

UC Irvine

UC Irvine Electronic Theses and Dissertations

Title

Reliability estimation as a normative principle in dynamic models of decision making

Permalink

<https://escholarship.org/uc/item/8856539t>

ISBN

9798189267413

Author

Khoudary, Ari

Publication Date

2026-08-31

Copyright Information

This work is made available under the terms of a Creative Commons Attribution-NonCommercial-ShareAlike License, available at <https://creativecommons.org/licenses/by-nc-sa/4.0/>

Peer reviewed|Thesis/dissertation

UNIVERSITY OF CALIFORNIA,
IRVINE

Reliability estimation as a normative principle in dynamic models of decision making

DISSERTATION

submitted in partial satisfaction of the requirements
for the degree of

DOCTOR OF PHILOSOPHY

in Cognitive Sciences

by

Ari Khoudary

Dissertation Committee:
Associate Professor Aaron M. Bornstein, Chair
Associate Professor Megan A. K. Peters
Professor Joachim Vandekerckhove

DEDICATION

To all the bridge-builders, academic and otherwise.

TABLE OF CONTENTS

	Page
LIST OF FIGURES	v
LIST OF TABLES	ix
ACKNOWLEDGMENTS	x
VITA	xii
ABSTRACT OF THE DISSERTATION	xiv
1 Introduction	1
2 Reliability-weighted integration explains time-varying effects of expectations on perceptual decisions	7
2.1 Introduction	7
2.2 Static effects of expectations on perceptual decisions: formal models and empirical data	8
2.3 Conceptual gap: role of memory retrieval dynamics	25
2.4 Dynamic reliability-weighted multi-source sequential sampling model	33
2.5 Dynamic reliability-weighted integration unifies time-varying effects of expectations	41
2.6 Discussion	70
2.7 Summary and Conclusion	83
3 Expectations dynamically modulate visual evidence accumulation	85
3.1 Introduction	85
3.2 Methods: experimental design	89
3.3 Methods: computational modeling	93
3.4 Results	101
3.5 Discussion	110
3.6 Conclusion	112
4 Reasoning goals and representational decisions in computational cognitive neuroscience: lessons from the drift diffusion model	114
4.1 Introduction	114

4.2	Philosophical tools for thinking about computational cognitive models	116
4.3	A philosophical introduction to the DDM	129
4.4	Reasoning goals and representational decisions in the collapsing bounds debate	136
4.5	Summary and Conclusion	152
5	Conclusion	153
5.1	Summary	154
5.2	Future Directions	155
	Bibliography	158
	Appendix A Supplementary Materials for Chapter 2	181
	Appendix B Supplementary Materials for Chapter 3	191

LIST OF FIGURES

	Page
2.1 Two dimensions of uncertainty that observers can have about expectations when making perceptual decisions. Source-dependent uncertainty (orange) captures uncertainty about <i>what</i> an expectation predicts, whereas deployment-dependent uncertainty (pink) captures uncertainty about <i>which</i> expectation to use for each decision. The relative degree of uncertainty within a dimension is represented by the opacity of its respective color, with highly uncertain conditions in lighter shades and less certain conditions in darker shades. Gray-text citations indicate existing findings of dynamic effects observed in each quadrant; we show that our model can capture each of these findings in Section 2.5.	32
2.2 Graphical depiction of single-trial dynamics in our model. Quantities corresponding to memory are plotted in blue, quantities corresponding to vision are plotted in yellow, and quantities corresponding the decision variable are plotted in green. Figures were generated using a vision sampling rate of 60Hz, a memory “sampling rate” of 6Hz ($\gamma = 10$), and an encoding noise ratio $\epsilon = 0.1$	37
2.3 Steady-state dynamic weights on memory (blue) and sensory (yellow) evidence. The rate at which memory evidence enters the decision process is indicated by different line styles (solid = equal rate as vision, dashed = high theta frequencies, dotted = low theta frequencies). (A) Weights when sensory evidence is totally uninformative ($\theta_{vis} = 0.5$). (B) Weights when sensory evidence is moderately informative ($\theta_{vis} = 0.65$). (C) Weights when sensory evidence is strongly informative ($\theta_{vis} = 0.8$). Shaded regions correspond to standard error computed across all levels of ϵ and the subset of γ corresponding to each linetype.	38
2.4 Dynamic reliability-weighted integration captures Hanks et al. (2011) dynamic effect of expectations. In all figures, memory evidence (expectations; Π) is plotted in blue, vision evidence is plotted in yellow, and the decision variable is plotted in green. (A) On the left, a stylized reproduction of the original model used by Hanks et al. (2011) to explain their “late effect” of expectations. On the right, a representative trial from our model demonstrating the same single-trial dynamics. (B) Model components involved in generating our representative trial above.	45

2.5	Dynamic reliability-weighted integration captures time-varying effect of expectations observed by Rungratsameetaweemana, Itthipuripat et al. (2018). (A) Original finding, left: violations of expectations increase frontal theta amplitude at later timepoints in the trial. Our simulation, right: weights on memory—interpretable as an incentive to sample evidence from memory—are greater for unexpected stimuli (red) relative to neutral (gray) and expected (blue). Crucially, the effect increases over time and is greatest at the end of the decision. Shaded areas reflect standard error across 15 simulated subjects. (B) Original finding, left: flicker rate and expectation manipulations induce canonical modulations of accumulation-based choice behavior. Slow flicker trials (purple), where fewer samples of visual evidence are presented per unit time, lead to slower reaction times and lower choice accuracy relative to fast flicker trials (green). Our simulation, right: choices and reaction times generated from our model replicate these behavioral effects, validating our choices for sensory signal strength θ_{vis} and threshold z . Errorbars in both (A) and (B) reflect standard error across $n = 18$ simulated “subjects”, each of which had a unique value of γ and z	50
2.6	Dynamic reliability-weighted integration captures time-varying effects from Bornstein et al. (2023). In all figures, memory evidence (expectations; Π) is plotted in blue, vision evidence is plotted in yellow, and the decision variable is plotted in green. Vertical lines correspond to timepoints when memory (blue) and sensory (yellow) evidence come “online” for the observer. Memory samples in both simulations were generated at a rate of 6Hz ($\gamma = 10$). (A) Our model captures key dynamics produced by Bornstein et al.’s (2023) multi-stage DDM (left). First, we show that expectations set the starting point for sensory evidence accumulation in a trial-by-trial manner (right, both figures). Then, we show that anticipatory sampling from a strong expectation expedites the decision process once weak sensory evidence becomes available (left panel under “our simulation”). Additionally, we show that sampling from memory continues throughout the anticipation period when an expectation carries no information about the likely choice outcome, such that choice dynamics are dominated by the strong sensory evidence once it comes online (right panel under “our simulation”). (B) Single-trial model components that generated the left and right panels for the “our simulation” portion of A. Components for the left panel—strong expectation with weak evidence—are plotted in the top row, and components for the right panel—weak expectation with strong evidence—are plotted in the bottom row.	59

2.7	Steady-state dynamics of evidence weights and source reliability estimates driving Bornstein et al.'s (2023) observation of uncertainty-adaptive anticipatory memory sampling. Quantities corresponding to memory are plotted in blue and quantities corresponding to vision are plotted in yellow. Solid lines correspond to simulations with a weak sensory signal ($\theta_{vis} = 0.55$), dashed lines correspond to simulations with a strong sensory signal ($\theta_{vis} = 0.8$). The vertical line indicates the onset of sensory evidence after a 4000ms anticipation period; dynamics for longer anticipation durations are plotted in Figure A.3. (A) Weights on each evidence source throughout a single decision. (B) Source reliability estimates throughout a single decision.	60
2.8	Dynamic reliability-weighted integration captures Hachisuka, Shor, Liu et al. (2026) "late" effect of expectations. (A) Experimental design in Hachisuka, Shor, Liu et al. (2026) . On each trial, human subjects were presented with an image for 2 seconds and asked to verbally report the content of the image. In the first block, participants viewed "degraded" images whose contents are difficult to discern in the absence of prior knowledge. Then, in a second block, participants were presented with denoised grayscale versions of the same images, facilitating the one-shot learning characteristic of Mooney images. Finally, in the third block, participants were presented with degraded images again. Crucially, 50% of the images had just been presented in full resolution during the learning block. This manipulation allows for discriminating neural timecourses as a function of pre- and post-learning. Image reproduced from Hachisuka, Shor, Liu et al. (2026) with colored borders added for clarity. (B) Original finding: location of implanted electrodes across all patients (left), average activity from frontoparietal (FPN) electrodes (right). Activity post-learning (orange) ramps up more quickly than activity pre-learning (blue), generating a "late" effect of one-shot learned expectations about ambiguous sensory input. Our simulation: decision variables averaged across trials reproduce the pre- and post-learning dynamics measured by FPN electrodes. Shaded areas reflect standard error across $n = 20$ simulated subjects.	66
2.9	Expectation-heterogeneity induces sequential effects on RTs, regardless of whether expectations are explicitly instructed (left) or learned through experience (right). Across all cue levels, switching cues across trials significantly slows responses, whereas repeating cues across trials significantly speeds responses. In all panels, individual lines reflect subject means, points are group-level marginal means estimated using fitted regression coefficients, errorbars reflect 95% confidence intervals around that estimate.	70

3.1	Task design. (A) Participants first learn that different colored borders make different predictions about the probability of observing one of two possible scene images. The predictive strength θ_{mem} of each cue is defined by the frequency with which it is followed by one of the two scene images. (B) After learning, participants are presented with stochastic visual evidence inside of the colored borders. Their task is to report which image is dominant (i.e., presented more frequently) in a 60Hz stream of visual evidence, followed by a confidence rating in that perceptual decision.	89
3.2	Response-coded model fits to simulated and human-subjects data. Shaded gray areas depict modal onsets and durations of the sensory noise periods, which followed a constant hazard rate across trials. Note the difference in y-axis values across Figures B-D. (A) RT distributions split between low and high threshold groups. (B) Fit of the response-coded model to data generated by a serial process. (C) Fit of the response-coded model to data generated by our parallel reliability-weighted process. (D) Fit of the response-coded model to human subjects data, split by fitted threshold value. *** = $p < .001$, **= $p < .01$, *= $p < .05$, · = $p < .1$	106
3.3	Fitted boundary values modulate model-free signatures of expectation integration. (A) RT distributions demonstrating that group-level differences in choice dynamics are conserved across correct and error response. Gray regions depict modal onsets and durations of sensory noise periods. (B) Observers with lower thresholds demonstrated greater accuracy gains as a function of θ_{cue} . (C) Confidence in incorrect decisions is more strongly modulated by θ_{cue} for observers with lower thresholds.	110
4.1	Graphical depictions of two standard forms of a diffusion model of speeded two-alternative choice. Standard interpretations of the key variables are provided in the main text. (A) The extended drift diffusion model. This form of the model features parameters whose values vary trial-by-trial but remain fixed within the course of a single trial. Starting point (yellow) and non-decision time (red) values are drawn from uniform distributions and drift rates (blue) are drawn from normal distributions. Threshold values (green) are constant within a trial, and can be allowed to vary as a function of experimental condition and/or participant. (B) The diffusion model with collapsing decision boundaries. This form of the model features parameters whose values are effectively fixed trial-by-trial and decision boundaries whose values decrease (“collapse”) over the course of a single trial. Drift rate and, in some models, starting point are permitted to vary as a function of experimental condition but are assumed to have fixed effects across trials within an experimental condition.	130

LIST OF TABLES

	Page
3.1 Calibrated threshold-coherence tuples used to generate data from both the serial and parallel models. Proportion correct responses averaged over 150,000 simulated trials.	100
3.2 Contrasts (error - correct) on estimated marginal mean RT (z-scored and log-transformed within subject) at each value of θ_{mem}	102
3.3 Contrasts (error - correct) on estimated marginal mean confidence at each value of θ_{mem}	103
3.4 Fitted drift rates from the response-coded model in each epoch, averaged across subjects. P-values reflect a one-sample t-test against 0; SEM = standard error of the mean (n=41).	103
3.5 No difference in fitted drift rates for the probability of making a cue-congruent response between congruent and incongruent trials.	104

ACKNOWLEDGMENTS

Science is social knowledge, and the production of scientific knowledge is reliably improved by being embedded in a rich social environment. All of the work in this dissertation was made possible by the people I have been so fortunate to have been supported by over the past several years.

I would first like to thank my co-advisors, Aaron Bornstein and Megan Peters, for bravely venturing into uncharted intellectual territory with me, exercising seemingly superhuman amounts of patience, and exploding the scope of what I deemed possible for myself and my career. I cannot thank you enough for your labor and support throughout these years. It has been a privilege to learn and grow with you both.

I thank my committee members, Joachim Vandekerckhove, Mark Steyvers, and Lauren Ross, for their careful engagement with my scholarship and vested interest in improving it. There are marks of your thinking speckled all over this body of work. I thank Felipe De Brigard, Maureen Ritchey, Richard Atkins, and the late Daniel Dennett for the same, albeit during the critical early years when obtaining this PhD was still very far from reality. Additionally, I thank Michael and Manuella Yassa for their boundless enthusiasm and consistent support of my development as a scholar.

Irvine is a weird place to live but a great place to do a PhD. I thank my labmates, cohort-mates, and more broadly-scoped academic peers for celebrating and commiserating with me throughout the highs and lows of grad school, and for making this weird little place feel like home. In particular, I am grateful to have been shaped by being in community with Ainsley May, Jú Dias, Alice Stewart, Aly Arends, Nora Harhen, Jungsun Yoo, Keiland Cooper, Alexa Booras, Bianca Leonard, Sharon Noh, Dale Zhou, Jerry Guo, Nidhi Banavar, Rochelle Kaper, Oğul Can Yurdukul, Adriana Chavez De La Peña, Kathleen Medriano, Raven Zhang, Angela Shen, Elle Giovanni, Pharrell Allen, and Aria Gastón-Panthaki. I'm also incredibly grateful to Kevin O'Neill and Ata Karagoz for their support from a distance.

I thank Aidan Goeschel, Rohin Palsule, and Vanessa Wu for infecting me with their curiosity and giving me the opportunity to start learning how to be a scientific mentor. I also thank Nadeja Jackson, Alison Krueger, Nurazmera Mossa, and Lauryn Chandler for the same, as well as for their logistical and emotional support with human subjects data collection.

Last, but most certainly not least, I thank my wonderful constellation of non-academic friends and family for their “outside” perspective and unconditional support. Lilac Friel, Bella and Spencer Pepke, Cleo Hyzy, Naseeb Nino, and Emma Arcos: thank you for taking on my problems as if they were your own and reminding me that I'm never as alone as I might feel. Mom, Dad, Anthony, Peter, and Theresa, I love you way more than I show it and I can't wait for the Khoudary Family Renaissance courtesy of UPenn. And to the most loyal and doting feline companion I could ever ask for—Ms. Bella aka Stink Khoudary—thank you for the sweetness, cuteness, levity, and emotional support. I truly couldn't have done this without you.

This work was supported by a fellowship from the Schneiderman-NIMH T32 training program in Learning and Memory (NIMH T32MH119049). The text in Chapter 4 of this dissertation is a reprint of the material as it appears in Khoudary, A., Peters, M. A. K.*, & Bornstein, A. M.* (2025). Reasoning goals and representational decisions in computational cognitive neuroscience: Lessons from the drift diffusion model. *The European Journal of Neuroscience*, 61(7), e70098. The co-authors listed in this publication directed and supervised research which forms the basis for the dissertation.

VITA

Ari Khoudary

EDUCATION

Doctor of Philosophy in Cognitive Sciences	2026
University of California, Irvine	<i>Irvine, CA</i>
Master of Science in Cognitive Neuroscience	2024
University of California, Irvine	<i>Irvine, CA</i>
Bachelor of Science in Psychology and Philosophy	2019
Boston College	<i>Chestnut Hill, MA</i>

RESEARCH EXPERIENCE

Graduate Research Fellow	2023–2026
University of California, Irvine	<i>Irvine, California</i>
Graduate Student Researcher	2021–2023
University of California, Irvine	<i>Irvine, California</i>
Research Assistant	2019–2021
Duke University	<i>Durham, NC</i>
NSF-REU Fellow	2018
New York University	<i>New York, NY</i>
Undergraduate Research Fellow	2017–2019
Boston College	<i>Chestnut Hill, MA</i>

TEACHING EXPERIENCE

Teaching Assistant	2021–2023
University of California, Irvine	<i>Irvine, CA</i>

REFEREED JOURNAL PUBLICATIONS

- Khoudary, A.**, Bornstein, A.M.*, Peters, M.A.K.* (2026) Blind ignorance or rational (in)attention? On the history and future of metabolic considerations for cognitive models (Commentary on Haueis & Colaço, 2026). *Behavioral and Brain Sciences*.
- Khoudary, A.**, Peters, M.A.K.*, Bornstein, A.M.* (2025) Reasoning goals and representational decisions in computational cognitive neuroscience: lessons from the drift diffusion model. *European Journal of Neuroscience*.

Miceli, K.*, Morales-Torres, R.*, **Khoudary, A.**, Faul, L., Parikh, N., & De Brigard, F. (2023). Perceived plausibility modulates hippocampal activity in episodic counterfactual thinking. *Hippocampus*.

Khoudary, A., O'Neill, K., Faul, L., Murray, S., Smallman, R., De Brigard, F. (2022). Neural differences between internal and external episodic counterfactual thoughts. *Philosophical Transactions of the Royal Society B*.

Khoudary, A., Hanna, E., O'Neill, K.G., Iyengar, V., Clifford, S., De Brigard, F.*, Cabeza, R.*, Sinnott-Armstrong, W.* (2022). A Functional Neuroimaging Investigation of Moral Foundations Theory. *Social Neuroscience*.

De Brigard, F., **Khoudary, M.**, & Murray, S. (2022). Times Imagined and Remembered. In Hoerl, C., McCormack, T., & Fernandes, A. (Eds.). *Temporal Asymmetries in Philosophy and Psychology*. Oxford University Press.

Mele, A., Nadelhoffer, T., **Khoudary, M.** (2021). Folk psychology and proximal intentions. *Philosophical Psychology*.

REFEREED CONFERENCE PUBLICATIONS

Wu, Y.*, **Khoudary, A.***, Peters, M.A.K. (2026) Characterizing time-varying effects of evidence volatility on perceptual confidence. *Extended Abstract at the 9th Conference on Cognitive Computational Neuroscience*.

Goeschel, A*, Palsule, R.A.*, Chen, J., Bornstein, A.M., **Khoudary, A.** (2026) Quantifying information gain in narrative stimuli using open-source large language models. *Extended Abstract at the 9th Conference on Cognitive Computational Neuroscience*.

Khoudary, A., Bornstein, A.M.*, Peters, M.A.K.* (2025) Investigating implicit and explicit expectations in perceptual decision making. *Proceedings of the 47th Annual Meeting of the Cognitive Science Society*.

Khoudary, A., Peters, M.A.K.*, Bornstein, A.M.* (2022) Precision-weighted evidence integration predicts time-varying influence of memory on perceptual decisions. *Proceedings of the 2022 Conference on Cognitive Computational Neuroscience*.

ABSTRACT OF THE DISSERTATION

Reliability estimation as a normative principle in dynamic models of decision making

By

Ari Khoudary

Doctor of Philosophy in Cognitive Sciences

University of California, Irvine, 2026

Associate Professor Aaron M. Bornstein, Chair

Normative models are powerful tools for advancing cognitive science. But where do their norms come from, and how well do they align with other constraints on the target system? This dissertation examines these questions within the family of sequential sampling models used to study the speed-accuracy tradeoff in decision-making. Across three chapters, we show how relaxing a common assumption generates normative solutions that require *within-trial* variability in the value of two key parameters (drift rate and decision threshold), and—critically—that the dynamics of this variability are governed by observers’ evolving belief about the reliability of evidence bearing on choice.

We begin by specifying a novel model formalizing the role of memory retrieval dynamics in biasing decisions when observers have uncertainty about prior probabilities, and show that it can reproduce dynamic effects of prior probability observed across different tasks, species, and neuroimaging modalities. Then, we report empirical data from a novel task designed to test the key prediction of our model—that choices ought to be most strongly biased by prior probability at the moments in time when sensory evidence is maximally uncertain—and report direct evidence in support of this prediction. Finally, we present a “philosophical toolkit” forming the conceptual foundations of this work and demonstrate its utility by using it to re-analyze a long-standing debate about time-varying decision thresholds. Taken

together, this work jointly illustrates how normative models offer principled starting points for developing and testing theories about neural and cognitive function, and how philosophy offers meta-scientific support toward that end.

Chapter 1

Introduction

There is no such thing as philosophy-free science; there is only science
whose philosophical baggage is taken on board without examination.

Daniel Dennett, *Darwin's Dangerous Idea*, 1995

When making decisions, humans and other agents not only have to determine *which* option to commit to but also *how long* to spend deliberating. Further, they must do so in environments that are stochastic and dynamic, where possible outcomes and/or evidence quality can change over the course of the decision. How do agents decide *when* to decide, and on what basis do they make this choice?

Decades of cross-disciplinary research has approached these questions by studying simple two-alternative decisions about low-level perceptual features like visual motion direction (Gold and Shadlen, 2007; Ratcliff and McKoon, 2008; Hanks and Summerfield, 2017) or auditory sound localization (Brunton et al., 2013; Pinto et al., 2022; Pagan et al., 2025). These simple tasks have become paradigmatic in the study of “complex” decision-making because they are exceptionally well-suited for careful experimental control and straightforward mapping onto theoretical models with varying levels of neurobiological detail. Indeed,

this subfield of cognitive neuroscience has a strong history of iterative model-based research, with early models offering a formal framework for interpreting neural observations and guiding the design of maximally-discriminatory follow-up experiments (e.g., Kiani et al., 2013; Thura et al., 2012).

At the heart of this research program is the theoretical framework of *sequential sampling*, which posits that the decision process involves sampling information from an evidence source, integrating those samples over time, and committing to one of the alternatives once the integrated evidence exceeds a threshold value (Wald, 1947; Bogacz et al., 2006; Shadlen and Shohamy, 2016; Ratcliff et al., 2016). A powerful theoretical feature of models in this family is that they conceptually dissociate the *mechanism* of accumulation (i.e., how information is encoded and integrated over time) from the *rule* determining how a choice is made on the basis of that accumulated evidence (i.e., where and how the threshold is set). Critically, cross-species research has shown that both the mechanism of accumulation *and* the decision rule accord with prescriptions of normative models (Bogacz et al., 2006; Brunton et al., 2013; Hanks and Summerfield, 2017), such that humans and other mammals appear to make optimal use of perceptual evidence when making decisions.

This dissertation uses theoretical, empirical, and philosophical research to critically examine—and further refine—the normative models used to ascribe optimality (or lack thereof) to choice behavior. We do this by examining the assumptions upon which normative models rest, identifying the limitations of those assumptions with respect to questions of scientific interest, and then discussing normative solutions on less-restrictive (or more realistic) assumptions. Chapters 2 and 3 use this approach to interrogate the mechanism whereby expectations—knowledge of prior probabilities—modify the mechanism of accumulation, whereas Chapter 4 considers normative models for the decision rule. Importantly, we conduct these investigations not simply for the sake of explicating when choice processes are “optimal” or not. Rather, we aim to advance the theoretical study of decision-making—and

its role in adaptive cognition more broadly—by incorporating principled constraints derived from the contributions of learning and memory systems to solving the problem of latent state inference under uncertainty.

In so doing, we also demonstrate the double-edged sword of model-based experimental design. Because models support understanding by idealizing components of the modeled system that lie “outside” the scope of a modelers’ intended application, tasks designed to meet assumptions of those models likewise abstract away, or experimentally control, variability presumed to be “external” to the scientific question under investigation. Without continual critical reflection on *what* the questions of scientific interest are and *how* they are translated into models and tasks, we run the risk of becoming drunkards chained to the streetlights created by highly-simplified normative models.

