $P > 0.05$ ]. This suggests that the frequency bias is not a multiplicative effect on the LPR as suggested by the a-DDM model, but rather an independent additive bias.



**Fig 5.17 The category that was sampled more frequently was chosen more often (i.e. a frequency bias).** The main effects of the LPR and the difference in samples between the winner and loser were both significant. The LPR (calculated after taking the last sample on each trial) had a significant positive effect on the likelihood of choosing the winner. The difference in the total number of samples taken from the winner minus the total number of samples taken from the loser exhibited a significant positive effect on the likelihood of choosing the winning category (i.e. a frequency bias). Thus, the category that was sampled more frequently was chosen more often, after taking the cumulative evidence into account. Bars are shown with 95% confidence intervals.

### 5.3.7 Q3. Consistency effect: switching between categories

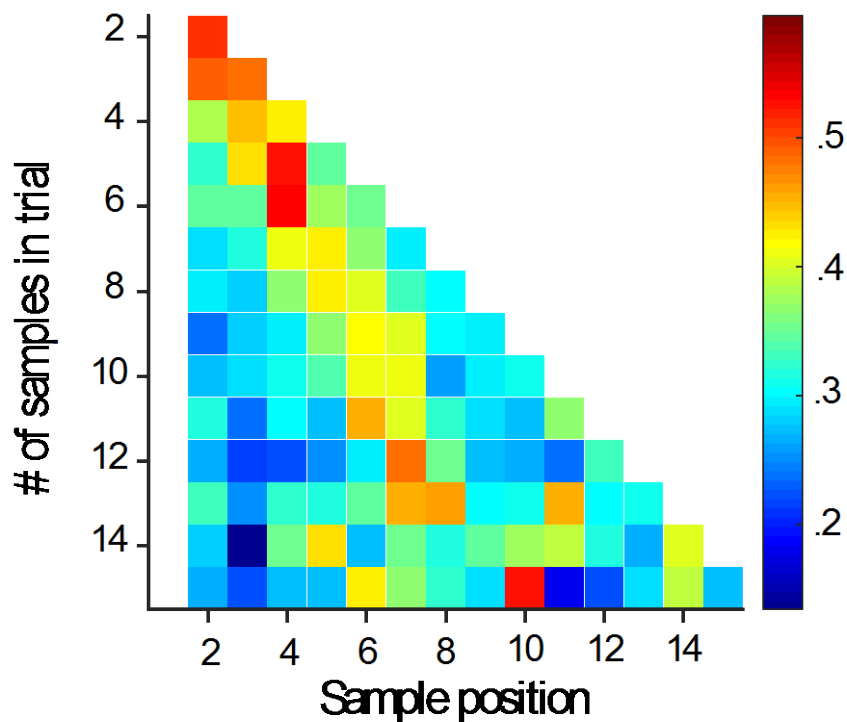
Finally, we identified samples on which observers switched between the two categories (e.g. colours) within a trial. If a sample *S* came from the same category as the previous sample (*S*-1), then sample *S* was denoted a 'stay' sample (e.g. Red-Red). However, if a sample *S* came from a different category as the previous sample (*S*-1), then sample *S* would be considered a 'switch' sample (e.g. Red – Blue). The first sample of a trial was excluded, as it did not have an *S*-1 sample. A Wilcoxon signed rank test indicated that participants were significantly less likely than chance to switch between the two categories: 32.2% of all fixations were switches ( $Z_{28} = -4.8$ ,  $p < 0.001$ ). These results provide a first piece of evidence that participants did not switch constantly between distributions (i.e. round-wise sampling), which should result in a 'switch' proportion significantly higher than 50% (only 4 out of 36 subjects had a switch proportion higher than 50%, the highest being 59.5%); nor did observers merely sample the element closest to their previous fixation, which would result in a switch proportion closer to 50%. This observation is particularly interesting because as we saw previously: on average not all elements are sampled (~ 8/12 elements are sampled). Thus, merely sampling the closest element might lead to being able to consider more elements in the decision before the response deadline. However, observers seem to switch less than one would expect by chance. **Table 5.18** shows the influence of evidence strength and reliability on the number of switches made by observers. Neither evidence strength nor evidence reliability significantly influenced the number of switches made by observers (**Table 5.18**).

**Table 5.18** Influence of evidence strength ( $\mu$ ) and reliability ( $\sigma$ ) on the number of switches.

| $\sigma$      | High $\mu$      | Low $\mu$       | Effect on switches     | Significance                       |
|---------------|-----------------|-----------------|------------------------|------------------------------------|
| High $\sigma$ | $2.27 \pm 0.11$ | $2.26 \pm 0.1$  | $\mu : \text{n.s.}$    | $F_{(1,35)} = 0,$<br>$P > 0.05$    |
| Low $\sigma$  | $2.23 \pm 0.11$ | $2.24 \pm 0.11$ | $\sigma : \text{n.s.}$ | $F_{(1,35)} = 0.95,$<br>$P > 0.05$ |

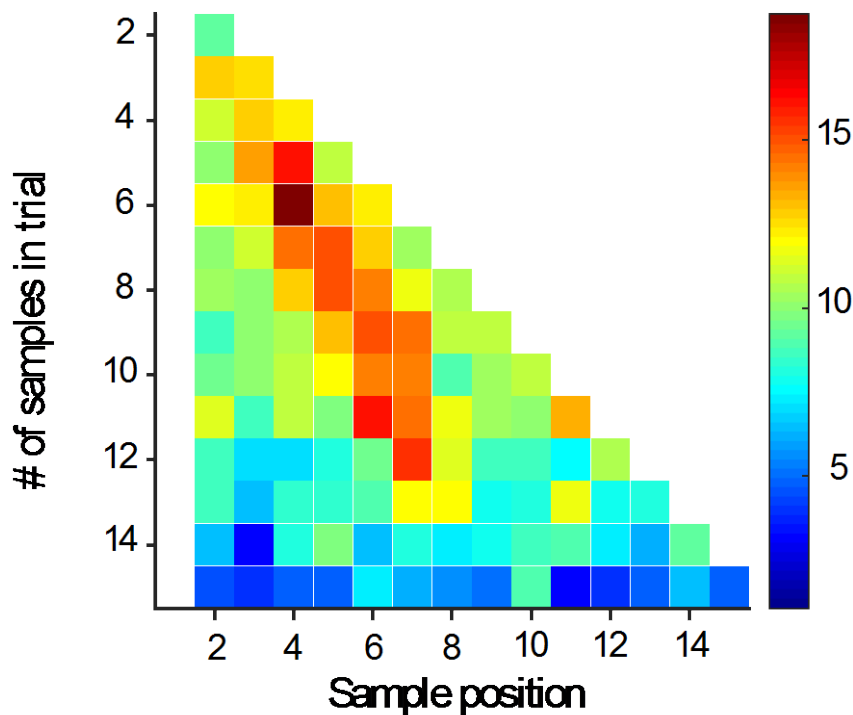
Mean number of switches and SD are shown per condition. Neither  $\mu$  nor  $\sigma$  had a significant effect on the number of switches.

Interestingly, a further piece of evidence supporting the idea that observers are avoiding to switch between categories is shown in **Fig. 5.19**. Here, the probability of a sample being a switch is plotted as a function of the sample position within the trial (x – axis), for different trial lengths (i.e. the number of samples taken; y - axis). Interestingly, the probability of switching was higher near the middle of the trial length indicated by the warmer colours (e.g. closer to red on a blue-red colour scale).



**Fig 5.19** The probability of a sample being a switch as a function of sample position, split by trial length. The probability of a sample being a switch is shown on a blue-green-red spectrum, with blue indicating a lower probability of the sample being a switch and red indicating a higher probability of the sample being a switch. Samples are more likely to be a switch near the middle of the trial length.

This effect is even more apparent when we multiply these probabilities with the number of subjects that have trials in each bin on the y – axis (**Fig. 5.20**). This figure is merely shown to illustrate that for some trial lengths only very few observers have data and should thus play a smaller role when considering the previous figure.



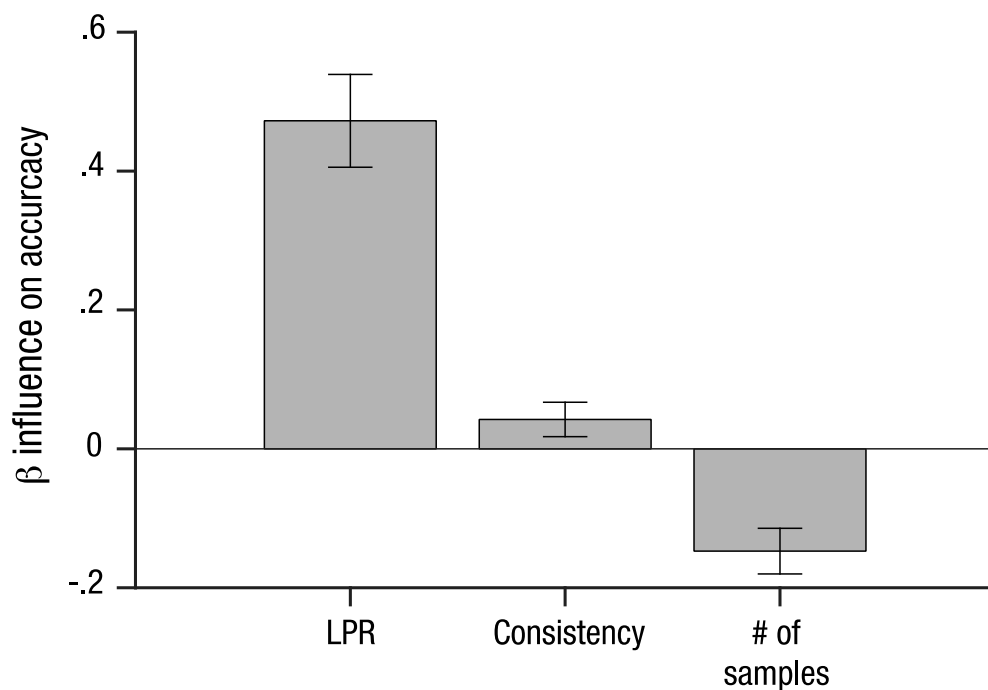
**Fig 5.20** The probability of a sample being a switch as a function of sample position, split by trial length and scaled by the number of subjects that have trials in each bin on the y - axis. The probability of a sample being a switch is shown on a blue-green-red spectrum, with blue indicating a lower probability of the sample being a switch and red indicating a higher probability of the sample being a switch. Samples are more likely to be a switch near the middle of the trial length.

Thus, there is evidence to support the hypothesis that in this dataset observers switch infrequently, and at the midpoint of their drawn samples, i.e. using a ‘summary sampling’ policy. Why might this be the case? A first hypothesis is that it is just simply more difficult to integrate the evidence for two categories simultaneously (i.e. similar to a switch-cost demonstrated when observers have to switch between two tasks (Jersild, 1927). This

hypothesis would receive support from the earlier finding that observers are more likely to switch near the middle of a trial sequence (e.g. first accumulating evidence from one category before turning to the other category, thereby minimising the potential ‘switch-cost’). Therefore, we aimed to investigate whether switching between categories reduced accuracy (i.e. provoked a switch cost). To do this, we calculated how ‘consistent’ the sampling was on each trial. Consistency was defined as the sum of samples that were sampled sequentially from the same category divided by the total number of samples on a given trial (e.g. the number of stay samples divided by the total number of stay + switch samples). Thus, more consistent trials consisted of proportionally less switches. We again used probit regression to predict observers’ accuracy (0 = incorrect, 1 = correct) as a function of three sets of predictors: (i) the cumulative evidence captured by the LPR after the last sample, (ii) the proportion of consistent samples out of all samples, and (iii) the number of elements sampled. We normalised the consistency by the number of samples in order to disentangle the consistency effect from any effect related to the number of samples, as more switches will be made on trials with more samples. The regression model we used was as follows, where in each case:

$$p(\text{correct}) = [b + w1 \cdot LPR_{last} + w2 \cdot (n_{\text{consistent}} / n_{\text{samples}}) + w3 \cdot (n_{\text{samples}})]$$

where  $\phi$  denotes the cumulative normal distribution,  $LPR_{last}$  denotes the LPR after the last sampled element before making a response,  $n\_consistent$  denotes the proportion of samples identified as consistent (e.g. 'stay') samples, and  $n\_samples$  denotes the number of elements that were sampled. The LPR again positively predicts the likelihood of being correct [ $t_{(35)} = 13.8$ ;  $P < 0.001$ ; **Fig. 5.21**]. More importantly, the consistency had a positive effect on accuracy (i.e. a consistency bias, [ $t_{(35)} = 3.5$ ;  $P < 0.01$ ]). Moreover, the number of samples had a negative effect on accuracy [ $t_{(35)} = -8.8$ ;  $P < 0.001$ ].



**Fig 5.21 The consistency effect on accuracy.** The probability of being correct goes up for more consistent sampling, indicated by the positive effect of consistency (e.g. the proportion of consistent or stay samples). Importantly, this effect is independent from the total number of samples taken, which shows a negative effect on accuracy. Bars are shown with 95% confidence intervals.

### **5.3.8 Computational modeling: recency, frequency, & consistency bias**

Next, we turned to computational modeling in order to investigate the underlying mechanisms of the recency, frequency, and consistency bias. We aimed to capture these effects with relatively simple models (**Appendix 1.1**). We used two different theoretical models: (i) a model that estimated the probability of either category being the winner using Bayesian inference implementing different variants of a memory leak term, and (ii) the adaptive gain model in which the gain of information processing adapts as a function of the previously sampled information (adaptation from (Cheadle et al., 2014)).

However, neither model was able to accurately capture all aforementioned aspects of the data (**Appendix 1.1**).

### **5.3.9 Q4. Drivers of the sampling policy: information selection**

So far we have only looked at descriptive aspects of the sampling policy and the effects it has on decision making. Are elements that are sampled closer to the decision weighted more heavily? Are observers more biased towards the overly sampled category? Does a sort of switch cost hamper evidence integration? However, additional questions remain. In this next section we will take a closer look at what the potential drivers of the sampling policy are when sampling information in a 2-AFC perceptual decision making task.

#### **5.3.9.1 Sampling policy: asymmetrical sampling and the equity policy**

First, we investigated the key factors that influence from which category (winner = 1 vs. loser = 0) the next item will be sampled. More specifically, we were mainly interested in two key factors that might influence which category

is sampled: (i) do sampling choices depend on the evidence accumulated so far (i.e. confirmation bias), and (ii) are sampling choices influenced by the relative difference of the number of samples taken from each category? Our previous finding that observers employed summary sampling (i.e. switches were more likely near the middle of the sample sequence) raises the suggestion that observers aim to minimize switches and to also sample approximately the same number of samples from each category (i.e. an 'equity' policy).

We used probit regression to predict the category of observers' next sample  $S+1$  (winning versus losing category) as a function of five sets of predictors: (1) the cumulative evidence captured by the LPR after the sample  $S$ , (2) the total number of samples after the sample  $S$ , (3) the total number of switches after sample  $S$ , (4) the difference in number of samples taken from the winner - loser after sample  $S$ , and (5) the interaction between the cumulative evidence and this difference in number of samples taken from each category. The cumulative evidence predictor in this regression captures a form of confirmation bias: a significant positive effect would indicate observers are more likely to sample elements that confirm their current hypothesis. The total number of samples / switches could signify a bias towards sampling the winner (positive effect) / loser (negative effect) as the observer samples more elements / switches more often. The total number of switches was included in the regression because perhaps participants were less likely to know the identity of the winner when they made proportionally more switches (e.g. due

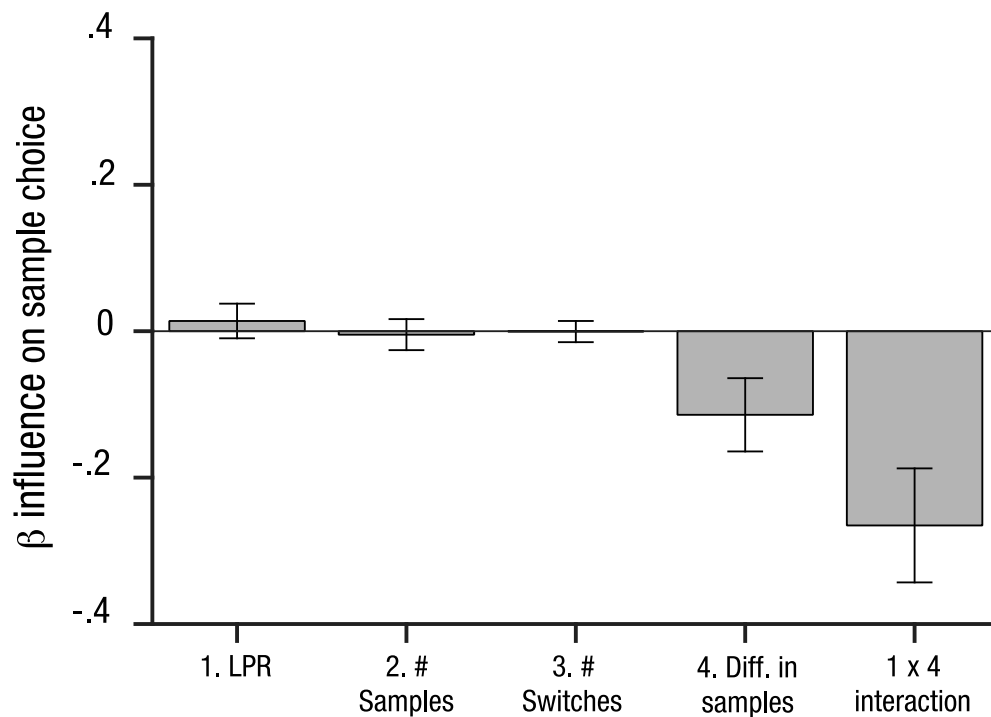
to the identified switch costs in this chapter). The difference in the number of samples taken from the winner - loser and its interaction with the cumulative evidence capture our second main variable of interest in this regression. The regression model we used was as follows, where in each case:

$$p(High_{S+1}) = \phi[b + w1 \cdot LPR_S + w2 \cdot (n\_samples_S) + w3 \cdot (n\_switches_S) + w4 \cdot (n\_winner_S - n\_loser_S) + w5 \cdot (LPR_S \cdot (n\_winner_S - n\_loser_S))]$$

where  $\phi$  denotes the cumulative normal distribution,  $LPR_S$  denotes the LPR after the sample  $S$  before sampling sample  $s + 1$ ,  $n\_samples_S$  denotes the total number of samples taken after sample  $S$ ,  $n\_switches_S$  denotes the total number of switches taken after sample  $S$ , and  $n\_winner_S - n\_loser_S$  denotes the difference in number of samples taken from the low and high category after sample  $S$ .

The evidence accumulated up to and including sample  $S$  ( $LPR_S$ ) did not significantly predict sampling choices at time point  $S+1$  [ $t_{(35)} = 1.15$ ;  $P > 0.05$ ;

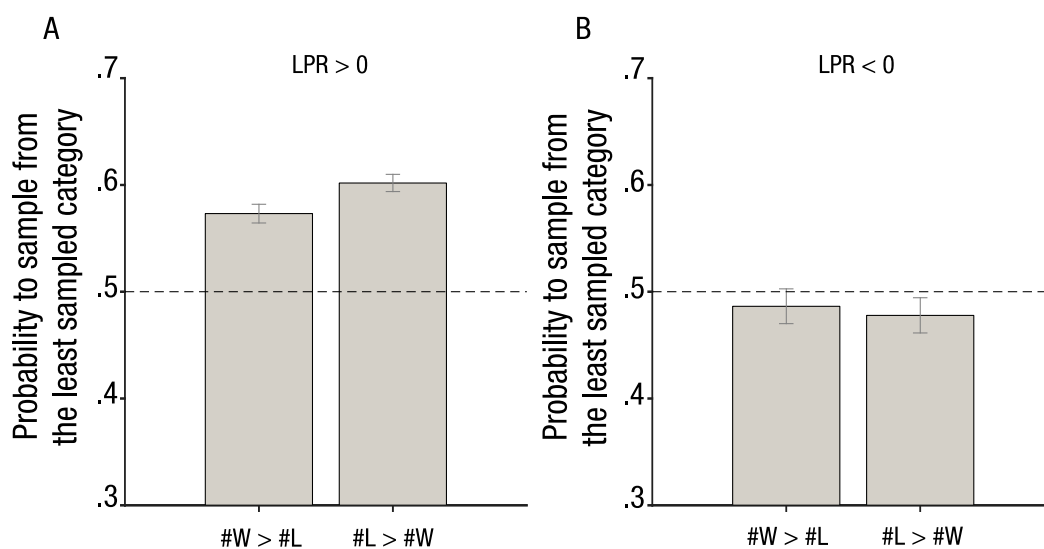
**Fig. 5.22].**



**Fig 5.22 Predicting the category of the next sample  $S + 1$ .** The difference in number of samples taken from the winning - losing category at sample  $S$  significantly predicted sampling choice at time  $S + 1$ . Moreover, the interaction between the cumulative evidence and the difference in samples taken from each category was significant. Bars are shown with 95% confidence intervals.

Neither the number of samples or switches taken up to and including sample  $S$  had a significant influence on sample choice of the next sample  $S + 1$  [ $t_{(35)} = -0.44$ ;  $P > 0.05$  and  $t_{(35)} = -0.07$ ;  $P > 0.05$ , respectively]. More importantly, the difference in number of samples taken from the winning – losing category at sample  $S$  significantly predicted sampling choice at time  $S + 1$  [ $t_{(35)} = -4.47$ ;  $P < 0.001$ ]. Moreover, the interaction between the LPR and the difference in samples taken from each category was significant [ $t_{(35)} = -6.7$ ;  $P < 0.001$ ].

Thus, the cumulative evidence on itself did not predict sampling choices, but its interaction with the difference in samples taken did. Thus, observers tend to preferentially sample elements from the category they have sampled the least from so far, and this effect is influenced by the cumulative evidence up to that time point. **Fig. 5.23** shows the probability of sampling from the winner (W) / loser (L) as a function of LPR (left vs. right plot) and the difference in samples taken from each category in order to better understand this significant interaction effect.



**Fig 5.23** The probability of sampling from the least sampled category as a function of the difference in samples taken from the winning and losing category. (A) The difference in samples effect is strongest when the LPR is positive (e.g. correct hypothesis about the identity of the winner). (B) The difference in samples effect is reduced and no longer significant when the LPR is negative (e.g. incorrect hypothesis about the identity of the winner). Bars are shown with 95% confidence intervals.

As **Fig. 5.23** shows, the degree to which the difference in samples taken from each category influences the sampling policy is strongest when the LPR is positive (e.g. you have strong evidence for selecting the winner), and is no longer significant when the LPR is negative.

In a subsequent step we aimed to see if we could recreate earlier aspects of the data by building in a model sampling policy. More specifically, we set out to recreate the tendency to switch categories near the middle of the sample sequence (e.g. near the middle of the trial length).

We wondered whether the equity policy (e.g. a preference to sample approximately the same number of samples from each category) shown by observers might actually be a consequence of another sampling mechanism. Taking the equity policy and the fact that observers tend to make switches near the middle of the sample sequence together might indicate that observers are sampling one category until they reach a set level of certainty about its estimate. Once this certainty criterion is reached, a switch is made and evidence sampling & integration for the other category starts.

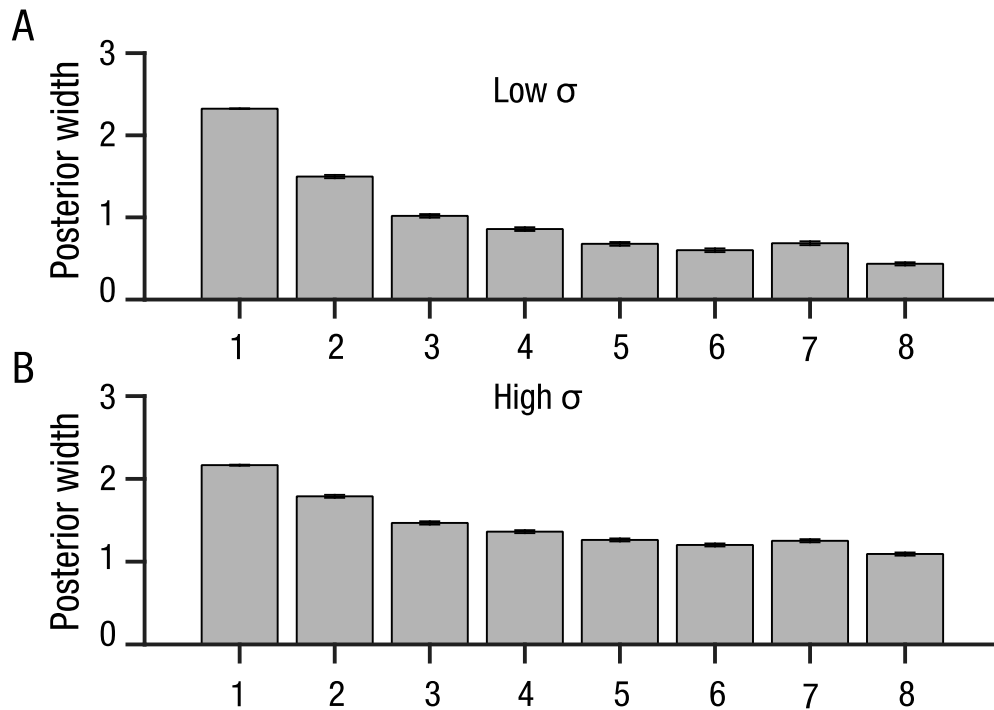
#### **5.3.10 Q5. Category uncertainty as a driver of the sampling policy**

How can we quantify the uncertainty surrounding a category estimate after each sample? To do this, we again turned to a Bayesian framework. The

decision environment in our experiment is characterised by the evidence strength ( $\mu$ ) and a given evidence reliability ( $\sigma$ ). For simplicity we assumed that observers' priors for each trial depend on the distribution of the evidence across the entire experiment, instead of having the prior evolve as the experiment goes on. More specifically, we employed a flat (uniform) prior across the range of numbers for different levels of  $\sigma$ . On each trial and after each sample belonging to category one, this prior is updated by a likelihood function representing the likelihood that the evidence strength of category one is  $X$  and the evidence reliability is  $Y$ , given sample  $S$ . Via Bayes-rule, this update results in the posterior of category one after sample  $S$ , and this posterior serves as the prior for the next sample of category one (so on and so forth). After each sample, we find the location at which the posterior is maximal (e.g. the most likely  $\mu$  and  $\sigma$ ). The most likely  $\sigma$  in this formulation, henceforth referred to as the 'posterior width', is used to quantify the uncertainty surrounding the category estimate after each new sample on a given trial.

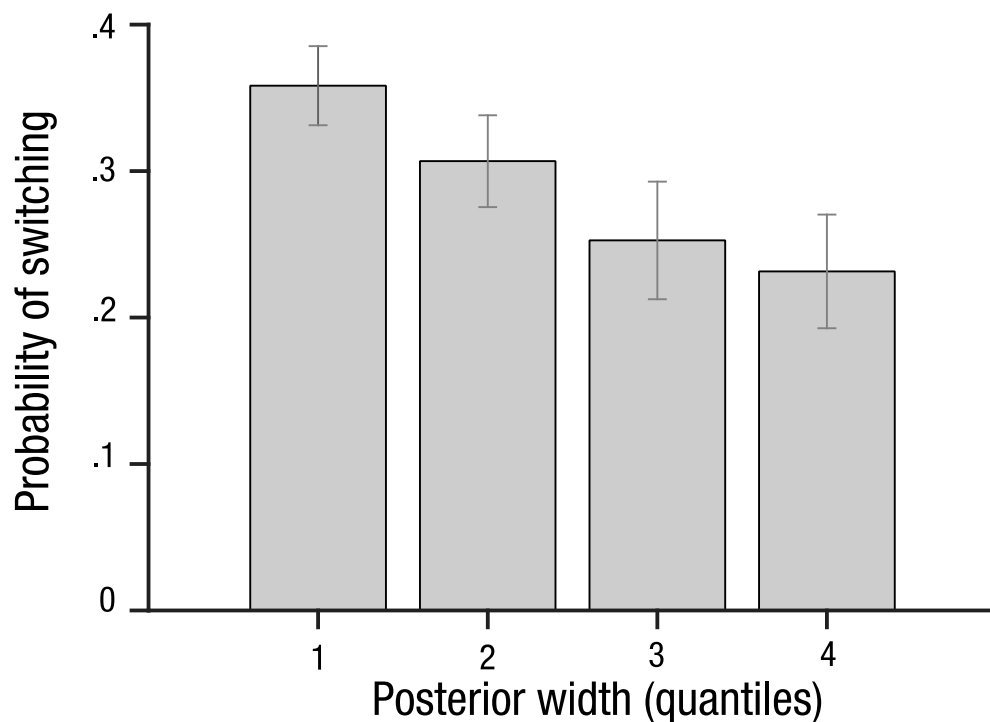
Next, we ran a computational simulation in which we simulate how this posterior width decreases as an observer samples more information, split for low and high evidence reliability trials. For this simulation we used the same environmental statistics used in the experiment (same  $\mu$ ,  $\sigma$ , range of stimuli...) but utilised a completely random sampling policy. **Fig. 5.24** shows how the posterior width (i.e. uncertainty) decreases as more samples are taken, and that this rate of decrease in uncertainty is faster in trials with more reliable

evidence (lower  $\sigma$ ) compared to trials with less reliable evidence (higher  $\sigma$ ).



**Fig 5.24 Evolution of the posterior width as more samples are taken (computational simulation).** The posterior width decreases as more samples are taken, and this rate of decrease in uncertainty is faster in trials with more reliable evidence (lower  $\sigma$ , top panel) compared to trials with less reliable evidence (higher  $\sigma$ , bottom panel).

If observers use an uncertainty criterion to guide their sampling policy (switch vs. stay), we would expect to see different probabilities of switching as a function of this posterior width. More specifically, one would expect a higher probability of switching when observers are more certain about their estimate about a category (i.e. smaller posterior width). This is indeed what we see [ $F(2.46, 86.01) = 70.32$ ,  $P < 0.001$ ] (**Fig. 5.25**).



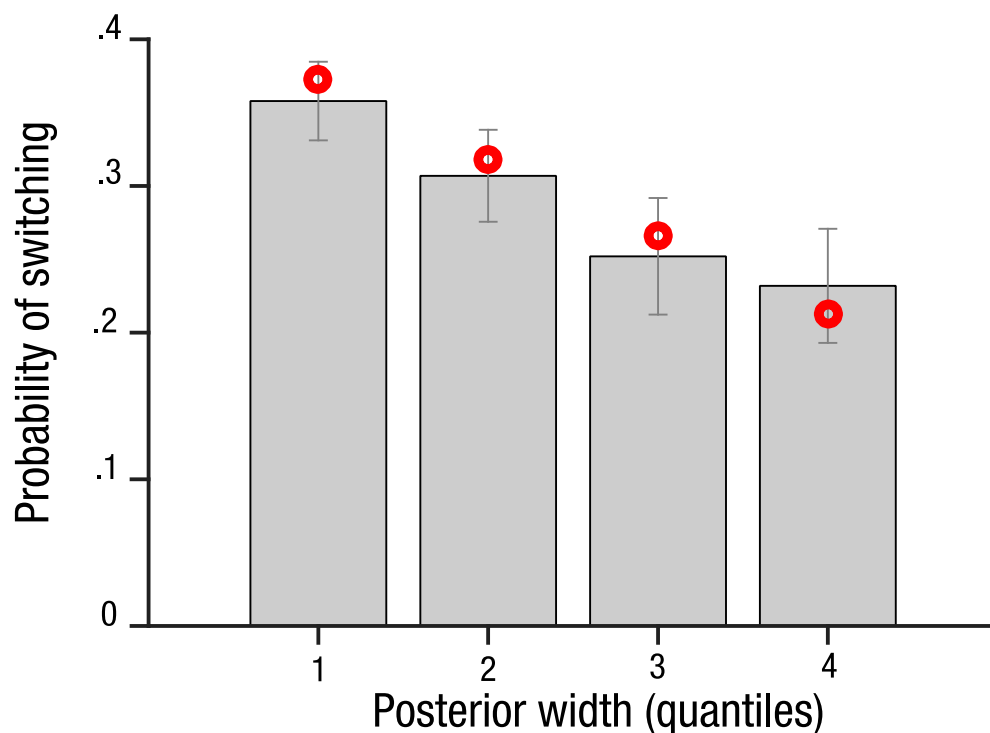
**Fig 5.25 The probability of switching as a function of posterior width.**

The probability of switching is higher when the posterior width is lower (e.g. lower uncertainty). Quantile one has the lowest uncertainty. Bars are shown with 95% confidence intervals.

In a next step, using probit regression we investigated whether the posterior width is a significant predictor of making a switch. We used a similar approach as when we predicted the category (winner = 1 vs. loser = 0) of the next sample. More specifically, we predicted whether the next sample was a switch (= 1) or not (= 0) as a function of the posterior width. As expected from figures 5.24 and 5.25, the posterior width had a significant negative effect on the likelihood of switching between categories [ $t_{(35)} = -10.6$ ;  $P < 0.001$ ]. Moreover,

this effect survived the inclusion of the number of samples viewed in the regression [ $t_{(35)} = -15.7$ ;  $P < 0.001$ ].

Next, we tested whether a simple computational model with two free parameters: 1) a threshold on the posterior width (e.g. switch when the posterior width dips below a threshold) and 2) a noise parameter modelling the trial by trial variability of this threshold. This simple model could predict the probability of switching as a function of the posterior width. **Fig 5.26** shows that such a model can indeed replicate the probability of switching as a function of posterior width analysis at both the subject and average level.



**Fig 5.26** The probability of switching as a function of posterior width as predicted by a threshold model. Model fits are presented in red. Bars are shown with 95% confidence intervals.

In this chapter we have looked at the effects of the sampling policy on choices and at potential drivers of the sampling policy. In these analyses, the sampling policy has been looked at from the perspective of information harvesting (e.g. which elements are sampled and its interplay with decision making). However, two other aspects are additionally important when considering the sampling policy in decision making: how much time do observers spend on sampling a particular stimulus (e.g. dwell times) and what drives observers to decide that they have sampled enough information and respond.

### **5.3.11 Sampling policy: dwell times?**

We investigated the key factors that influence how long observers viewed a specific stimulus (e.g. the dwell time). More specifically, we were mainly interested in three key factors that might influence dwell times: (i) does the sampling latency depend on the absolute accumulated evidence so far (i.e.  $|LPR|$ ); absolute LPR was used because we hypothesized that dwell times would be affected by the amount of sampled evidence, irrespective of whether this evidence was in favour of the correct response. (ii) is the sampling latency influenced by the relative difference of the number of samples taken from each category. (iii) is the sampling latency influenced by the estimated uncertainty at the moment of sampling.

We used linear regression to predict observers' sample latency of sample  $S$  as a function of six predictors: (1) the absolute cumulative evidence captured

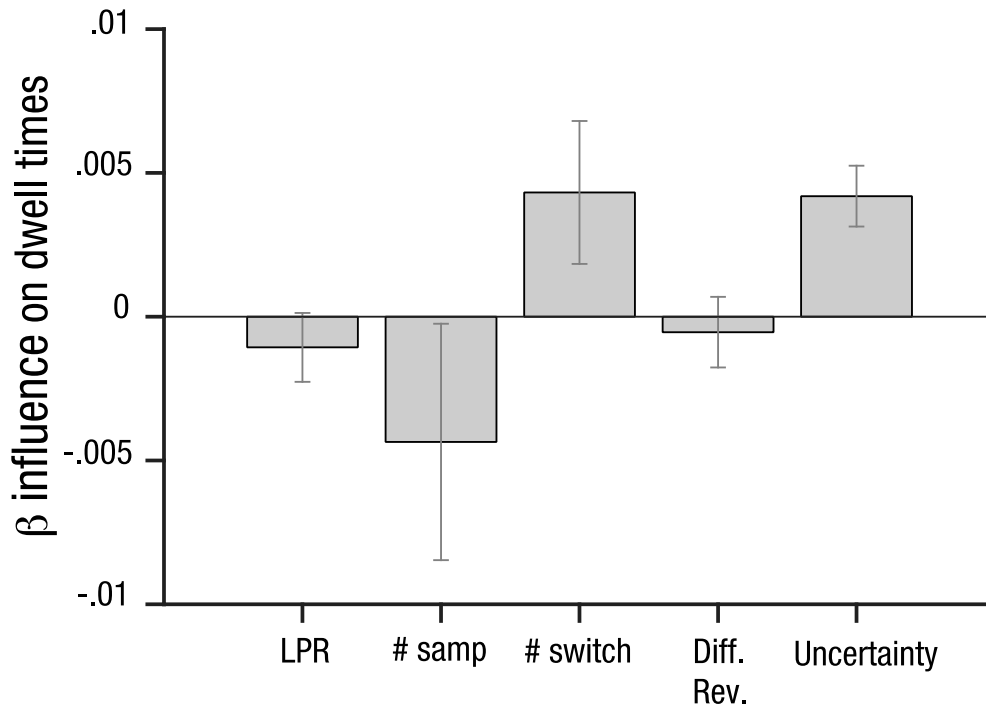
by the LPR after the sample  $S$ , (2) the total number of samples after the sample  $S$ , (3) the total number of switches after sample  $S$ , (4) the difference in number of samples taken from the winner – loser after sample  $S$ , and (5) the posterior width of the category that was sampled.

The regression model we used was as follows, where in each case:

$$p(Dwell_S) = \phi[b + w1 \cdot |LPR_S| + w2 \cdot (n\_samples_S) + w3 \cdot (n\_switches_S) + w4 \cdot (n\_winner_S - n\_loser_S) + w5 \cdot (Posterior_S)]$$

where  $\phi$  denotes the cumulative normal distribution,  $|LPR_S|$  denotes the absolute LPR after the sample  $S$ ,  $n\_samples_S$  denotes the total number of samples taken after sample  $S$ ,  $n\_switches_S$  denotes the total number of switches taken after sample  $S$ , and  $n\_winner_S - n\_loser_S$  denotes the difference in number of samples taken from the winning and losing category after sample  $S$ .  $Posterior_S$  denotes the estimated width of the posterior after sample  $S$ .

The dwell time was significantly affected by the number of switches: dwell time increased when more switches were made [ $t_{(35)} = 3.3$ ;  $P < 0.001$ ; **Fig. 5.27**]. Additionally, the uncertainty about the sampled category had a significant positive effect on dwell time [ $t_{(35)} = 7.8$ ;  $P < 0.001$ ]. Thus, when observers were more uncertain they dwelled on the evidence for longer. This effect was significant even when including the total number of samples, which had a significant negative effect on dwell time [ $t_{(35)} = -2.1$ ;  $P < 0.05$ ].



**Fig 5.27 Predicting dwell times.** The dwell time at sample  $S$  was significantly affected by number of switches made, the category uncertainty (e.g. the posterior width), and the number of samples taken after sample  $S$ . Bars are shown with 95% confidence intervals.

### 5.3.12 Sampling policy: when to stop sampling?

Finally, we investigated the key factors that influence when observers decide to stop sampling additional information and respond. We were interested in the same set of predictors used to predict the dwell times. We used probit regression to predict whether observers' would respond at sample  $S$ , as a function of six predictors: (1) the absolute cumulative evidence captured by the LPR after the sample  $S$ , (2) the total number of samples after the sample  $S$ , (3) the total number of switches after sample  $S$ , (4) the difference in

number of samples taken from the winner - loser after sample  $S$ , and (5) the posterior width of the category that was sampled on sample  $S$ .