A sizable body of existing work has taken exactly this approach to develop sequential sampling models that offer normative guarantees on less-restrictive starting premises. This work has focused primarily on the decision rule—or threshold—because this model component determines how decision-makers solve the *speed-accuracy tradeoff*, which is precisely the problem for which normative models provide optimal solutions. For several early models, the optimal solution was given by the Sequential Probability Ratio Test (SPRT; Wald, 1945), which minimizes the average amount of time needed to make decisions for a fixed accuracy level, or maximizes average accuracy for a desired average reaction time. The SPRT achieves this guarantee by assuming that the observer knows how informative the sensory evidence will be *before* it is presented to them, and combines that information with their desired error rate to optimally trade off sampling time with choice accuracy.

While there are certainly some decision contexts where observers have the requisite knowledge to implement the SPRT (e.g., crossing the road at a familiar intersection), many decisions involve deliberation upon sensory evidence whose properties are *not* known in advance. How should agents decide when to decide in this more general case? Converging findings

have shown that when evidence strength (or informativeness) is not known ahead of time, the optimal decision rule corresponds to a threshold value that *changes* over the course of deliberation (Drugowitsch et al., 2012; Moran, 2015; Malhotra et al., 2018). The precise form of the boundary depends on properties of the decision environment (e.g., the mixture of signal strengths, task pacing, etc.), with highly noisy and uncertain environments requiring more aggressively-changing boundaries relative to environments that are more stably informative (or *reliable*). In other words, the *dynamics* of the optimal decision boundary are defined as a function of the *reliability* of sensory evidence forming the basis of choice.

Within this collection of existing normative models, it has also been shown that expectations—knowledge of the prior probability of choice outcomes—can be integrated into decisions by biasing the initial value of the accumulator toward the threshold corresponding to the more likely outcome in proportion to the log-odds of the prior belief (Edwards, 1965; Bogacz et al., 2006). However, all existing findings supporting this mechanism—including its more complex extension to uncertain environments (Moran, 2015)—*also* assume that observers have perfect knowledge of what the true prior probability is for each decision. In doing so, they tacitly assume that integrating this information into the decision process occurs instantaneously, e.g., that the neurons encoding accumulated evidence are already pre-potentiated toward the outcome with a higher prior probability. Again, while it is certainly the case that observers meet these conditions in some contexts, a normative solution to the more general case of *uncertainty about expectations* remains lacking from the literature.

Chapter 2 of this dissertation explicitly addresses this gap by synthesizing theoretical principles of learning, memory, and cue combination to propose a novel sequential sampling model formalizing the role of uncertainty about expectations in the decision process. At the heart of our model are the ideas that retrieving information from memory takes time, and that the dynamics of that retrieval process ought to modulate effects of expectations on perceptual decisions. We capture these ideas by modeling expectations as their own dynamic source of

evidence that are sampled from *in parallel* to sensory information. Then, we propose that evidence from both sources (memory and vision) are combined according to a time-varying estimate of the *relative* reliability of each evidence source, such that expectations are weighted most strongly at the moments in time when sensory evidence is most uncertain. We show that this single model can reproduce time-varying effects of expectations observed in different species performing different tasks while undergoing different neuroimaging procedures, suggesting that it offers a unifying explanation for these previously-disparate findings.

In Chapter 3, we present results from a novel experiment that puts this model to empirical test. Human observers learned probabilistic cue-image associations through experience, and then had those associations cued trial-by-trial in a perceptual decision task. On each trial, we keep the informativeness of both the expectation and sensory evidence constant, but dynamically manipulate their *relative* reliability by inserting stochastic epochs of sensory noise. The key prediction of our theory, contra existing normative models, is that the decision variable should *continue accumulating* toward the option predicted by the expectation at a rate proportional to the predictiveness of that expectation. We tested this prediction by applying a generalized drift diffusion model with a time-varying drift rate to the human subjects data, and found that human subjects indeed exhibit the signature of dynamic integration predicted by our theory. In line with the theoretical fact that expectations most strongly impact decisions when time is limited, we found that this effect was most prominent in observers who adopted a lower threshold on the decision task.

Finally, in Chapter 4, we outline a “philosophical toolkit” that formed the foundation for this dissertation work, and show how it can diffuse a long-standing debate about decision thresholds in sequential sampling models. We review cross-disciplinary research from philosophy of modeling, computational neuroscience, and cognitive modeling to demonstrate how *reasoning goals* necessarily structure the process and products of model-based research. Because all models are incomplete—and because that incompleteness is how models support expla-

nation and understanding—model builders are forced to make representational decisions: deciding *what* to represent in their model and *how* to represent it (Harvard and Winsberg, 2022). By examining the history of representational decisions leading to the development of time-varying threshold models, and contrasting it to the development of the “extended DDM” (Ratcliff et al., 2016), we show how the two models reflect different quantitative and epistemic goals across sub-communities of model users. We conclude by showing how reasoning goals shape both the form and content of formal model selection methods, and, accordingly, propose that model users can solve this specific model selection problem (fixed vs. time-varying boundaries) by considering how the commitments of each model correspond with their own goals for reasoning about a target through a model.

...

This body of work is the product of our reflections on the commitments of normative models of speeded decision-making. Where do their norms come from, and how well do they capture our desired theoretical commitments? We showed that for both the decision rule and the accumulation process, relaxing the assumption of perfect knowledge of environmental statistics—the informativeness of sensory and mnemonic evidence, respectively—produces a normative solution that involves a time-varying adjustment to the decision process. Critically, the dynamics of that adjustment are governed by the uncertainty, or reliability, of the evidence bearing on the decision. Reliability estimation thus appears to serve as a normative principle unifying time-varying models of speeded decision-making.

Chapter 2

Reliability-weighted integration explains time-varying effects of expectations on perceptual decisions

2.1 Introduction

In order to stay alive, humans and other animals must make quick decisions about how to act on the basis of noisy, ambiguous, and limited sensory information. Expectations—prior knowledge about statistical regularities—are a powerful source of information that decision-makers can use to reduce uncertainty and make more efficient decisions (e.g., Seriès and Seitz, 2013; Summerfield and de Lange, 2014; de Lange et al., 2018).

Precisely how expectations enter into the decision process is a topic of active research. Early theoretical, behavioral, and neural evidence suggested that expectations modulate decisions by biasing the deliberation process before the onset of sensory evidence (Edwards, 1965; Carpenter and Williams, 1995; Ratcliff and Rouder, 1998; Simen et al., 2009; Rao et al.,

2012). More recently, however, it has been shown that expectations can exert biases that *vary* over the time it takes to make a single choice (Hanks et al., 2011; Bornstein et al., 2023; Afacan-Seref et al., 2018; Rungratsameetaweemana, Itthipuripat et al., 2018; Diaz, Pisauero et al., 2024). These *time-varying* effects of expectations have been observed across different species, behavioral tasks, and neuroimaging modalities, but—critically—have been explained by different formal and/or neural processes in each instance.

Here, we demonstrate that these disparate observations of dynamic effects of expectations can be explained by a principled *uncertainty-reduction* process operating over parallel “streams” of memory and sensory evidence. To do this, we first review “static” models of expectations in decision-making and highlight a common assumption made by all of them: that observers have no uncertainty about their expectations. We then discuss why uncertainty about expectations ought to induce time-varying effects on expectations by appealing to the literature on *memory sampling* in two-alternative choice, before presenting a mathematical model specifying a principled process whereby observers combine uncertain expectations with sensory evidence that is itself uncertain. Finally, we demonstrate our model’s ability to capture several seemingly disparate observations of dynamic expectation effects across species and measurement modalities (behavior and single-unit recordings in both monkeys and humans, EEG, and fMRI), suggesting that it offers a unifying explanation of many previously unconnected phenomena.

2.2 Static effects of expectations on perceptual decisions: formal models and empirical data

To contextualize our model, we review the origins and development of static models of expectations in perceptual decision-making. We first describe the sequential sampling framework within which all models in this article are specified, and then introduce the diffusion decision

or drift-diffusion model (DDM). We focus on this model because it is the foundation upon which starting point models of expectations were initially developed. Accordingly, we then discuss the form, assumptions, and limitations of early starting point models, before turning to review some of the empirical evidence supporting them. Finally, we discuss the subset of drift rate models that also posit static, or time-constant, effects of expectations on the drift rate of accumulation processes, as well as empirical evidence supporting them. Throughout this section, we highlight the common principle of *uncertainty-weighting* as it appears in different formally-specified decision processes.

2.2.1 Sequential sampling and the diffusion decision model (DDM) of two-alternative choice

The most common framework for studying perceptual decisions based on dynamically-evolving sensory evidence is *sequential sampling*, often termed evidence accumulation (Gold and Shadlen, 2007; Forstmann et al., 2016; Ratcliff et al., 2016). This conceptual framework for studying the psychology and neuroscience of (mainly) two-alternative decisions is based on the mathematical framework of sequential analysis in statistics (Barnard, 1946; Wald, 1945; Gold and Shadlen, 2002). Both frameworks posit that decisions are made by continuously sampling information from an evidence source, extracting and integrating decision-relevant information over time, and committing to one of the options once the accumulated evidence surpasses a threshold value. The formal framework of sequential analysis yields normative solutions to the accumulation process, and thus can be used to build models that optimize the tradeoff between decision accuracy and deliberation time (the *speed-accuracy tradeoff*).

Importantly, not all sequential sampling models are normative. Some can be proven to non-optimally trade off speed and accuracy (e.g., race models; Vickers, 1970; Bogacz et al., 2006), whereas the normativity of others cannot be evaluated at all because of their mathematical form (e.g., the extended diffusion decision model; Ratcliff and Tuerlinckx, 2002; Bogacz et

al., 2006). In an influential paper, Bogacz et al. (2006) demonstrated that the earliest form of the *diffusion decision model* (or drift-diffusion model; DDM) optimally trades off speed with accuracy under different formalizations of that trade-off. Additionally, they showed that five other prominent models become equivalent to the DDM when their parameter settings optimize those same criteria (Bogacz et al., 2006), providing strong evidence for the normativity of the process specified by the DDM. Because early starting point models were designed to retain the optimality of the DDM while still incorporating a bias toward one option (Edwards, 1965; Link, 1975; Bogacz et al., 2006), we turn next to describing the DDM.

The first application of a sequential sampling modeling to the study of human decision processes is attributed to Stone (1960), who posited that the dynamics of two-alternative deliberation could be captured by a Gaussian random walk:

$$dx = A dt + c dW, \quad x(0) = 0 \tag{2.1}$$

where $x(0)$ represents the starting point of the decision variable, dx represents a change in the decision variable x over a unit of time dt , A represents the internal signal strength (corresponding to the drift rate of the process), and $c dW$ represents noise/diffusion in the accumulation process which follows a normal distribution with mean 0 and variance c^2 (Bogacz et al., 2006). The process ends once the decision variable surpasses a threshold value z , at which point the decision maker commits to the choice corresponding to the threshold value reached ($+z$ or $-z$). Throughout, we will call this model the “original” DDM to distinguish it from the “extended” DDM which is a probabilistic extension designed to maximize the model’s ability to fit behavior across a wide range of experimental conditions (Ratcliff and Rouder, 1998; Ratcliff and Tuerlinckx, 2002; Ratcliff et al., 2016).

The optimality of the original DDM comes from its equivalence with the sequential probability ratio test (SPRT; Bogacz et al., 2006; Moran, 2015). The SPRT gives the optimal procedure for determining when enough evidence has been gathered to terminate a two-alternative, discrete sampling process (Wald, 1945; Barnard, 1946; Wald and Wolfowitz, 1948). Optimality in the SPRT (hereafter “SPRT-optimality”) is defined as minimizing the average amount of time needed to make a decision at a fixed level of accuracy, or maximizing the proportion of correct responses for a fixed decision time (Wald and Wolfowitz, 1948; Bogacz et al., 2006). We discuss this form of optimality and its relationship to other forms at length in Khoudary et al. (2025b); a comprehensive overview is also provided by Bogacz et al. (2006).

Finally, it is worth mentioning that, in both the original and extended DDM, the drift rate A and prior probability Π terms are assumed to be functions of the decision environment. The drift rate term A has been robustly shown to be proportional to the signal-to-noise ratio in sensory stimuli (often termed the *coherence* of evidence; Palmer et al., 2005a; Gold and Shadlen, 2007; Ratcliff et al., 2016). The prior probability term Π likewise has been shown to scale with the frequency at which each choice outcome is presented (often termed the *bias* in the environment; Palmer et al., 2005b; Simen et al., 2009; Moran, 2015). The diffusion noise term c^2 is almost always treated as a fixed parameter during model fitting and is not commonly linked to properties of the decision environment (though see Zylberberg et al., 2016). The threshold term z is *optimally* defined as a function of the decision environment, but *psychologically* corresponds to individual differences in risk tolerance (Bogacz et al., 2006; Ratcliff et al., 2016), and thus is commonly allowed to vary across individuals during model fitting. In what follows, we focus on normative considerations for the drift rate A and prior probability Π terms.

2.2.2 Starting point models of expectations

The most prominent approach to incorporating expectations into sequential sampling models—and the DDM specifically—is adjusting the starting point $x(0)$ of the decision process in proportion to the magnitude of the expectation (Edwards, 1965; Bogacz et al., 2006; Simen et al., 2009; Gold and Shadlen, 2007; Ratcliff et al., 2016). This section reviews the normative principles motivating this convention, the assumptions upon which its optimality rests, and empirical evidence supporting starting point formalizations of expectations.

2.2.2.1 Normative grounding for starting point models

Edwards (1965) proposed one of the first extensions to the original DDM: adjusting the starting point of the decision process $x(0)$ to capture any existing biases that an agent might have toward deciding in favor of one of the two outcomes. Let Π denote the probability with which the choice option corresponding to $+z$ is correct on each trial. An *environment* is considered biased if $\Pi \neq 0.5$; i.e., if one of the two options is more frequently correct or rewarded relative to the other. Conversely, an *observer* is inferred to be biased toward one of the two options if $\Pi = 0.5$ but their choices are best described by a model where $x(0) \neq 0$. Edwards (1965) showed that the optimality of the original DDM can be retained if, in a biased environment, observers adjust their starting point $x(0)$ according to:

$$x_0 = \frac{c^2}{2A} \log \frac{\Pi}{1 - \Pi} \quad (2.2)$$

where A is the drift rate and c is the variance of a Gaussian diffusion process (Bogacz et al., 2006). A highly similar model was proposed by Link (1975), who instead posited that the log-odds ratio be scaled by a factor of 4. These differences stem from different formalizations of the speed-accuracy tradeoff, which we discuss in the next section. The *similarities* in the models, however, reveal two general principles about how expectations

are optimally integrated into perceptual decisions. First, the starting point should be offset in proportion to the log-odds ratio of choice outcomes before observing any evidence. Second, this offset should be scaled by the *ratio* of noise (c^2) to signal (A) in the decision environment (Bogacz et al., 2006). Critically, this scaling factor leads to exponentially greater weighting of expectations in decisions as the signal-to-noise ratio of evidence approaches zero (Figure A.1). In other words, expectations ought to be weighted more strongly as the uncertainty of sensory evidence increases.

2.2.2.2 Assumptions and limitations of the normative starting point model

The small difference in the forms of Edwards’ (1965) and Link’s (1975) models stem from different choices about how to formalize the speed-accuracy tradeoff as an optimality criterion. Edwards 1965 used the standard definition of SPRT-optimality (Section 2.2.1), whereas Link’s (1975) model minimizes error rate for a fixed decision threshold. Although these different definitions of optimality led to nearly identical solutions in this case, different definitions of optimality lead to qualitatively different normative models (and, by extension, distinct theories of optimal decision processes). A core argument of this article is that the normativity of the starting point model—along with its empirical utility—has overly simplified formal theories of expectation-setting and integration. Specifically, we propose that formal models and experimental settings alike have largely obscured the role that *uncertainty* about expectations has in shaping how they are integrated with sensory evidence that is itself uncertain (Figure 2.1). Incorporating uncertainty about expectations is important not only for capturing real-world scenarios within which observers make decisions, but also for interpreting the dynamics of information transfer across brain regions in more complex decision-making settings (e.g., Hachisuka, Shor, Liu et al., 2026).

To demonstrate these points, we begin by discussing the two assumptions that models must meet in order to be considered SPRT-optimal, and how these assumptions are violated in settings of considerable scientific interest. Assumption 1 describes conditions the decision

environment must meet, whereas Assumption 2 describes conditions that the observer’s knowledge must meet in order to execute the SPRT-optimal procedure. In this way, SPRT-optimality can be thought to be comprised of conditions on both external (Assumption 1) and internal (Assumption 2) factors of the decision process.

2.2.2.2.1 Assumption 1: homogeneous decision environment The first assumption that any SPRT-optimal model must make is that properties of the decision environment are *fixed*, *constant*, or *homogeneous*. This means that the drift rate A , diffusion magnitude c^2 , and prior probability Π *should* take on one scalar value that is constant across all decisions made in that context. On this assumption, any trial-level variability in the accumulation process is attributed solely to the noise term c^2 . When experimental conditions are deliberately constructed such that this first assumption of SPRT-optimality is met, behavior in humans and non-human animals appears largely SPRT-optimal (Link, 1975; Bogacz et al., 2006; Simen et al., 2009; Brunton et al., 2013).

However, the utility of SPRT-optimality for reasoning normatively about human behavior is limited for at least two reasons. First, real-world choice environments hardly—if ever—are homogeneous on all the dimensions that SPRT-optimality requires. Second, experimental tasks are also rarely homogeneous in an SPRT-compatible way. Unless an experiment explicitly sets out to test a prediction of an SPRT-optimal model, it is likely to manipulate the signal strength or difficulty of a decision at least pseudorandomly across trials. An important implication, then, is that arguments for starting point models that appeal to optimality are only valid for behavioral data collected in homogeneous decision environments. Once some trial-level variability in the signal strength A , evidence volatility c^2 , or prior probability Π is introduced, the starting point model is no longer the optimal procedure for integrating expectations about choice outcomes into the decision process (Moran, 2015).

2.2.2.2.2 Assumption 2: full access to properties of the environment The second assumption required for SPRT-optimality is that decision-makers have access to the fixed, true values of the environmental properties corresponding to the parameters A , c^2 , and Π : sensory signal strength, sensory signal volatility, and the prior probability of either choice option. This assumption follows directly from the procedure on which SPRT optimizes the speed-accuracy tradeoff. An observer who has access to the true values of A , c^2 , and Π , along with their desired proportion of correct responses, can, in theory, compute the precise threshold value z that leads to SPRT-optimal behavior (minimizing the average amount of time needed to respond at a desired level of accuracy).

A core limitation on the utility of this assumption is that humans and other animals rarely—if ever—have direct access to this information when making perceptual decisions in the real world. Instead, multiple learning and memory systems work together to store and flexibly utilize information about statistical regularities, both about the decision-maker and their environment, acquired across different timescales of experience (Seriès and Seitz, 2013; de Lange et al., 2018; Wang et al., 2021; Bakkour et al., 2019; Yoo and Bornstein, 2024; Gläscher et al., 2010; Bornstein and Daw, 2013, 2012; Noh et al., 2023). A more realistic starting premise, then, is that decision-makers *learn* about the values of A , c^2 , and Π through experience and—critically—treat these sparsely-informative expectations as appropriately uncertain. Accordingly, observers must rely on their learning processes to optimize behavior, while *also* incorporating their uncertainty about the learned values into the optimization process (e.g., Hanks et al., 2011; Khodadadi et al., 2014; Findling et al., 2025). We expand upon these points in Sections 2.2.3.2.2 and 2.3.2.1.

2.2.2.3 Empirical support for starting point models

A number of empirical findings support the idea that expectations bias perceptual decisions by modifying the starting point of the evidence accumulation process (Carpenter and Williams, 1995; Ratcliff and McKoon, 2008; Simen et al., 2009; van Ravenzwaaij et al., 2012;

Mulder et al., 2012; de Lange et al., 2013; Bornstein et al., 2023). Early work initially identified starting point biases via fits to human behavior (Carpenter and Williams, 1995, Ratcliff and McKoon, 2008, Simen et al., 2009, van Ravenzwaaij et al., 2012), with later work linking starting point effects to neural activity in humans (de Lange et al., 2013, Mulder et al., 2012, Bornstein et al., 2023) and non-human primates (Rao et al., 2012). This section briefly reviews these findings, which demonstrate that (1) humans do appear to optimally integrate expectations with sensory evidence and (2) expectations exert effects on the starting point even when the conditions for SPRT-optimality are not met.

One of the few experiments that explicitly created an environment that meets the assumptions for SPRT-optimality was conducted by Simen et al. (2009). The authors fixed signal strength A across trials, and manipulated the prior probability of choice outcomes Π across blocks of trials. In support of optimal behavior, the authors found selective effects of prior probability Π on the starting point of the decision process: starting point values were selectively greater in the biased blocks than in the neutral blocks (Simen et al., 2009). The authors additionally manipulated the amount of trials—and the time elapsed between them—across blocks, such that some blocks contained several trials presented quickly in succession, whereas other blocks contained fewer trials with greater amounts of time elapsed between them. Combining this manipulation with that of the prior probability makes a somewhat surprising prediction: as the amount of elapsed time between trials becomes sufficiently low, and the prior probability Π becomes sufficiently high, observers should opt for a *non-integrative* decision procedure: responding immediately with the choice that has a higher Π (Simen et al., 2009).

This prediction follows directly from (1) the optimality of lowering thresholds for faster-paced choice environments with more opportunities to make a decision and (2) the optimality of biasing the starting point toward the more likely option in proportion to Π .¹ RT distributions

¹Prediction (1) here derives from the *reward-rate* formalization of the speed-accuracy tradeoff, which posits that observers’ objective function is to maximize their total amount of reward (i.e., correct choices) per

across conditions demonstrated that observers indeed increasingly adopted non-integrative strategies: at the highest levels of Π and task-pacing, most responses were made immediately at the onset of sensory evidence (Simen et al., 2009). Conversely, as the pacing of the task slowed down—and observers had fewer opportunities to make correct responses—the proportion of non-integrative responses *decreased*, an effect that is particularly pronounced in the most biased environments ($\Pi = 0.9$; Simen et al., 2009). Taken together, these findings suggest that humans can optimally modify their decision processes to account for temporal and statistical properties of different choice environments.

A study by Mulder et al. (2012) used functional neuroimaging to further investigate the neural correlates of starting point effects. The authors contrasted two types of expectations: one consisting of information about the likely outcome (i.e., prior probability) and another consisting of information about the outcome more likely to be rewarded (i.e., potential payoff). To do this, they incorporated a cue at the start of each trial of a motion discrimination task. The cues consisted of a left- or right-pointing arrow and either a “%” (indicating prior probability) or “€” (indicating potential payout). Although this trial-level manipulation of expectations makes the SPRT no longer optimal, Mulder et al. (2012) found that both types of cues induced effects on the starting point. Further, they found that starting point offsets returned by the model parametrically modulated BOLD activity in frontoparietal regions and anterior cingulate gyrus, suggesting these areas may encode information about prior probability during decision-making (Mulder et al., 2012).

Finally, a study by Rao et al. (2012) revealed neural evidence of a starting point effect in rhesus macaques performing a motion discrimination task. The authors obtained electrophysiological measurements from neurons in measurements from neurons in cortical areas previously shown to dissociably encode sensory and decisional information (middle tempo-

unit time (e.g., Drugowitsch et al., 2012). In homogeneous environments, reward-rate and SPRT-optimality are equivalent (Moran, 2015); the reward-rate formalization simply adds the ability to make predictions about how task timing ought to impact optimal choice procedures.

ral (MT) and lateral intraparietal cortex (LIP), respectively). Rao et al., 2012 found that trial-varying cues induce directional activity in LIP, such that firing rates are increased for neurons selective to the cued motion direction and decreased for neurons selective to the alternative. The authors did not find an effect of expectations on firing rates in MT, which led them to conclude that expectations modulate decisions by selectively biasing the starting point of evidence.