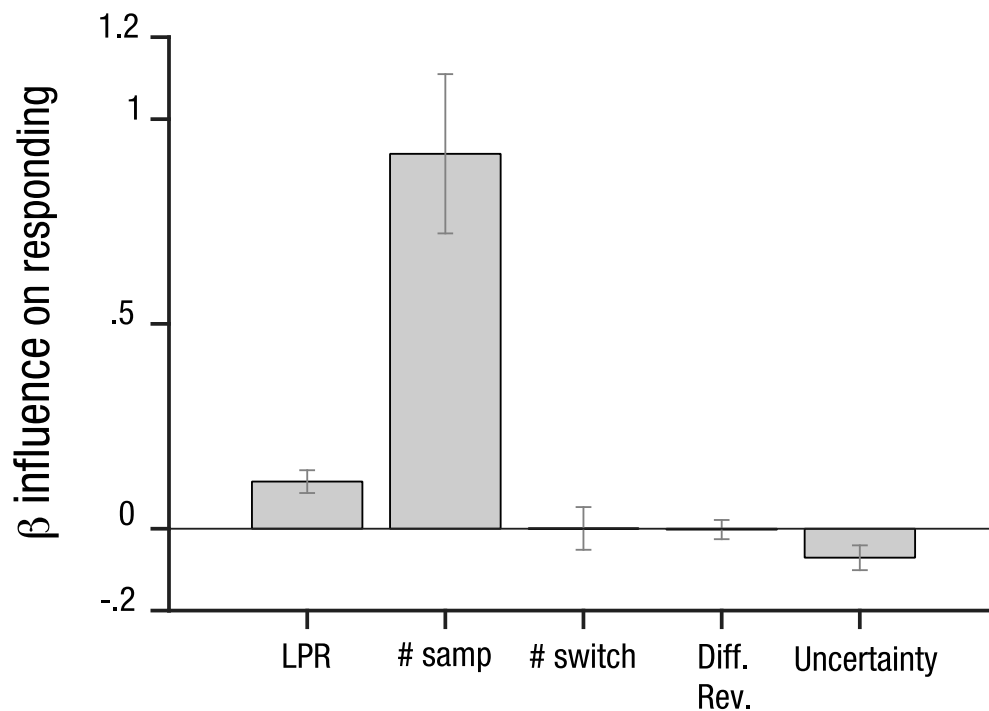
The regression model we used was as follows, where in each case: |

$$p(\text{respond}_S) = \phi[b + w1 \cdot |LPR_S| + w2 \cdot (n\_samples_S) + w3 \cdot (n\_switches_S) \\ + w4 \cdot ((n\_winner_S - n\_lower_S) + w5 \cdot (Posterior_S)]$$

where  $\phi$  denotes the cumulative normal distribution,  $|LPR_S|$  denotes the absolute LPR after the sample  $S$ ,  $n\_samples_S$  denotes the number of samples taken after sample  $S$ ,  $n\_switches_S$  denotes the number of switches taken after sample  $S$ , and  $n\_winner_S - n\_lower_S$  denotes the difference in number of samples taken from the winning and losing category after sample  $S$ .  $Posterior_S$  denotes the estimated width of the posterior after sample  $S$  for the sampled category.

Observers were more likely to respond when the evidence itself was higher [ $t_{(35)} = 8.1$ ;  $P < 0.001$ ]. The number of samples taken also significantly predicted the likelihood of responding: observers were more likely to respond as the number of samples taken increased [ $t_{(35)} = 9.2$ ;  $P < 0.001$ ]. The number of switches made and differences in the number of elements sampled between the two categories had no significant effect on the tendency to respond [ $P > 0.05$ ]. Finally, the uncertainty about each category significantly influenced the tendency to respond. As expected, observers were less likely to respond when they were more uncertain about the category they had just

sampled from [ $t_{(35)} = -4.58$ ;  $P < 0.001$ ; **Fig. 5.28**].



**Fig 5.28 Predictors of the tendency to respond.** The absolute LPR and the number of samples taken had a significant positive effect, while the uncertainty had a significant negative effect on the tendency to respond. Bars are shown with 95% confidence intervals.

## 5.4 Discussion

The purpose of experiment one was to delve deeper into the characteristics and determinants of the sampling policy when deciding between two categories or options, and how the sampling policy affects these decisions themselves. More specifically, we have demonstrated effects of the sampling policy on the weighting of information and on the decision making process itself (e.g. a recency effect, frequency bias, and a consistency bias). It is

important to note that the frequency bias in this experiment is demonstrated in a paradigm in which the decision value of each sample can be accurately quantified (unlike in most value based decision making tasks). For example, value based decision making tasks often employ a 'rating' session beforehand in which participants rate the stimuli that they will see in the experiment. However, the context in which the stimuli appear could potentially influence the value of a stimulus at a given point during the experiment (Khaw, Glimcher, & Louie, 2017). Our modified paradigm thus allowed us to more accurately separate out the choice effects due to the frequency bias and choice effects due to the amount of evidence sampled from each category.

Furthermore, our results suggest that this frequency bias is not a multiplicative effect as suggested by the attentional DDM. Thus, rather than multiplicatively scaling the DV of the fixated option compared to the non-fixated option, an additive bias might better describe the frequency bias (at least in our gaze contingent paradigm).

Furthermore, we have shown that switching between sampling from one category to another incurs some kind of cost, demonstrated by a consistency effect. Observers in our task seem to adapt their sampling policy in order to minimize this switch cost, evidenced by the low proportion of switches (e.g. 32%) and the tendency to switch near the middle of the sample sequence. However, the mechanisms underlying this sampling behaviour are still poorly understood. For example, Hills & Hertwig described two distinct sampling policies in their study: round-wise and summative sampling (Hills & Hertwig,

2010). Even though the demonstrated consistency effect could be facilitating a summative sampling policy, we have provided no evidence for a causal relationship. Furthermore, it is unclear whether participants employ their sampling strategy regardless of the decision environment they are operating under. Or would a summative sampler be able to adapt their sampling policy based on the type of decision information available or based on other task characteristics.

While we have tried to model the underlying mechanisms of these biases, the implemented models were unable to capture the full extent of our data.

A neurological investigation of these biases could shed light on their underlying mechanisms. For example, researchers utilising fMRI have shown that the value representations of stimuli in the orbito-frontal cortex (OFC) can be modulated by fixation location (McGinty et al., 2016). It would be interesting to investigate whether these value representations would be influenced by for example a frequency or consistency bias. For example, would value representations on 'switch fixations' be less accurate?

Moreover, researchers have found evidence for the ventromedial prefrontal cortex (vmPFC) to track the value of the default / attended option, while the dorsal anterior cingulate cortex (dACC) tracks the alternative / unattended option (Boorman, Rushworth, & Behrens, 2013; Kolling, Behrens, Mars, &

Rushworth, 2012). Therefore, would dACC activity ramp up towards a switch? Similarly, would it be possible to predict the sampling behaviour of participants 'on-line' based on neural activity in the dACC? Furthermore, other methodological techniques have been used to capture selective attention. For example, Nunez and colleagues have identified independent EEG components of selective attention (Nunez, Vandekerckhove, & Srinivasan, 2017). Alternatively, frequency tagging stimuli (e.g. oscillating stimuli at different frequencies) is another methodology that has been used in order to track selective attention with EEG (Müller, Malinowski, Gruber, & Hillyard, 2003; Toffanin, de Jong, Johnson, & Martens, 2009).

Finally, we have identified various distinct drivers of the sampling policy.

(1) The difference in samples taken from each category influences which category is sampled from next, in the sense that observers are more likely to sample from the category they sampled the least from. Interestingly, this effect is modulated by the cumulative evidence: the magnitude of this effect is bigger when the cumulative evidence favours the correct category.

(2) The uncertainty about the estimates of each category significantly predicts the likelihood of making a switch and the dwell time.

(3) Finally, we have taken a brief look at what makes people decide when to stop sampling for information and make their choice. Observers were more likely to decide as they sampled more evidence and as the absolute LPR increased (e.g. the amount of evidence for either hypothesis). Furthermore, observers were less likely to respond when they were more uncertain about the sampled category.

## **Chapter 6: Voluntary information sampling when deciding between two categories (experiment two)**

### **6.1 Introduction**

The interesting results of chapter 5 drove us to take this project further. More specifically, in the next experiment we wanted to allow participants to sample an infinite amount of evidence from each category (in chapter 5 the total number of samples is constrained by nature of the design: only a finite number of samples can be present on the computer screen). Moreover, we wanted to replicate our results from chapter 5 in a design that is more suitable for the collection of neural and pupil size data. Urai, Braun, & Donner have recently demonstrated that pupil size can encode decision uncertainty (Urai, Braun, & Donner, 2017). However, pupil size measurement is inaccurate in designs where observers are not required to maintain central fixation (Scheepers & Crocker, 2004). This inaccuracy is due to changes in the geometry of the recorded pupil as a function of gaze position (Atchison, Smith, & Smith, 2000) and manufacturers of eye-tracking systems have acknowledged this issue. However, do note that researchers have developed paradigm specific methodologies to try and remedy this issue (Gagl, Hawelka, & Hutzler, 2011). Our goals for this second experiment are accompanied with the side-effect that for certain analyses we will have less power due to the smaller amount of trials / samples taken. However, this is exactly why we ran experiment 2: can we replicate our findings from the first experiment in a design that is better suited for the collection of neural and pupil size data?

## 6.1 Current study

*Sample.* The experiment took place in the Department of Experimental Psychology at the University of Oxford. Eighteen subjects volunteered (7 males and 11 females, age between 18 and 30 years old) to participate in experiment 2. Subjects reported normal or corrected to normal vision and had no history of any psychological disorders. Participants provided written consent before the experiment. The experiment was approved by the local ethics committee (approval no. MSD-IDREC-C1-2009-1).

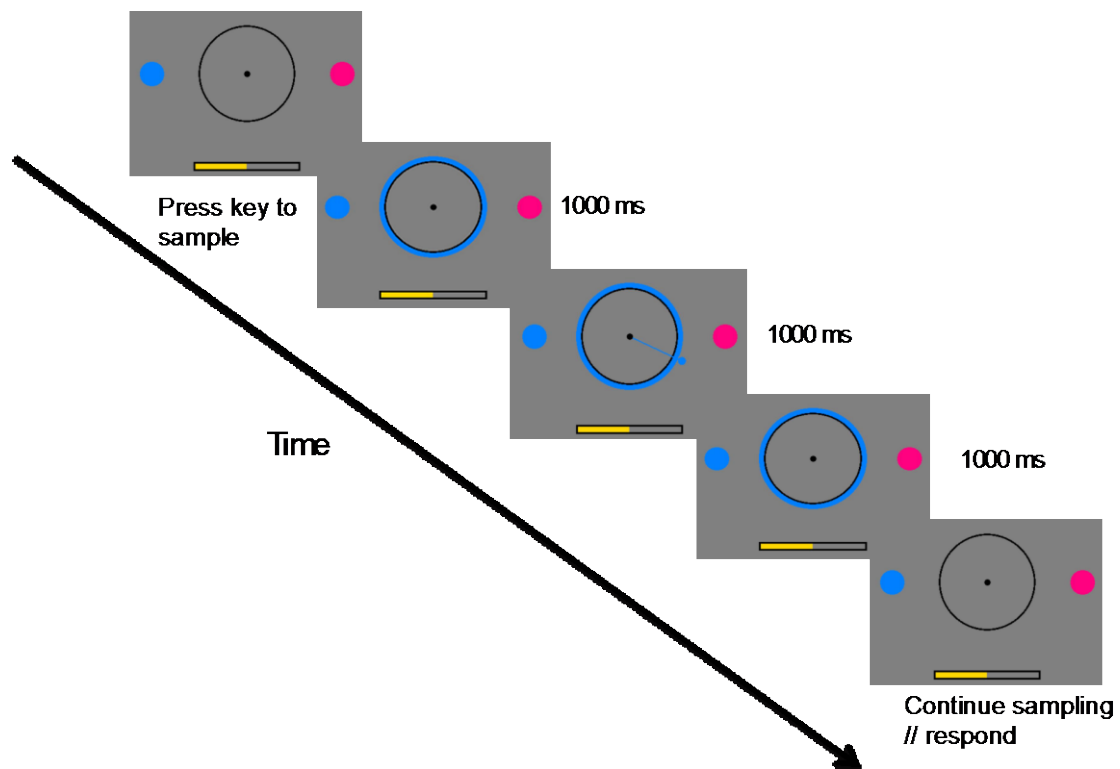
*Stimuli.* Stimuli were generated using PsychToolbox ([psychtoolbox.org](https://psychtoolbox.org)) for MATLAB (Mathworks) and presented on a 17-inch screen (resolution, 1024 x 768) viewed from a distance of 60 cm. In this task, observers had to indicate which of two categories was oriented most clockwise and could sample evidence (i.e. angle orientations) from both categories (**Fig 6.1**). We chose to use angle orientations over number stimuli as we were hoping to take this paradigm further and collect fMRI data if we were able to replicate our findings from chapter 5.



### Which category is more clockwise?

**Fig 6.1 Two example stimulus displays.** Observers had to indicate which of two distributions (pink or blue) is most likely to be oriented more clockwise. In this case the red sample is oriented more clockwise than the blue sample. Observers could sample an infinite number of stimuli (e.g. angles) from each distribution (i.e. category). Note that observers had to infer which generative distribution was oriented most clockwise, and thus could not simply sample one element from each category and respond with 100% accuracy. The stimulus display is enlarged for readability.

Each trial started with a black circle (diameter 320 pixels) and a black fixation dot (diameter 5 pixels) in the centre of the screen (**Fig 6.2**).



**Fig 6.2 Demonstration of the sequence of taking one sample.** The trial starts with the standard display: a small black central fixation dot, surrounded by a large black circle. A coloured circle accompanies this black circle on each side, indicating the two categories observers could sample from. Whenever the standard display was present, observers could choose to sample or respond. Whenever a sample was requested, the large central black circle was surrounded by the colour of the requested category for 1000 ms (i.e. the category information display). Subsequently the stimulus, a coloured line and coloured circle on the edge of the large black central circle, was displayed for 1000 ms. The category information screen again followed the stimulus display for 1000ms. After this interval, the category information was removed from the central black circle and observers could either choose to sample more evidence or to respond.

Smaller coloured circles (diameter 50 pixels) were present on each side of the central black circle, in order to indicate the possible categories the stimuli would be sampled from and to convey the response mapping. More specifically, the locations of these coloured circles were spatially congruent with the keys on the keyboard that were assigned to sample information from each category (i.e. in the case of **Fig. 6.2**: left for samples from the blue category).

Observers could sample evidence from either category by pressing one of two keys whenever the large central circle was black. Whenever a sample was requested, the large central black circle was surrounded by the colour of the requested category for 1000 ms (i.e. the 'category information display', blue in this case). Subsequently the stimulus, a coloured line and coloured circle on the edge of the large black central circle, was displayed for 1000 ms. The decision evidence was conveyed by the angle constructed by this coloured line connecting the small coloured circle on the edge of the large black central circle with the fixation dot. The category information screen again followed the stimulus display for 1000ms. After this interval, the category information was removed from the central black circle and observers could either choose to sample more evidence or to respond. These long periods between stimulus presentations were employed for two reasons. First, the orientation judgment paradigm was constructed for fMRI use. Therefore, a sufficient delay between samples was necessary. Second, the inter stimulus interval could have drastic effects on the way observers solve the task. For example the possibility to

rapidly sample stimuli after one another might have allowed subjects to solve the task visually rather than weighting and integrating each sample.

The bar on the bottom of the screen signaled the probability of earning a bonus reward. The bar started out half-full at the start of each block and would grow or shrink on a trial by trial basis as a function of performance. Incorrect responses were more influential on the probability of reward than correct responses: an incorrect response would decrease the bar by two units and a correct response would increase the bar by one unit. We pre-determined the accuracy (85%) subjects need to achieve in order to get the maximum reward. On the contrary, if the subjects' performance is at chance, the bar will be empty and they would receive no reward. Rewards again ranged from £5 (performance at chance level) to 15£. Implementing this reward bar necessitates observers to strike a balance between sampling many items (leading to ~100% accuracy) and solving many trials in order to maximise their reward.

The response mapping and the colours (pink and blue) were counterbalanced across subjects. Visual feedback on accuracy was available on each trial: the central fixation dot would either turn green (i.e. a correct response) or red (i.e. an incorrect response) for 100 ms. The inter trial interval (ITI) consisted of a fixation dot that was presented in the middle of the screen for a total duration of 1000 ms.

Experiment two consisted of **10 blocks** each lasting approximately 5 minutes. The number of trials per block was variable due to observers being able to sample evidence for each category for as long as they would like. The experiment was preceded by five practice trials.

*Design.* On each trial, a reference angle  $R$  was drawn randomly from a uniform distribution (1 - 360 degrees). The angles were drawn from two Gaussian distributions with mean  $\mu_1$  and  $\mu_2$  and standard deviation  $\sigma$ . Observers were asked to choose which of the two distributions (pink or blue) was most clockwise. The difficulty varied between trials:  $\mu_2 - \mu_1 = \pm 8$  or  $\mu_2 - \mu_1 = \pm 12$  in low and high  $\mu$  trials, respectively, and  $\sigma = 8, 14$  or  $20$  in low, medium, and high  $\sigma$  trials, respectively. Three levels of  $\sigma$  were employed as this allowed us to test a wider range of decision environments for the potential follow up study.

Note that these values represent the 'real' evidence strength and reliability of the two distributions. However, since infinite samples can be drawn from these distributions the actual evidence strength and reliability can be different from trial to trial. Accuracy and feedback was computed based on the underlying distributions instead of the sampled evidence. While this is exactly the same as experiment one, note that in experiment one the sampling was constrained and the 'objective' evidence would always lead observers to the

correct choice if they would have sampled all elements. Each of the twelve trial types occurred equally often, and their presentation order within a block was random.

## 6.3 Results

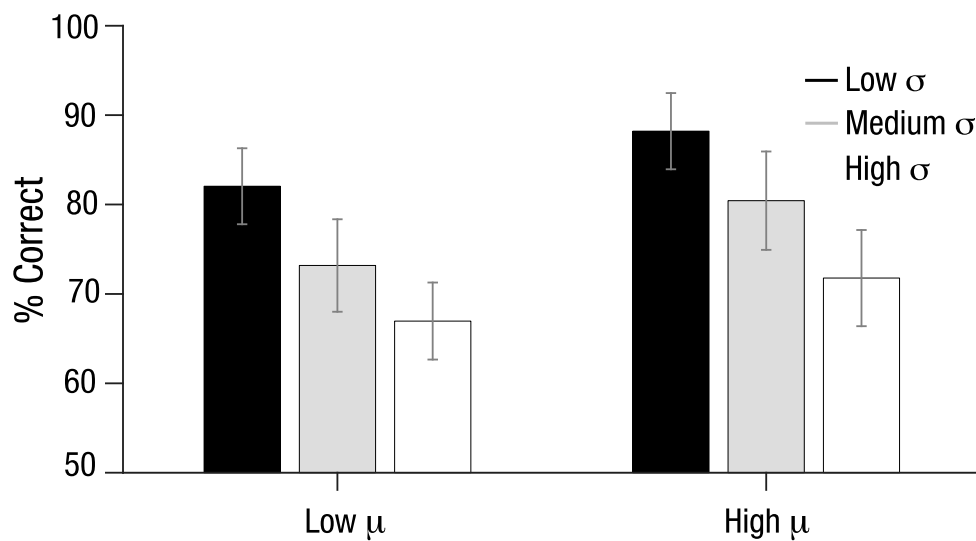
We excluded the data from the 5 practice trials in every analysis reported in this section. First, we will present descriptive results on accuracy and sampling rates in order to give the reader a more accurate picture of how subjects perform the task. Next, we will focus on replicating the findings from experiment one, and again focus on two major aspects of the sampling policy: selective weighting of information and information selection.

The overall mean accuracy was 76.75% ( $n = 18$ ). An average of 7.2 stimuli were revealed on each trial ( $SD = 3.35$ ). Note that on average participants sampled less stimuli in experiment 2 compared to experiment 1, even though participants were able to sample an infinite number of samples.

### 6.3.1 Effects of evidence strength $\mu$ and evidence reliability ( $\sigma$ )

Performance again depended on both the mean ( $\mu$ ) and variance of the distribution from which the numbers were drawn (**Fig. 6.3.**). Accuracy increased on high  $\mu$  trials compared with low  $\mu$  trials [ $F(1.91, 32.44) = 30.68$ ,  $P < 0.001$ ] and decreased on high  $\sigma$  (more variable evidence) trials compared

with low  $\sigma$  trials [ $F(1, 17) = 10.8, P < 0.01$ ], replicating previous work and the findings from experiment 1 (De Gardelle & Summerfield, 2011; Michael et al., 2013; Vandormael et al., 2017). However, the interaction between  $\mu$  and  $\sigma$  was not significant [ $F(1.89, 32.1) = 0.24, P > 0.05$ ].



**Fig 6.3 Influence of evidence strength ( $\mu$ ) and reliability ( $\sigma$ ) on accuracy.**

Performance increased for higher evidence strength and higher evidence reliability. Bars are shown with 95% confidence intervals.

The influence of evidence strength ( $\mu$ ) and evidence reliability ( $\sigma$ ) on the sampling rate is depicted in **Table 6.4**. Observers sampled less elements when the evidence strength was higher [ $F_{(1,17)} = 5.24; P < 0.05$ ]. Moreover, the interaction between evidence strength and reliability was significant [ $F_{(1.88,31.93)} = 4.28$ ]. This interaction was driven by the fact that observers

sampled the least number of samples during the high evidence reliability / strength trials.

**Table 6.4 Influence of evidence strength ( $\mu$ ) and reliability ( $\sigma$ ) on the number of samples.**

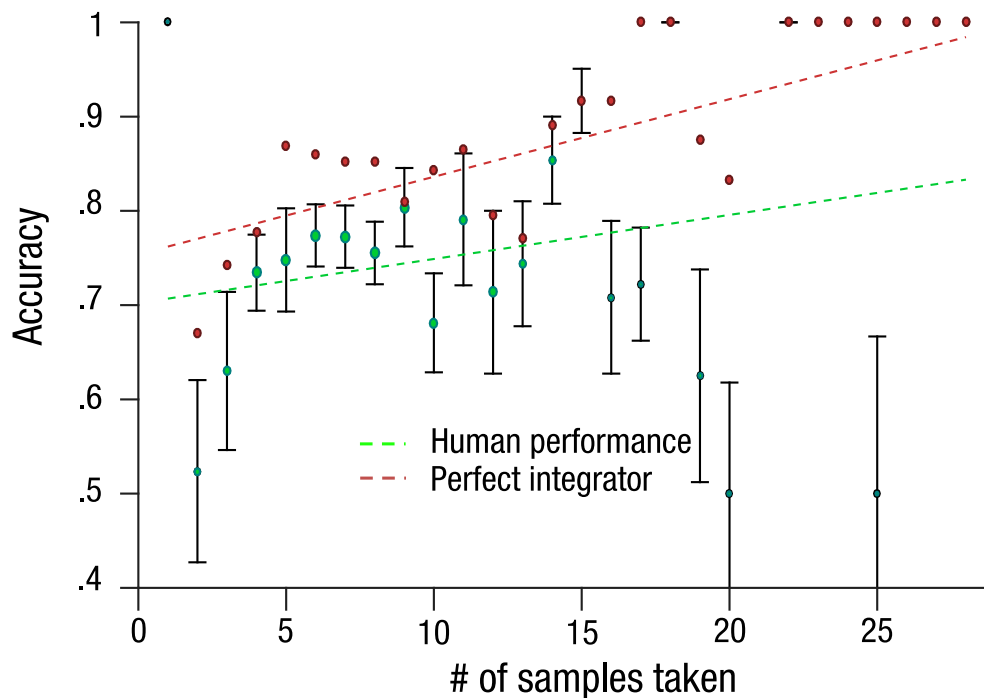
| $\sigma$      | High $\mu$      | Low $\mu$       | Effect on # of samples                                       | Significance                              |
|---------------|-----------------|-----------------|--------------------------------------------------------------|-------------------------------------------|
| High $\sigma$ | $7.51 \pm 0.74$ | $7.56 \pm 0.69$ | $\mu \uparrow \rightarrow \# \text{ samp} \downarrow$        | $F_{(1,17)} = 5.24$ ,<br>$P < 0.05$       |
| Med $\sigma$  | $7.46 \pm 0.5$  | $7.53 \pm 0.63$ | $\sigma : \text{n.s.}$                                       | $F_{(2,33.96)} = 0.97$ ,<br>$P > 0.05$    |
| Low $\sigma$  | $6.72 \pm 0.61$ | $7.79 \pm 0.73$ | High $\mu * \text{Low } \sigma : \# \text{ samp} \downarrow$ | $F_{(1.88,31.93)} = 4.28$ ,<br>$P < 0.05$ |

Mean number of samples and SD are shown per condition. Higher values of  $\mu$  lead observers to take fewer samples. This effect was particularly prominent in the high evidence strength and high evidence reliability trials.

In summary, when the evidence strength was high (high  $\mu$ ), observers were more accurate and took fewer samples. When the evidence reliability was high (low  $\sigma$ ) observers were more accurate. Evidence reliability had no significant influence on the number of samples taken. However, the interaction between the evidence reliability and the evidence strength was significant, with observers sampling less during high  $\mu$  / low  $\sigma$  trials.

### 6.3.2 Q1. Sampling for a prolonged period of time and the effect on decision accuracy

**Fig. 6.5** shows evidence for an asymptotic level of performance in our task: performance levels out in the 76% range and does not increase as more samples are taken (e.g. the number of samples taken is not a significant predictor of performance [ $t_{(17)} = -1.75$ ;  $P > 0.05$ ]). Note that the human performance curve could be seen as quadratic. However, given the low number of trials with number of samples at the extremes, we have chosen to focus on those trials for which the majority of subjects have data when interpreting this graph. Furthermore, regression techniques described later on in the chapter will be more sensitive for capturing the effect of number of samples taken on performance than this aggregate view.

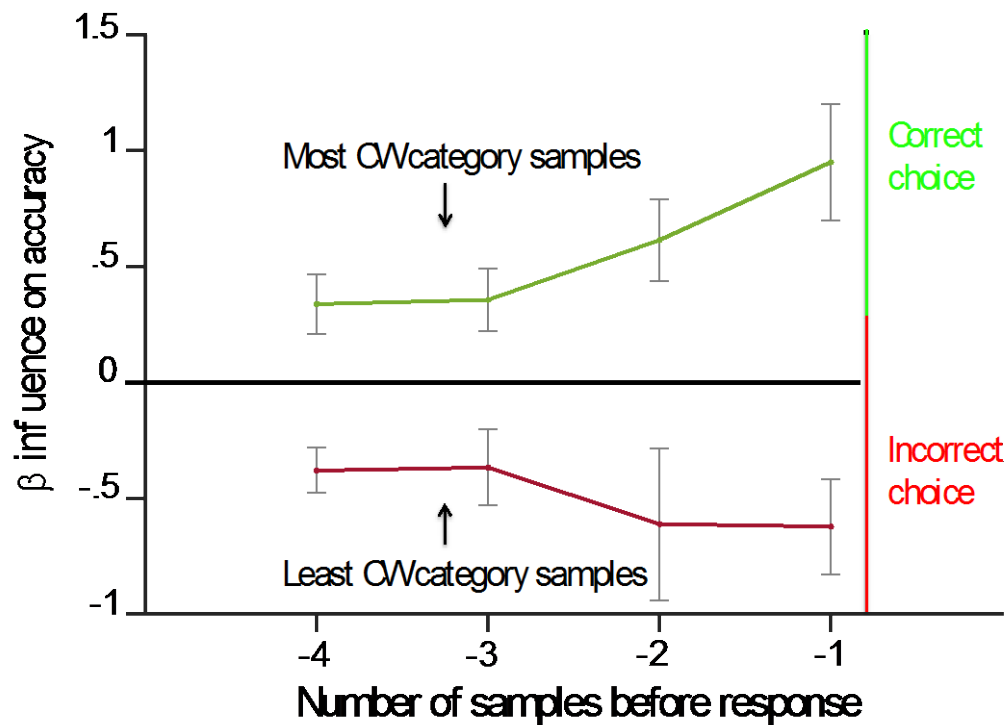


**6.5 Asymptotic level of performance.** Human performance is shown in green, with the size of the bubbles reflecting the number of subjects per number of samples taken (i.e. most subjects sample between 4-9 times). Performance of a perfect integrator without loss is shown in red. Note that the perfect integrator sees exactly the same decision information as the human observers. Weighted least-squares regression lines are shown via the dotted lines.

### 6.3.3 Sampling policy and the weighting of information: recency

Next, we again turned to logistic regression in order to investigate the possible sources of this asymptotic level in performance. More specifically, we regressed choices (i.e. chosen the winner / most clockwise = 1, chosen the loser / least clockwise = 0) on the decision information (i.e. the sampled angles) as a function of when they were sampled relative to when the decision was made (**Fig. 6.6**). The sampled decision information is normalized with

respect to the reference angle in order to combine trials with different levels of evidence strength and reliability. We controlled for individual variability in response times (due to the free-response paradigm) by aligning samples relative to when the response occurred instead of when the stimulus appeared. The decision information was binned based on when a sample was revealed relative to when the response was made (i.e. 1 sample, 2 samples, 3 samples or 4 samples before the response). Moreover, we created these regressors separately for the most clockwise (e.g. the eventual ‘winner’) and the least clockwise (e.g. the eventual loser), resulting in 8 regressors. **Fig. 6.6** illustrates that samples closer to the response exhibited a larger influence on choice than samples further away from the response [ $t_{(17)} = 2.8$ ;  $P < 0.05$  and  $t_{(17)} = -4.6$ ;  $P < 0.001$  for samples from the least CW category and most CW category respectively]. Furthermore, samples from the least CW category made observers more likely to choose the least CW category (i.e. negative  $\beta$  weights,  $p < 0.001$ ), while samples from the most CW category increased the likelihood of choosing the most CW category (i.e. positive  $\beta$  weights;  $p < 0.001$ ). Thus in the orientation paradigm, observers also exhibited a recency effect: information sampled closer to the response was weighted more strongly than information sampled earlier on the trial.



**Fig 6.6** Information sampled closer to the response has a stronger influence on choices (i.e. a recency effect). The bottom four regressors correspond to samples from the least clockwise (CW) category while the top four regressors correspond to samples from the most clockwise category. Most CW category samples are accompanied by positive  $\beta$  weights (e.g. higher likelihood of choosing the high category) while least CW category samples are accompanied by negative  $\beta$  weights (e.g. higher likelihood of choosing the least CW category). Importantly, information sampled near the end of the trial was more influential (larger  $\beta$  weights) than information sampled at the beginning of the trial. Bars are shown with 95% confidence intervals.

### 6.3.4 Q2. Frequency bias: preference for the overly sampled category

Next, we looked at whether we could replicate the frequency bias found in chapter 5, e.g. a preference for the overly sampled category (Armél & Rangel, 2008; Krajbich et al., 2010; Krajbich & Rangel, 2011; Shimojo et al., 2003).

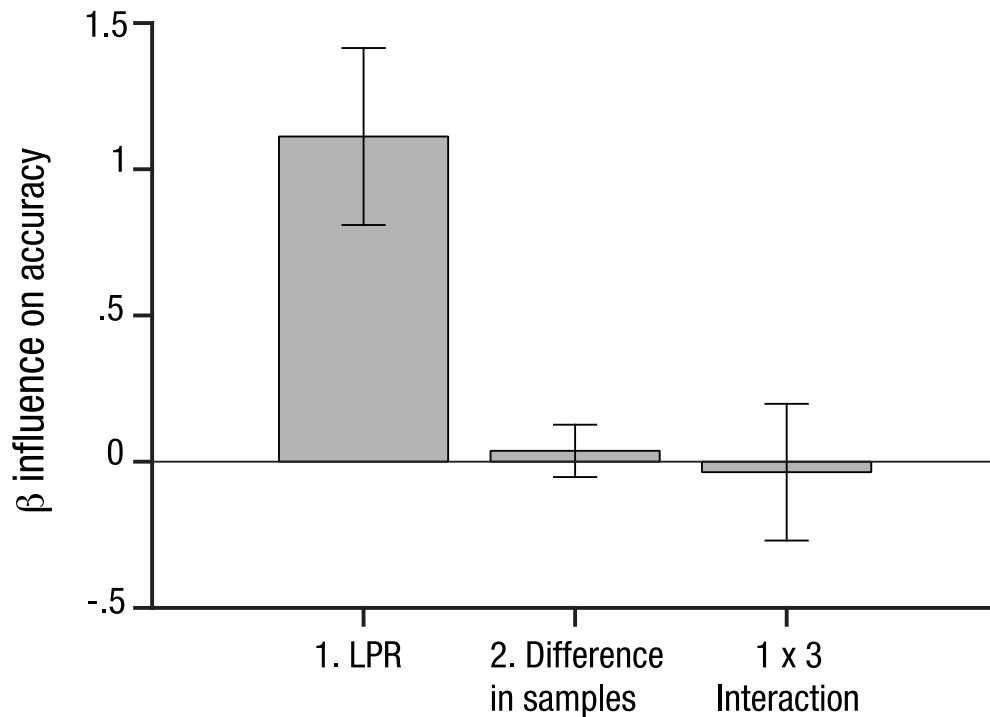
In order to investigate whether observers were biased towards the overly sampled category (i.e. a ‘frequency bias’), we again turned to a probit regression. We used probit regression to predict observers’ accuracy (most clockwise versus least clockwise category) as a function of three sets of predictors: (i) the cumulative evidence captured by the LPR after the last sample, (ii) the frequency bias measured by the difference between the number of items sampled from the most CW category minus the number of items sampled from the least CW category, and (iii) the interaction between these two quantities. The regression model that we used was defined as follows:

$$p(\text{Most}) = \phi[b + w1 \cdot LPR_{last} + w2 \cdot (n_{most} - n_{least}) + w3 \cdot (LPR \cdot (n_{most} - n_{least}))]$$

where  $\phi$  denotes the cumulative normal distribution,  $LPR_{last}$  denotes the LPR after the last sampled element before making a response,  $n_{most}$  denotes the number of elements that were sampled from the most clockwise category, and  $n_{least}$  denotes the number of elements that were sampled from the least clockwise category.

Remember that this regression model allows us to separate out the influence of the decision information (i.e. the LPR) and the frequency bias (i.e. the

number of samples taken from the least clockwise category minus the samples taken from the most clockwise category). Our results replicate the effect that the LPR positively predicts the likelihood of choosing the most clockwise category, providing support for the LPR as an approximation of the cumulative evidence in the orientation paradigm [ $t_{(17)} = 7.2$ ;  $P < 0.001$ ; **Fig. 6.7**]. However, there was no evidence for a frequency bias [ $t_{(17)} = 0.8$ ;  $P > 0.05$ ]. The interaction of the cumulative evidence (LPR) and the frequency bias was again not significant [ $t_{(17)} = -0.3$ ;  $P > 0.05$ ]. Thus, we could not replicate the earlier frequency bias (e.g. the overly sampled category was *not* chosen more often in the orientation paradigm).



**Fig 6.7 Frequency bias not replicated in the orientation paradigm.** The main effect of the LPR was significant providing additional evidence that the sampling policy is not stochastic. The LPR (calculated after taking the last sample on each trial) had a significant positive effect on the likelihood of choosing the most clockwise category. However, the difference in samples between the most clockwise and least clockwise category was not significant. Thus, we could not replicate the frequency bias in the orientation paradigm. Bars are shown with 95% confidence intervals.

### 6.3.5 Q3. Consistency effect: switching between categories

Next, we returned to the analyses investigating switching behaviour, in order to replicate the consistency effect. Switches were defined using the same methodology used in chapter 5. If a sample  $S$  came from the same category

as the previous sample (S-1), then sample S would be considered a consistent ('stay') sample (e.g. Red-Red). However, if a sample S came from a different category as the previous sample (S-1), then the sample S would be considered a switch sample (e.g. Red-Blue → the blue sample is identified as an inconsistent / switch). The first sample of each trial was again excluded, as these samples cannot have an S-1 sample.

Interestingly, the sampling behaviour in the orientation paradigm was markedly different compared to the sampling behaviour in experiment one. The Wilcoxon signed rank test indicated that observer's tendency to switch between the two categories did not significantly differ from chance: 47.18% of all samples were classified as switches ( $Z_{79} = -0.28$ ,  $p > 0.05$ ). Interestingly, two types of observers could be identified when looking at the individual subject level: 7 out of 18 observers showed switch rates higher than 50%, with 5 of these observers having switch rates higher than 70%. Thus, these observers could be classified as 'round-wise samplers' as defined in chapter 5. 11 of 18 subjects showed switch rates lower than 50% with 7 of these subjects having switch rates lower than 30%. Thus, these observers could be classified as 'summary samplers', showing similar behaviour to the observers in the gaze contingent paradigm. Note that in chapter 5 only 4 out of 36 subjects showed switch rates higher than 50%.

**Table 6.8** shows the influence of evidence strength and reliability on the number of switches made by observers. In experiment one, neither evidence strength nor evidence reliability significantly influenced the number of

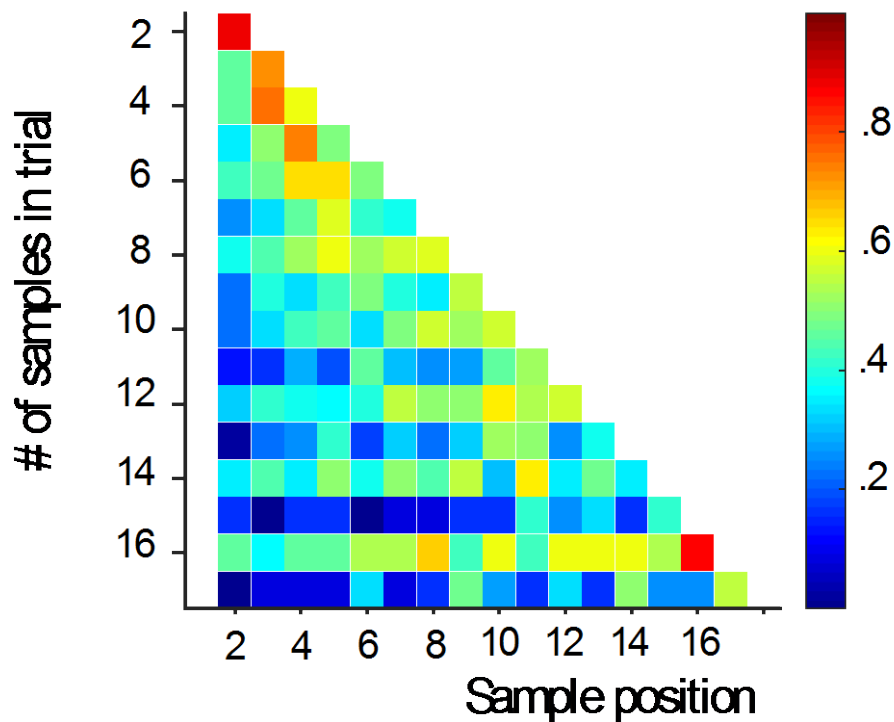
switches made by observers. However in experiment two, observers made less switches when the evidence strength was high [ $F_{(1,17)} = 15.1$ ;  $P < 0.001$ ]. Note however, that observers also took fewer samples in high  $\mu$  trials (see above). Moreover, the effect of evidence strength was no longer significant when analysing the proportion of switches [ $F_{(1,17)} = 15.1$ ;  $P < 0.001$ ].

| $\sigma$      | High $\mu$      | Low $\mu$       | Effect on # of samples                                    | Significance                             |
|---------------|-----------------|-----------------|-----------------------------------------------------------|------------------------------------------|
| High $\sigma$ | $2.68 \pm 0.5$  | $2.78 \pm 0.46$ | $\mu \uparrow \rightarrow \# \text{ switches} \downarrow$ | $F_{(1,17)} = 15.1$ ,<br>$P < 0.001$     |
| Med $\sigma$  | $2.86 \pm 0.39$ | $3.08 \pm 0.57$ | $\sigma : \text{n.s.}$                                    | $F_{(1.74,29.54)} = 2.5$ ,<br>$P > 0.05$ |
| Low $\sigma$  | $2.56 \pm 0.54$ | $3.15 \pm 0.57$ | $\mu * \sigma : \text{n.s.}$                              | $F_{(1.93,32.77)} = 2.2$ ,<br>$P > 0.05$ |

**Table 6.8 Influence of evidence strength ( $\mu$ ) and reliability ( $\sigma$ ) on the number of switches.**

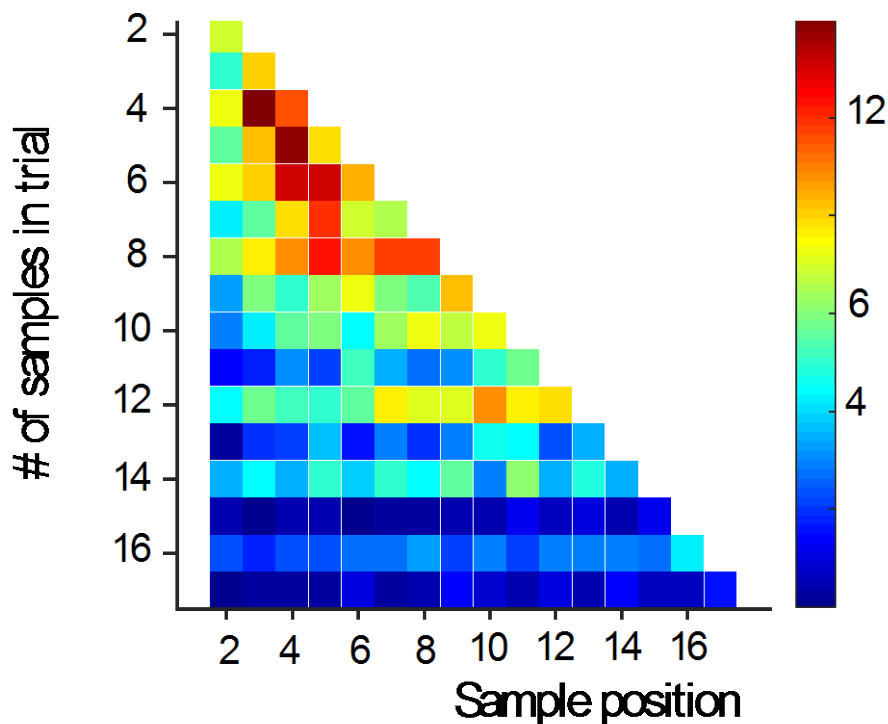
Mean number of switches and SD are shown per condition. Neither  $\mu$  nor  $\sigma$  had a significant effect on the number of switches.