Taken together, the work reviewed in this section lends empirical support for a starting point effect of expectations. Model fitting to human behavior reveals that starting point biases describe effects of expectations well, further suggesting that human perceptual decision-making is largely optimal. Neuroimaging work has linked these starting point effects to neural activity in frontoparietal regions, suggesting that expectations enter into the decision process as an additional source of evidence that’s integrated with sensory information to optimize the decision process. However, a growing body of work suggests that expectations might *also* exert effects on the drift rate of the decision process (Hanks et al., 2011; Moran, 2015; Dunovan and Wheeler, 2018; Diaz et al., 2024; Kelly et al., 2021; Bornstein et al., 2023). We discuss these findings and how they relate to optimal decision models in the next section.

2.2.3 Drift rate models of expectations

Whereas starting point models posit that expectations simply shift the initial value of the decision process $x(0)$, drift rate models posit that the time-evolving evidence itself is accumulated more efficiently for outcomes that have a higher prior probability Π . Conceptually, drift rate models posit that expectations *amplify* the signal of expectation-congruent evidence and *attenuate* the signal of expectation-incongruent evidence. This type of effect on decisions is qualitatively similar to models that posit a sharpening of sensory representations

as a function of expectations (e.g., Kok et al., 2012; Afacan-Seref et al., 2018), but changes the locus of the effect from the sensory stage to the decisional stage.

Before discussing the theoretical and empirical evidence supporting drift rate effects of expectations, it is worth briefly discussing the distinction between *static* and *dynamic* effects of expectations and how they have been mapped to different models in the literature. In this article, we use the term “static” to refer to an effect of expectations that remains constant across the amount of time it takes to complete a single trial. Static effects can be contrasted against “dynamic” effects of expectations, which we use to mean an effect that *changes* over the time it takes to make a choice. The literature often maps the static vs. dynamic distinction onto effects on the starting point vs. drift rate of the accumulation process (e.g., Dunovan et al., 2014; Moran, 2015; Malhotra et al., 2018). This mapping reflects the general interpretation that starting point effects modify the *initial conditions* of the process, whereas drift rate effects modify its *dynamics*. Importantly, however, effects on drift rate can *themselves* be dynamic, such that the rate of accumulation changes over the time it takes to make a single choice (e.g., Afacan-Seref et al., 2018; Hanks et al., 2011). Accordingly, we will use the terms *static drift* and *dynamic drift* to refer to these two different classes of drift rate effects. Finally, our usage of “dynamic” is intended to apply only to effects that occur within a single trial, rather than effects that change over the course of an experimental session (e.g., Alister and Evans, 2026).

2.2.3.1 Normative grounding for static drift models of expectations

Moran (2015) combined formal decision theory with numerical simulations to investigate how the original DDM can maximize reward rate in an environment that fails to meet the assumptions of SPRT. This was achieved by introducing one simple modification to the formal decision problem: permitting the sensory signal strength to vary across trials. Adding trial-level variability in the sensory signal relaxes the assumption of a homogeneous decision environment (Section 2.2.2.2.1), which in turn generates uncertainty in the observer’s knowl-

edge of the statistical properties of the environment (Section 2.2.2.2.2). When an observer does not know in advance how decisive the sensory evidence will be, they cannot appropriately set their threshold in accordance with SPRT, and thus the original DDM fails to provide a normative solution for trading off speed and accuracy in heterogeneous environments (Moran, 2015; Drugowitsch et al., 2012). Despite this, Moran (2015) used numerical simulations to investigate how the process specified by the DDM might approximate optimality (i.e., maximize reward rate) in environments that are both heterogeneous *and* biased (i.e., $\Pi \neq 0.5$). The author found that an additive bias on *both* the starting point *and* drift rate parameters were necessary for maximizing reward rate in this environment, in contrast to previous findings suggesting that no offset on drift was required (van Ravenzwaaij et al., 2012; see Section 4.3.2 of Khoudary et al. (2025b) for an extended discussion of these two findings).

2.2.3.2 Assumptions and limitations of the normative drift model

The results of Moran (2015) thus indicate that, in biased and heterogeneous environments, a static drift model—coupled with a starting point offset—maximizes reward rate (and thus optimizes the decision process). While these findings advance theoretical knowledge of optimal behavior in more realistic environments, they necessarily rest on simplifying assumptions about the environment and the observer. We discuss each of these assumptions—and their conceptual limitations—before turning to the theoretical framework and computational model we developed to address the gaps left by Moran’s 2015 analysis.

2.2.3.2.1 Assumption 1: *expectation*-homogeneous decision environment The first assumption underlying Moran’s (2015) findings is that there is one fixed level of prior probability Π (or bias) in the environment, but that the sensory evidence strength (and, by extension, the drift rate A) varies unpredictably across trials. This corresponds to an environment that is *expectation*-homogeneous but *evidence*-heterogeneous: observers use the

same expectation for all decisions made in the environment, but the difficulty of the decision varies from trial to trial. We will use the term *fully heterogeneous* to describe environments that have trial-level variability in *both* the sensory signal strength and the prior probability (which we shall denote Π_{cue} to differentiate from the block-level expectation Π).

As with Assumption 1 for starting point models, basing the normative solution on expectation-homogeneous but evidence-heterogeneous environments restricts the scope of Moran’s (2015) findings to behavior measured in these environments. Experimentally, this corresponds to block-level manipulations of prior probability Π with trial-level manipulations of sensory signal strength. While the lion’s share of early work on expectations in perceptual decisions used exactly this type of decision environment (Carpenter and Williams, 1995; Ratcliff and McKoon, 2008; Simen et al., 2009; Hanks et al., 2011; van Ravenzwaaij et al., 2012), it is becoming increasingly common to measure behavior in fully heterogeneous environments (Kok et al., 2013; Dunovan et al., 2014; Dunovan and Wheeler, 2018; Kelly et al., 2021; Bornstein et al., 2023; Diaz et al., 2024; Khoudary et al., 2025a). To our knowledge, a normative model for this type of behavioral problem has yet to be specified.

2.2.3.2.2 Assumption 2: perfect knowledge of Π The second assumption underlying Moran’s (2015) findings is that observers have perfect knowledge of the prior probability Π , but imperfect (or uncertain) knowledge of the signal strength (and corresponding drift rate A) on each trial. It is precisely this trial-level uncertainty in sensory signal strength that distinguishes Moran’s findings from the SPRT, and that justifies an effect of expectations on the drift rate in addition to the starting point.

Unlike the assumption of expectation-homogeneity, experimental conditions only sometimes relax this assumption of perfect prior knowledge. Indeed, the dominant approach has been to explicitly instruct humans on prior probabilities, regardless of whether they are manipulated at the block- or trial-level (e.g., Carpenter and Williams, 1995; Ratcliff and McKoon, 2008;

Simen et al., 2009; Hanks et al., 2011; Dunovan et al., 2014; Dunovan and Wheeler, 2018; Diaz et al., 2024). In studies using non-human animals, observers are trained on tens of thousands of trials until the experimenters are confident that the animals have sufficiently learned environmental biases (Hanks et al., 2011; Rao et al., 2012), such that there is little to no uncertainty in their knowledge of either Π or Π_{cue} . As discussed in the context of starting point models (§2.2.2.2), the assumption that observers have near-perfect knowledge of prior probabilities completely elides the questions of *where* expectations come from and *how* the source of expectations might impact their integration with sensory evidence. The theoretical framework and computational model we present in the following sections aim to address this gap directly, positing a central role of *memory retrieval dynamics* in shaping the integration of uncertain expectations with sensory evidence that is itself uncertain.

2.2.3.3 Empirical support for static drift models of expectations

Before detailing our framework and model, it is worth briefly reviewing some of the recent work supporting static effects of expectations on the drift rate. One of the first studies to report an effect of expectations on drift was conducted by Dunovan et al. (2014), who measured both behavior and neural activity in fully heterogeneous environments where cues unambiguously informed observers about Π_{cue} on a trial-by-trial basis. The authors found that a “multi-stage model” where both the starting point and drift rate varied as a function of trial-level expectation cue outperformed several simpler models on both quantitative model comparisons and qualitative match to behavioral data. A later study by Dunovan and Wheeler (2018) linked these effects to BOLD activity in category-selective regions of inferotemporal cortex. The authors showed that cue-evoked BOLD timecourses were systematically modulated by the predictiveness of a cue, whereas sensory-evoked timecourses were more strongly modulated by the *mismatch* between a cue’s prediction and the subsequent contents of perception. An EEG study by Rungratsameetaweemana, Itthipuripat et

al. (2018) provided corroborating evidence for this “mismatch” effect, which we show that our model captures in Section 2.5.2.

A study by Kelly et al. (2021) further revealed effects of expectations on drift rate that are modulated by the difficulty of the decision task. Across three blocks, Kelly et al. (2021) manipulated choice difficulty by manipulating the signal strength and amount of time allotted to make each choice. Expectations were manipulated trial-by-trial, with predictive probabilities learned during the practice phase of the task (Kelly et al., 2021). The authors then built a “neurally-informed” sequential sampling model that used EEG measurements to constrain the threshold, non-decision, and starting point terms, and compared the performance of this model to that of the extended DDM (eDDM). In addition to showing that the neurally-informed model outperformed the eDDM on measured and simulated data, the authors reported positive biases on drift rates as a function of expectations across all difficulty blocks (Kelly et al., 2021). Further, these biases were linked to changes in the amplitude of the centroparietal positivity (CPP), an event-related EEG component that has been well-established to co-vary with evidence strength (O’Connell et al., 2012; van Vugt et al., 2019). Specifically, expectation-based effects on the CPP were strongest when observers (1) had a limited amount of time to make a choice and (2) there was a mismatch between the cue’s prediction and the subsequent content of perception (Kelly et al., 2021), further corroborating the findings discussed above.

Finally, a recent study by Liu et al. (2025) showed that effects of expectations on the drift rate and starting point can also be observed in rodents. The authors investigated brain-wide recordings from rodents who learned expectations within a decision environment, and used several formal models to characterize the effects of learned expectations on neural activity. Liu et al. (2025) found that learned expectations modulate the initial value and gain on firing rates both for single neurons and populations thereof, but only if those firing rates appeared to encode accumulated evidence or action initiation rather than sensory evidence.

These findings further support the idea that expectations can modify decision—but not sensory-related—processing by modulating both the starting point and drift rate of evidence accumulation on a trial-by-trial basis (Liu et al., 2025).

2.2.4 Section summary

This section reviewed the formal and empirical evidence supporting a *static* effect of expectations on perceptual decisions. Again, our usage of static in this article pertains to *how* the effect of expectations is modeled; a static effect corresponds to an offset on the decision process that does not change over the course of a single decision. On this definition, expectations can—and should—exert static effects on both the starting point and drift rate (Wald and Wolfowitz, 1948; Bogacz et al., 2006; Moran, 2015). We reviewed literature demonstrating behavioral and neural evidence that human and non-human observers optimally incorporate expectations into the decision process, either as a static offset solely on the starting point or on both the starting point and drift rate (Simen et al., 2009; Mulder et al., 2012; Rao et al., 2012; Moran, 2015; Dunovan et al., 2014; Dunovan and Wheeler, 2018; Kelly et al., 2021; Liu et al., 2025).

Mounting evidence, however, suggests that effects of expectations might *vary* within the time it takes to make a choice (Hanks et al., 2011, Rungratsameetaweemana, Itthipuripat et al., 2018, Bornstein et al., 2023, Hachisuka, Shor, Liu et al., 2026). Static starting point models categorically cannot capture such *time-varying* effects of expectations. Static drift rate models can approximate time-varying effects, but are limited in positing effects that are strictly linear in time (Moran, 2015). In the following sections, we suggest—and then show—that time-varying effects of expectations might be driven by the dynamics of *memory retrieval* processes that function to reduce uncertainty during latent state inference processes (Wang et al., 2021; Shadlen and Shohamy, 2016). To do this, we first introduce sequential sampling approaches for modeling the role of memory retrieval in action selection

in Section 2.3.1. Then, we discuss how *uncertainty about expectations* has been omitted from the theoretical study of perceptual decision-making despite its ubiquity in real-world settings, and delineate two dimensions of expectation-uncertainty that have been obscured by assumptions of normative models (Section 2.3.2; Figure 2.1). We then present a novel sequential sampling model formalizing the most general construal of our framework (Section 2.4) and use simulations to show that it reproduces dynamic effects of expectations observed across several different species, tasks, and neuroimaging modalities (Section 2.5). Finally, we present a secondary analysis of behavioral data to empirically validate a novel prediction generated by our framework of uncertainty-driven memory sampling to further demonstrate its utility for understanding the dynamics of decision-making under uncertainty.

2.3 Conceptual gap: role of memory retrieval dynamics

All mathematical modeling—and normative modeling in particular—requires idealizing properties of the system being modeled (i.e., the *target*; Khoudary et al., 2025b). Idealizations simplify aspects of target systems by *omitting* and/or *distorting* their properties, resulting in formal representations that “selectively attend” to aspects of the system about which model users aim to reason (Portides, 2021). Omitting uncertainty about expectations from formal models—and the experiments used to test them—is a powerful idealization that has led to great progress in understanding the neural and computational mechanisms of optimal decision-making. In doing so, however, the field has systematically ignored the role that *memory retrieval* might play in shaping the dynamics of expectation-guided perceptual decisions.

The relevance of memory for expectation-guided decisions becomes obvious when considering the question of where expectations come from *outside* of laboratory settings; one answer

could be previous outcomes of similar decision problems. Because human memory systems store information about past experiences at multiple levels of detail (Tarder-Stoll et al., 2026), precisely what information about the past comes to mind during decision-making—and in what form—is an area of active research (Zeithamova, Schlichting, Preston, 2012; Aronowitz, 2019; Biderman et al., 2020; Wang et al., 2021; Shushruth et al., 2022; Bakkour et al., 2019). For present purposes, we focus on a relatively simple type of information: probabilities of choice outcomes conditioned on a specific context. These “contextual expectations” (Seriès and Seitz, 2013) are acquired through repeated experience with a decision environment (de Lange et al., 2018), and have been shown to facilitate various types of perceptual judgments such as object recognition (Oliva and Torralba, 2007), visual search (Zhou and Geng, 2024), and two-alternative discrimination (Bornstein et al., 2023). Expectation-setting in the real world thus involves retrieving information from memory, but the role of memory retrieval dynamics in shaping perceptual decisions remains largely underexplored. Sequential sampling, or evidence accumulation, provides a common formalism for building theories about how agents trade off speed and accuracy when making decisions based on *both* external (sensory) and internal (memory) information. Combining this formalism with theoretical insights from learning and memory can help advance understanding of the general principles that support flexible and adaptive decision-making in humans, which is what we aim to achieve with the framework and model presented in this article.

2.3.1 Memory as a dynamic source of internal evidence

Retrieving information from memory takes time. Formalizing the relationship between the strength of a memory and the amount of time needed to retrieve it was the goal of the earliest form of the extended drift diffusion model (Ratcliff, 1978), which is now commonly used to analyze behavior in value-based decision-making tasks where the relevant “evidence” used to make a choice is strictly internal to decision-makers (Ratcliff et al., 2016; Krajbich et al., 2010). The *memory sampling* framework explains the ability of accumulator models

to fit this class of data by positing that observers sequentially sample information from memory to make predictions about the outcome of each choice, terminating deliberation once the predicted value of one choice sufficiently outweighs the predicted value of the other (Shadlen and Shohamy, 2016; Bakkour et al., 2019; Wang et al., 2021; Biderman et al., 2020; Zylberberg et al., 2024).

Crucially, evidence supporting memory sampling theories of action selection comes from experiments where memory strength is permitted to vary across trials (Wang et al., 2021; Shadlen and Shohamy, 2016; Bakkour et al., 2019). This variability derives both from structured manipulations of experimental stimuli (e.g., Nicholas and Mattar, 2026) and the stochasticity of memory retrieval processes both within and across individuals (Wang et al., 2021; Bornstein and Pickard, 2020). Measuring and manipulating memory as a source of decisional evidence has revealed that memory-based decisions obey the same dynamics governing sensory-based decisions: stronger memories lead to faster and more accurate decisions, whereas weaker memories lead to slower and less accurate decisions (Shadlen and Shohamy, 2016; Biderman et al., 2020; Wang et al., 2021).

More recently, it has been shown that humans can adaptively control this memory sampling process, engaging in prolonged sampling only when the expected utility of doing so outweighs the costs of investing more time and effort into the retrieval process (Bornstein et al., 2023; Callaway et al., 2024). The framework of sequential sampling—and its formal equivalence with partially observable Markov decision problems (Rao, 2010; Drugowitsch et al., 2012; Huang et al., 2012)—permits interpreting this adaptive control of memory retrieval in terms of minimizing uncertainty during latent state inference. When sampling information from an evidence source, observers can take three possible actions: choose latent state $+$, choose latent state $-$, or keep sampling. These actions a are determined by the distance of the decision variable x_t from the threshold z at each point in time as follows:

$$a_t = \begin{cases} \text{choose+} & \text{if } x_t \geq z_t \\ \text{choose-} & \text{if } x_t \leq -z_t \\ \text{sample} & \text{if } x_t < |z_t| \end{cases} \quad (2.3)$$

Interpreting the decision variable x_t as reflecting observers’ probabilistic, time-varying belief about the true state of the world (+ or −) further permits interpreting SAMPLE actions as decisions to continue acquiring information to reduce uncertainty in x_t (Drugowitsch et al., 2012; Bornstein et al., 2023; Rao, 2010; Kaelbling et al., 1998). Accordingly, prolonged sampling from memory—measured via response times and/or neural activity—can be interpreted as reflecting greater uncertainty about the true value of the data-generating distribution (i.e., the outcome predicted by the set of relevant associative and/or episodic memories). The relationship between sampling time and uncertainty is well-established in the perceptual decision literature, where agents of all species demonstrate robust, systematic relationships between sensory evidence uncertainty and sampling or deliberation time (Gold and Shadlen, 2002, 2007; Hanks et al., 2011; Fetsch et al., 2014; Kiani et al., 2014; Hanks and Summerfield, 2017).

2.3.2 Two dimensions of uncertainty about expectations

Again, a core argument of this article is that the assumptions required for optimal decision processes have obscured the role of uncertainty-driven sampling from memory. To help address this gap, we distinguish two dimensions along which uncertainty about expectations can vary (Figure 2.1). First, we posit that the *source* of an expectation—or how it is acquired and stored in memory—entails uncertainty about *what* the expectation predicts. Second, we posit that environmental demands on the *deployment* of expectations—or how they are used to guide future decisions—entails uncertainty about *which* expectation to use across consecutive decisions. As shown in the gray text of Figure 2.1, neural and behavioral

evidence for time-varying effects of expectations have been observed in experimental settings corresponding to each of the four quadrants delineated by these dimensions. Section 2.5 shows that our sequential sampling model formalizing this framework reproduces each of the time-varying effects reported by these studies, suggesting that uncertainty-driven memory sampling offers a unifying explanation for time-varying effects of expectations. In support of this ultimate claim, the rest of this section elaborates upon the dimensions of uncertainty picked out by our framework. For each dimension, we discuss how it relaxes the assumptions of previous normative models, along with the theoretical and empirical motivations for doing so.

2.3.2.1 Source-related uncertainty

The dimension of source-related uncertainty (orange axis in Figure 2.1) is intended to capture uncertainty associated with how expectations are acquired and subsequently stored (or represented) in memory. Accordingly, it relaxes the assumptions of perfect knowledge of Π or Π_{cue} that are required for normative static effects of expectations (Sections 2.2.2.2.1 and 2.2.3.2.2). Relaxing this assumption is especially important because it has obscured investigation into *where* expectations come from and *how* the uncertainty associated with an expectation ought to impact its integration with sensory evidence during decision-making.

An observer can be said to have source-related uncertainty about their expectations if they acquire those expectations on the basis of implicit, experience-based learning rather than explicit instruction.² This type of uncertainty is thus about the true prior probability of stimulus occurrence in the decision environment (Π) or the predictive probability of a learned cue with respect to its associated stimulus (Π_{cue}). A core idea behind our framework is that when prior probabilities are not immediately accessible or directly known by the observer, they must *estimate* these probabilities on the basis of past experience. Doing so

²This dimension of uncertainty roughly corresponds to the description-experience gap as used in the literature on risky, value-based choice (Hertwig and Erev, 2009; Heilbrunner and Hayden, 2016).

requires retrieving relevant information from memory in order to construct an estimate of the prior probability. This process should be jointly modulated by (1) the strength of the true prior probability and (2) the amount of information an observer has stored in memory (or, the precision of their belief). Accordingly, more certain probabilities (closer to 0 or 1) should take fewer samples to estimate, whereas less certain probabilities (closer to 0.5) require prolonged sampling for observers to obtain steady-state estimates. As an observer gains more experience—either with the relevant associations comprising the memory or with estimating probabilities based on a specific set of associations—their belief becomes more precise, and there should be less overall uncertainty in the estimate regardless of its true value. “Structural” expectations—reflecting general statistical information that is preserved across contexts (Seriès and Seitz, 2013)—can be thought to reflect an asymptotic form of these expectations learned through experience.

As mentioned above, the majority of experimental work has ignored this dimension of uncertainty about expectations, opting instead to give human observers explicit knowledge of prior probabilities or extensively training non-human animals such that any uncertainty about Π_{cue} is minimized. One exception to this trend comes from a small handful of experiments where observers implicitly learn environmental biases while performing the main decision task. Behavior and neural activity in both humans (Rungratsameetaweemana, Itthipuripat et al., 2018; Wert et al., 2025) and rodents (Findling et al., 2025) has been shown to be sensitive to these latent environmental biases, with observers of both species demonstrating evidence of learning within as few as 20 trials. Yet how these learned biases manifest in neural representations and the decision process remains underexplored. In humans, for example, these implicitly learned expectations appear to selectively impact decisional, rather than sensory, EEG components (Rungratsameetaweemana, Itthipuripat et al., 2018). Rodent studies, in contrast, have shown that information about the latently-learned environmental bias can be decoded from up to 30% of the entire brain, spanning early sensory to higher-level cortical areas (Findling et al., 2025) rather than being confined to decisional circuitry. These

findings thus shed some light on the neural mechanisms involved in integrating uncertain expectations with sensory evidence, but leave open the question of which computations drive this integration. In humans, there is the additional question of *which* representations these computations operate over; the framework and model we develop in this article suggests that this representation is a dynamic, probabilistic “stream” of evidence samples retrieved from associative memory.

2.3.2.2 Deployment-related uncertainty

The dimension of deployment-related uncertainty (purple axis in Figure 2.1) captures uncertainty associated with how expectations are cued in advance of decisions. Namely, an observer can be said to have deployment-related uncertainty if they are making decisions in an environment where the prior probability Π_{cue} differs across decisions in a manner unknown to them. This dimension of uncertainty relaxes the assumption of *expectation-homogeneity* that is ubiquitous among normative models of expectations in decision-making (Edwards, 1965; Bogacz et al., 2006; Simen et al., 2009; Moran, 2015; Malhotra et al., 2018), along with foundational experiments designed to test their predictions (Carpenter and Williams, 1995; Ratcliff and McKoon, 2008; Simen et al., 2009). As shown in Figure 2.1, deployment-related uncertainty is much lower in experiments that manipulate expectations on the level of blocks of trials, such that observers make hundreds to thousands of consecutive decisions using the same expectation.

Relaxing the assumption of expectation-homogeneity is important for two reasons. First, the assumption itself strongly constrains the space of real-world problems to which existing normative models can be applied. Human observers are rarely tasked with making more than a handful of consecutive perceptual decisions using the same exact expectation, much less hundreds or thousands as in block-level experimental manipulations. Second, and relatedly, it is becoming increasingly common to measure behavior and neural activity of observers making decisions in expectation-heterogeneous environments (Dunovan et al., 2014; Dunovan

and Wheeler, 2018; Aitken and Kok, 2022; Kok et al., 2012; Kok et al., 2013; Bornstein et al., 2023; Diaz, Pisauro et al., 2024). Although the prescriptions of existing normative models can offer some guidance for interpreting findings from these environments, a principled model of how deployment-related uncertainty ought to impact decision processes would allow for finer-grained prescriptions guidance for future experimental designs.

Two dimensions of uncertainty about expectations

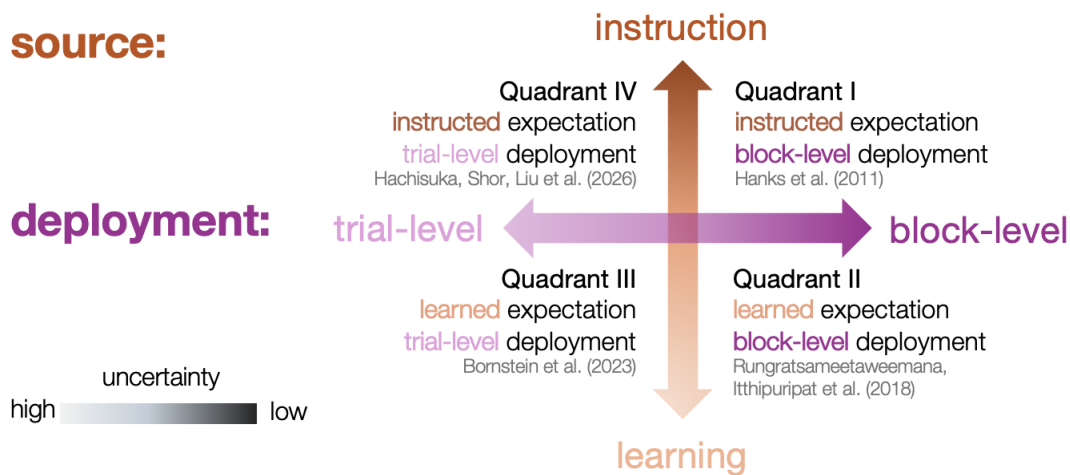


Figure 2.1: Two dimensions of uncertainty that observers can have about expectations when making perceptual decisions. Source-dependent uncertainty (orange) captures uncertainty about *what* an expectation predicts, whereas deployment-dependent uncertainty (pink) captures uncertainty about *which* expectation to use for each decision. The relative degree of uncertainty within a dimension is represented by the opacity of its respective color, with highly uncertain conditions in lighter shades and less certain conditions in darker shades. Gray-text citations indicate existing findings of dynamic effects observed in each quadrant; we show that our model can capture each of these findings in Section 2.5.

2.3.3 Dynamic reliability estimation as a unifying normative principle

A final motivation for relaxing the assumption of expectation-homogeneity comes from the form of previous normative solutions to decision problems in heterogeneous sensory environments. Multiple converging studies have shown that optimal decision-making in evidence-heterogeneous environments—where observers have uncertainty about the strength or du-

ration of sensory evidence—requires some dynamically-changing component of the decision process (Frazier and Yu, 2007; Drugowitsch et al., 2012; Huang et al., 2012; Deneve, 2012; Glaze et al., 2015; Malhotra et al., 2018; Kilpatrick et al., 2019; Barendregt et al., 2022; Callaway et al., 2024; Fang et al., 2026). The intuition behind the need for a dynamic component is the same across all models: when an observer doesn’t know in advance how informative the evidence will be, they must *estimate* its informativeness per unit time—or its *reliability*—over the course of evidence sampling. This evidence reliability estimate is then used to optimize the speed-accuracy tradeoff by dynamically modulating the decision process in real-time. Specifics of how the dynamic adjustment is incorporated vary across model applications and objective functions (see Kilpatrick et al., 2019 for a review), but the general premise of modifying the decision process in proportion to the estimated reliability of evidence remains constant. Our work extends this principle to the integration of expectations into perceptual decisions. We show that multiple disparate empirical observations of time-varying effects of expectations can be explained by a single model that (1) treats expectations as arising from a series of samples from memory and (2) dynamically adjusts their weighting in the decision process according to a time-evolving belief about the reliability of sensory evidence *relative* to that of the expectation.

2.4 Dynamic reliability-weighted multi-source sequential sampling model

This section presents a sequential sampling model formalizing the conceptual framework developed above. One of its core features is modeling expectations as an independent source of probabilistic information that observers sample from and accumulate *in parallel* to sensory evidence. This process thus equips observers with two sources of probabilistic evidence that they have to combine to make decisions. We draw on well-established theoretical and empirical findings in multisensory integration to propose that these two evidence sources—

memory and sensation—are combined according to their *relative* reliability (e.g., Fetsch et al., 2011; Landy et al., 2011; Angelaki et al., 2009). We model the reliability estimation process as dynamically updating with each new sample of noisy evidence from each source (Figure 2.2); this eliminates the need for using elapsed time as a proxy variable, and encompasses cases where proxies such as time diverge from evidence reliability. The following section shows how this dynamic, relative reliability-weighted process captures time-varying effects of expectations observed by Hanks et al. (2011), Rungratsameetaweemana, Itthipuripat et al. (2018), Hachisuka, Shor, Liu et al. (2026) , and Bornstein et al. (2023).

2.4.1 Model specification

In our model, simulated agents sequentially sample evidence from two sources $s = \{memory, vision\}$. Samples for both sources are generated by a Bernoulli distribution with parameter θ_s , which we assume to be unknown by the observer (Figure 2.2A, left). Additionally, we assume that these samples are noisily encoded at a proportion defined by ϵ , such that:

$$o_{s_t} \sim \text{Bernoulli}((1 - \epsilon)\theta_s + \epsilon(1 - \theta_s)) \quad (2.4)$$

Next, we posit that observers use these samples to form and update a Beta-distributed belief $g_s(t)$ about the value of θ_s for each evidence source (Figure 2.2, middle). On this specification, θ corresponds to the “signal strength” for each evidence source: the true θ_{vis} value is identical to signal strength A for sensory evidence, and the true θ_{mem} value is identical to predictive probability Π_{cue} . Each new sample from each source updates the belief $g_s(t)$ according to:

$$g_s(t) \equiv p(\theta_s | o_{s_t}, o_{s_{t-1}}, \dots, o_{s_{t_0}}) \sim \text{Beta}(\alpha_s, \beta_s) \quad (2.5)$$

$$g_s(t) \equiv \begin{cases} \text{Beta}(\alpha_{s_{t-1}} + 1, \beta_{s_{t-1}}), & \text{if } o_{s_t} > 0 \\ \text{Beta}(\alpha_{s_{t-1}}, \beta_{s_{t-1}} + 1), & \text{if } o_{s_t} < 0 \end{cases} \quad (2.6)$$

The central tendency of $g_s(t)$ captures the observer's posterior belief about each source's likely signal strength at each moment in time, and the variance of $g_s(t)$ captures the observer's uncertainty about the inferred signal strength. We define internal evidence reliability λ_{s_t} as a term that captures both the central tendency and variance of $g_s(t)$ (Figure 2.2A, right):

$$\lambda_{s_t} = \sum (g_s(t) \cdot H(X))^{-1} \quad (2.7)$$

where X is a vector of probability values ranging from 0 to 1 and H is the information entropy function. Each source's dynamic reliability estimate thus corresponds to the dot product of the observer's belief about that source's signal strength θ_s with the information entropy function $H(X)$. This specification creates a reliability estimate that quantifies how informative an ideal observer ought to believe a particular evidence source is, based on how informative it has been up to this point in the trial. Then, we define the optimal weighting of each evidence source w_{s_t} using the standard ideal-observer model of cue integration (Figure 2.2B, left):

$$w_{memory_t} = \frac{\lambda_{memory_t}}{\lambda_{vision_t} + \lambda_{memory_t}}; \quad w_{vision_t} = \frac{\lambda_{vision_t}}{\lambda_{vision_t} + \lambda_{memory_t}} \quad (2.8)$$

Observers weight each observation from each evidence source by these source-specific weights, creating a decision variable x_t that is a time-varying reliability-weighted sum of samples from independent streams of memory and visual evidence (Figure 2.2B, right):

$$x_t = x_{t-1} + o_{memory_t} w_{memory_t} + o_{vision_t} w_{vision_t} \quad (2.9)$$

The decision process terminates according to

$$|x_t| > z \tag{2.10}$$

where z is a scalar defining the boundary separation (threshold) value.

In addition to the integration of x_t over time, our model posits independent reliability-weighted accumulators for memory and vision that represent the unimodal evidence sampled from each source (Figure 2.2B, right):

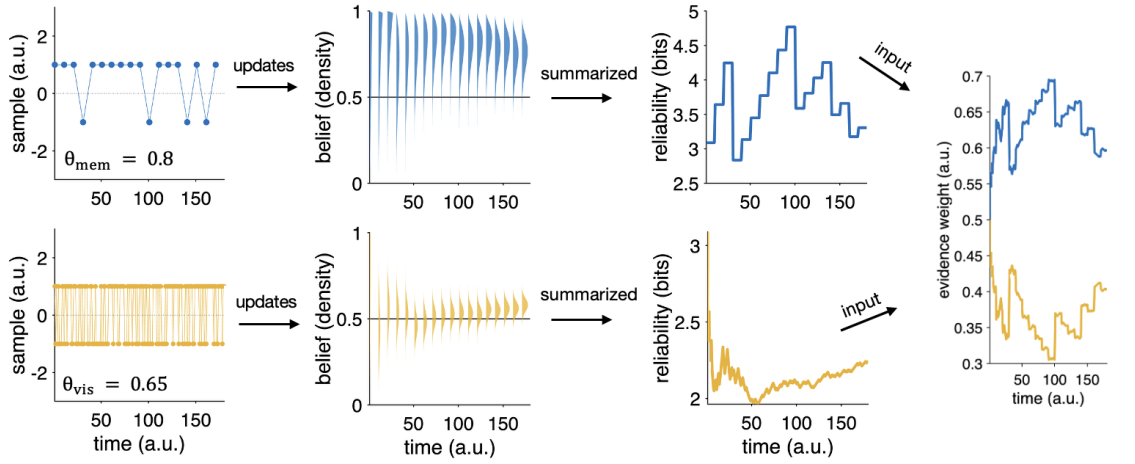
$$\begin{aligned} x_{memory_t} &= x_{memory_{t-1}} + o_{memory_t} w_{memory_t} \\ x_{vision_t} &= x_{vision_{t-1}} + o_{vision_t} w_{vision_t} \end{aligned}$$

Additionally, the model assumes that memory samples are generated at a fraction of the rate that visual samples are received by the observer, such that the total number of samples generated by each evidence source (memory and vision) is given by:

$$n_{memory} = n_{vision} \cdot \gamma^{-1} \tag{2.11}$$

where γ is a scalar that, in this model, we treat as a fixed, exogenously-set value. Unless otherwise specified, all findings in this article assume that γ scales the “rate” of memory evidence generation such that memory samples are delivered along the range of theta oscillatory frequencies that have been associated with task-induced memory retrieval in primates (2-12Hz; Jacobs, 2014; Kragel et al., 2020; Vivekananda et al., 2021; Seger et al., 2023). For the majority of simulations, this corresponded to $\gamma \in \{30, 20, 15, 12, 10, 8, 7, 6, 5\}$. In so doing, we demonstrate that our findings are robust to many possible settings of γ within a physiologically-observed range.

(A) Reliability estimation process



(B) Dynamic reliability-weighted integration process

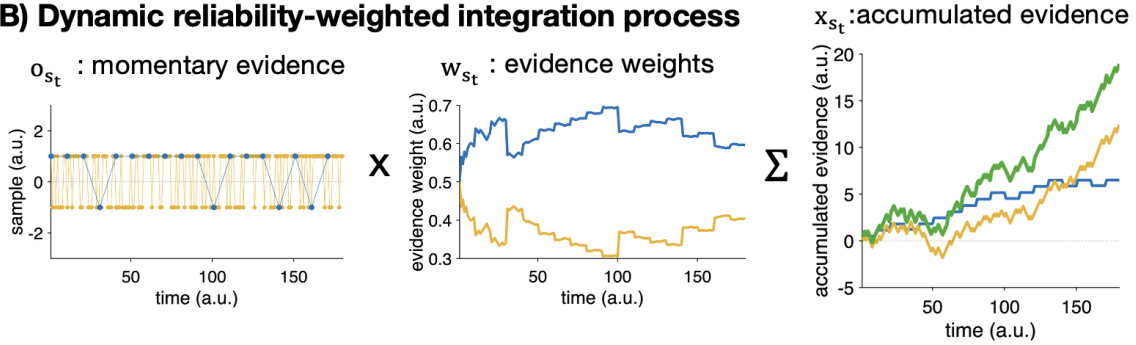


Figure 2.2: Graphical depiction of single-trial dynamics in our model. Quantities corresponding to memory are plotted in blue, quantities corresponding to vision are plotted in yellow, and quantities corresponding to the decision variable are plotted in green. Figures were generated using a vision sampling rate of 60Hz, a memory “sampling rate” of 6Hz ($\gamma = 10$), and an encoding noise ratio $\epsilon = 0.1$.

(A) Model components involved in the dynamic reliability estimation process specified by our model. The momentary evidence o_{s_t} is generated by a Bernoulli distribution with $p = \theta_s$, and some proportion of the encoded evidence is assumed to be corrupted by Bernoulli noise ϵ (Equation 2.4). The true value of θ_s defines how informative each evidence source s is with respect to the two-alternative choice. Each new sample o_{s_t} updates a Beta-distributed belief g_{s_t} about the value of θ . We assume that the observer’s internal estimate of each source’s reliability λ_{s_t} is equivalent to the inverse of the dot product of g_{s_t} with the information entropy function. The relative magnitude of each of these reliability estimates λ_{s_t} then defines how an observer weights each source of evidence at each moment in time (w_{s_t}). (B) Model components depicting the reliability-weighted evidence integration process. We assume that observers weight each evidence source by the dynamic relative reliability estimate computed above, and sum these samples over time to create a decision variable (green) that is a reliability-weighted blend of memory (blue) and sensory (yellow) evidence.

2.4.2 Steady-state model dynamics

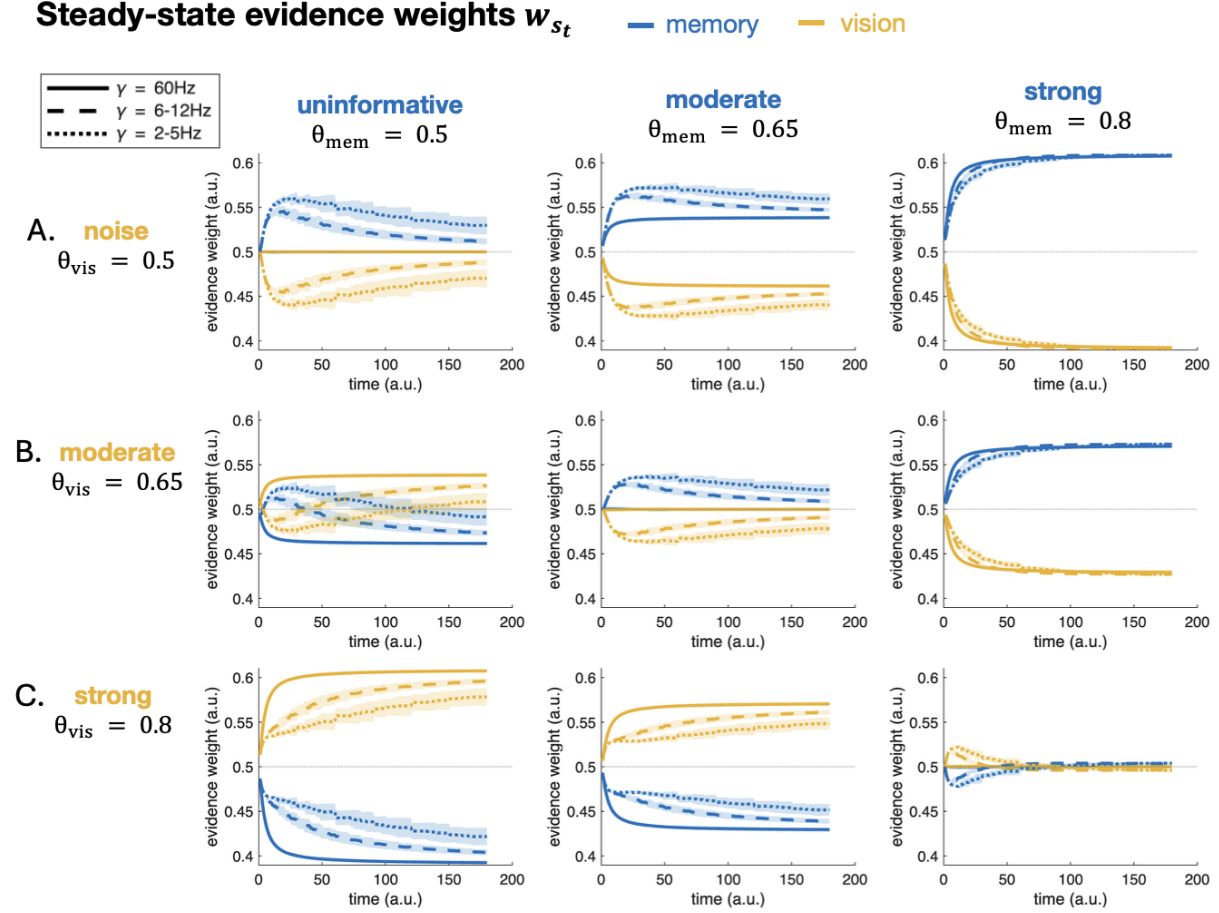


Figure 2.3: Steady-state dynamic weights on memory (blue) and sensory (yellow) evidence. The rate at which memory evidence enters the decision process is indicated by different line styles (solid = equal rate as vision, dashed = high theta frequencies, dotted = low theta frequencies). (A) Weights when sensory evidence is totally uninformative ($\theta_{vis} = 0.5$). (B) Weights when sensory evidence is moderately informative ($\theta_{vis} = 0.65$). (C) Weights when sensory evidence is strongly informative ($\theta_{vis} = 0.8$). Shaded regions correspond to standard error computed across all levels of ϵ and the subset of γ corresponding to each linetype.

To further elucidate the dynamics generated by our model, we display steady-state timeseries of the relative evidence weights generated across different combinations of memory and sensory evidence strength (Figure 2.3). Steady-state timeseries were obtained by averaging across 5,000 simulations of 100 decision trials for each combination of memory evidence strength θ_{mem} , sensory evidence strength θ_{vis} , encoding noise ϵ , and memory “sampling rate” γ . We simulated the process with $\theta_s \in \{0.5, 0.65, 0.8\}$, $\epsilon \in \{0.01, 0.06, 0.11, 0.15, 0.2\}$,

and $\gamma \in \{30, 20, 15, 12, 10, 8, 7, 6, 5\}$, and $z = \infty$. For simulations where $\theta_{mem} > 0.5$, we included only the subset of trials where memory and sensory evidence were in agreement about the correct decision (i.e., cues expectations were “valid” or “congruent” with sensory evidence).

Figure 2.3 demonstrates how the steady-state timeseries change across different combinations of θ_s and γ . Importantly, we display timeseries with $\gamma = 1$ —corresponding to identical memory and vision “sampling rates”—in solid lines to aid in interpreting the effects of γ on our process. These timeseries, along the diagonal of Figure 2.3 where $\theta_{vis} = \theta_{mem}$, validate that there ought to be equal weight on both evidence sources when they are equally informative *and* enter the decision process at the same rate. The dashed and dotted lines on the diagonal plots, however, show that slowing the rate of memory evidence sampling relative to vision generates a transient *upweighting* in one of the sources relative to the other. Specifically, for sufficiently weak sensory evidence ($\theta_{vis} < 0.8$), expectations are upweighted early on relative to vision, with the degree of upweighting scaling inversely with θ_{vis} (top left and middle panels of Figure 2.3). When sensory evidence is sufficiently strong, however, *vision* is transiently upweighted before they quickly collapse to being identical (bottom right panel of Figure 2.3). Further, the degree to which equally-informative expectations are upweighted relative to vision scales with how quickly memory samples become available for the observer: memory is more strongly upweighted for evidence generated at the low theta frequency (dotted lines) relative to the high theta frequency (dashed lines).

Panels *off* the diagonal in Figure 2.3 further display dynamics when memory and sensory evidence differ in their informativeness. The top row of Figure 2.3 shows that, when sensory evidence is completely uninformative, memory remains consistently upweighted throughout the decision process. Comparing across moderate (top row, middle) and strong (top row, right) expectations shows that both the magnitude and rate at which memory is upweighted scales as a function of its informativeness. The middle row of Figure 2.3 shows

that moderately strong sensory evidence elicits heterogeneous dynamics across differently-informative expectations. When expectations carry no information (middle row, left) *and* they become available at a slower rate than sensory information (dashed and dotted lines), they are again transiently upweighted relative to sensory information. However, as the reliability estimates λ_s converge on their steady-state values over the course of sampling, vision becomes *upweighted* relative to memory. And again, the *rate* with which this “crossover” occurs depends on how quickly memory evidence becomes available to the observer (dashed vs. dotted lines). When expectations are *more* informative than sensory evidence, however, memory is quickly and persistently upweighted relative to vision (right panel in Figure 2.3B). Finally, the bottom row of Figure 2.3 shows that when sensory evidence is sufficiently strong (θ_{vis}), it *always* ought to be weighted more than memory evidence that is uninformative (left panel of Figure 2.3C), moderately informative (middle panel of Figure 2.3C), or even equally informative but arriving at a slower rate (right panel of Figure 2.3, dotted and dashed lines).

Taken together, the dynamics in Figure 2.3 capture core theoretical ideas in our model. When an evidence source has continuously provided reliable information throughout a decision (i.e., when θ_s is large), that source should increasingly bias one’s decision over time. However, when an evidence source has continuously provided weak or unreliable information throughout a decision (i.e., when θ_s is small or null), the *alternative* evidence source should increasingly be relied on over time. That these effects are modulated by the relative *rate* at which samples become available to the observer further underscores the importance of considering the temporally-extended nature of memory retrieval, and how those dynamics shape the integration of uncertain expectations with sensory evidence that is itself uncertain.

2.5 Dynamic reliability-weighted integration unifies time-varying effects of expectations

This section simulates the model specified above under four different experimental conditions, each corresponding to a unique quadrant in Figure 2.1. In doing so, we demonstrate that the conceptual framework of uncertainty-based memory sampling—and the dynamic reliability-weighted model it inspired—offers a unifying explanation for time-varying effects of expectations on perceptual decisions. Crucially, each of the findings simulated below come from different research groups, experimental designs, and neuroimaging modalities (local field potentials, EEG, fMRI, and intracranial recordings), and were accordingly explained by different neurocomputational processes. The ability of our model to capture latent dynamics across this broad range of explanatory targets suggests that it identifies a core computational process driving the integration of prior knowledge with incoming sensory information.

All simulations were carried out in MATLAB 2023B and 2025A using scripts available on Github. Details about each simulation are contained in its corresponding subsection below.

2.5.1 Quadrant I: block-level instructed expectations

The first quadrant of Figure 2.1 corresponds to instances when we predict time-varying effects of expectations to be most subtle: expectations come from a source with high certainty (direct instruction) and there is no uncertainty about which expectation to use across decisions (block-level deployment). Intriguingly, however, one of the first and most influential observations of time-varying effects of expectations were reported from a task using exactly these settings (Hanks et al., 2011).

2.5.1.1 Study design & key findings

Hanks et al. (2011) measured behavior and local field potentials from humans and macaques performing a standard motion discrimination task. Humans were explicitly instructed about the block-level prior probabilities, and macaques were extensively trained until their behavior demonstrated successful learning of block-level reward probabilities, effectively equipping them with instructed expectations. Sensory evidence coherence, or signal strength, was allowed to vary across trials, generating an *evidence-heterogeneous* but *expectation-homogeneous* decision environment in each block. With this design, Hanks et al. (2011) reported behavioral, neural, and computational evidence indicating that observers nonlinearly increase their weighting of the prior probability Π as a function of elapsed decision time (Figure 2.4A, left). Further, this effect was greatest on trials where sensory evidence was maximally uncertain (Hanks et al., 2011).