Next, we investigated whether switches were more likely near the middle of the sample sequence. **Fig. 6.9** plots the probability of a sample being a switch as a function of the sample position within the trial (x – axis), for different trial lengths (i.e. the number of samples taken; y - axis). In contrast to experiment one, the probability of switching was not higher near the middle of the trial length.



**Fig 6.9** The probability of a sample being a switch as a function of sample position, split by trial length. The probability of a sample being a switch is shown on a blue-green-red spectrum, with blue indicating a lower probability of the sample being a switch and red indicating a higher probability of the sample being a switch..

Multiplying these probabilities with the number of subjects that have trials in each bin on the y – axis does not affect the non-replication (**Fig. 6.10**). Thus, the non – replication is not affected by the fact that for some trial lengths only very few observers have data in the orientation paradigm.



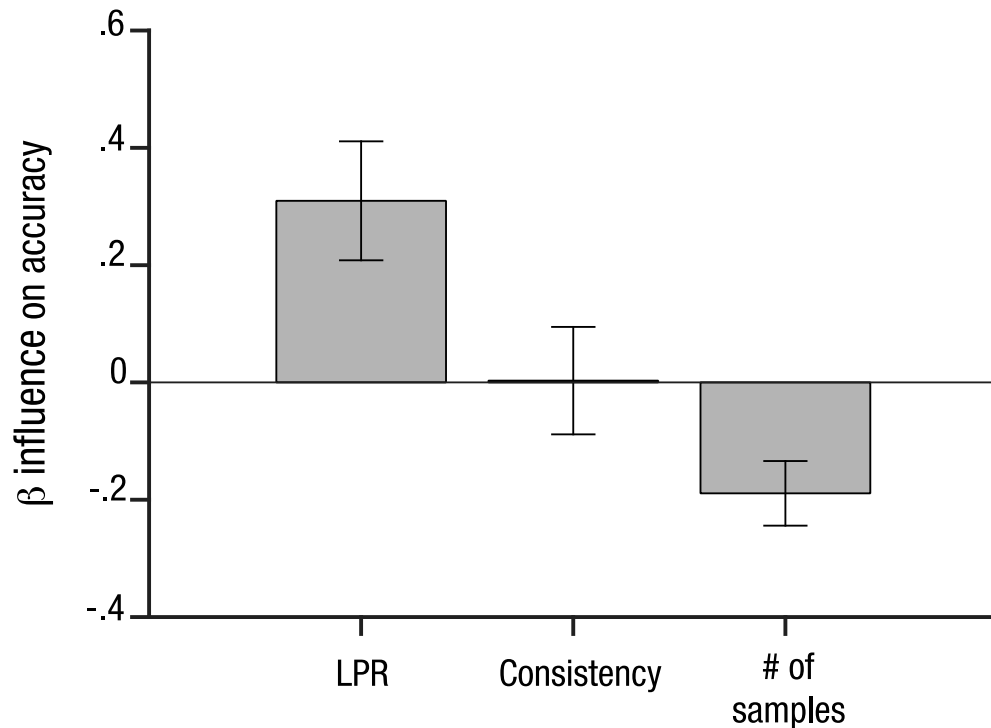
**Fig 6.10** The probability of a sample being a switch as a function of sample position, split by trial length and scaled by the number of subjects that have trials in each bin on the y - axis. The probability of a sample being a switch is shown on a blue-green-red spectrum, with blue indicating a lower probability of the sample being a switch and red indicating a higher probability of the sample being a switch. In contrast to experiment one, samples were not more likely to be a switch near the middle of the trial length.

Next, we investigated whether we could replicate the switch cost highlighted in experiment one. To do this, we again calculated how 'consistent' the sampling was on each trial. Consistency was defined as the sum of samples that were sampled sequentially from the same category divided by the total number of samples on a given trial (e.g. the number of stay samples divided by the total number of stay + switch samples). Thus, more consistent trials

consisted of proportionally less switches. We again used probit regression to predict observers' accuracy (most clockwise = 1 versus least clockwise = 0) as a function of three sets of predictors: (i) the cumulative evidence captured by the LPR after the last sample, (ii) the proportion of consistent samples out of all samples, and (iii) the number of elements sampled. Remember that we normalised the consistency by the number of samples in order to disentangle the consistency effect from any effect related to the number of samples, as more switches will be made on trials with more samples. The regression model we used was as follows:

$$p(Most) = \phi[b + w1 \cdot LPR_{last} + w2 \cdot (n_{consistent} / n_{samples}) + w3 \cdot (n_{samples})]$$

where  $\phi$  denotes the cumulative normal distribution,  $LPR_{last}$  denotes the LPR after the last sampled element before making a response,  $n_{consistent}$  denotes the number of samples identified as consistent ('stay') samples, and  $n_{samples}$  denotes the number of angles that were sampled. The LPR again positively predicted the likelihood of choosing the high category providing extra support for the LPR as an approximation of the cumulative evidence [ $t_{(17)} = 5.99$   $P < 0.001$ ; **Fig. 6.11**]. However, consistency had no significant effect on accuracy (in contrast to the consistency bias found in experiment one, [ $t_{(17)} = 0.7$ ;  $P > 0.05$ ]). The number of samples had a negative effect on accuracy, indicating that observers performed worse when they sampled more elements, replicating the result from chapter 5 [ $t_{(17)} = -6.72$ ;  $P < 0.001$ ].



**Fig 6.11 No consistency effect on accuracy in the orientation paradigm.**

The probability of choosing the most clockwise category is not affected by consistent sampling in the orientation paradigm (unlike the significant consistency effect shown in chapter 5). Bars are shown with 95% confidence intervals.

In summary, we have been able to replicate the asymptotic performance and the recency effect. However, we have not been able to replicate the frequency- and the consistency bias in the orientation judgment paradigm. Possible explanations for these failed replications vary from the experimental paradigm itself to the use of orientation stimuli.

First, in this paradigm observers are afforded a  $\pm 1000$  ms interval before the stimulus is presented after sampling from the desired category. Researchers

have previously reported a switch cost when voluntarily switching between two tasks and have reported that the size of the switch cost depends on the length of the Response-Stimulus-Interval (RSI; (Arrington & Logan, 2004; Rogers & Monsell, 1995). For example, Arrington & Logan showed an increased switch cost for short (100ms) compared to long RSI's (1000ms). The interval between requesting a sample and sample presentation in our tasks is comparable to traditional RSI's. Importantly, this interval is approximately eight times longer in the orientation judgment paradigm compared to the gaze contingent paradigm. Thus, this difference in 'RSI' could potentially explain why we do not see evidence for a switch cost in the orientation paradigm.

It is noteworthy that the sampling policy is more heterogeneous in the orientation paradigm compared to the gaze contingent paradigm. The two types of observers (Summative samplers / round-wise samplers) present in our task illustrate this heterogeneity.

Second, orientation is a class of stimuli with more perceptual ambiguity (encoding noise of 2-6 degrees; (Solomon & Morgan, 2006) compared to number stimuli. Similarly, it seems likely that it is easier to maintain an accurate representation of sampled numbers in working memory. Thus, under limited resources and when performing a task with more ambiguous stimuli, observers might opt for a different strategy: instead of estimating the average orientation of each category semi-sequentially (i.e. summary sampling),

observers instead switch continuously between both distributions in order to make one by one comparisons and count the number of comparisons each category has won (i.e. round-wise sampling).

### **6.3.6 Drivers of the sampling policy: information selection**

In this next section we again turn to the potential drivers of the sampling policy when sampling information in a 2-AFC perceptual decision making task.

#### **6.3.6.1 Sampling policy: asymmetrical sampling and the equity policy**

We again started off with assessing the key factors that influence from which category (i.e. most clockwise = 1 vs. least clockwise = 0) the next item will be sampled. More specifically, we were mainly interested in two key factors that might influence which category is sampled: (i) do sampling choices depend on the evidence accumulated so far (i.e. confirmation bias), and (ii) are sampling choices influenced by the relative difference of the number of samples taken from each category? Remember that which category was sampled next in chapter 5 was mainly predicted by the difference in the number of elements sampled from each category and its interaction with the accumulated evidence. In the gaze contingent paradigm, we found evidence for an ‘equity’ policy in which observers aim to sample approximately the same number of samples from each category. However, this effect could potentially be driven by the fact that observers are limited to sampling a finite number of pieces of information present on the screen. Since observers were allowed to sample

infinitely from both categories in the orientation judgment paradigm, we could test this alternative hypothesis explicitly.

We again used probit regression to predict observers' next sample  $S+1$  (most clockwise versus least clockwise category) as a function of five sets of predictors: (1) the cumulative evidence captured by the LPR after the sample  $S$ , (2) the total number of samples after the sample  $S$ , (3) the total number of switches after sample  $S$ , (4) the difference in number of samples taken from the most – least clockwise category after sample  $S$ , and (5) the interaction between the cumulative evidence and this difference in number of samples taken from each category. The regression model we used was as follows:

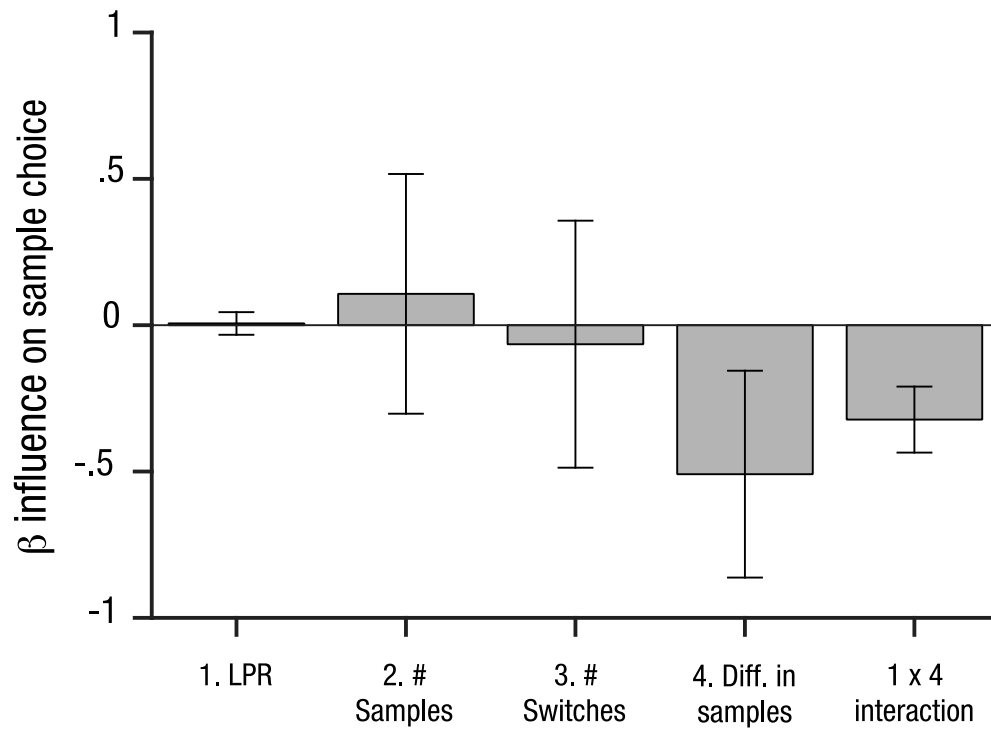
$$p(Most_{S+1}) = \phi[b + w1 \cdot LPR_S + w2 \cdot (n\_samples) + w3 \cdot (n\_switches) + w4 \cdot (n\_most - n\_least) + w5 \cdot LPR_S \cdot (n\_most - n\_least)]$$

where  $\phi$  denotes the cumulative normal distribution,  $LPR_S$  denotes the LPR after the sample  $S$  before sampling sample  $s + 1$ ,  $n\_samples$  denotes the total number of samples taken after sample  $S$ ,  $n\_switches$  denotes the total number of switches taken after sample  $S$ , and  $n\_most - n\_least$  denotes the difference in number of samples taken from the most and least clockwise category after sample  $S$ .

The evidence accumulated up to and including sample  $S$  ( $LPR_S$ ) did not

significantly predict sampling choices at time point  $S+1$  [ $t_{(17)} = 0.30$ ;  $P > 0.05$ ;

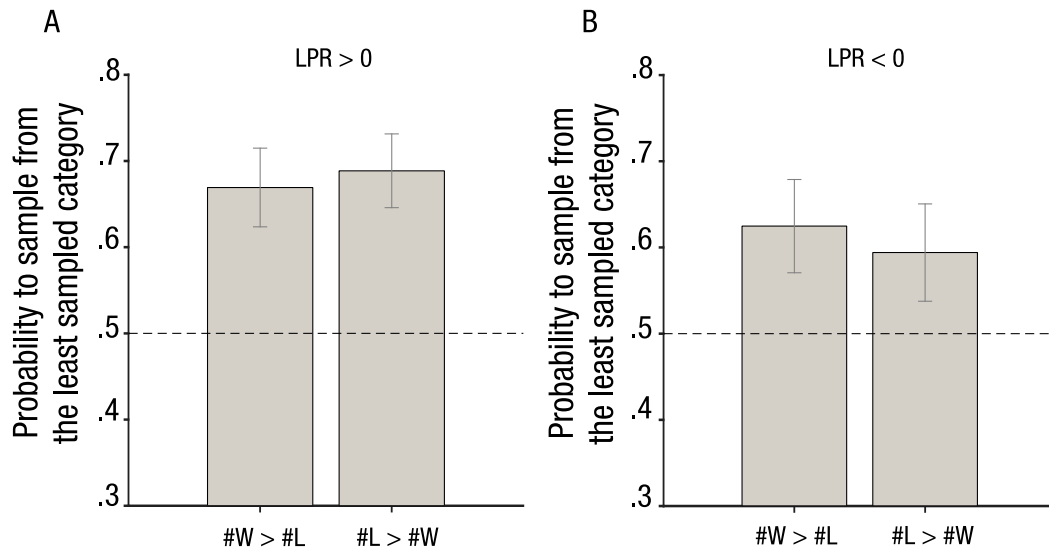
**Fig. 6.12**].



**Fig 6.12 Predicting the category of the next sample  $S + 1$ .** Replicating chapter 5, the difference in number of samples taken from the most – least clockwise category at sample  $S$  significantly predicted sampling choice at time  $S + 1$ . Moreover, the interaction between the cumulative evidence and the difference in samples taken from each category was significant. Bars are shown with 95% confidence intervals.

Neither the number of samples or switches taken up to and including sample  $S$  had a significant influence on sample choice of the next sample  $S + 1$  [ $t_{(17)} = 0.51$ ;  $P > 0.05$  and  $t_{(17)} = -0.30$ ;  $P > 0.05$ , respectively]. More importantly, the difference in number of samples taken from the most – least clockwise

category at sample  $S$  significantly predicted sampling choice at sample  $S + 1$  [ $t_{(17)} = -2.8$ ;  $P < 0.001$ ]. Moreover, the interaction between the cumulative evidence and the difference in samples taken from each category was again significant [ $t_{(17)} = -5.6$ ;  $P < 0.001$ ]. Thus, the cumulative evidence on itself did not predict sampling choices, but its interaction with the difference in samples taken did. Thus, observers tend to preferentially sample elements from the category they have sampled the least from so far, and this effect is influenced by the current cumulative evidence at that time point (i.e. replicating the findings of chapter 5). **Fig. 6.13** shows the probability of sampling from the most (M) and least (L) clockwise category as a function of LPR (left vs. right plot) and the difference in samples taken from each category in order to better understand this significant interaction effect.



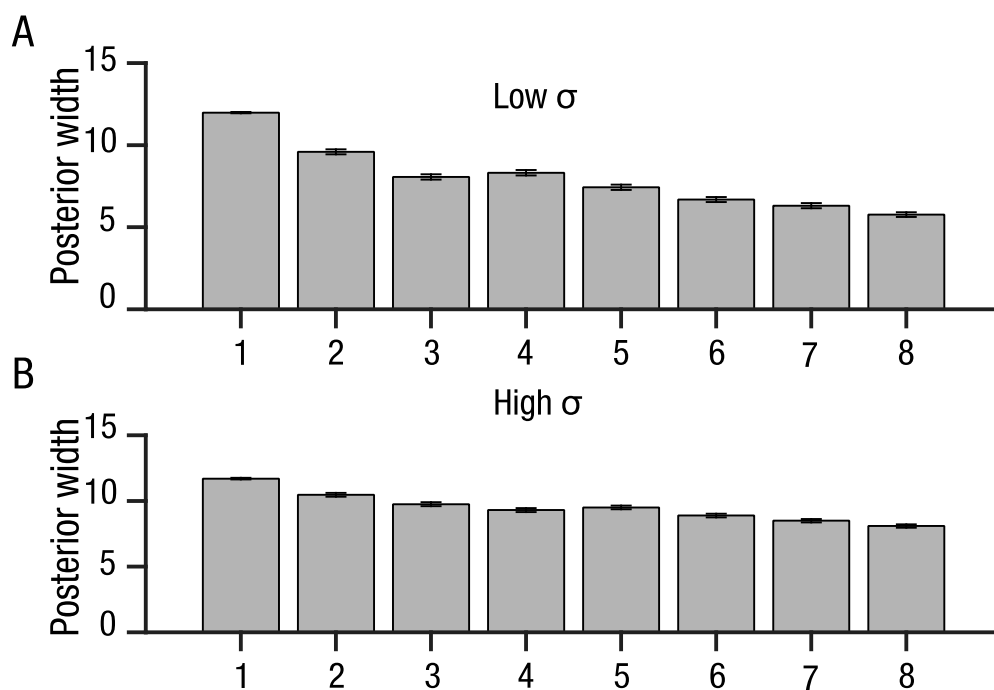
**Fig 6.13** The probability of sampling from the high / low category as a function of the difference in samples taken from each. (A) The difference in samples effect is strongest when the LPR is positive (e.g. correct hypothesis about the identity of the most clockwise category). (B) The difference in samples effect is reduced but still significant when the LPR is negative (e.g. incorrect hypothesis about the identity of the most clockwise category). Bars are shown with 95% confidence intervals.

As **Fig. 6.13** shows, the degree to which the difference in samples taken from each category influences the sampling policy is strongest when the LPR is positive (e.g. you have strong evidence for selecting the winner), and is reduced but still significant when the LPR is negative.