2.5.1.2 Computational modeling

The authors explained these findings by theorizing that observers learn a *time-dependent accuracy* (TDA) function as they make decisions in evidence-heterogeneous environments. Learning such a function is behaviorally advantageous because it allows observers to optimize their speed-accuracy tradeoffs without having full knowledge of the statistics of the environment. Instead, they can leverage the systematic—and even idiosyncratic—association between sensory evidence strength and the amount of time needed to commit to a decision to *infer* signal strength A based on elapsed time (Hanks et al., 2011). A *dynamic bias signal* (DBS) can then be defined based on the ratio of the prior probability to time-dependent accuracy. Formally,

$$DBS(t) = \frac{\log(\frac{k \cdot \Pi}{1 - k \cdot \Pi})}{\log(\frac{TDA(t)}{1 - TDA(t)})} z(t) \quad (2.12)$$

where Π is the prior probability of a motion direction, z is the height of decision boundary, and k is a constant that scales Π to accommodate cases where primate observers misestimate the true prior probability. The time-dependent accuracy function (TDA) was computed based on observers’ behavior in the unbiased blocks, and then was used to fit behavior in the biased blocks (Hanks et al., 2011).

Although not explicitly described as such, Hanks et al.’s (2011) dynamic bias signal implements a dynamic reliability-weighted process very similar to the one we specify in Section 2.4.1. Rather than estimate evidence reliability directly, however, Hanks et al., 2011 suggest that evidence reliability is inferred based on elapsed decision time. The inference about evidence reliability is thus given by the association of longer deliberation time with lower coherence stimuli, which also generally are associated with a lower degree of choice accuracy. However, as we will show next, these same dynamics can be generated by our reliability-weighted process which does not assume or require that observers learn time-dependent accuracy functions.

2.5.1.3 Simulation methods

The goal of our simulations was to show that our relative reliability-weighted process generates the same “late effect” observed by Hanks et al. (2011): greater weighting of expectations *late* in the decision process, especially when sensory evidence is weakly informative. To do so, we generated timeseries data using $\theta_{vis} = 0.50$, $\theta_{mem} = 0.8$, $\epsilon = 0.11$, and $\gamma = 10$ with a visual presentation rate of 60 Hz (effective memory “sampling rate” = 6Hz). These parameter settings correspond to the experimental conditions where the late effect was most prominent (Hanks et al., 2011). We discretized the continuous-time simulation such that each timestep corresponded to one “frame” of presented visual evidence; each trial’s three second simulation thus consisted of $n_{vision} = 180$ and $n_{memory} = 18$. Importantly, our explanation for the effects observed by Hanks et al. (2011) (Figure 2.4) are independent of threshold z , which was set at an arbitrary value of $z = 250$.

2.5.1.4 Simulation results

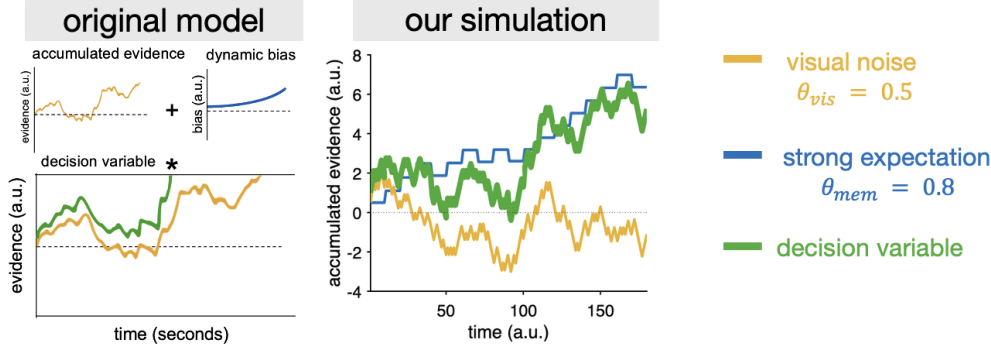
Figure 2.4A demonstrates that, at a single-trial level, our model (right) produces the same dynamics as Hanks et al.’s (2011) dynamic bias model (left). Crucially, however, our model suggests that the “dynamic bias signal” is comprised of an independent *memory accumulator* that samples and integrates evidence generated by memory in parallel to the sensory evidence presented by the observer.

Figure 2.4B displays the corresponding model components for this representative trial. The model components generating these dynamics are displayed in (Figure 2.4B). Just as each sample of visual evidence o_{vision_t} contributes to the observer’s estimate of visual evidence reliability g_{vision_t} , so too does each “sample” of memory evidence o_{memory_t} update the observer’s estimate of memory evidence reliability g_{memory_t} . Because the signal strength of vision is so low relative to the evidence provided by memory, the expectation is continuously estimated as being more reliable (or informative) throughout the trial, leading to consistently higher weights on memory than sensory evidence (Figure 2.4B, rightmost plots). When these dynamic, relative-reliability defined weights are used to integrate the sensory and memory evidence over time, they generate the timeseries plotted in Figure 2.4A (right), where the decision variable becomes increasingly biased toward the expectation over the course of a single decision. Further, the rightmost panel in Figure 2.3A further demonstrates that the increased weighting over time is present in the steady-state dynamics of our model averaged across several plausible values of γ and z .

Taken together, these simulation findings demonstrate that the relative reliability-weighted process we specified to account for the role of uncertainty about expectations can capture a canonical time-varying effect of expectations: increased weighting as a function of elapsed decision time (Hanks et al., 2011). Crucially, these effects were observed in experimental settings where we predict uncertainty-driven time-varying effects to be the most subtle (Quadrant I of Figure 2.1), and thus serve as a strong test of the plausibility of the frame-

work and model that we propose. The following simulation findings further demonstrate the ability of our model to capture time-varying effects observed in even more uncertain circumstances.

(A) Single-trial dynamics capture Hanks et al. (2011) dynamic bias effect



(B) Single-trial model components

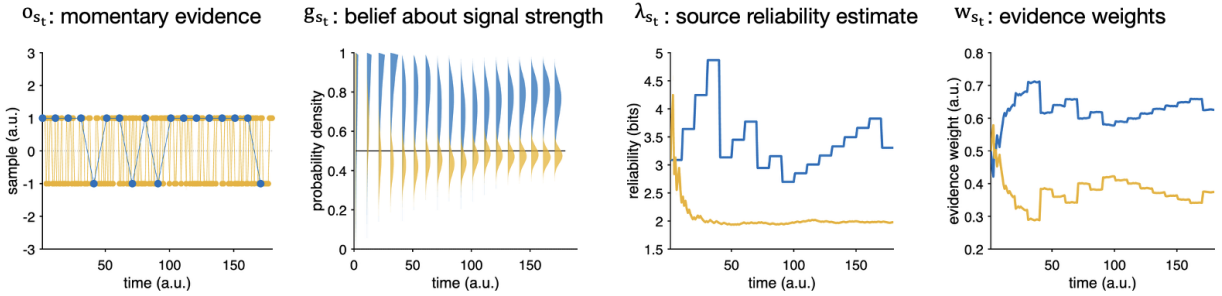


Figure 2.4: Dynamic reliability-weighted integration captures Hanks et al. (2011) dynamic effect of expectations. In all figures, memory evidence (expectations; Π) is plotted in blue, vision evidence is plotted in yellow, and the decision variable is plotted in green. (A) On the left, a stylized reproduction of the original model used by Hanks et al. (2011) to explain their “late effect” of expectations. On the right, a representative trial from our model demonstrating the same single-trial dynamics. (B) Model components involved in generating our representative trial above.

2.5.2 Quadrant II: block-level learned expectations

An active debate in the perceptual decision-making literature is whether expectations can modulate sensory responses themselves (Kok et al., 2012; Kok et al., 2013; Afacan-Seref et al., 2018) or only “post-sensory” decision processes (Summerfield and de Lange, 2014; Rungratsameetaweemana and Serences, 2019). Specifically, it has been argued that trial-level

expectation cues can elicit *attentional* effects on sensory responses that are conceptually distinct from those driven by prior knowledge of stimulus probabilities (Summerfield and de Lange, 2014). To address this question, Rungratsameetaweemana, Itthipuripat et al. (2018) designed an experiment where observers had to learn block-level expectations *during* a perceptual decision task. This in-context learning of prior probabilities bypasses the concern about attentional effects elicited by cue-based manipulations of expectations, such that any effects on sensory responses must be attributed to learned expectations about prior probabilities.

2.5.2.1 Study design & key findings

Rungratsameetaweemana, Itthipuripat et al. (2018) report behavioral and EEG data from $n=17$ human subjects performing the task. Sensory stimuli were comprised of small red and blue bars whose locations and orientations were refreshed at slow (33.33 Hz) and fast (50 Hz) refresh rates. Across all decisions, the signal-to-noise ratio of sensory evidence was fixed at 0.685. The authors independently manipulated expectations about color, orientation, and motor response across blocks comprised of 60 trials each, but fixed the magnitude of the prior probability to $\Pi = 0.7$ across all blocks. This means that in biased blocks, $n = 42$ trials were presented with the dominant stimulus feature whereas $n = 22$ trials were presented with the other stimulus feature. Crucially, neutral blocks presented both stimulus features equally, and subjects completed 4 blocks of 60 trials under neutral expectations as well.

With this design, the authors showed that expectations can exert robust effects on choice behavior in the absence of modulations to sensory processing as measured by EEG. Specifically, Rungratsameetaweemana, Itthipuripat et al. (2018) showed that expectations had no effect on the early visual negative potential, an event-related EEG component that indexes sensory information processing. However, the authors did observe effects of expectations on the centroparietal positivity (CPP), a different event-related EEG component that is strongly associated with the decision variable in accumulator models (O’Connell et al., 2018), as well

as effects of expectations on the amplitude of parietal alpha and frontal theta measurements. Crucially, all of these effects occurred *after* each component reached its peak value for the decision, *and* were driven by trials where the sensory evidence was *unexpected* (or *incongruent* with the learned expectation for that block). In what follows, we show that our dynamic reliability-weighted integration process can capture the dynamics observed in frontal theta amplitude.

2.5.2.2 Simulation methods

Simulations were run with $\theta_{vis} = 0.685$, $\theta_{mem} = 0.7$ for biased blocks, and $\theta_{mem} = 0.5$ for neutral blocks. Although the signal-to-noise ratio of sensory evidence was held constant across trials, the authors manipulated the rate of sensory evidence presentation randomly across trials. Accordingly, we ran simulations with both the low (33.33 Hz) and high (50 Hz) flicker rates used in the original study. This simulation’s aim of capturing neural activity in the theta frequency further aided in constraining the range of γ values used in the low and high flicker rate conditions. We defined the range of γ values such that memory evidence was generated across the range of theta activity observed in humans (1.5-8Hz; Rungratsameetaweemana, Itthipuripat et al., 2018; Jacobs, 2014). Simulations with low flicker rates used $\gamma \in \{4, 6, 8, 10, 12, 20\}$ corresponding to activity at 8.25, 5.5, 4.125, 3.3, 2.5, and 1.5Hz, and simulations with high flicker rates used $\gamma \in \{4, 6, 8, 10, 12, 20\}$ to generate memory samples at 8.3, 6.25, 4.16, 3.5, 2.5, and 1.5Hz, respectively.

With these free parameters of the model constrained by the original study, we then defined the range of threshold values $z \in \{3, 4, 5\}$ based on the similarity between trial-averaged behavioral data produced by the model in simulation and by human observers in the experiment. Importantly, this trial-averaged behavior was generated by simulated $n = 18$ “subjects”, each with a unique combination of memory sampling rate γ and threshold z . Each “subject” contributed the same number of trials as completed by human subjects in the experiment: 360 biased trials and 120 neutral trials at each flicker rate (720 biased and

240 neutral total). We additionally varied evidence encoding noise ϵ at the whole-sample level, such that each run used $\epsilon \in \{0.01, 0.06, 0.11, 0.15, 0.2\}$. Figure 2.5 displays results generated with $\epsilon = 0.11$, and Figure A.2A displays results averaged across all values of ϵ . As in the original experiment, the first 1000ms of each trial were comprised of sensory noise ($\theta_{vis} = 0.5$). Then, sensory evidence was presented for 850ms at either the low or high flicker rate, before being replaced by sensory noise for 600ms (Rungratsameetaweemana, Itthipuripat et al., 2018). Each timestep in the simulation corresponded to one frame of evidence, such that slow flicker trials had $n = 82$ timesteps and fast flicker trials had $n = 124$ timesteps.

The timeseries plotted in Figure 2.5A (right) used a normalization procedure to align memory weights generated in the low and high flicker simulations onto a common numerical space. For each subject, weights were averaged across trials in each of the expectation x flicker rate conditions, and then interpolated onto a common grid comprised of $n=100$ timesteps. Once normalized, weights were then averaged across flicker levels within each participant, and then averaged across participants to obtain condition-specific group means and standard errors. Group-averaged timeseries for each of the low and high flicker conditions in their native numeric spaces are displayed in Figure A.2B.

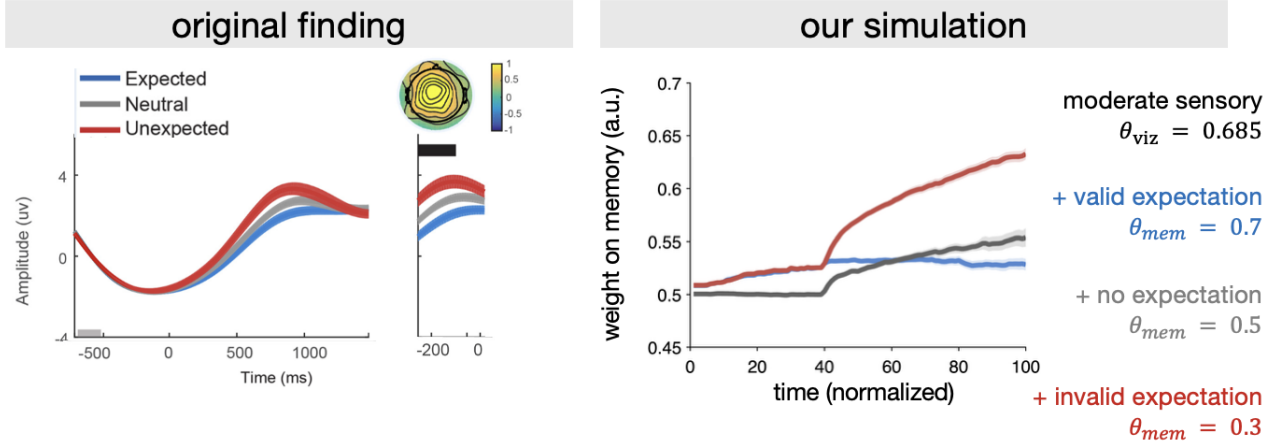
2.5.2.3 Simulation results

Figure 2.5A shows that our model’s time-evolving weights on memory capture the “late” effect that Rungratsameetaweemana, Itthipuripat et al. (2018) observed on frontal theta amplitude. After the initial period of sensory noise, weights on memory slowly increase over time in all three expectation conditions. However, the rate of increase is largest on trials when the expectation is *invalid* with respect to sensory evidence (red line), relative to trials where the expectation carries no information about sensory evidence (gray line) or valid/congruent information about sensory evidence (blue line).

This effect emerges naturally from positing that observers use the same time-constant boundary to terminate their decisions across different experimental conditions (low/high flicker, expected/neutral/unexpected). Because evidence weights in our model are estimated dynamically over the course of sampling, trials where observers take longer to make a choice naturally generate weights on memory that continue to increase relative to trials where observers terminate their decisions earlier. Crucially, this only holds if the expectation magnitude θ_{mem} and sensory evidence θ_{vis} are held constant across trials, which is precisely the case when comparing valid and invalid trials. Neutral trials appropriately generate weights whose magnitudes are intermediate between valid and invalid. However, because there is no true information in the neutral expectation signal, its timeseries is quantitatively indistinguishable from valid expectations that don't need to be sampled because a choice has already been made.

Figure 2.5B additionally shows that the ranges of γ and z used in our simulation capture the key behavioral effects, validating the link between our model's latent processes and the EEG measurements in the original study. Crucially, in doing so, we demonstrate that our model can reproduce both behavior and neural dynamics in environments corresponding to Quadrant II in Figure 2.1: learned expectations manipulated on the level of blocks. Our findings further raise an alternative interpretation of the frontal theta effects observed by Rungratsameetaweemana, Itthipuripat et al. (2018) : rather than reflecting executive control processes, the late increase in theta amplitude may reflect prolonged sampling of learned expectations on unexpected trials relative to neutral and expected.

(A) Trial-averaged weights on memory capture Rungratsameetaweemana, Itthipuripat et al. (2018) “late” effect on frontal theta amplitude



(B) Simulation parameters capture behavioral effects

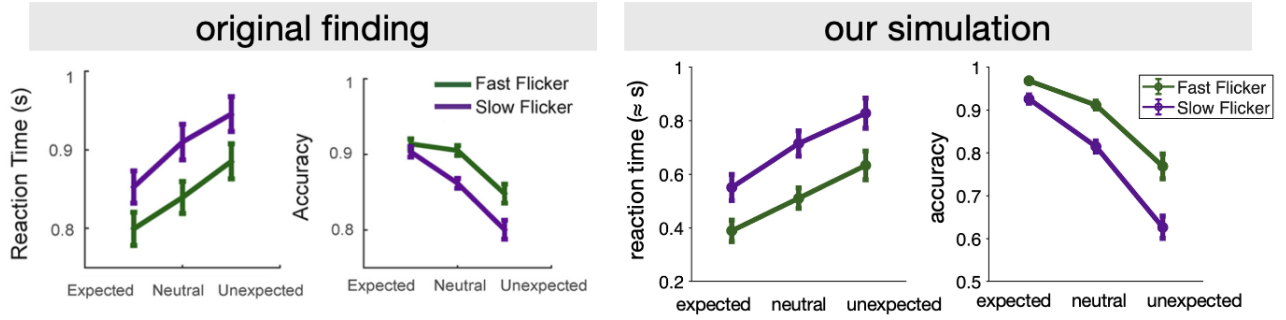


Figure 2.5: Dynamic reliability-weighted integration captures time-varying effect of expectations observed by Rungratsameetaweemana, Itthipuripat et al. (2018). (A) Original finding, left: violations of expectations increase frontal theta amplitude at later timepoints in the trial. Our simulation, right: weights on memory—interpretable as an incentive to sample evidence from memory—are greater for unexpected stimuli (red) relative to neutral (gray) and expected (blue). Crucially, the effect increases over time and is greatest at the end of the decision. Shaded areas reflect standard error across 15 simulated subjects. (B) Original finding, left: flicker rate and expectation manipulations induce canonical modulations of accumulation-based choice behavior. Slow flicker trials (purple), where fewer samples of visual evidence are presented per unit time, lead to slower reaction times and lower choice accuracy relative to fast flicker trials (green). Our simulation, right: choices and reaction times generated from our model replicate these behavioral effects, validating our choices for sensory signal strength θ_{vis} and threshold z . Errorbars in both (A) and (B) reflect standard error across $n = 18$ simulated “subjects”, each of which had a unique value of γ and z .

2.5.3 Interim summary

Taken together, Sections 2.5.1 and 2.5.2 show that dynamic, relative reliability-weighted integration of independent memory and sensory accumulators can reproduce dynamic effects of expectations observed in multi-modal, cross-species neural measurements. Both of these simulations focused on dynamics observed in *expectation-homogeneous* environments, where observers have no uncertainty about which expectation to use across decisions. The key factor differentiating the two simulations is whether those expectations were acquired through instruction (Section 2.5.1) or learned through experience with the decision environment (Section 2.5.2). Both types of block-level expectations exerted time-varying effects that occurred *late* in the decision process: with the Hanks et al. (2011) increasing nonlinearly over elapsed time, and the Rungratsameetaweemana, Itthipuripat et al. (2018) effect occurring after the peak of the ERP signal on each trial. We turn next to demonstrating that our model can capture existing findings corresponding to the remaining two quadrants in Figure 2.1, where observers have uncertainty about *which* expectation to use for each decision.

2.5.4 Quadrant III: trial-level learned expectations

On our framework, experimental settings corresponding to Quadrant III are those where we expect time-varying effects of expectations to be most pronounced: there is source-related uncertainty about *what* the true value of the expectation is, as well as deployment-related uncertainty about *which* expectation will be used on every trial. A recent study by Bornstein et al. (2023) measured human behavior and neural activity in precisely these settings and found evidence for adaptive, uncertainty-mediated time-varying effects of expectations that occur *early* in the decision process, before any sensory evidence has been presented. We expand upon this finding below, and show that our model offers the first formal explanation for it.

2.5.4.1 Study design & key findings

Bornstein et al. (2023) utilized a two-phase task design, where observers first learn probabilistic cue-stimulus relationships in a learning phase, and then make cue-guided perceptual decisions in a decision-making phase. The expectation cues consisted of colored fractals, and the perceptual stimuli were grayscale scene and face images. The decision task consisted of two-alternative discrimination between two images from the same category, e.g., reporting whether the dominant stimulus on each trial was Face A or Face B. Accordingly, observers learned that each fractal makes a unique prediction about which image is more likely. Crucially, images from the same category (faces vs. scenes) were equally predicted by their corresponding fractal, such that, for example, faces were strongly predicted by their fractals ($\Pi_{cue} = 0.8$) whereas scenes were weakly predicted by their fractals ($\Pi_{cue} = 0.6$). These learning conditions thus created both item- and category-specific predictions, such that, for example, the correct answer for face decisions was more predictable than the correct answer for scene decisions. Finally, during learning, participants only had a small amount of exposure to each cue-stimulus pairing (~ 30 trials) to ensure that there was learning-related uncertainty around each expectation.

Sensory evidence during the decision-making phase was presented at an effective rate of 30Hz: frames were updated at 60Hz, but each frame of signal was forward- and backward-masked by noise (phase-scrambled superpositions of the to-be-discriminated images). The difficulty of the decision was manipulated by varying the proportion of frames containing the target image, such that trials with closer to equal frames of target and lure images were more difficult to resolve. Each trial in the decision phase began with a brief (750ms) presentation of a fractal image followed by an interstimulus interval that varied uniformly across 4000, 6000, and 8000ms. Once the ISI elapsed, sensory evidence appeared on screen and observers had up to 3000ms to make their response. Crucially, the authors yoked decision difficulty to the category membership of the sensory evidence; e.g., in a given block, face evidence

would always have a low signal-to-noise ratio whereas scene evidence would always have a high signal-to-noise ratio. Combined with the category-level expectations created during learning, this manipulation created *joint* about sensory evidence on each trial: a probabilistic expectation about the correct answer (defined by Π_{cue}) and a *deterministic* expectation about the difficulty of the upcoming decision (learned through experience in the decision phase).

Because each cue gives observers deterministic information about the relative signal strength of the upcoming decision, this design allowed Bornstein et al. (2023) to identify behavioral effects in line with predictions of normative starting point models. Specifically, on as many as a third of trials with a high Π_{cue} that also predicted a low upcoming signal strength A , observers opted to make a response *before* observing any sensory evidence at all. This prediction follows directly from normative starting point models, which show that as Π approaches 1 and decisions approach maximum difficulty, the optimal procedure is to respond immediately on the basis of the expectation (Simen et al., 2009).