### 6.3.7 Q5. Category uncertainty as a driver of the sampling policy

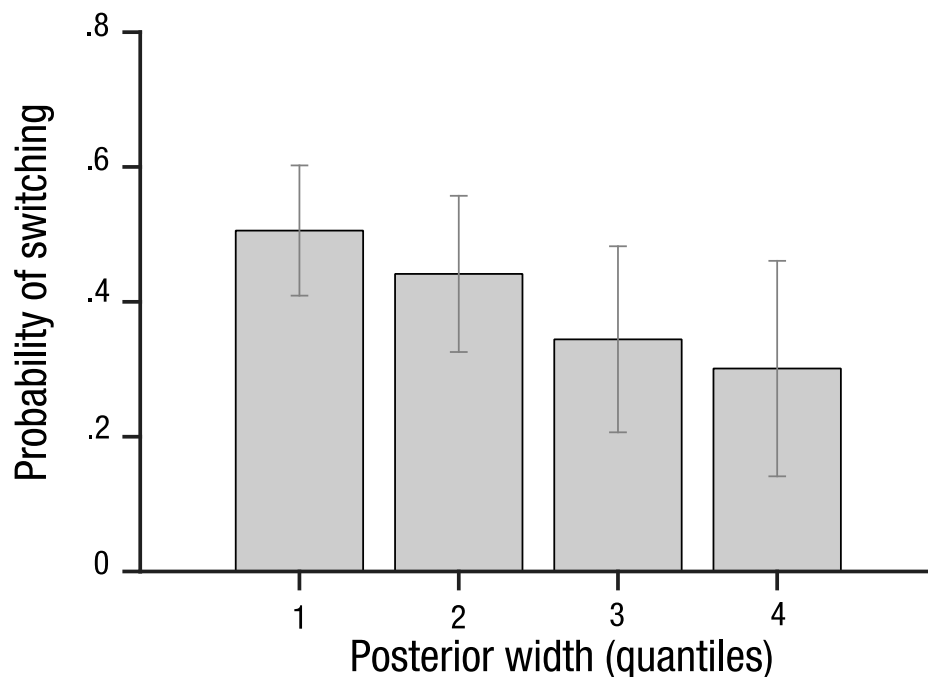
Next, we aimed to replicate the effects of category uncertainty as a driver of the sampling policy (chapter 5). We replicated the original computational

simulation in which we simulate how this posterior width decreases as an observer samples more information, split for low and high evidence reliability trials. Importantly, for this simulation we now used the environmental statistics from the orientation judgment paradigm (same  $\mu$ ,  $\sigma$ , range of stimuli...) and again utilised a completely random sampling policy. **Fig. 6.14** replicates our findings from chapter 5 and shows how the posterior width (i.e. uncertainty) decreases as more samples are taken, and that this rate of decrease in uncertainty is faster in trials with more reliable evidence (lower  $\sigma$ ) compared to trials with less reliable evidence (higher  $\sigma$ ).



**Fig 6.14 Evolution of the posterior width as more samples are taken (computational simulation).** The posterior width decreases as more samples are taken, and this rate of decrease in uncertainty is faster in trials with more reliable evidence (lower  $\sigma$ , top panel) compared to trials with less reliable evidence (higher  $\sigma$ , bottom panel).

Similarly, we were able to replicate a similar pattern of different probabilities of switching as a function of this posterior width (**Fig. 6.15**). More specifically, observers were again more likely to switch when they were more certain about their estimate about a category (i.e. smaller posterior width; [ $F(1.24,21.1) = 10.75$ ,  $P < 0.05$  ]).



**Fig 6.15** The probability of switching as a function of posterior width.

The probability of switching is higher when the posterior width is lower (e.g. lower uncertainty). Quantile one has the lowest uncertainty. Bars are shown with 95% confidence intervals.

Next, we again used probit regression to investigate whether the posterior

width is a significant predictor of making a switch (as was the case in chapter 5). We again predicted whether the next sample was a switch ( $= 1$ ) or not ( $= 0$ ) as a function of the posterior width. As expected from figure 6.15, the posterior width had a significant negative effect on the likelihood of switching between categories [ $t_{(17)} = -5.2$ ;  $P < 0.001$ ].

Finally, we looked at whether the same factors from experiment one were able to predict when observers would stop sampling in the orientation paradigm. We did not replicate the dwell times/sample latency analyses since the sample latency is fixed for each orientation stimulus in the current paradigm.

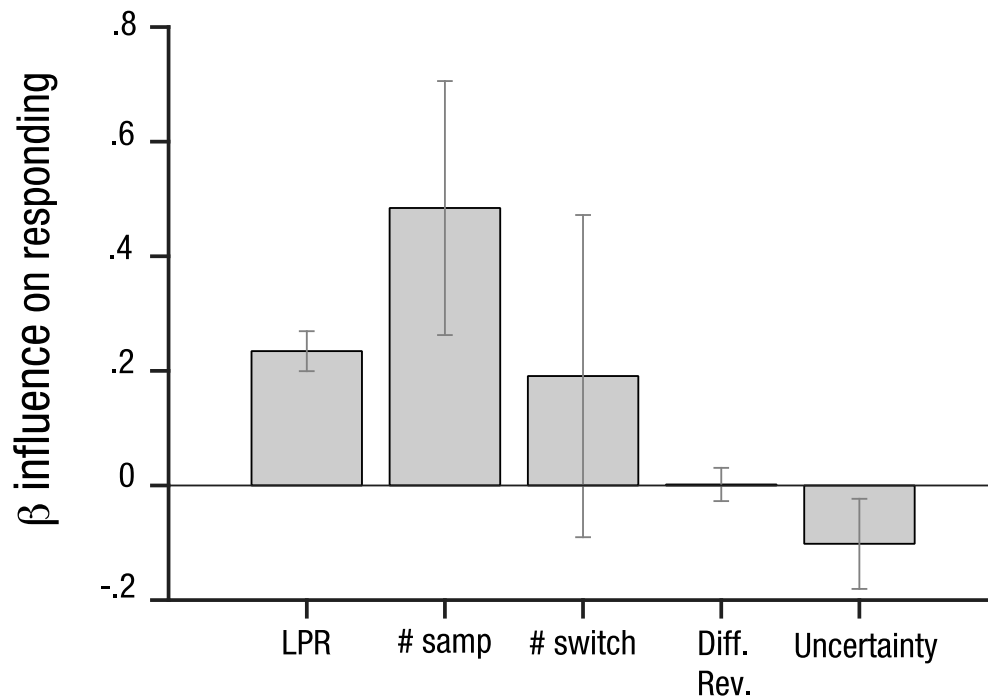
### **6.3.8 Sampling policy: when to stop sampling?**

We used linear regression to predict observers' sample latency as a function of the same six predictors: (1) the absolute cumulative evidence captured by the LPR after the sample  $S$ , (2) the number of samples after the sample  $S$ , (3) the number of switches after sample  $S$ , (4) the difference in number of samples taken from the most – least clockwise category after sample  $S$ , and the posterior width of the category that was sampled at sample  $S$ .

The regression model we used was as follows, where in each case:

$$p(\text{respond}_S) = \phi[b + w1 \cdot |LPR_S| + w2 \cdot (n\_samples_S) + w3 \cdot (n\_switches_S) \\ + w4 \cdot ((n\_most_S - n\_least_S) + w5 \cdot (Posterior_S))]$$

where  $\phi$  denotes the cumulative normal distribution,  $|LPR_S|$  denotes the absolute LPR after the sample  $S$ ,  $n\_samples_S$  denotes the total number of samples taken after sample  $S$ ,  $n\_switches_S$  denotes the total number of switches taken after sample  $S$ , and  $n\_Most_S - n\_Least_S$  denotes the difference in number of samples taken from the most and least clockwise category after sample  $S$ .  $Posterior_S$  denotes the estimated width of the posterior after sample  $S$  for the category that was sampled on sample  $S$ . The same factors from chapter 5 were able to predict when observers stop sampling in experiment two (**Fig. 6.16**). First, observers were more likely to respond when the evidence for either hypothesis was stronger [ $t_{(17)} = 13.1$ ;  $P < 0.001$ ]. The number of samples taken also significantly predicted the likelihood of responding: observers were more likely to respond as the number of samples taken increased [ $t_{(17)} = 4.3$ ;  $P < 0.001$ ]. The number of switches made and differences in the number of elements sampled between the two categories had no significant effect on the tendency to respond [ $P > 0.05$ ]. Finally, observers were less likely to respond when they were more uncertain about the category they had just sampled from [ $t_{(17)} = -2.5$ ;  $P < 0.05$ ; **Fig. 6.16**].



**Fig 6.16 Predictors of the tendency to respond.** The absolute LPR and the number of samples taken had a significant positive effect, while the uncertainty had a significant negative effect on the tendency to respond. Bars are shown with 95% confidence intervals.

## 6.4 Discussion

While we were able to replicate the asymptotic performance and recency findings, we were unable to replicate the consistency and frequency biases in a design that is more suitable for pupil size investigation and neural data acquisition such as EEG or fMRI.

Interestingly, the analyses investigating drivers of the sampling policy were broadly consistent between paradigms. In both paradigms, the identity of the next sample was predicted by the difference in the number of samples taken from each category: observers were more likely to sample elements from the

category they had sampled least from. Interestingly this effect was strongest when the LPR was positive, i.e. when the evidence that they had sampled was in favour of the correct option. Observers showed evidence for an ‘equity’ policy in both tasks, whereby they preferred to sample an equal amount of elements from each category. Importantly, this effect was replicated in chapter 6, showing that the tendency to sample elements from the category that was sampled the least was not due to a finite number of samples present on the screen.

The tendency to switch was again predicted by the uncertainty associated with the sampled category. As observers became more certain about their estimate, they were more likely to switch and start sampling the other category. Moreover, the policy that determines whether an observer would respond was not only influenced by the passage of time and the cumulative sampled evidence, but also by the uncertainty.

The failed replications could be influenced by the longer interval between requesting a new sample of information and sample presentation in the orientation paradigm compared to the gaze contingent paradigm or because the orientation stimuli are associated with higher perceptual ambiguity. While this longer interval would be necessary for fMRI data acquisition, it is not a necessity for EEG or pupillometry studies.

Moreover, the sampling policy was more heterogeneous in the orientation paradigm compared to the gaze contingent paradigm. While almost every observer in the gaze contingent paradigm used a summative sampling

strategy, an extensive proportion of observers used a round-wise sampling strategy in the orientation paradigm.

It remains an open question whether the emergence of a round-wise strategy in our sample compared to the strategies used in chapter 5, is due to the paradigm itself (e.g. the longer interval between requesting a stimulus and receiving a stimulus), the orientation stimuli, or individual characteristics of our participants. This could easily be tested in a follow-up experiment using number stimuli in which the length of the interval between requesting the stimulus and receiving the stimulus is manipulated.

However, it is clear that the context or decision environment can heavily influence the sampling policy. For example, while some paradigms expose strong primacy effects (Ludwig & Evens, 2016), we have demonstrated a recency effects. Similarly, while observers mainly employed summative sampling in the gaze contingent experiment, a number of observers employed round-wise sampling in the orientation paradigm.

It would be interesting to investigate individual differences between these two types of participants. For example, differences in working memory capacity could influence the strategy that observers employ. For example, observers with a lower working memory capacity might opt for round-wise sampling strategies. Similarly, it would be interesting to see whether observers who employ one strategy (e.g. round-wise sampling) would perform worse if they were forced to use the other strategy (e.g. summary sampling). Moreover, recent research has started to highlight individual differences in gaze biases

and their effects on behaviour (Thomas, Molter, Krajbich, Heekeren, & Mohr, 2017)

The degree of flexibility of the sampling policy still remains an open question. Chapter 5 and 6 have provided clear evidence for the fact that the sampling policy is influenced by on-going cognitive processing, suggesting that the sampling policy could be flexible and adaptive. Similarly, in chapter 4 we have shown that the preference to sample inliers is an adaptive strategy because it will lead to an improved accuracy for participants due to an interplay between their own cognitive constraints and the decision environment.

However, recent research has shown that the frequency bias was applied even when it lead to more inconsistent choices (Thomas et al., 2017). This opens up an interesting question: are sampling policies static within people or do they change depending on the context? Moreover, does this depend on the particular sampling policy (e.g. frequency bias versus summative/round-wise sampling)?

## **Chapter 7: Conclusion**

In this thesis, we have emphasised the importance of considering information sampling as an integral component of the decision making process. To do this, we have used behavioural and eye-tracking studies in combination with computational modelling. In this chapter, I am going to highlight the important findings from previous chapters.

### **7.1 Identified drivers of the sampling policy**

We have provided evidence for a variety of drivers of the sampling policy.

#### **7.1.1 Confirmation bias at the perceptual feature level**

In chapter three, when the decision value of stimuli was defined by the combination of two perceptual characteristics (e.g. colour and shape), the sampling policy was driven by a colour and shape confirmation bias. The decision values themselves did not affect the sampling policy. Thus, when the decision value of stimuli was defined by a combination of perceptual feature values, the sampling biases were situated at the level of perceptual features. However, it remained unclear whether the absence of an effect of the decision information on the sampling policy was due to the experimental design (e.g. decision information level biases cannot occur within a short enough timeframe to drive eye-movements) or due to the stimulus (e.g. the decision information is constructed by the combination of each perceptual feature value).

### **7.1.2 Confirmation bias at the decision information level**

Thus, in chapter four we adapted the paradigm of chapter three, and used two-digit numbers as the stimuli in order to disambiguate these possibilities. Thus, the decision information was not carried by the perceptual information, but rather by the symbolic meaning of the stimulus. Furthermore, the perceptual characteristics of these two-digit numbers were not as salient as the shape and colour feature values of the stimuli in chapter three. The experiments in chapter four provided evidence for two distinct effects of how the decision values of the sampled elements can influence which element is sampled next. First, we provided evidence for a confirmation bias driven by the decision values: observers were more likely than chance to sample information that confirmed their current assessment of the stimulus array. Thus, decision values could drive the sampling policy even when information is sampled very rapidly.

### **7.1.3 Robust sampling**

Furthermore, we showed that observers were robust samplers: observers were more likely than chance to sample inlying (i.e. elements close to the category boundary) compared to outlying elements (i.e. elements further away from the category boundary). Moreover, we showed that the original robust averaging effect was at least in part due to robust sampling.

#### **7.1.4 Summary sampling**

Chapter six showed that when integrating evidence from two categories in a gaze-contingent design, observers tend to sample as consistent as possible (e.g. minimise the number of switches), a strategy referred to as 'summary sampling'. Observers focussed on sampling an option extensively before switching to the other option.

#### **7.2.5 Equity sampling**

Observers were also more likely to sample evidence from the category they had sampled the least from, and this effect was influenced by the accrued evidence. Importantly, this effect was not due to the finite number of samples present on the stimulus display, as the effect was replicated in the orientation paradigm mentioned in chapter six. In this orientation paradigm observers could sample an infinite number of stimuli. Interestingly however, we were not able to replicate all of the effects from chapter five in this orientation paradigm. Moreover, we distinguished two distinct types of samplers in our group of participants: summary samplers and round-wise samplers. Thus, the sampling policy of observers is not only determined by the decision environment but also by individual differences between observers.

## 7.2 Effects of the sampling policy on decision-making

We have not only investigated drivers of the sampling policy in this thesis. We have also shown that the sampling policy can affect decision making, even when information is sampled in rapid succession.

Chapter three showed a benefit of voluntary resampling compared to no resampling and random resampling.

Chapter four showed that the robust sampling bias could be beneficial under certain circumstances, and showed that the original robust averaging effect was at least in part due to robust sampling.

Chapter 5 provided evidence for two distinct biases of the sampling policy on accuracy. (1) A consistency bias, whereby observers were more likely to be correct when their sampling was more consistent (e.g. when they made proportionally fewer switches). And (2) a frequency bias whereby observers were more likely to choose the category from which they sampled the most, independent of the sampled evidence itself. Furthermore, our results suggest that this frequency bias is not a multiplicative effect as suggested by the attentional DDM, but rather an additive bias (at least in our gaze contingent paradigm).

### **7.3 Development of the landscape approach to analyse eye-movement data**

Finally, we have developed a novel analysis approach for eye-tracking data that jointly estimates the size of the sampling window and the decision information harvested therein. Unlike other methods, the landscape approach defines the width of the information-processing spotlight within the same task, independently from the decision values. The approach is thus a valuable alternative to the commonly used retinal sensitivity mapping. This is particularly relevant because recent research has shown that the use of retinal sensitivity mapping can sometimes lead to the overestimation of the degree to which information is processed around each fixation (Morvan & Maloney, 2012). Instead of requiring an additional retinal sensitivity mapping session, researchers could thus instead rely on a proxy of the ‘utilised’ retinal sensitivity as estimated from gaze behaviour on the task by the landscape approach.

Furthermore, the landscape approach allows researchers ample flexibility. For example, researcher could use soft filters, where the information falling further from the point of gaze has a weaker influence on choices. This can be a useful tool when designing experimental designs, as peripheral discrimination thresholds are different for certain for different feature dimensions such as colour and shape. Alternatively, they could use hard filters in which information within the boundaries of the filter is sampled equally. Similarly, the approach is immune to otherwise arbitrary choices made by the experimenter

during more standard fixation analyses (e.g. criteria for determining fixations, size and shape of the fixation windows...).

Finally, and perhaps most importantly, the approach allows the researcher to not only quantify the degree to which an element is sampled (e.g. the width of the processing spotlight), but allows the researcher to measure the sampled value of any construct related to the sampled stimulus. For example in chapter four, we constructed prototypicality landscapes in which the height of the filters was determined by how in- or out-lying the sampled stimulus was. Similarly, we constructed evidence landscapes in which the height of the filters was determined by the decision value of the sampled stimulus. The landscape approach is thus ideally placed for combining eye tracking methods with computational models of decision making.

#### **7.4 The importance of acknowledging the sampling policy as an integral stage of the decision making process.**

We have provided evidence for drivers of the sampling policy as well as evidence for the direct effects this sampling policy can have on the decision making process. While the experimental paradigms used in this thesis are particularly suited for eliciting overt sampling behaviours, this is not necessarily required in order for the sampling policy to influence the decision making process.

For example, Poletti et al. recently provided evidence for selective attention within the fovea (Poletti, Rucci, & Carrasco, 2017). Thus, even in the

traditional decision making paradigms in which observers are instructed to maintain central fixation and stimuli are presented within the foveal locus, the sampling policy might be affecting the decision making process.

## **Appendix.**

### **Appendix 1.1 Capturing the recency, frequency, and consistency bias**

We aimed to capture the recency, frequency, and consistency bias using two relatively simple computational models. We used two different theoretical models: (1) a model that estimated the probability of either category being the winner using Bayesian inference implementing different variants of a memory leak term, and (2) the adaptive gain model in which the gain of information processing adapts as a function of the previously sampled information.

#### **Appendix 1.1 A Bayesian model implementing different implementations of a leaky process.**

One possible reason for the recency, frequency, and consistency bias is that human observers are limited capacity observers. The asymptotic performance, recency effect, and switch cost led us to believe that participants struggle to integrate all the available information due to capacity limitation. As a result, observers might try to adopt a sampling strategy as a compensatory strategy for these limited capacity constraints. In the Bayesian model we will describe next, we try to simulate this limited capacity by having different instantiations of leak on the accumulated evidence.

Our Bayesian model assumes that the state space in which observers make estimates about the true mean ranges from -14 to 14 in this experiment (e.g.

the numbers in experiment one were always distributed  $\pm 14$  around the reference value of that trial). We assumed the prior belief over the state space to be uniform at the start of each trial (i.e. observers are not biased towards a particular category colour:  $P(R_M) = P(B_M)$ ). As observers take more samples, this prior belief will be updated. Every sample itself is characterised by a normal probability density function centered on the observed value (e.g. the sampled number) with an uncertainty  $\sigma_s$ . Thus each sample is associated with a normal probability density function (the likelihood function) reflecting the probability of observing that sample, given the (unknown) true mean of its category ( $P(S_i|R_M)$  or  $P(S_i|B_M)$ ). Implementing Bayes rule yields

$$P(R_M|S_i) = \frac{P(S_i|R_M) \times P(R_M)}{\sum_M P(S_i|R_M) \times P(R_M)}$$

where  $P(R_M|S_i)$  is the posterior probability that the true mean of the sampled category (e.g. red in this case) is equal to M, given the value of the sampled number  $S_i$ . This posterior probability is normalised via the denominator  $\sum_M P(S_i|R_M) \times P(R_M)$  after each sample. Crucially, this posterior is influenced by a free parameter of the model: leak value  $\lambda$ . The leak was implemented by simply adding  $\lambda$  to the width of the posterior probability after sample  $S_i$ , resulting in a higher uncertainty of the estimate for higher values of  $\lambda$ . This posterior probability will form the prior for the next sample of the same category and this process continues for subsequent samples. Thus under this model, observers will have obtained two separate estimates for the mean of each category given the sampled evidence at the end of each trial. Trial by

trial model responses are estimated by calculating the ROC (via the trapezoid rule) based on the cumulative posterior probabilities of each category. When the area under the curve is 0.5, the model is equally likely to choose low or high. Values higher (lower) than 0.5 indicate a preference for the high (low) distribution.