An important goal of Bornstein et al.’s (2023) study, however, was to show that the start-point setting process is itself one of evidence accumulation. To do this, they computed a metric termed the *reinstatement index*, the integrated evidence of stimulus reinstatement over the cue period, measured at each timepoint as the correlation between multi-voxel patterns measured when subjects directly perceived each target image with those observed when subjects engaged in putative cue-based retrieval of that target. Reinstatement indices were computed for neural activity in two regions of interest only—the fusiform face area (FFA) and the parahippocampal place area (PPA)—whose activity is known to be specialized for discriminating similar faces and scenes, respectively. Importantly, the authors theorized that the reinstatement index should be *inversely* related to associative memory strength: weaker associations require more sampling from memory, whereas stronger associations can be retrieved after a smaller number of memory samples (Bornstein et al., 2023; Bornstein and Daw, 2013). Consistent with this hypothesis, Bornstein et al. (2023) found that the

reinstatement index was lower on trials with a stronger predictive probability and increased as a function of time between onset of cue and response time. Interestingly, the authors also found that the reinstatement index was greater on trials with a long cue-stimulus ISI *only* when a cue signaled weak upcoming sensory evidence. The authors offered no formal explanation for this phenomenon, which we term *uncertainty-adaptive anticipatory memory sampling*.

2.5.4.2 Computational modeling

The authors did, however, test variants of a diffusion model to investigate their hypothesis that expectation-setting proceeds via evidence accumulation. Their core model (multi-stage DDM; MSDDM) was comprised of *two* diffusion processes with independent drift rates, one corresponding to an expectation-setting phase and the other corresponding to the sensory accumulation phase. Importantly, the terminal value of the decision variable in the first process sets the starting point for the sensory accumulation process on a trial-by-trial basis, effectively positing a two-stage process whereby memory dynamically biases sensory evidence accumulation on a trial-by-trial basis (Bornstein et al., 2023; Srivastava et al., 2017; (Figure 2.6A, left)). The model assumes one common threshold for the whole process, and was tested against two other models that selectively disabled key features of the MSDDM: a standard DDM with a constant drift rate across trials (1DDM) and two different DDMs that were fit independently to the expectation-setting and sensory-accumulation phases of each trial (2DDM), critically *without* information passing from the first stage to the next. These comparison models thus allowed for testing the necessity of (1) a time-varying drift rate and (2) a direct link between expectation-setting and sensory-accumulation processes, respectively.

To link the reinstatement index with predictions of the MSDDM, Bornstein et al. (2023) quantified the relationship between the reinstatement index and response times for choices made after the onset of sensory information. First, the authors found that the reinstatement

index predicted faster response times even when controlling for cue strength or evidence coherence, indicating that this metric captures meaningful trial-level variability in choice behavior. Next, Bornstein et al. (2023) found that this speeding was unique to trials where the cue validly predicted the choice outcome, consistent with a hypothesis that the reinstatement index corresponds to a start-point setting process. Finally, the authors found that the effects of reinstatement index were strongest for valid trials with weak sensory evidence. Taken together, these findings strongly support the idea that humans use their expectations in an uncertainty-adaptive manner: drawing more strongly on internal evidence when sensory evidence is insufficient to make the decision, but not bothering to do the extra effort of memory retrieval when decisions are likely to be easy to resolve with sensory evidence alone.

2.5.4.3 Simulation methods

We used both trial-level and steady-state simulations to demonstrate our model’s ability to capture the uncertainty-adaptive anticipatory sampling reported by Bornstein et al. (2023). Sensory evidence was generated at 60 Hz, and memory evidence was generated at frequencies ranging from 2-12 Hz ($\gamma \in \{30, 20, 15, 12, 10, 8, 7, 6, 5\}$ and with evidence encoding noise $\epsilon \in \{0.01, 0.06, 0.11, 0.15, 0.2\}$. We fixed the threshold to an arbitrarily large value of $z = 250$ because, in a fixed-period task, the dynamics we aimed to capture are largely independent of threshold. The authors tested expectation levels $\Pi_{cue} = \{0.5, 0.6, 0.7, 0.8\}$ and sensory evidence levels individually calibrated to two separate strengths corresponding to low (65%) and high (85%) decisional accuracy in the absence of any expectation (Bornstein et al., 2023). We focus on capturing the two dynamics that comprise the uncertainty-adaptive anticipatory sampling finding: a strong expectation predicting a low accuracy (i.e., weak evidence) trial, and a neutral expectation predicting a high accuracy (i.e., stronger evidence) trial. The corresponding steady-state timeseries were generated using the same process described in Section 2.4.2. As in Bornstein et al. (2023), we fixed the duration of sensory evidence presentation to 3000ms and ran separate simulations for each of the short, medium, and

long anticipation durations (4000, 6000, and 8000ms) used in the study. Main text findings report simulations from the shortest cue-stimulus interval (4000ms); steady-state dynamics for the longer two intervals can be found in Supplemental Figure A.3.

Additionally, to capture the joint nature of cues in Bornstein et al. (2023), we added *anticipated* coherence to the relative reliability computations occurring during the cue-stimulus interval. Rather than set sensory evidence weights to zero, we assume that each new sample of evidence from memory gives observers information how reliable the upcoming sensory signal will be. Specifically, for each timestep where a memory sample is drawn before the onset of sensory evidence, simulated agents receive one noisy sample from a Bernoulli distribution where θ_{vis} is defined by the sensory signal reliability on that trial. Importantly, this sample is used only to update the belief about upcoming signal reliability $g(t)_{viz}$, the time-varying reliability of each evidence source $\lambda(t)_s$, and the time-varying weights on evidence $w(t)_s$ defined by the relative values of λ_s at each point in time; it does not enter into either the sensory accumulator or the decision variable.

We do, however, assume that the decision variable is continuously updated with each new piece of information from memory during the period of time between the onset of the cue and the onset of sensory evidence. Specifically, we update the decision variable as follows:

$$x(t) = \begin{cases} x_{t-1} + w_{memory_t} * o_{memory_t} + w_{vision_t} * \text{Bernoulli}(0.5) & \text{if } t < t_{stim} \\ x_{t-1} + w_{memory_t} * o_{memory_t} + w_{vision_t} * o_{vision_t} & \text{if } t \geq t_{stim} \end{cases} \quad (2.13)$$

where t_{stim} denotes the timepoint at which sensory evidence becomes available.

2.5.4.4 Simulation results

Figure 2.6A shows that single-trial dynamics of our model (right) capture key properties of Bornstein et al.’s (2023) MSDDM (left). First, both panels in Figure 2.6A (right) show that

our model creates a process where the memory accumulator (blue lines) sets the starting point for the vision accumulator (yellow lines) on a trial-by-trial basis. When a strong expectation about choice outcome *also* predicts weak upcoming sensory evidence (first panel, right side of Figure 2.6A), the decision variable (green line) is increasingly biased toward the choice predicted by the expectation as a function of elapsed time between cue onset (vertical blue line) and sensory evidence onset (yellow vertical line). However, when an expectation carries no information about likely choice outcome, but signals strong upcoming sensory evidence (second panel, right side of Figure 2.6A), the decision variable is much less biased by memory during the interval between cue and sensory evidence.

Figure 2.6B displays the corresponding model components for these single-trial dynamics. When memory evidence is consistently more reliable than sensory evidence (top row), the model continuously places more weight on memory than vision throughout the trial (last panel, top row). However, when memory is uninformative, the model switches to placing more weight on sensory evidence *during* the anticipation period (last panel, bottom row). If we consider weights on each evidence source as corresponding to observers’ incentive to sample information from that source (as we did in Section 2.5.2), then these dynamics can be interpreted as reproducing Bornstein et al.’s (2023) observation uncertainty-adaptive anticipatory sampling effect on the single-trial level.

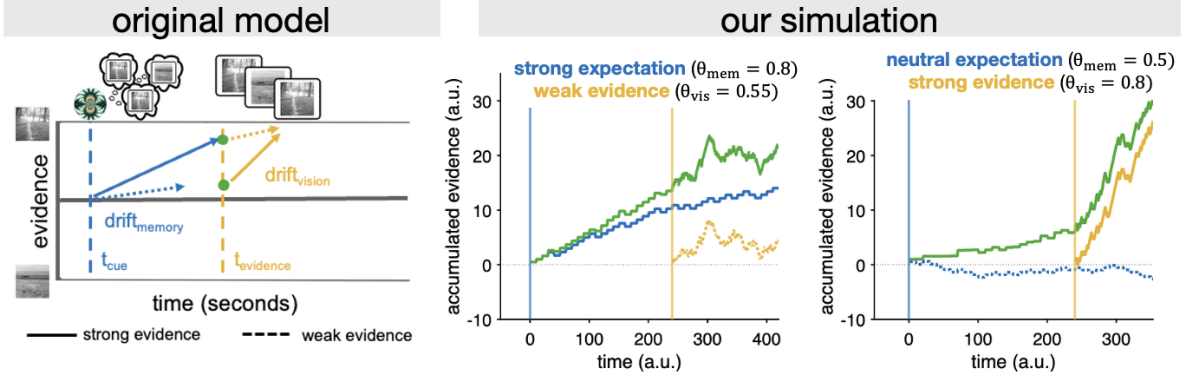
Figure 2.7 shows that these dynamics are also present in steady-state model simulations. Specifically, comparing the solid and dashed lines in the right panel of Figure 2.7A displays the uncertainty-adaptive sampling effect: weights on informative expectations increase with elapsed time between cue and stimulus onset, and the magnitude of these weights is greater when sensory evidence is anticipated to be highly noisy (dashed line). The steady-state evidence weights for the inverse case of an uninformative expectation paired with strong sensory evidence (left panel of 2.7A) are interestingly non-monotonic. In the case where a cue makes no prediction about the likely choice outcome, but signals that upcoming sensory

evidence will be strong (solid lines), increasingly more weight is placed on vision relative to memory during the cue-stimulus interval. However, once the sensory evidence becomes available (vertical line), the model slightly and briefly *downweights* vision for immediately after evidence onset, before returning to the same dynamics of increasing the weight on vision relative to memory. This transient downweighting is also observed—albeit to a lesser degree—for longer cue-stimulus interval durations (Figure A.3).

The steady-state source reliability estimates displayed in Figure 2.7B elucidate the time-evolving quantities driving these different dynamics in evidence weights. In the case of a weak cue predicting strong sensory evidence (solid lines, left panel), vision is consistently estimated as being more reliable than memory throughout the trial. Crucially, however, these reliability estimates are formed on the basis of an *expected* value about sensory evidence strength that is estimated at the frequency of the memory sampling process. These slower evidence generation dynamics lead to inflated source reliability estimates, as discussed in Section 2.4.2 and found in other work where small numbers of memory samples drive behavior (Bornstein and Daw, 2013; Vul et al., 2014). Once sensory evidence actually comes online and presents the observer with far more samples per unit time, the reliability estimate λ_{vis} quickly reaches its asymptotic value, which is less than the original estimates based on memory samples. This same principle explains the inflated reliability estimate λ_{mem} for $\theta_{mem} = 0.8$ relative to λ_{vis} for $\theta_{vis} = 0.8$.

Taken together, both the single-trial and steady-state simulations in this section demonstrate that our model can capture the core time-varying effect observed by Bornstein et al. (2023): uncertainty-adaptive anticipatory sampling of learned expectations on a trial-by-trial basis. Additionally, we showed that our parallel process generates the same trial-level start-point setting specified by Bornstein et al.’s 2023 MSDDM, suggesting that our model might be a more general case of the two-stage process they identify.

(A) Single-trial dynamics capture Bornstein et al. (2023) effects



(B) Single-trial model components

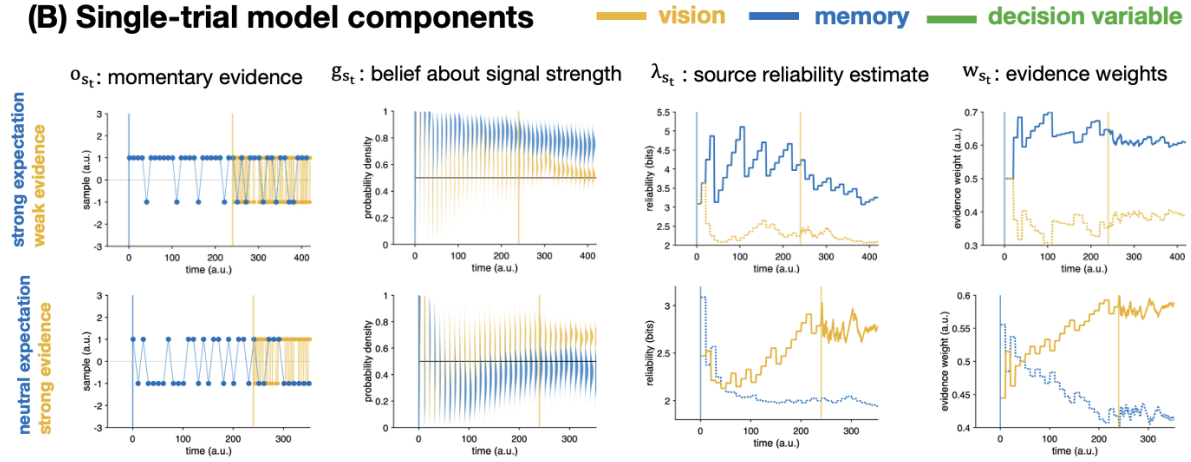
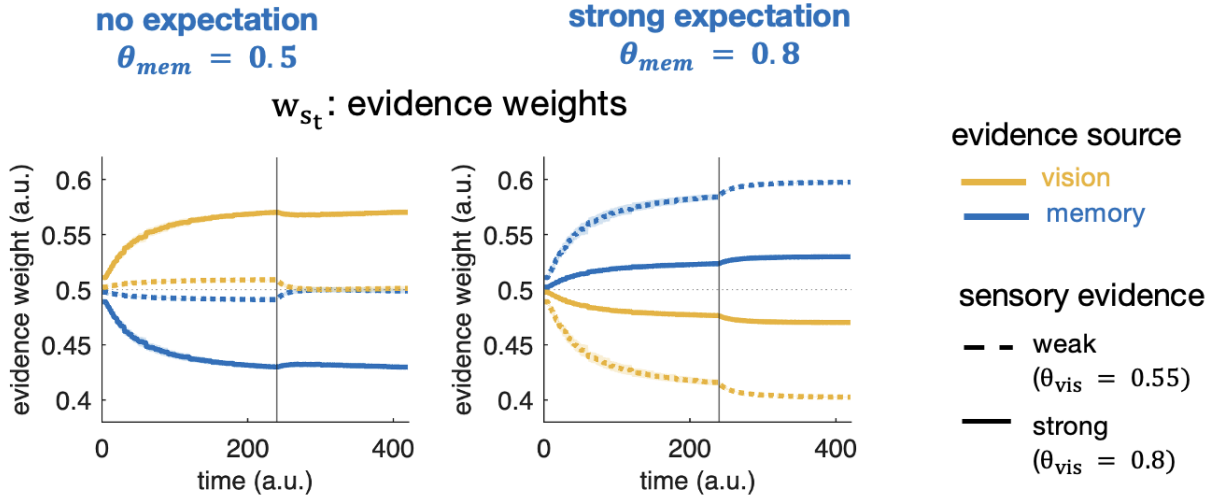


Figure 2.6: Dynamic reliability-weighted integration captures time-varying effects from Bornstein et al. (2023). In all figures, memory evidence (expectations; Π) is plotted in blue, vision evidence is plotted in yellow, and the decision variable is plotted in green. Vertical lines correspond to timepoints when memory (blue) and sensory (yellow) evidence come “online” for the observer. Memory samples in both simulations were generated at a rate of 6Hz ($\gamma = 10$). (A) Our model captures key dynamics produced by Bornstein et al.’s (2023) multi-stage DDM (left). First, we show that expectations set the starting point for sensory evidence accumulation in a trial-by-trial manner (right, both figures). Then, we show that anticipatory sampling from a strong expectation expedites the decision process once weak sensory evidence becomes available (left panel under “our simulation”). Additionally, we show that sampling from memory continues throughout the anticipation period when an expectation carries no information about the likely choice outcome, such that choice dynamics are dominated by the strong sensory evidence once it comes online (right panel under “our simulation”). (B) Single-trial model components that generated the left and right panels for the “our simulation” portion of A. Components for the left panel—strong expectation with weak evidence—are plotted in the top row, and components for the right panel—weak expectation with strong evidence—are plotted in the bottom row.

(A) Steady-state evidence weights capture uncertainty-adaptive anticipatory sampling in Bornstein et al. (2023)



(B) Steady-state source reliability estimates

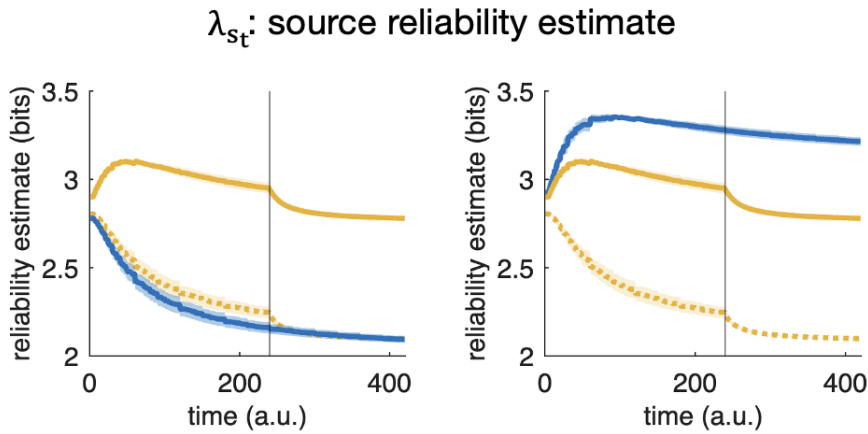


Figure 2.7: Steady-state dynamics of evidence weights and source reliability estimates driving Bornstein et al.’s (2023) observation of uncertainty-adaptive anticipatory memory sampling. Quantities corresponding to memory are plotted in blue and quantities corresponding to vision are plotted in yellow. Solid lines correspond to simulations with a weak sensory signal ($\theta_{vis} = 0.55$), dashed lines correspond to simulations with a strong sensory signal ($\theta_{vis} = 0.8$). The vertical line indicates the onset of sensory evidence after a 4000ms anticipation period; dynamics for longer anticipation durations are plotted in Figure A.3. (A) Weights on each evidence source throughout a single decision. (B) Source reliability estimates throughout a single decision.

2.5.5 Quadrant IV: trial-level instructed expectations

A recent study by Hachisuka, Shor, Liu et al. (2026) measured human behavior and neural activity in experimental settings that correspond to the fourth and final quadrant in Figure 2.1: an expectation-heterogeneous environment where observers have uncertainty about *which* expectation to use for each consecutive choice, but little to no uncertainty about *what* each expectation predicts.

2.5.5.1 Study design & key findings

Hachisuka, Shor, Liu et al. (2026) combined high resolution fMRI, intracranial recordings, and deep neural network modeling to investigate the neural and computational mechanisms of one-shot perceptual learning in humans. To do this, they used “Mooney images”: photographic stimuli that initially appear ambiguous or abstract, but are actually just degraded versions of images that are instantly recognizable when presented without degradation (examples are shown in Figure 2.8A). Crucially, the apparent ambiguity of Mooney images disappears after just one exposure to the full-resolution image, such that observers cannot “unsee” the true contents of degraded image after being exposed to its full-resolution counterpart. This property of Mooney images makes them well-suited for studying the neurocomputational mechanisms involved in one-shot perceptual learning.

Hachisuka, Shor, Liu et al. (2026) developed a task where neural activity is recorded upon initially viewing a Mooney image (“pre-Mooney”), then viewing the full-resolution corresponding image in grayscale (“gray”), and then finally viewing the Mooney images again after having seen their full-resolution counterparts (“post-Mooney”; Figure 2.8A, right). On our framework, this corresponds to an initial “no expectation” block (blue border, Figure 2.8A), followed by a “one-shot learning” block (gray border, Figure 2.8A), and then finally a “test” block (orange border, Figure 2.8A). In each block, trials began with a fixation cross presented for 1000-2000ms followed by presentation of a Mooney or grayscale image for

2000ms (Figure 2.8A, left). Patients with intracranial electrodes were prompted to answer the question “Can you name the object hidden in the image?” with Yes/No keyboard presses made using different hands, whereas non-patients were prompted to verbally respond to the question “What did you see?”. Crucially, the test block following grayscale image presentation contained Mooney images observed in the first block (now *post*-Mooney) as well as brand new “pre”-Mooney images that had not been previously presented to the observer. Thus, the “post” images in the test block are those for which observers have an unambiguous expectation, whereas there is no such expectation for the “pre” images.

The authors report cross-modal evidence suggesting that one-shot perceptual expectations are encoded in higher-level visual cortex (Hachisuka, Shor, Liu et al., 2026). The main finding of interest for the present analysis, however, comes from intracranial recordings from the frontoparietal network (FPN), which encompasses several key regions implicated in perceptual decision-making (Hanks and Summerfield, 2017; Polanía et al., 2014). As shown in the left panel Figure 2.8B (“original finding”), FPN neurons exhibited accumulation-like responses to the Mooney images, with responses to post-Mooney images (orange line) increasing nonlinearly over time relative to pre-Mooney images (blue line). The authors suggest that this difference in slope later in the trial “may be related to recognition-triggered decision-related activity” (Hachisuka, Shor, Liu et al., 2026, p.7), which is precisely what our model predicts ought to occur. To test the ability of our model to capture these dynamics, we conducted the simulations described below.

2.5.5.2 Simulation methods

The original timeseries data were generated from approximately 500 FPN electrodes implanted across $n = 19$ patients. Each patient completed at least 30 trials of the pre/post Mooney task, and timeseries were averaged across trials in each condition for each electrode in each subject (Hachisuka, Shor, Liu et al., 2026). Because the form of our model is designed to capture individual-level processes, we treated each run of the simulation as the

average of all electrodes implanted in a single patient. Further, we assumed that each subject contributed data using some unique contribution of relative memory “sampling rate” γ and threshold z . For all subjects, we defined the strength of sensory and memory evidence as follows: $\theta_{vision} = 0.545$, $\theta_{pre} = 0.55$, and $\theta_{post} = 0.9$. We chose values slightly greater than the minimum of 0.5 for sensory evidence and pre-Mooney expectation to capture the effects of low-level visual signals and repetition-based familiarity, respectively (Hachisuka, Shor, Liu et al., 2026). As in the original study, sensory evidence was generated at a rate of 75Hz and presented for 2000 simulated ms ($n = 150$ maximum timesteps per trial).

Given the rate of sensory evidence presentation, we generated data using $\gamma = \{6, 7, 8, 9, 10, 12, 15, 18, 21, 30\}$ such that memory samples were generated at rates spanning 2.5-12.5Hz across all subjects. The decision process was terminated using threshold values $z = \{4, 6\}$. As in Section 2.5.2, we manipulated evidence encoding noise ϵ on the whole-sample level, running different simulations for $\epsilon \in \{0.01, 0.06, 0.11, 0.15, 0.2\}$. As with all other findings, we report results from the simulation with $\epsilon = 0.11$; timeseries averaged across all ϵ values are displayed in Figure A.4. For each cue type (pre- vs. post-Mooney), we simulated 30 trials using each unique combination of γ and z , which created a sample size of $n = 20$ simulated subjects. We averaged the decision variable of our model across runs of the simulation within each subject, and then display the group averages in Figure 2.8B, which also contains shaded regions depicting the standard error across simulated subjects.

2.5.5.3 Simulation results

We focused our simulation only on the pre- and post-Mooney timeseries (blue and orange lines in Figure 2.8B), since these recordings are the most relevant to our theory. The pre-Mooney timeseries correspond to accumulation of weak sensory evidence with only a weak expectation about its content, whereas the post-Mooney timeseries correspond to accumulation of the same weak sensory evidence but with a strong expectation about its content. Importantly,

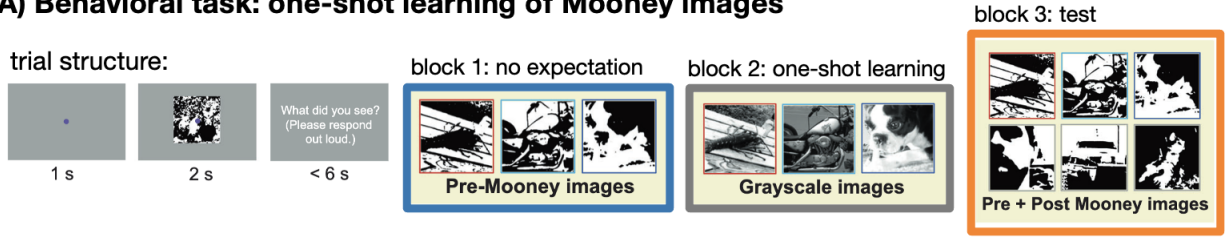
observers are not told *which* grayscale image to use as their expectation for any given trial in the test block; they simply observe the Mooney image and then make a report.

Our sampling-based theory of expectation retrieval, setting, and integration predicts that time-varying effects of memory retrieval on choice dynamics ought to be observed in precisely this kind of experimental setting. Figure 2.8B quantitatively validates that prediction by showing that our dynamic reliability-weighted process can reproduce similar dynamics as those observed by Hachisuka, Shor, Liu et al. (2026) in neural recordings. When observers have learned a strong expectation for a particular sensory input, but are not instructed ahead of time to deploy that expectation for an upcoming choice, then their choices become increasingly biased by that expectation as a function of elapsed decision time (orange lines in Figure 2.8B). However, when an observer does *not* have the relevant expectation for a particular sensory input, the decision variable increases at a lower and more constant rate (blue lines in Figure 2.8B). Figure A.4 further demonstrates that these effects are present in steady-state dynamics of our model as well.

Hachisuka, Shor, Liu et al. (2026) speculated that these time-varying dynamics might be driven by recognition-related neural activity. Although we do not explicitly include recognition-driven retrieval dynamics in our model, our ability to reproduce the “late” effect observed by Hachisuka, Shor, Liu et al. (2026) lends support to the authors’ original interpretation of memory-related dynamics. Further, we offer a formal account for *why* these dynamics occur *when* they do: expectation-setting takes time, and uncertainty about *which* expectation to use across decisions will increase the amount of time it takes to set & integrate the relevant expectation into the decision process. These recent findings from Hachisuka, Shor, Liu et al. (2026) nicely complement the late effect of expectations first reported by Hanks et al. (2011). Whereas the late effect in Hanks et al. (2011) was driven by heterogeneity in the sensory evidence strength across trials, the late effect in Hachisuka, Shor, Liu et al. (2026) is driven by heterogeneity in the *expectation* across trials. While it is possible that

some Mooney images are more easily discernible than others in the pre-learning stage, this incidental variability in sensory evidence strength is unlikely to drive robust effects across observers. Taken together, the Hanks et al. (2011) and Hachisuka, Shor, Liu et al. (2026) findings lend converging support for the core theoretical argument of this paper: dynamic effects of expectations are driven by the time-varying *relative* uncertainty of memory and sensory evidence.

(A) Behavioral task: one-shot learning of Mooney images



(B) Trial-averaged decision variables capture Hachisuka, Shor, Liu et al. (2026) “late” effect in frontoparietal electrodes

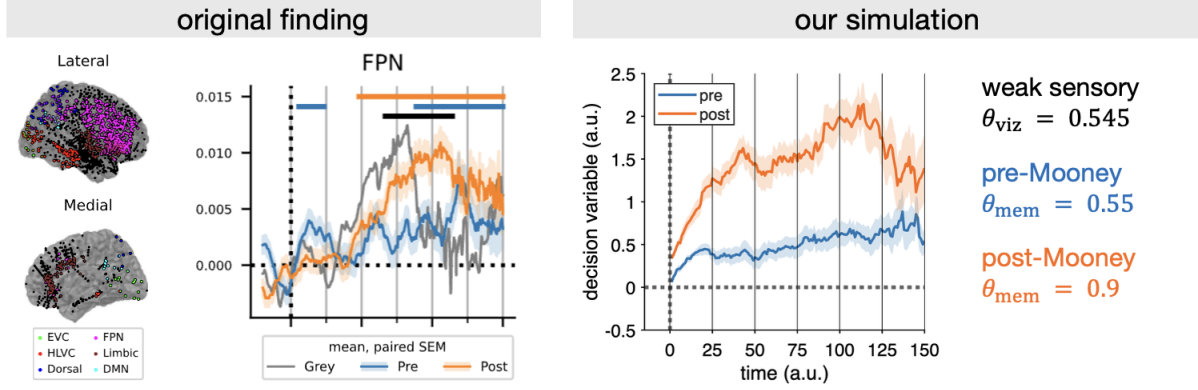


Figure 2.8: Dynamic reliability-weighted integration captures Hachisuka, Shor, Liu et al. (2026) “late” effect of expectations. (A) Experimental design in Hachisuka, Shor, Liu et al. (2026). On each trial, human subjects were presented with an image for 2 seconds and asked to verbally report the content of the image. In the first block, participants viewed “degraded” images whose contents are difficult to discern in the absence of prior knowledge. Then, in a second block, participants were presented with denoised grayscale versions of the same images, facilitating the one-shot learning characteristic of Mooney images. Finally, in the third block, participants were presented with degraded images again. Crucially, 50% of the images had just been presented in full resolution during the learning block. This manipulation allows for discriminating neural timecourses as a function of pre- and post-learning. Image reproduced from Hachisuka, Shor, Liu et al. (2026) with colored borders added for clarity. (B) Original finding: location of implanted electrodes across all patients (left), average activity from frontoparietal (FPN) electrodes (right). Activity post-learning (orange) ramps up more quickly than activity pre-learning (blue), generating a “late” effect of one-shot learned expectations about ambiguous sensory input. Our simulation: decision variables averaged across trials reproduce the pre- and post-learning dynamics measured by FPN electrodes. Shaded areas reflect standard error across $n = 20$ simulated subjects.

2.5.6 Interim summary

The previous two subsections (2.5.4 and 2.5.5) show that our model can reproduce dynamic effects of expectations observed in *expectation-heterogeneous* environments, where observers face deployment-related uncertainty about which expectation will be used on each trial (left side of the quadrants in Figure 2.1). First, we reproduced the uncertainty-adaptive anticipatory sampling finding reported by Bornstein et al. (2023), where observers increase their sampling of memory over elapsed time before evidence onset *only* when a predictive cue signals weak upcoming sensory evidence (Figure 2.6A, right; Figure 2.7A, right). Then, we generated timeseries that reproduce the difference in pre- and post-learning activity observed in frontoparietal electrodes while humans discriminated Mooney images (Hachisuka, Shor, Liu et al., 2026), lending formal support to the original authors’ interpretation of memory-related processing.

Combined with earlier findings that our model captures captures dynamic effects in *expectation-homogeneous* environments (Sections 2.5.1-2.5.2), this section has shown that our principled model can reproduce time-varying effects of expectations on perceptual decisions observed across different experiments, species, neuroimaging modalities, and research groups. Each of these findings were initially explained using distinct cognitive and/or computational processes: learning time-dependent accuracy functions (Hanks et al., 2011), conflict monitoring (Rungratsameetaweemana, Itthipuripat et al., 2018), two-stage evidence accumulation (Bornstein et al., 2023), and memory recognition (Hachisuka, Shor, Liu et al., 2026). What we have shown, however, is that these findings can be unified through our principled theory of uncertainty-driven memory sampling and integration (Figures 2.1 and 2.2).

Importantly, our results do not necessarily imply that previous explanations are incorrect. Indeed, most of them already invoke different theoretical dimensions of our model: uncertainty-driven dynamic weighting on expectations (Hanks et al., 2011), expectation-setting via evidence accumulation (Bornstein et al., 2023), and memory retrieval/reinstatement as a pro-

cess that contributes evidence to the perceptual decision variable (Hachisuka, Shor, Liu et al., 2026). Our model can thus be interpreted as *synthesizing* existing theories about time-varying effects into a single principled computational framework, rather than invalidating the previous findings or explanations.

2.5.7 Uncertainty about expectations induces sequential effects on RTs

In addition to unifying previously-disparate observations into a common principled process, the broader framework of uncertainty-driven memory sampling generates a testable empirical prediction: if the process of *setting* expectations itself involves evidence accumulation over time, then *switching* expectations across consecutive decisions ought to *slow* reaction times relative to when expectations remain unchanged across decisions. Further, this effect should be observed independently of other behavioral and cognitive factors known to modulate reaction time. We tested this prediction using linear regressions on publicly-available data from two studies that measured human behavior in expectation-heterogeneous environments, but differed in whether observers used learned cues through instruction (Diaz, Pisauro et al., 2024) or experience (Bornstein et al. (2023)); details about each study can be found in Section A.2.

Regressions were fit using the `lm()` function in R 4.4.2. Estimated marginal means, 95% confidence intervals, and contrast coefficients were computed using the `emmeans` package. All reaction times were log-transformed and z-scored within subject before entered into the regression, and thus we did not include any random effects for subjects. All models contained fixed effects of predictors of interest (previous cue and previous target), as well as for factors known to affect reaction time (cue validity, cue level, choice accuracy, sensory evidence strength, accuracy of choice on immediately preceding trial, and total amount of time on task). Sensory evidence strength and time on task (defined as raw trial number) were

modeled as continuous regressors, whereas all other variables were modeled categorically so as not to enforce a linearly-increasing effect of cue probability. Data and code to reproduce these results can be obtained from Github.

Figure 2.9 shows that both instructed and learned expectations induce sequential effects on reaction times. For both the Diaz, PISAURO et al. (2024) and Bornstein et al. (2023) datasets, RTs were significantly slower when cues switched across two consecutive trials relative to when they were repeated ($\beta_{instructed} = 0.029$, $t_{16687} = 3.858$, $p < .001$; $\beta_{learned} = 0.055$, $t_{4678} = 3.172$, $p = .002$). Crucially, we did not observe any significant effects of switching or repeating choice targets across trials ($\beta_{instructed} = 0.003$, $t_{16687} = 0.516$, $p = 0.606$; $\beta_{learned} = -0.007$, $t_{4678} = -0.363$, $p = .717$) in either dataset, suggesting that the observed effects are driven solely by switching expectations—but not sensory choice outcomes—across trials. Full summaries of fitted coefficients can be found in Section A.2.

Taken together, these findings lend empirical support to the core theoretical motivation behind our model: uncertainty about expectations ought to induce effects on the dynamics of perceptual choice. We focused our analyses on behavior measured in expectation-heterogeneous environments, where observers have uncertainty about *which* expectation they will need to use for each decision. The framework of uncertainty-driven memory sampling predicts that behavior should exhibit “switch costs” related to the time it takes to adjust expectations across trials with different cues, which is precisely what we observed in both the Diaz, PISAURO et al. (2024) and Bornstein et al. (2023) datasets. Further, these effects were observed independently of canonical experience-related effects on RT, providing evidence for effects of uncertainty at both the trial- and task-level. It is worth noting that these predictions are not generated directly by our model, which assumes that reliability estimates are constructed anew on each trial and thus predicts uniform effects on choice dynamics regardless of the preceding choice. Incorporating a mechanism specifying how the reliability estimates are learned over the course of sampling and then retrieved at the moment of

choice (Section 2.6.3.3 is a promising direction for future work to formally model these serial effects. These findings demonstrate the theoretical utility of considering how uncertainty about expectations modulates the dynamics of perceptual choice.

Switching **instructed** and **learned** expectations across decisions slows reaction times

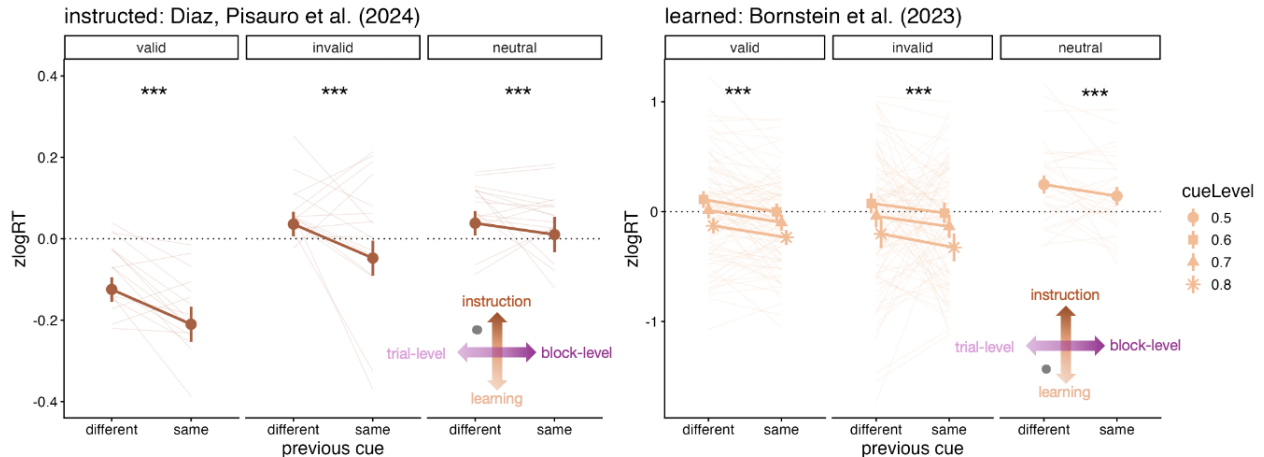


Figure 2.9: Expectation-heterogeneity induces sequential effects on RTs, regardless of whether expectations are explicitly instructed (left) or learned through experience (right). Across all cue levels, switching cues across trials significantly slows responses, whereas repeating cues across trials significantly speeds responses. In all panels, individual lines reflect subject means, points are group-level marginal means estimated using fitted regression coefficients, errorbars reflect 95% confidence intervals around that estimate.

2.6 Discussion

We have shown that uncertainty-driven memory sampling can explain time-varying effects of expectations on perceptual decision-making. A core theme throughout the article has been interrogating the assumptions behind extant models specifying the optimal procedure for integrating beliefs about prior probability (or expectations) into evidence accumulation processes. In doing so, we identified *uncertainty about expectations* as a critical feature that has been omitted from existing normative models. We then used the framework of sequential sampling to develop a verbal theory and formal model articulating how uncertainty about expectations ought to impact evidence accumulation. Regression analyses on existing

behavioral data provided empirical support for a novel prediction generated by our verbal theory, and simulation results demonstrated that our model can reproduce time-varying effects of expectations observed under various experimental conditions and neuroimaging modalities.

This section further articulates the contributions that our framework and model make to broader literature on expectations in perceptual decision-making. To do this, we first discuss how our model is related to—and distinct from—neighboring models in the literature. Then, we discuss insights our model offers for neural investigations into evidence accumulation. Next, we discuss the simplifying assumptions upon which our model rests, articulating both how those assumptions constrain the scope of phenomena the model can capture and how they might be relaxed to capture more complex phenomena of interest. Finally, we conclude with novel predictions that our model makes about choice behavior and neural dynamics, further demonstrating its utility for advancing research on memory-guided visual behavior.

2.6.1 Related models

The sequential sampling family of models is broad and structurally heterogeneous, with each model specifying a slightly different process whereby evidence is accumulated to commit to a choice. Further, different models within this family have been developed to support different approaches to scientifically studying decision making. We have previously shown how considering the difference in reasoning goals motivating model specification and application can resolve apparent tensions in model comparison and selection, especially concerning the inclusion of components that vary as a function of time within the trial (Khoudary et al., 2025b).

In this work, we use sequential sampling to specify a formal theory about a data-generating process whereby two sources of dynamic, probabilistically-evolving evidence are combined based on first principles from statistics and psychology (i.e., reliability-weighted cue combi-

nation). We structure our discussion of related models along this axis of use cases, beginning first by addressing other simulation-based findings of time-varying diffusion models. Then, we discuss “multi-stage” diffusion models that are often used to approximate continuously-changing processes when models are applied to data. Finally, we discuss “psychometric” diffusion models, a term we use to refer to models designed specifically for purposes of latent variable estimation.

2.6.1.1 Dynamic diffusion models

The first class of models we survey are those where the drift rate and/or decision threshold change continuously as a function of time (Deneve, 2012; Huang et al., 2012; Malhotra et al., 2017). One of the most similar previously-specified models was proposed by Deneve (2012), who investigated the question of how the dynamic bias signal identified by Hanks et al. (2011) might be implemented via Bayesian neural inference. In addition to inferring the probability of the correct choice conditional on all observed evidence thus far, Deneve’s (2012) model also infers the sensory evidence strength on a trial-by-trial basis. The model then uses this time-varying estimate to scale the decision thresholds and weighting of evidence within the course of a single decision. This results in a process where both sensory weights and decision thresholds *increase* with time on trials with highly certain sensory evidence, but *decrease* with time on trials with highly *uncertain* sensory evidence (Deneve, 2012, Figure 3). Deneve (2012) further shows that this process also generates time-varying weights on expectations, such that they become more strongly weighted in the decision process as a function of elapsed time on a single choice. Finally, Deneve (2012) shows how these computations can be implemented using both spike trains of single motion-selective neurons as well as populations of neurons that share motion-selective properties, offering proof-of-principle for the feasibility of (1) jointly inferring uncertainty and motion direction and (2) using these quantities to dynamically modulate decision boundaries and drift rates (i.e., weighting of evidence at each moment in time).

Complementary analyses by Huang et al. (2012) and Malhotra et al. (2017) used partially observable Markov decision processes (POMDPs) to further demonstrate that the dynamic bias signal reported by Hanks et al. (2011) can be generated via time-varying decision boundaries. As discussed in Section 2.3, POMDPs permit framing the evidence accumulation process as one of action selection: at each moment in time, an observer chooses whether to continue sampling or commit to one of the two choice outcomes based on the evidence they have accumulated thus far (Equation 2.3). The policy of a POMDP defines which action the observer takes at each timestep, and optimal policies are those that maximize expected future reward (Rao, 2010; Huang et al., 2012; Malhotra et al., 2017). Huang et al. (2012) showed that the optimal policy for a POMDP in an expectation-homogeneous, evidence-heterogeneous environment corresponds to decision boundaries that (1) decrease exponentially over time and (2) are shifted in proportion to the prior probability Π , effectively implementing a starting point offset. A more recent study from Malhotra et al. (2018) built upon these findings to show that the *rate* of bound collapse depends jointly on the prior probability and the mixture of sensory evidence strength in the decision environment. The authors first showed that decision boundaries only exhibit exponential collapse in environments where one of the signal strengths is sufficiently low; otherwise, reward rate is optimized by using boundaries that are constant over time (Malhotra et al., 2018). Then, they showed that adding a strong prior probability *linearizes* the collapse of boundaries over time (Figure 9 in Malhotra et al., 2018), such that for sufficiently high prior probability (in this case, $\Pi = 0.7$) the optimal boundaries remain parallel but decrease monotonically over time. Accordingly, increasingly more evidence is required to overcome the prior as deliberation time increases—precisely the dynamic effect observed by Hanks et al. (2011) and reproduced by our model.

Taken together, the simulation-based findings of dynamic diffusion models strongly suggests that a time-varying modification to the decision process is needed when observers have uncertainty about sensory evidence (Drugowitsch et al., 2012; Deneve, 2012; Huang et al., 2012; Malhotra et al., 2018). Although our model does not explicitly incorporate time-

varying decision boundaries, we have shown that it reproduces both the original finding these models aimed to explain (Hanks et al., 2011) in addition to explaining several other time-varying effects of expectations on perceptual decisions (Section 2.5). Critically, our model suggests that several of the dynamics predicted by expectation-induced effects on decision boundaries could also be generated by a dynamic reliability-weighted process integrating samples from parallel streams of memory and sensory evidence. Because time-varying drift rates can often mimic time-varying thresholds, carefully designed experiments coupled with high-resolution neuroimaging will likely be required to adjudicate these formally similar—but conceptually different—theories of how expectations modify the dynamics of choice. Further, as Deneve (2012) suggests, it could be the case that expectations exert time-varying effects on *both* the drift rate and decision thresholds. One possibility is that expectations with low source-related uncertainty act primarily on decision thresholds, whereas expectations with high source-related uncertainty act primarily on the drift rate, owing to this latter parameter’s weaker dependence on observers’ explicit knowledge or decisional strategies.

2.6.1.2 Multi-stage diffusion models

Multi-stage diffusion models correspond to a stochastic process where the drift rate, diffusion coefficient, and threshold change discretely—rather than continuously—over time (Srivastava et al., 2017). We focus specifically on the class of multi-stage *drift* models, which generate “switches” in the slope and direction of the decision variable at specific points in time (Ratcliff, 1980; Diederich and Busemeyer, 2006; Diederich and Oswald, 2014, 2016; Srivastava et al., 2017).

2.6.1.2.1 Two-stage models of expectations A study by Diederich and Busemeyer (2006) was among the first to apply this kind of model to the question of how asymmetries in response outcomes—defined as differences in expected payoffs—modulate choice dynamics. The authors found that a two-stage model where payoffs were initially processed with one

drift rate, and then stimuli were processed with a different drift rate outperformed models that incorporated the payoffs as offsets on both the starting point and drift rate (Diederich and Busemeyer, 2006). This finding is a conceptual predecessor to the results of Bornstein et al. (2023), who used a model developed by Srivastava et al. (2017) to show that brain and behavioral data were best captured by a two-stage process whereby evidence accumulation from memory sets the starting point for sensory evidence accumulation on a trial-by-trial basis.

The right panel of Figure 2.6A shows that the decision variable of our model effectively implements a time-continuous version of the two-stage model used by Bornstein et al. (2023): anticipatory accumulation from memory biases the starting point of the sensory evidence accumulator jointly as a function of expectation strength and sampling time. Two key features distinguish our model from the multi-stage models developed by Srivastava et al. (2017) and Diederich and Busemeyer (2006). First, our model posits two independent accumulators for each evidence source, and a decision variable that is a dynamic reliability-weighted combination of samples from each source. This leads to the second difference, which is that we posit an effect of memory that *continues on* once the sensory evidence has been presented. Rather than suggest that the decision process involves discrete “switches” based on sampling one source versus another, we propose that both sources of evidence are continuously sampled and integrated throughout the deliberation process. We discuss this assumption of our model, and possible extensions beyond it, in later parts of the Discussion.

2.6.1.2.2 Opposing process theory Although not expressed using the formalism of sequential sampling, Press et al.’s (2020) “opposing process theory” can be thought to correspond to a process with a multi-stage drift rate. Press et al. (2020) developed opposing process theory to resolve an apparent paradox between “Bayesian” and “cancellation” theories of how expectations modulate the gain of sensory information. Whereas Bayesian theories—prominent in sensory perception research—assert that expected information ought

to have a higher gain, cancellation theories—prominent in action research—assert that *unexpected* information ought to have higher gain. Opposing process theory resolves this tension by positing a two-stage process whereby sensory evidence for expected outcomes initially is processed with a higher gain, but that surprising inputs (quantified using the Kullback-Leibler divergence of an observer’s posterior relative to their prior) generate increased gain later on during processing, potentially via phasic release of catecholamine (Press et al., 2020).

The findings we display in Figures 2.5A and 2.6B suggest that our model is capable of formalizing the dynamics predicted by opposing process theory using a time-continuous modeling framework. The right panel of Figure 2.5A shows that our model generates greater weights on unexpected, followed by *unexpectedable*, sensory evidence as a function of elapsed time. The rightmost panel of Figure 2.6B (top) additionally shows that strong expectations generate early weights that prioritize predictions from memory relative to contents of perception. Indeed, the conjunction of these dynamics—early upweighting of expected information and later upweighting of unexpected information—was recently observed via EEG decoding of a task specifically designed to test opposing process theory (Rittershofer et al., 2026). Given that our model can reproduce both of these dynamics in independent cases, it follows that it should be able to capture their conjunction as well.