In order to capture the asymptotic performance, recency, frequency and consistency biases, we implemented four distinct implementations of a leaky process. The model could either have no leak ( $\lambda = 0$ ), a leak proportionate to the number of elements sampled (more elements sampled leads to a higher leak  $\lambda$ ), only leak on the non-sampled category (e.g. memory decays on the inactive category), or leak on both categories regardless of the sampled category. The two free parameters uncertainty  $\sigma_s$  and leak  $\lambda$  were estimated by fitting model accuracy to human accuracy (split by conditions or split by fixation lengths). The log-likelihood was used as our measure of goodness of fit:

$$\ln L(model|\sigma, \lambda) = - \sum_i ((T_i^{corr} * \log(ROC_i)) + (T_i^{inc} * \log(1 - ROC_i)))$$

where  $i$  represents the various conditions or fixation lengths.  $T^{corr}$  represents the number of trials where the participant was correct and  $T^{inc}$  the number of trials where the participant was incorrect.  $ROC_i$  is the probability of the model choosing the correct response for condition  $i$ . These negative log-likelihoods are calculated for all combinations of  $\sigma$  and  $\lambda$  in our grid search and the parameter combination with the lowest negative log-likelihood is chosen as

the best fit for our data (gradient descent algorithms yielded the same best fitting parameter values).

All three Bayesian models implementing some sort of non-zero leak  $\lambda$  could capture the asymptotic performance and recency bias. However, none of the four implementations could replicate the frequency and consistency bias.

| <i><b>Bayesian leak</b></i>   |                |
|-------------------------------|----------------|
| <b>Asymptotic performance</b> | Replicated     |
| <b>Recency bias</b>           | Replicated     |
| <b>Frequency bias</b>         | Not replicated |
| <b>Consistency bias</b>       | Not replicated |

**Table 1. Bayesian models implementing different variations of a leak parameter.** The asymptotic performance and recency bias were replicated by all four implementations. However, none of these models were able to replicate the frequency and consistency bias.

## **Appendix 1.2 An adaptive gain-field model**

The adaptive gain-field model is an adaptation of the adaptive gain control model described earlier (Cheadle et al., 2014). The key idea underlying the adaptive gain control model is that the information processing rapidly adapts as a function of the environment. Adaptive gain control has been shown to occur very rapidly (Cheadle et al., 2014; Li et al., 2017). Furthermore, Michael, de Gardelle, & Summerfield (2014) showed rapid adaptation to variance in a priming experiment, providing evidence for an adaptation process that can happen in under 100 ms. Thus, it is entirely possible that subsequent samples in our experiment could be affected by previous samples when scanning the display. The adaptive gain control model has been used to explain recency effects and to show that information that is similar to previous information is more influential on choice (Cheadle et al., 2014). Therefore, our hypothesis was that observers in our task might be sampling information in a way to maximize gain-allocation within a single trial in order to cope with their limited capacity.

In this adaptive gain-field model adaptation each sampled number is passed through a Gaussian gain-field representing the likelihood of particular features across the number (decision) space (GF). We assumed that at the beginning of the trial, the gain field starts off with a global sensitivity towards the category boundary (i.e. where the reference lies). This gain field adapts according to the local environmental statistics i.e. seen samples. The gain

field ( $GF$ ) towards the category boundary can be denoted as a probability density function of the normal distribution:

$$GF_n^0 = N(f_n, 0, \sigma_g)$$

Where  $N(f_n, 0, \sigma_g)$  corresponds to the probability density of a Gaussian evaluated at each possible feature  $n$  ( $n \in -14, -13 \dots 13, 14$ ). Here, the Gaussian probability density function has a mean of 0, which represents the mean-centered reference of the trial, and a standard deviation  $\sigma_g$ , which is a free parameter that represents the global sensitivity towards the category boundary. In other words, the lower  $\sigma_g$  is, the greater sensitivity towards the trial reference (indicated by the heightened Gaussian peak at where the reference lies).

This gain field is then updated based on the local environment. On each sample, the “sample gain field” updates the previous gain field:

$$GF_n^i = N(f_n, S^i, \sigma_s)$$

Where  $N(f_n, S^i, \sigma_s)$  corresponds to the probability density of a Gaussian with a mean at where the sample  $i$  ( $S^i$ ) lies, and a standard deviation of  $\sigma_s$ , evaluated at each possible feature  $n$  ( $n \in -14, -13 \dots 13, 14$ ). The width of the probability density function  $\sigma_s$ , is another free parameter of the model, which corresponds to the sensitivity towards each sample. A low  $\sigma_s$  would mean that samples would receive higher gain. For each sample  $i$ , the gain

field (GF) is updated by taking the product of the previous GF with the sample gain field.

In order to get the model estimate of the cumulative evidence from each sample, we first assume uncertainty or noise towards sample  $i$  by introducing random noise from a zero-mean Gaussian distribution with a standard deviation of  $\xi_{\text{early}}$  to the sample. The parameter  $\xi_{\text{early}}$  is the third free parameter of the model, controlling the uncertainty over the samples. Lastly the model estimated sample ( $\theta_i$ ) is computed as follows:

$$\theta_i = X_i \cdot S_i$$

where  $X_i$  is the gain field value at where the noisy sample lies. The model estimated cumulative evidence for sample  $i$  ( $\theta_i$ ) is computed by summing the weighted product of the noisy sample's corresponding gain field value and its noisy sample decision variable (i.e.  $S_i$ ). The model response is simply calculated by comparing the cumulative tallies of cumulative evidence for each category while corrupting this estimate by late zero-mean Gaussian noise  $\xi_{\text{late}}$  (the fourth free parameter). Goodness of fit was evaluated in the same way as for the Bayesian leak models.

We have tried four different implementations of this adaptive gain control model. First, the GF update could happen either before or after the integration of a sample. Second, the GF update could be centered on the running average of the sampled information or on the fixated number. However, none

of the different implementations were able to predict the recency and consistency bias (**Table 2**).

|                               | <i>Bayesian leak</i> | <i>Adaptive gain field model</i> |
|-------------------------------|----------------------|----------------------------------|
| <b>Asymptotic performance</b> | Replicated           | Replicated                       |
| <b>Recency bias</b>           | Replicated           | Not replicated                   |
| <b>Frequency bias</b>         | Not replicated       | Replicated                       |
| <b>Consistency bias</b>       | Not replicated       | Not replicated                   |

**Table 2.** Neither the Bayesian nor adaptive gain field model were able to accurately describe all four behavioural effects.

## Appendix 2.1 List of the exploratory analyses in this thesis.

The analysis plan was always formulated a priori for each experiment.

However, I have indicated the exploratory analyses (e.g. those analyses for which we did not have any specified predictions) in the table below.

| <b><u>Exploratory analyses</u></b>                                                                                                                                                                                                                                                                                                                                                                       |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b><u>Chapter 3.</u></b><br><ul style="list-style-type: none"> <li>- While we expected evidence for a confirmation bias, we did not have any predictions regarding the colour/shape confirmation bias.</li> </ul>                                                                                                                                                                                        |
| <b><u>Chapter 4.</u></b><br><ul style="list-style-type: none"> <li>- Investigation of the relationship between performance and sampling policy.</li> </ul>                                                                                                                                                                                                                                               |
| <b><u>Chapter 5 &amp; 6.</u></b><br><ul style="list-style-type: none"> <li>- The influence of the number of samples on accuracy.</li> <li>- The influence of stimuli that were sampled early (primacy) versus late (recency) on accuracy.</li> <li>- Predicting the category of the next sample.</li> <li>- Predicting dwell times.</li> <li>- Predicting when observers would stop sampling.</li> </ul> |

## References

- Ahmad, S., & Yu, A. J. (2013). Active sensing as bayes-optimal sequential decision making. *arXiv preprint arXiv:1305.6650*.
- Allais, M. (1953). L'extension des théories de l'équilibre économique général et du rendement social au cas du risque. *Econometrica, Journal of the Econometric Society*, 269-290.
- Anderson, B. A., Laurent, P. A., & Yantis, S. (2011). Learned value magnifies salience-based attentional capture. *PloS one*, 6(11), e27926.
- Armel, K. C., & Rangel, A. (2008). The impact of computation time and experience on decision values. *The American Economic Review*, 98(2), 163-168.
- Arrington, C. M., & Logan, G. D. (2004). The cost of a voluntary task switch. *Psychological science*, 15(9), 610-615.
- Atchison, D. A., Smith, G., & Smith, G. (2000). Optics of the human eye.
- Baldwin, A., Baker, D., & Hess, R. (2016). What do contrast threshold equivalent noise studies actually measure? Noise vs. nonlinearity in different masking paradigms. *PloS one*, 11(3), e0150942.
- Baldwin, J., Burleigh, A., Pepperell, R., & Ruta, N. (2016). The Perceived Size and Shape of Objects in Peripheral Vision. *i-Perception*, 7(4).
- Ballard, D. H., Hayhoe, M. M., & Pelz, J. (1995). 15 Memory Limits in Sensorimotor Tasks. *Models of information processing in the basal ganglia*.
- Barkan, R., Zohar, D., & Erev, I. (1998). Accidents and decision making under uncertainty: A comparison of four models. *Organizational Behavior and Human Decision Processes*, 74(2), 118-144.
- Barlow, H. B. (1961). Possible principles underlying the transformations of sensory messages.
- 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.
- Bastardi, A., & Shafir, E. (1998). On the pursuit and misuse of useless information. *Journal of personality and social psychology*, 75, 19-32.
- Betz, T., Kietzmann, T. C., Wilming, N., & Koenig, P. (2010). Investigating task-dependent top-down effects on overt visual attention. *Journal of Vision*, 10(3), 15-15.
- Birch, E. E., Shimojo, S., & Held, R. (1985). Preferential-looking assessment of fusion and stereopsis in infants aged 1-6 months. *Investigative ophthalmology & visual science*, 26(3), 366-370.
- 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.
- Boorman, E. D., Rushworth, M. F., & Behrens, T. E. (2013). Ventromedial prefrontal and anterior cingulate cortex adopt choice and default reference frames during sequential multi-alternative choice. *Journal of Neuroscience*, 33(6), 2242-2253 %@ 0270-6474.
- Brezis, N., Bronfman, Z. Z., & Usher, M. (2015). Adaptive spontaneous transitions between two mechanisms of numerical averaging. *Scientific reports*, 5, 10415.

- Britten, K. H., Shadlen, M. N., Newsome, W. T., & Movshon, J. A. (1992). The analysis of visual motion: a comparison of neuronal and psychophysical performance. *The Journal of Neuroscience*, 12(12), 4745-4765.
- Broadbent, D. E. (1958). The effects of noise on behaviour.
- Burgess, A. E., & Colborne, B. (1988). Visual signal detection. IV. Observer inconsistency. *JOSA A*, 5(4), 617-627.
- Buswell, G. T. (1935). *How people look at pictures*: University of Chicago Press Chicago.
- Camerer, C. (2000). Iterated Reasoning in Dominance-solvable Games: ch.
- Carandini, M., Demb, J. B., Mante, V., Tolhurst, D. J., Dan, Y., Olshausen, B. A., . . . Rust, N. C. (2005). Do we know what the early visual system does? *Journal of Neuroscience*, 25(46), 10577-10597 %@ 10270-16474.
- Cassey, T. C., Evens, D. R., Bogacz, R., Marshall, J. A., & Ludwig, C. J. (2013). Adaptive sampling of information in perceptual decision-making. *PloS one*, 8(11), e78993.
- Cerf, M., Frady, E. P., & Koch, C. (2009). Faces and text attract gaze independent of the task: Experimental data and computer model. *Journal of Vision*, 9(12), 10-10.
- Cheadle, S., Wyart, V., Tsetsos, K., Myers, N., De Gardelle, V., Castañón, S. H., & Summerfield, C. (2014). Adaptive gain control during human perceptual choice. *Neuron*, 81(6), 1429-1441.
- Chelazzi, L., Perlato, A., Santandrea, E., & Della Libera, C. (2013). Rewards teach visual selective attention. *Vision research*, 85, 58-72.
- De Gardelle, V., & Summerfield, C. (2011). Robust averaging during perceptual judgment. *Proceedings of the National Academy of Sciences*, 108(32), 13341-13346.
- Dijksterhuis, A. (2004). Think different: the merits of unconscious thought in preference development and decision making. *Journal of personality and social psychology*, 87(5), 586.
- Dijksterhuis, A., Aarts, H., & Smith, P. K. (2005). The power of the subliminal: On subliminal persuasion and other potential applications. *The new unconscious*, 1.
- Ditterich, J. (2006). Stochastic models of decisions about motion direction: behavior and physiology. *Neural Networks*, 19(8), 981-1012.
- Dreher, J.-C., & Tremblay, L. (2016). *Decision Neuroscience: An Integrative Perspective*: Academic Press.
- Drugowitsch, J., Moreno-Bote, R., Churchland, A. K., Shadlen, M. N., & Pouget, A. (2012). The cost of accumulating evidence in perceptual decision making. *Journal of Neuroscience*, 32(11), 3612-3628.
- Durgin, F. H., Doyle, E., & Egan, L. (2008). Upper-left gaze bias reveals competing search strategies in a reverse Stroop task. *Acta Psychologica*, 127(2), 428-448.
- Ebbinghaus, H. (1913). *Memory: A contribution to experimental psychology*: University Microfilms.
- Eckstein, M. P. (2011). Visual search: A retrospective. *Journal of Vision*, 11(5), 14.
- Eckstein, M. P., Schoonveld, W., Zhang, S., Mack, S. C., & Akbas, E. (2015). Optimal and human eye movements to clustered low value cues to increase decision rewards during search. *Vision research*, 113, 137-154.

- Einhäuser, W., & König, P. (2003). Does luminance-contrast contribute to a saliency map for overt visual attention? *European Journal of Neuroscience*, 17(5), 1089-1097.
- Erev, I. (1998). Signal detection by human observers: a cutoff reinforcement learning model of categorization decisions under uncertainty. *Psychological review*, 105(2), 280.
- Erev, I., & Barron, G. (2005). On adaptation, maximization, and reinforcement learning among cognitive strategies. *Psychological review*, 112(4), 912.
- Faisal, A. A., Selen, L. P. J., & Wolpert, D. M. (2008). Noise in the nervous system. *Nature reviews. Neuroscience*, 9(4), 292.
- Fantz, R. L. (1964). Visual experience in infants: Decreased attention to familiar patterns relative to novel ones. *science*, 146(3644), 668-670.
- Feigenson, L., Dehaene, S., & Spelke, E. (2004). Core systems of number. *Trends in cognitive sciences*, 8(7), 307-314.
- Gagl, B., Hawelka, S., & Hutzler, F. (2011). Systematic influence of gaze position on pupil size measurement: analysis and correction. *Behavior research methods*, 43(4), 1171-1181.
- Glaholt, M. G., & Reingold, E. M. (2009). Stimulus exposure and gaze bias: A further test of the gaze cascade model. *Attention, Perception, & Psychophysics*, 71(3), 445-450.
- Glaholt, M. G., & Reingold, E. M. (2009). The time course of gaze bias in visual decision tasks. *Visual Cognition*, 17(8), 1228-1243.
- Glaholt, M. G., Wu, M.-C., & Reingold, E. M. (2010). Evidence for top-down control of eye movements during visual decision making. *Journal of Vision*, 10(5), 15-15.
- Glaze, C. M., Kable, J. W., & Gold, J. I. (2015). Normative evidence accumulation in unpredictable environments. *Elife*, e08825.
- Gold, J. I., & Shadlen, M. N. (2007). The neural basis of decision making. *Annu. Rev. Neurosci.*, 30, 535-574.
- Gottlieb, J., Hayhoe, M., Hikosaka, O., & Rangel, A. (2014). Attention, reward, and information seeking. *The Journal of Neuroscience*, 34(46), 15497-15504.
- Green, D. M., & Swets, J. A. (1974). Signal detection theory and psychophysics (Rev. ed.). *Huntington, NY: RF Krieger*.
- Groopman, J. E., & Prichard, M. (2007). *How doctors think* (Vol. 82): Springer.
- Haberman, J., & Whitney, D. (2010). The visual system discounts emotional deviants when extracting average expression. *Attention, Perception, & Psychophysics*, 72(7), 1825-1838.
- Hanks, T., & Summerfield, C. (2017). Perceptual Decision Making in Rodents, Monkeys, and Humans. *Neuron*, 93(1).
- Hansen, T., Pracejus, L., & Gegenfurtner, K. R. (2009). Color perception in the intermediate periphery of the visual field. *Journal of Vision*, 9(4), 26-26.
- Hayhoe, M., & Ballard, D. (2005). Eye movements in natural behavior. *Trends in cognitive sciences*, 9(4), 188-194.
- Hayhoe, M., Mennie, N., Sullivan, B., & Gorgos, K. (2005). *The role of internal models and prediction in catching balls*.
- He, P., & Kowler, E. (1989). The role of location probability in the programming of saccades: Implications for "center-of-gravity" tendencies. *Vision research*, 29(9), 1165-1181.

- Henderson, J. M., & Ferreira, F. (1990). Effects of foveal processing difficulty on the perceptual span in reading: implications for attention and eye movement control. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 16(3), 417-429.
- Henderson, J. M., & Hollingworth, A. (1999). High-level scene perception. *Annual review of psychology*, 50(1), 243-271.
- Hertwig, R., Barron, G., Weber, E. U., & Erev, I. (2006). The role of information sampling in risky choice. *Information sampling and adaptive cognition*, 72-91.
- Hertwig, R., & Erev, I. (2009). The description–experience gap in risky choice. *Trends in cognitive sciences*, 13(12), 517-523.
- Hickey, C., Chelazzi, L., & Theeuwes, J. (2010). Reward changes salience in human vision via the anterior cingulate. *Journal of Neuroscience*, 30(33), 11096-11103.
- Hills, T. T., & Hertwig, R. (2010). Information search in decisions from experience: do our patterns of sampling foreshadow our decisions? *Psychological science*, 21(12), 1787-1792.
- Howes, A., Warren, P. A., Farmer, G., El-Deredy, W., & Lewis, R. L. (2016). Why contextual preference reversals maximize expected value. *Psychological review*, 123(4), 368.
- 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), 90-98.
- Hyde, D. C., & Spelke, E. S. (2011). Neural signatures of number processing in human infants: evidence for two core systems underlying numerical cognition. *Developmental science*, 14(2), 360-371.
- Inhoff, A. W., & Rayner, K. (1986). Parafoveal word processing during eye fixations in reading: Effects of word frequency. *Perception & Psychophysics*, 40(6), 431-439.
- Itti, L., & Baldi, P. (2009). Bayesian surprise attracts human attention. *Vision research*, 49(10), 1295-1306.
- Itti, L., & Koch, C. (2000). A saliency-based search mechanism for overt and covert shifts of visual attention. *Vision research*, 40(10), 1489-1506.
- Itti, L., Koch, C., & Niebur, E. (1998). A model of saliency-based visual attention for rapid scene analysis. *IEEE Transactions on Pattern Analysis & Machine Intelligence*(11), 1254-1259.
- Jansen, L., Onat, S., & König, P. (2009). Influence of disparity on fixation and saccades in free viewing of natural scenes. *Journal of Vision*, 9(1), 29-29.
- Jersild, A. T. (1927). Mental set and shift. *Archives of psychology*.
- Kahneman, D. (2011). *Thinking, fast and slow*: Macmillan.
- Kahneman, D., & Tversky, A. (1979). Prospect theory: An analysis of decision under risk. *Econometrica: Journal of the econometric society*, 263-291.
- Kahneman, D., & Tversky, A. (1984). Choices, values, and frames. *American psychologist*, 39(4), 341.
- Khaw, M. W., Glimcher, P. W., & Louie, K. (2017). Normalized value coding explains dynamic adaptation in the human valuation process. *Proceedings of the National Academy of Sciences*, 201715293 %@ 201710027-201718424.

- Kienzle, W., Franz, M. O., Schölkopf, B., & Wichmann, F. A. (2009). Center-surround patterns emerge as optimal predictors for human saccade targets. *Journal of Vision*, 9(5), 7-7.
- Klauer, K. C., Voss, A., Schmitz, F., & Teige-Mocigemba, S. (2007). Process components of the Implicit Association Test: a diffusion-model analysis. *Journal of personality and social psychology*, 93(3), 353.
- Klayman, J., & Ha, Y.-W. (1987). Confirmation, disconfirmation, and information in hypothesis testing. *Psychological review*, 94(2), 211.
- Klayman, J., & Ha, Y.-w. (1989). Hypothesis testing in rule discovery: Strategy, structure, and content. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 15(4), 596.
- Koch, C., & Ullman, S. (1987). Shifts in selective visual attention: towards the underlying neural circuitry *Matters of intelligence* (pp. 115-141): Springer.
- Kolling, N., Behrens, T. E. J., Mars, R. B., & Rushworth, M. F. S. (2012). Neural mechanisms of foraging. *science*, 336(6077), 95-98 %@ 0036-8075.
- Krajbich, I., Armel, C., & Rangel, A. (2010). Visual fixations and the computation and comparison of value in simple choice. *Nature neuroscience*, 13(10), 1292-1298.
- Krajbich, I., & Rangel, A. (2011). Multialternative drift-diffusion model predicts the relationship between visual fixations and choice in value-based decisions. *Proceedings of the National Academy of Sciences*, 108(33), 13852-13857.
- Kunst-Wilson, W. R., & Zajonc, R. B. (1980). Affective discrimination of stimuli that cannot be recognized. *science*, 207(4430), 557-558.
- Kusunoki, M., Gottlieb, J., & Goldberg, M. E. (2000). The lateral intraparietal area as a salience map: the representation of abrupt onset, stimulus motion, and task relevance. *Vision research*, 40(10), 1459-1468.
- Lakatos, P., Karmos, G., Mehta, A. D., Ulbert, I., & Schroeder, C. E. (2008). Entrainment of neuronal oscillations as a mechanism of attentional selection. *science*, 320(5872), 110-113.
- Land, M., Mennie, N., & Rusted, J. (1999). The roles of vision and eye movements in the control of activities of daily living. *Perception*, 28(11), 1311-1328.
- Land, M. F. (2006). Eye movements and the control of actions in everyday life. *Progress in retinal and eye research*, 25(3), 296-324.
- Li, V., Castañón, S. H., Solomon, J. A., Vandormael, H., & Summerfield, C. (2017). Robust Averaging Protects Decisions from Noise in Neural Computations. *PLoS Comput Biology*, 13(8), e1005723.
- Lichtenstein, S., & Slovic, P. (1973). Response-induced reversals of preference in gambling: An extended replication in Las Vegas. *Journal of Experimental Psychology*, 101(1), 16.
- Lim, S.-L., O'Doherty, J. P., & Rangel, A. (2011). The decision value computations in the vmPFC and striatum use a relative value code that is guided by visual attention. *Journal of Neuroscience*, 31(37), 13214-13223.
- 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.
- Luce, R. D. (1977). The choice axiom after twenty years. *Journal of mathematical psychology*, 15(3), 215-233.

- Ludwig, C. J. H., Davies, J. R., & Eckstein, M. P. (2014). Foveal analysis and peripheral selection during active visual sampling. *Proceedings of the National Academy of Sciences*, 111(2), E291-E299.
- Ludwig, C. J. H., & Evens, D. R. (2016). Information foraging for perceptual decisions.
- Manohar, S. G., Chong, T. T. J., Apps, M. A. J., Batla, A., Stamelou, M., Jarman, P. R., . . . Husain, M. (2015). Reward pays the cost of noise reduction in motor and cognitive control. *Current Biology*, 25(13), 1707-1716.
- Manohar, S. G., & Husain, M. (2013). Attention as foraging for information and value. *Frontiers in human neuroscience*, 7.
- Maris, E., & Oostenveld, R. (2007). Nonparametric statistical testing of EEG-and MEG-data. *Journal of neuroscience methods*, 164(1), 177-190.
- Markant, D. B., Settles, B., & Gureckis, T. M. (2016). Self-Directed Learning Favors Local, Rather Than Global, Uncertainty. *Cognitive science*, 40(1), 100-120.
- McGinty, V. B., Rangel, A., & Newsome, W. T. (2016). Orbitofrontal Cortex Value Signals Depend on Fixation Location during Free Viewing. *Neuron*, 90(6), 1299-1311.
- McKoon, G., & Ratcliff, R. (2012). Aging and IQ effects on associative recognition and priming in item recognition. *Journal of memory and language*, 66(3), 416-437.
- Meier, K. M., & Blair, M. R. (2013). Waiting and weighting: Information sampling is a balance between efficiency and error-reduction. *Cognition*, 126(2), 319-325.
- Michael, E., de Gardelle, V., Nevado-Holgado, A., & Summerfield, C. (2013). Unreliable Evidence: 2 Sources of Uncertainty During Perceptual Choice. *Cerebral Cortex (New York, NY)*, 25(4), 937-947. doi:10.1093/cercor/bht287
- Milstein, D. M., & Dorris, M. C. (2007). The influence of expected value on saccadic preparation. *Journal of Neuroscience*, 27(18), 4810-4818.
- Milstein, D. M., & Dorris, M. C. (2011). The relationship between saccadic choice and reaction times with manipulations of target value. *Frontiers in neuroscience*, 5.
- Moreland, R. L., & Zajonc, R. B. (1977). Is stimulus recognition a necessary condition for the occurrence of exposure effects? *Journal of personality and social psychology*, 35(4), 191.
- Moreland, R. L., & Zajonc, R. B. (1982). Exposure effects in person perception: Familiarity, similarity, and attraction. *Journal of Experimental Social Psychology*, 18(5), 395-415.
- Morvan, C., & Maloney, L. T. (2012). Human visual search does not maximize the post-saccadic probability of identifying targets. *PLoS Comput Biol*, 8(2), e1002342.
- Müller, M. M., Malinowski, P., Gruber, T., & Hillyard, S. A. (2003). Sustained division of the attentional spotlight. *Nature*, 424(6946), 309-312 %@ 0028-0836.
- Najemnik, J., & Geisler, W. S. (2005). Optimal eye movement strategies in visual search. *Nature*, 434(7031), 387.
- Najemnik, J., & Geisler, W. S. (2008). Eye movement statistics in humans are consistent with an optimal search strategy. *Journal of Vision*, 8(3), 4.

- Navalpakkam, V., & Itti, L. (2007). Search goal tunes visual features optimally. *Neuron*, 53(4), 605-617.
- Navalpakkam, V., Koch, C., Rangel, A., & Perona, P. (2010). Optimal reward harvesting in complex perceptual environments. *Proceedings of the National Academy of Sciences*, 107(11), 5232-5237.
- Navarro, D. J., & Perfors, A. F. (2011). Hypothesis generation, sparse categories, and the positive test strategy. *Psychological review*, 118(1), 120.
- Nelson, J. D. (2005). Finding useful questions: On Bayesian diagnosticity, probability, impact, and information gain. *Psychological review*, 112(4 %@ 0033-295X).
- Newell, B. R., & Rakow, T. (2007). The role of experience in decisions from description. *Psychonomic Bulletin & Review*, 14(6), 1133-1139.
- Nickerson, R. S. (1998). Confirmation bias: A ubiquitous phenomenon in many guises. *Review of general psychology*, 2(2), 175.
- Niebur, E., & Koch, C. (1996). *Control of selective visual attention: Modeling the "where" pathway*.
- Nørretranders, T. (1998). The User Illusion: Cutting Consciousness Down to Size, trans. Jonathan Sydenham (New York: Viking, 1998), 187.
- Nothdurft, H.-C. (2000). Saliency from feature contrast: additivity across dimensions. *Vision research*, 40(10), 1183-1201.
- Nunez, M. D., Vandekerckhove, J., & Srinivasan, R. (2017). How attention influences perceptual decision making: Single-trial EEG correlates of drift-diffusion model parameters. *Journal of mathematical psychology*, 76, 117-130 %@ 0022-2496.
- Ossmy, O., Moran, R., Pfeffer, T., Tsetsos, K., Usher, M., & Donner, T. H. (2013). The timescale of perceptual evidence integration can be adapted to the environment. *Current Biology*, 23(11), 981-986.
- Payne, J. W., & Braunstien, M. L. (1978). Risky choice: An examination of information acquisition behavior. *Memory & Cognition*, 6(5), 554-561.
- Peck, C. J., Jangraw, D. C., Suzuki, M., Efem, R., & Gottlieb, J. (2009). Reward modulates attention independently of action value in posterior parietal cortex. *Journal of Neuroscience*, 29(36), 11182-11191.
- Pelli, D. G. (1981). *Effects of visual noise*. (PH.D.), University of Cambridge, Cambridge.
- Pettibone, J. C., & Wedell, D. H. (2007). Testing alternative explanations of phantom decoy effects. *Journal of Behavioral Decision Making*, 20(3), 323-341.
- Pica, P., Lemer, C., Izard, V., & Dehaene, S. (2004). Exact and approximate arithmetic in an Amazonian indigene group. *science*, 306(5695), 499-503.
- Poletti, M., Rucci, M., & Carrasco, M. (2017). Selective attention within the foveola. *Nature neuroscience*, 20(10), 1413 %@ 1546-1726.
- Purcell, B. A., Schall, J. D., Logan, G. D., & Palmeri, T. J. (2012). From saliency to saccades: multiple-alternative gated stochastic accumulator model of visual search. *Journal of Neuroscience*, 32(10), 3433-3446.
- Ratcliff, R. (1978). A theory of memory retrieval. *Psychological review*, 85(2), 59.
- Ratcliff, R., & McKoon, G. (2008). The diffusion decision model: theory and data for two-choice decision tasks. *Neural computation*, 20(4), 873-922.
- Ratcliff, R., & Rouder, J. N. (1998). Modeling response times for two-choice decisions. *Psychological science*, 9(5), 347-356.

- Ratcliff, R., Smith, P. L., Brown, S. D., & McKoon, G. (2016). Diffusion decision model: current issues and history. *Trends in cognitive sciences*, 20(4), 260-281.
- Ratcliff, R., Thompson, C. A., & McKoon, G. (2015). Modeling individual differences in response time and accuracy in numeracy. *Cognition*, 137, 115-136.
- Ratcliff, R., Van Zandt, T., & McKoon, G. (1999). Connectionist and diffusion models of reaction time. *Psychological review*, 106(2), 261-300.
- Rayner, K. (1998). Eye movements in reading and information processing: 20 years of research. *Psychological bulletin*, 124(3), 372.
- Rehder, B., & Hoffman, A. B. (2005). Eyetracking and selective attention in category learning. *Cognitive psychology*, 51(1), 1-41.
- Rieskamp, J., & Hoffrage, U. (2008). Inferences under time pressure: How opportunity costs affect strategy selection. *Acta Psychologica*, 127(2), 258-276.
- Rogers, R. D., & Monsell, S. (1995). Costs of a predictable switch between simple cognitive tasks. *Journal of experimental psychology: General*, 124(2), 207.
- Rubinstein, A. (2003). "Economics and psychology"? The case of hyperbolic discounting. *International Economic Review*, 44(4), 1207-1216.
- Saito, T., Otani, M., & Kinjo, H. (2015). Does the Gaze Cascade Effect Occur in Various Judgments Other Than the Preference Judgment. *Cognitive Studies*, 22(3), 463-472.
- Scheepers, C., & Crocker, M. W. (2004). Constituent order priming from reading to listening: A visual-world study. *The on-line study of sentence comprehension: Eyetracking, ERP, and beyond*, 167-185.
- Schkade, D. A., & Johnson, E. J. (1989). Cognitive processes in preference reversals. *Organizational Behavior and Human Decision Processes*, 44(2), 203-231.
- Schotter, E. R., Berry, R. W., McKenzie, C. R. M., & Rayner, K. (2010). Gaze bias: Selective encoding and liking effects. *Visual Cognition*, 18(8), 1113-1132.
- Schulte-Mecklenbeck, M., Kühberger, A., & Ranyard, R. (2011). The role of process data in the development and testing of process models of judgment and decision making. *Judgment and Decision Making*, 6(8), 733.
- Schütz, A. C., Braun, D. I., & Gegenfurtner, K. R. (2011). Eye movements and perception: A selective review. *Journal of Vision*, 11(5), 9-9.
- Scolari, M., & Serences, J. T. (2009). Adaptive allocation of attentional gain. *The Journal of Neuroscience*, 29(38), 11933-11942.
- Shadlen, M. N., & Kiani, R. s. (2013). Decision making as a window on cognition. *Neuron*, 80(3), 791-806.
- Shadlen, M. N., & Roskies, A. L. (2012). The neurobiology of decision-making and responsibility: reconciling mechanism and mindedness. *Frontiers in neuroscience*, 6.
- Shimojo, S., Simion, C., Shimojo, E., & Scheier, C. (2003). Gaze bias both reflects and influences preference. *Nature neuroscience*, 6(12), 1317-1322.
- Simion, C., & Shimojo, S. (2007). Interrupting the cascade: Orienting contributes to decision making even in the absence of visual stimulation. *Attention, Perception, & Psychophysics*, 69(4), 591-595.
- Simonson, I. (1989). Choice based on reasons: The case of attraction and compromise effects. *Journal of consumer research*, 16(2), 158-174.

- Skov, R. B., & Sherman, S. J. (1986). Information-gathering processes: Diagnosticity, hypothesis-confirmatory strategies, and perceived hypothesis confirmation. *Journal of Experimental Social Psychology*, 22(2), 93-121.
- Slovic, P., & Lichtenstein, S. (1983). Preference reversals: A broader perspective. *The American Economic Review*, 73(4), 596-605.
- Slowiaczek, L. M., Klayman, J., Sherman, S. J., & Skov, R. B. (1992). Information selection and use in hypothesis testing: What is a good question, and what is a good answer? *Memory & Cognition*, 20(4), 392-405.
- Solomon, J. A., & Morgan, M. J. (2006). Stochastic re-calibration: contextual effects on perceived tilt. *Proceedings of the Royal Society of London B: Biological Sciences*, 273(1601), 2681-2686.
- Spiegel, M. F., & Green, D. M. (1981). Two procedures for estimating internal noise. *The Journal of the Acoustical Society of America*, 70(1), 69-73.
- Spitzer, B., Blankenburg, F., & Summerfield, C. (2016). Rhythmic gain control during supramodal integration of approximate number. *Neuroimage*, 129, 470-479.
- Spitzer, B., Waschke, L., & Summerfield, C. (2017). Selective overweighting of larger magnitudes during noisy numerical comparison. *Nature Human Behaviour*, 1(8), s41562-41017-40145.
- Starr, M. S., & Rayner, K. (2001). Eye movements during reading: Some current controversies. *Trends in cognitive sciences*, 5(4), 156-163.
- Summerfield, C., & Tsetsos, K. (2015). Do humans make good decisions? *Trends in cognitive sciences*, 19(1), 27-34.
- Swensson, R. G. (1972). The elusive tradeoff: Speed vs accuracy in visual discrimination tasks. *Perception & Psychophysics*, 12(1), 16-32.
- Takikawa, Y., Kawagoe, R., Itoh, H., Nakahara, H., & Hikosaka, O. (2002). Modulation of saccadic eye movements by predicted reward outcome. *Experimental brain research*, 142(2), 284-291.
- Tanner Jr, W. P., & Swets, J. A. (1954). A decision-making theory of visual detection. *Psychological review*, 61(6), 401.
- Tatler, B. W. (2009). Current understanding of eye guidance. *Visual Cognition*, 17(6-7), 777-789.
- Tatler, B. W., Hayhoe, M. M., Land, M. F., & Ballard, D. H. (2011). Eye guidance in natural vision: Reinterpreting salience. *Journal of Vision*, 11(5), 5-5.
- Tatler, B. W., & Vincent, B. T. (2009). The prominence of behavioural biases in eye guidance. *Visual Cognition*, 17(6-7), 1029-1054.
- Thomas, A. W., Molter, F., Krajbich, I., Heekeren, H. R., & Mohr, P. N. C. (2017). Gaze Bias Differences Capture Individual Choice Behavior. *bioRxiv*, 228825.
- Thompson, K. G., & Bichot, N. P. (2005). A visual salience map in the primate frontal eye field. *Progress in brain research*, 147, 249-262.
- Toffanin, P., de Jong, R., Johnson, A., & Martens, S. (2009). Using frequency tagging to quantify attentional deployment in a visual divided attention task. *International Journal of Psychophysiology*, 72(3), 289-298 %@ 0167-8760.
- Torrallba, A., Oliva, A., Castelhana, M. S., & Henderson, J. M. (2006). Contextual guidance of eye movements and attention in real-world scenes: the role of global features in object search. *Psychological review*, 113(4), 766.

- Towal, R. B., Mormann, M., & Koch, C. (2013). Simultaneous modeling of visual saliency and value computation improves predictions of economic choice. *Proceedings of the National Academy of Sciences*, 110(40), E3858-E3867.
- Treisman, A. M., & Gelade, G. (1980). A feature-integration theory of attention. *Cognitive psychology*, 12(1), 97-136.
- 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.
- Tsetsos, K., Moran, R., Moreland, J. C., Chater, N., Usher, M., & Summerfield, C. (2016). Reply to Davis-Stober et al.: Violations of rationality in a psychophysical task are not aggregation artifacts. *Proceedings of the National Academy of Sciences*, 201608989.
- Tsetsos, K., Usher, M., & Chater, N. (2010). Preference reversal in multiattribute choice. *Psychological review*, 117(4), 1275.
- Tversky, A., & Simonson, I. (1993). Context-dependent preferences. *Management science*, 39(10), 1179-1189.
- Urai, A. E., Braun, A., & Donner, T. H. (2017). Pupil-linked arousal is driven by decision uncertainty and alters serial choice bias. *Nature Communications*, 8.
- Usher, M., & McClelland, J. L. (2001). The time course of perceptual choice: the leaky, competing accumulator model. *Psychological review*, 108(3), 550.
- Vandormael, H., Castañón, S. H., Balaguer, J., Li, V., & Summerfield, C. (2017). Robust sampling of decision information during perceptual choice. *Proceedings of the National Academy of Sciences*, 114(10), 2771-2776.
- Vö, M. L.-H., & Henderson, J. M. (2010). The time course of initial scene processing for eye movement guidance in natural scene search. *Journal of Vision*, 10(3), 14.
- Von Neumann, J., & Morgenstern, O. (1947). *Theory of games and economic behavior*, 2nd rev.
- Voss, A., Nagler, M., & Lerche, V. (2013). Diffusion models in experimental psychology. *Experimental psychology*.
- Wald, A., & Wolfowitz, J. (1948). Optimum character of the sequential probability ratio test. *The Annals of Mathematical Statistics*, 326-339.
- Weber, E. U., Shafir, S., & Blais, A.-R. (2004). Predicting risk sensitivity in humans and lower animals: risk as variance or coefficient of variation. *Psychological review*, 111(2), 430.
- Wilson, T. D. (2004). *Strangers to ourselves*: Harvard University Press.
- Wolfe, J. M., Cave, K. R., & Franzel, S. L. (1989). Guided search: an alternative to the feature integration model for visual search. *Journal of Experimental Psychology: Human perception and performance*, 15(3), 419.
- Wyart, V., De Gardelle, V., Scholl, J., & Summerfield, C. (2012). Rhythmic fluctuations in evidence accumulation during decision making in the human brain. *Neuron*, 76(4), 847-858.
- Wyart, V., Myers, N. E., & Summerfield, C. (2015). Neural mechanisms of human perceptual choice under focused and divided attention. *Journal of Neuroscience*, 35(8), 3485-3498.
- Yang, S. C.-H., Lengyel, M., & Wolpert, D. M. (2016). Active sensing in the categorization of visual patterns. *Elife*, 5, e12215.

- Yantis, S., Anderson, B. A., Wampler, E. K., & Laurent, P. A. (2012). Reward and attentional control in visual search *The Influence of Attention, Learning, and Motivation on Visual Search* (pp. 91-116): Springer.
- Yarbus, A. L. (1967). *Eye movements during perception of complex objects*: Springer.
- Yarbus, A. L. (1967). *Eye movements during perception of complex objects*: Springer.
- Yechiam, E., & Busemeyer, J. R. (2006). The effect of foregone payoffs on underweighting small probability events. *Journal of Behavioral Decision Making*, 19(1), 1-16.
- Zajonc, R. B. (1968). Attitudinal effects of mere exposure. *Journal of personality and social psychology*, 9(2p2), 1.
- Zellman, E., Kaye-Blake, W., & Abell, W. (2010). Identifying consumer decision-making strategies using alternative methods. *Qualitative Market Research: An International Journal*, 13(3), 271-286.
- Zhao, Q., & Koch, C. (2011). Learning a saliency map using fixated locations in natural scenes. *Journal of Vision*, 11(3), 9-9.