2.6.1.3 Psychometric diffusion models

The final class of models we consider are those that we term *psychometric* models, deriving from their development for purposes of latent variable estimation across a variety of experimental settings (Khoumary et al., 2025b; Batchelder, 2010). The first model we discuss, the attentional DDM (aDDM), posits a time-varying offset onto the decision variable as a function of fixation location (Krajbich et al., 2010; Tavares et al., 2017; Yang and Krajbich, 2023). The second, the extended DDM (eDDM), posits trial-level variability in the drift rate and starting point terms of the original DDM (Ratcliff and McKoon, 2008; Ratcliff et al.,

2016). Our work in this article motivates alternative applications and/or interpretations of each model, which we discuss in turn.

2.6.1.3.1 The attentional DDM (aDDM) An increasingly prominent time-varying drift rate model is the “attentional DDM” (aDDM), which formalizes the effects of eye movements on the dynamics of value-based choice (2010). The aDDM posits that fixations on a specific option induce transient biases on the drift rate, such that evidence is accumulated more quickly for an option when an observer is overtly attending to it with their gaze. More recently, a multi-attribute generalization has been developed which further permits estimating the weighting of specific attributes—as opposed to entire outcomes—on the drift rate (Yang and Krajbich, 2023). A core feature distinguishing our model from the aDDM is *how* the dynamic effects on drift are generated: the aDDM uses measurements of fixations to *infer* time-varying biases on the drift rate, whereas our model posits that drift rate biases are a function of the time-evolving evidence itself. Further, the aDDM posits that the magnitude of drift modulations is governed by the subjective value of the attended choice option, whereas we suggest that the relative evidence reliability governs how strongly drift rates are biased.

These differences can be understood as reflecting the different cognitive processes—operationalized via different experimental tasks—that the models are designed to explain. In contrast to perceptual decision tasks that present one dynamically-evolving stimulus on each trial, stimuli in value-based decision tasks are comprised of two static, full-color images. Further, perceptual decision tasks have an objectively correct answer (defined via experimental manipulations), whereas value-based tasks are designed to elicit stable differences in subjective estimations of stimulus value. An interesting direction for future work would be developing a single model that unifies both types of decision processes into a common framework, perhaps by considering the role of uncertainty in scaling dynamic effects both of memory and attention on deliberation.

Further, given the close empirical link between eye movements and memory retrieval (Barker et al., 2025; Kragel et al., 2020; Hannula et al., 2010), as well as memory retrieval and dynamic construction of subjective value (Shadlen and Shohamy, 2016; Wang et al., 2021), it would be fruitful to consider the role of memory retrieval dynamics in guiding fixations during both value-based and perceptual choice. In particular, it would be interesting to consider how memory sampling may dynamically *guide* eye movements during sensory evidence accumulation, perhaps by directing gaze to areas that in space that are maximally informative for discriminating whether a trial is congruent or incongruent with what is predicted by memory. Another possible direction for future research is using the aDDM to identify moments in time when observers’ attention switches from external sensory information to internal information generated by memory (Verschooren et al., 2026; Kumle et al., 2025), a process that has been linked to differential recruitment of the hippocampus by subcortical and cortical brain regions, respectively (Poskanzer and Aly, 2023).

2.6.1.3.2 The extended DDM (eDDM) The extended DDM (eDDM; Ratcliff and Rouder, 1998; Ratcliff and Tuerlinckx, 2002; Ratcliff and McKoon, 2008) builds upon the original DDM (Equation 2.1) in two ways. First, it specifies an additional *non-decision time* term T_0 that additively offsets the first-passage time of the Gaussian walk to account for delays in RT attributable to putatively non-decision related processing, such as sensory evidence encoding and motor command execution. Second, it defines the drift rate, starting point, and non-decision time terms as *probability distributions* rather than constant values, such that the decision process can exhibit structured variability across trials. Both of these modifications enhance behavioral goodness-of-fit across several tasks, but a principled theory about *why* trial-level variability in these parameters (and of their specific probabilistic form) is still largely lacking.

Anticipatory, simultaneous, and/or post-decisional memory sampling may offer some insight into this question. Specifically, trial-level variability in the starting point may reflect indi-

vidual differences in how quickly memory information comes to mind and/or is weighted relevant to the upcoming decision. Drift rate variability could also be explained by periodic “bursts” of internal evidence generated by memory that direct observers’ attention away from their external environment, contribute internal evidence bearing on the externally-oriented decision, or some combination thereof. As such, it may be the case that the eDDM improves goodness-of-fit not because it better captures facts about the data-generating process, but rather that it more flexibly accommodates stable individual differences that are not accounted for by simpler models. Although hierarchical Bayesian approaches to parameter estimation handle concerns about individual differences in a statistically elegant and principled manner, they do not on their own provide answers to the deeper theoretical question of *why* individuals exhibit the differences that they do. Process models based on first principles—coupled with tasks designed to tease apart structured deviations from their predictions—are powerful complementary tools toward this end.

2.6.2 Relationship to neural data

Principled process models have been highly useful for advancing research into the neural bases of evidence accumulation, and in linking neural activity with observed behavior across species (Brody and Hanks, 2016; Gold and Shadlen, 2007; Forstmann et al., 2016; Hanks and Summerfield, 2017). Foundational research in non-human primates demonstrated signatures of evidence accumulation in single-cell and population-level activity (Newsome et al., 1989; Britten et al., 1996; Gold and Shadlen, 2002), and mapped different computational steps of the formal process onto distinct cortical regions comprising a decision-making circuit (Gold and Shadlen, 2007). More recent research using large-scale recordings in rodents has both challenged and supported the early image developed based on primate work, primarily by demonstrating that signatures of evidence accumulation can be observed across the entire brain instead of a specialized decision circuit (Liu et al., 2025; Findling et al., 2025; Bondy et al., 2025).

The framework, model, and findings we develop here contribute to this line of research by identifying the specific circumstances in which memory retrieval dynamics are most likely to modulate the dynamics of sensory evidence accumulation. Further, our model permits making specific predictions about what a candidate timeseries might look like under different assumptions about the relative reliability of both evidence sources, the observer’s uncertainty about that reliability, and the relative rate at which information is retrieved from memory. And because we aimed to keep our model as minimal as possible, future work can refine its predictions by incorporating neural observations about the timescales of activity in different brain regions (e.g., Shi et al., 2025), which is a core conceptual idea motivating our model.

2.6.3 Model assumptions and extensions

A core motivation for our model is identifying a principled process for integrating expectations with sensory evidence under less-idealized assumptions. Taking a formal approach, however, requires that we still make use of simplifying assumptions in order to make the modeling tractable. This section discusses these assumptions and suggests directions that future work can take to address them.

2.6.3.1 Independence of memory and sensory evidence accumulation

First, our model assumes that each source generates and accumulates evidence independently of the other. That is, the generating distribution for memory samples remains constant over the course of a single decision and does not change as a function of the content of sensory information. We justified this assumption on the basis of our goal to model simple pattern completion and/or associative retrieval processes, rather than a more complicated memory search process. We believe this assumption of a static generating distribution for memory—or the independence of memory and sensory evidence samples—is warranted for quick perceptual decisions on the order of a few seconds. However, it is likely that more prolonged perceptual deliberation might involve sequentially searching through relevant in-

formation in memory (Wang et al., 2021; Callaway et al., 2024). Future work can investigate this possibility by dynamically manipulating the content of sensory evidence and/or the relevant expectation within the course of a single trial (e.g., Kiani et al., 2013; Kilpatrick et al., 2019), as well as considering cases where observers have more richly structured expectations about the dynamic contents of sensory evidence (e.g., Chen and Bornstein, 2024; Kaper and Peters, 2026).

2.6.3.2 Memory “sampling rate” γ

Next, to capture putative differences in the timescales of sensory and mnemonic evidence generation, our model assumes that memory evidence is generated at a fraction of the rate that sensory evidence becomes available to the observer, captured by the γ parameter. For simulation findings aimed at re-producing empirical time-varying effects (Figures 2.5 and 2.8), we incorporated a wide range of memory “sampling rates” for two reasons. First, to show that our findings are robust across a range of plausible ratios in sampling rate, and second, to capture putative heterogeneity in this ratio across individuals and species (2014, 2020). We did, however, restrict the set of possible ratios to the broadest range of empirically-observed theta-band oscillatory activity (2-12Hz; Seger et al., 2023). While it is possible that the type of memory sampling we invoke here occurs at other relative “sampling rates”, we chose to restrict this set of findings to the theta frequency because of robust neural evidence implicating theta-band activity in task-evoked memory retrieval, particularly in the context of memory-guided visual behavior (Ter Wal et al., 2021; Zhang et al., 2018; Kragel et al., 2020; Kerrén et al., 2018; Zubair et al., 2026; Senoussi et al., 2025). An interesting direction for future research is to consider how memory “sampling rate” might be *controlled* by observers in a task- or uncertainty-dependent manner, or if it is better understood as a trait-level property that more so functions as a “bottleneck” on the memory sampling process.

2.6.3.3 Evidence reliability estimation

Finally, our model makes a handful of assumptions about the reliability estimation process. First, we assume that observers dynamically estimate the reliability of each evidence source on every single decision; essentially, that observers treat the evidence-generating distributions as nonstationary with unknown time variation. This is a simplifying assumption because the evidence-generating distributions are commonly stationary at the single-trial level, so it is possible that reliability estimates may carry over from trial-to-trial. Indeed, the secondary behavioral analyses conducted in Section 2.5.7 corroborate the idea that reliability estimates are “cached” across trials, as evidenced by the RT advantage for repeating learned cues relative to switching them (Figure 2.9). Modeling and measuring how observers learn to estimate the reliability of time-evolving evidence—perhaps via a process akin to Kalman filtering (Roweis and Ghahramani, 1998)—is a promising direction for future research.

Second, and relatedly, we assume that observers estimate the reliability of evidence *online* during the decision process. This assumption follows directly from our aim of considering how expectation-heterogeneity—or uncertainty about *which* expectation to use on each trial—ought to impact the dynamics of perceptual choice: when the environment dictates which expectation to retrieve on a trial-by-trial basis, any estimation of mnemonic evidence reliability must be done during the retrieval process itself. It would be interesting, however, to consider cases where observers themselves determine which information to retrieve from memory, either in anticipation of or simultaneously with perceiving ambiguous sensory evidence. Our model predicts that observers ought to retrieve information that maximally reduces uncertainty in their deliberation process, but leaves open the question of *how* observers determine which information is useful toward this end. Tasks that utilize multidimensional expectations, along with sensory evidence that evolves along longer timescales, can help shed light on these questions (e.g., Finn et al., 2018; Nicholas and Mattar, 2026; Mazor et al., 2026). Further, it would be interesting to investigate whether the same dynamic ef-

fects emerge when evidence reliability is used to *guide* memory retrieval, rather than being estimated online as a function of it.

Finally, our model assumes that there is no noise or error in observers’ reliability estimation process. This assumption follows directly from Equation 3.3, which states that the Beta-distributed belief $g_s(t)$ about source reliability θ_s is perfectly updated with each new sample of sensory and memory evidence. This is one of the strongest simplifying assumptions of our model, since it is unlikely that observers place equal weight on every sample at every timestep when constructing their belief about θ_s . A simple extension could be introducing stochasticity into the belief-updating process, perhaps by incorporating a softmax function into the update rule defined in Equation 3.3. Another possibility could be that observers differentially attend to or weight evidence based on its *congruence* with the other source, and that this affects the accuracy with which observers construct their belief $g_s(t)$. Investigating principled approaches to incorporating asymmetry in belief updating is a promising direction for understanding both the basic mechanisms of reliability-weighted evidence integration, and how they are altered in psychopathology.

2.7 Summary and Conclusion

We have proposed that dynamic reliability-weighted integration of parallel streams of memory and sensory evidence offers a unifying explanation for time-varying effects of expectations on perceptual decisions. To do this, we first discussed how expectations—beliefs about the prior probability of observing a particular outcome—are typically modeled as exerting *static* effects on perceptual decision processes by biasing either the starting point or drift rate of the decision process before any sensory evidence is observed (Section 2.2). We ended our review by discussing recent findings of *dynamic* effects of expectations on perceptual choice that cannot be explained by static models, and then presented a theoretical framework and computational model suggesting that dynamic effects could be driven by reliability-weighted

integration of evidence dynamically retrieved from memory (Sections 2.3-2.4). Finally, we showed that our model can reproduce dynamic effects observed across four different experimental tasks, research groups, and neuroimaging modalities, with each effect corresponding to a unique subspace of our broader theoretical framework (Figure 2.1; Section 2.5).

These findings suggest that the dynamic reliability-weighted process specified by our model offers a unifying explanation for time-varying effects expectations on perceptual decisions. However, direct empirical evidence supporting our model over multi-stage and/or static alternatives is still lacking, in large part because of how expectation-uncertainty has been omitted from normative models of perceptual decision-making (Section 2.3). Future work can combine the framework and model we developed in this article with emerging methods for studying decision-making in dynamic and uncertain environments to test the core prediction of our model: the *rate* of evidence accumulation is driven by fluctuations in *relative* reliability of memory and sensory information.

A promising candidate measurement toward this end is the centroparietal positivity (CPP), an event-related EEG signal whose slope has been shown to scale with the strength of both sensory and memory evidence independently (Tsvinev et al., 2025; van Vugt et al., 2019). The conceptual framework we develop in Section 2.3 predicts that effects of memory sampling should be strongest in fully heterogeneous environments, i.e., those where observers have uncertainty over the true predictive probability of expectation cues, where the cues are manipulated across trials, and where the sensory evidence strength also varies across trials. As suggested above, a behavioral task that dynamically manipulates the *relative* reliability of memory and sensory evidence within a single trial will be necessary to distinguish our parallel integrated model from the serial two-stage models discussed above.

Chapter 3

Expectations dynamically modulate visual evidence accumulation

3.1 Introduction

How do expectations shape perceptual decisions? Decades of neural, behavioral, and theoretical evidence have shown that expectations—beliefs about the prior probability of choice outcomes—enter the decision process by setting the starting point of evidence accumulation (Edwards, 1965; Link and Heath, 1975; Carpenter and Williams, 1995; Ratcliff and Rouder, 1998; Bogacz et al., 2006; Gold and Shadlen, 2007; Ratcliff and McKoon, 2008; Simen et al., 2009; van Ravenzwaaij et al., 2012; Moran, 2015; Bornstein et al., 2023). More recent work, however, suggests that expectations might also modulate the *rate* with which evidence is accumulated toward the more likely choice (Hanks et al., 2011; Deneve, 2012; Moran, 2015; Afacan-Seref et al., 2018; Malhotra et al., 2018; Bornstein et al., 2023; Rungratsameetawee- mana, Itthipuripat et al., 2018; Diaz, Pisauro et al., 2024). While both theories agree that expectations contribute evidence bearing on the perceptual decision, they disagree about

whether that evidence *statically* or *dynamically* modulates the sensory accumulation process, respectively.

Although static theories are grounded in the optimality of a starting-point offset (Edwards, 1965; Simen et al., 2009), more recent work has shown that dynamic effects are *also* necessary for optimizing the speed-accuracy tradeoff on less restrictive assumptions. Both Moran (2015) and Malhotra et al. (2018) showed that when sensory evidence strength (or coherence *c*) varies across decisions *and* observers have knowledge of unequal prior probabilities between the choice outcomes (i.e., expectations), the process that maximizes reward rate is one in which the effect of expectations changes over the course of sensory evidence sampling. These findings lend formal theoretical support to an earlier proposal by Hanks et al. (2011) that dynamic effects of expectations are driven by implicit computation of a time-dependent accuracy function (TDA) which is learned over the course of decision-making in uncertain environments. The core idea behind the TDA function is that, in an evidence accumulation framework, elapsed decision time can serve as a proxy for sensory evidence uncertainty: the longer it takes to make a choice, the more likely it is that the evidence forming the basis of choice is highly uncertain (Hanks et al., 2011). This theory was developed to explain the authors’ findings that both human behavior and primate neural recordings showed evidence of an increased weighting of expectations *late* in the decision process (Hanks et al., 2011). Such “late” effects on the decision process directly contradict prescriptions of starting point models, which predict that the influence of the expectation should *decrease* with sensory evidence sampling time.

Critically, however, both Moran’s (2015) and Malhotra’s (2018) findings assume that observers have perfect knowledge of the prior probabilities. Although this assumption is justified for experimental tasks where observers are explicitly instructed on prior probabilities—as in Hanks et al. (2011)—it fails to capture the experience-dependent uncertainty that characterizes real-world expectations generated by retrieving information about past experiences

from memory (e.g., Wang et al., 2021; Bornstein et al., 2023; Biderman et al., 2020; Tarder-Stoll et al., 2026; Ladyka-Wojcik et al., 2025). The systematic omission of expectation-uncertainty in the theoretical study of perceptual decision-making reflects a broader historical trend whereby research in perception is conducted in isolation from research on learning and memory, a gap that is actively being bridged by research into how the two systems function in concert with each other (Aly et al., 2026; Martin and Barense, 2023; Press et al., 2020). The conceptual framework and mathematical model we propose in Khoudary et al. (2026) contributes directly to this effort by developing a formal framework for theorizing about the role of uncertainty-driven memory sampling in perceptual inference. Specifically, we show that reliability-weighted integration of parallel “streams” of memory and sensory evidence can reproduce several previously-disparate neural and behavioral observations of dynamic effects of expectations (Hanks et al., 2011; Rungratsameetaweemana, Itthipuripat et al., 2018; Bornstein et al., 2023; Hachisuka, Shor, Liu et al., 2026), and that uncertainty about expectations induces sequential effects on reaction times (RTs) in perceptual decision tasks (Khoudary et al., 2026).

We developed this model based on first principles deriving from both theoretical and empirical research on memory and perception. First, drawing on decades of research demonstrating the active, temporally-extended nature of memory retrieval (Schacter and Addis, 2007), and more recent work linking this process to deliberation during value-guided decision-making (Bakkour et al., 2019; Shadlen and Shohamy, 2016; Wang et al., 2021), we posit that expectations are themselves a dynamic “stream” of internal evidence that observers sample *in parallel* to sensory information when making perceptual decisions. Then, building off rich theoretical work on optimal cue combination in multisensory perception (Angelaki et al., 2009; Landy et al., 2011), we suggest that memory and sensory evidence are combined according to a dynamically-updating estimate of the *relative reliability* of the two evidence sources. Although this model can reproduce the dynamic bias effect reported by Hanks et al. (2011), it does not require that observers use elapsed time as a proxy for evidence

uncertainty. Instead, we suggest that observers automatically update their beliefs about the reliability of each evidence source, and use those beliefs to determine how much weight to place on memory vs. sensation at different points during the deliberation process.

Our theory thus diverges from Hanks et al.’s 2011 TDA model by positing that weights on expectations can be *non-monotonic* in time, ebbing and flowing with the dynamics of sensory evidence rather than strictly increasing with elapsed time on task. Further, it differs from starting-point (van Ravenzwaaij et al., 2012) or multi-stage models (Bornstein et al., 2023; Diederich and Busemeyer, 2006) in assuming that expectations continue to be sampled throughout the duration of the decision, rather than ceasing to be consulted once sensory evidence becomes available. Here, we develop a behavioral task that allows us to test a key prediction of our theory: perceptual decisions should be most strongly influenced by expectations during moments in time when sensory evidence is maximally uncertain. Further, the degree of this influence should be modulated by the uncertainty of the expectation itself, with more certain expectations inducing stronger effects than less certain ones.

Observers first learned probabilistic cue-image associations (Figure 3.1A), which we then used to manipulate expectations trial-by-trial in a perceptual decision task (Figure 3.1B). We ensured memory information was persistently cued by presenting sensory evidence *inside* the learned, probabilistic prior cues (colored borders), and used a calibration procedure to ensure sensory evidence reliability was constant across trials and choice outcomes. Then, we dynamically manipulated their *relative* reliability by stochastically inserting epochs of pure noise into the sensory evidence stream (Figure 3.1B). Because the onset and duration of these noise periods followed a fixed hazard rate across trials (Maniscalco and Peters, 2024), elapsed time did not provide observers with information about the timing with which sensory evidence would transition from signal to noise. The key test of our hypothesis, then, is whether the decision variable during periods of sensory noise exhibits a significantly

positive slope toward the choice option predicted by the expectation, and if the magnitude of this slope is modulated by the strength of the expectation itself.

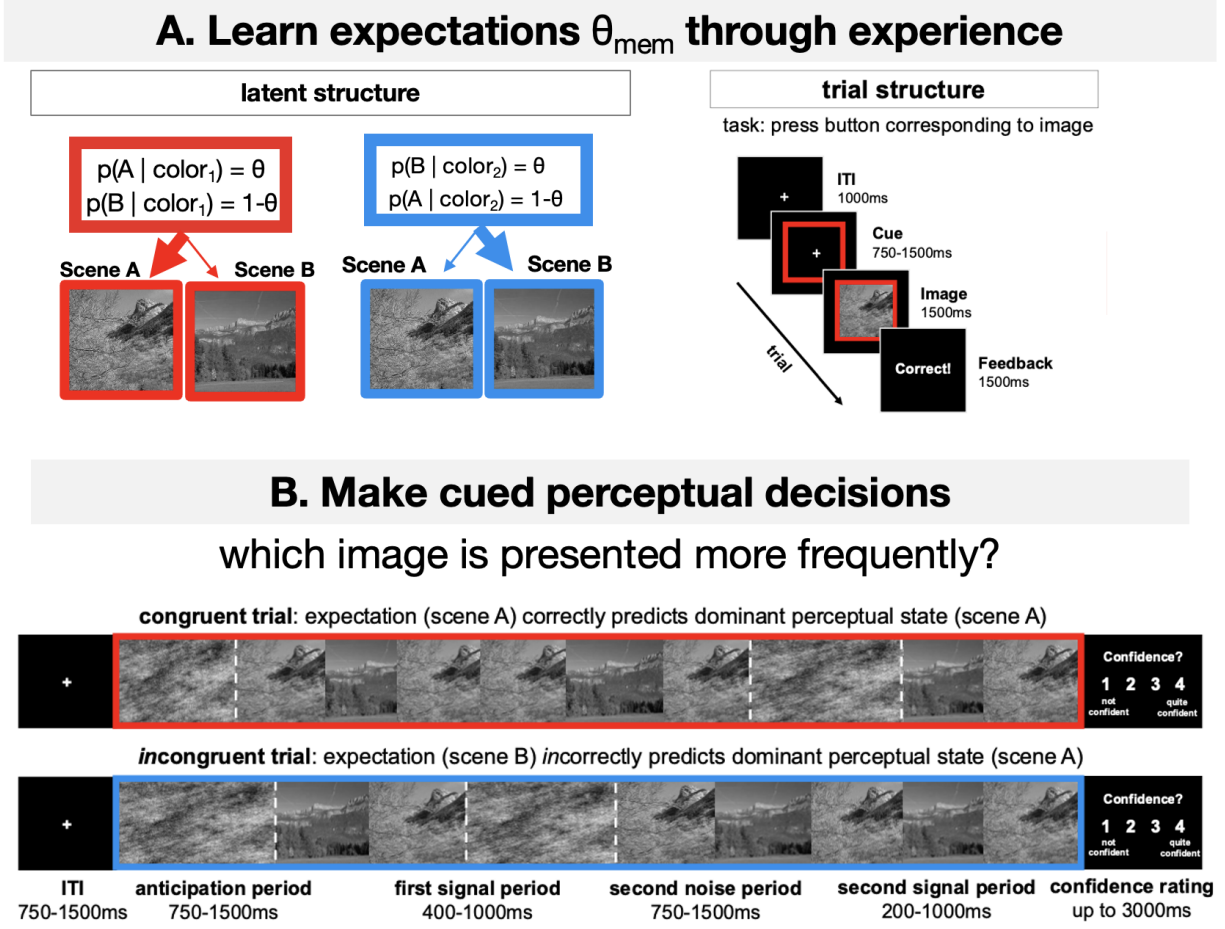


Figure 3.1: Task design. (A) Participants first learn that different colored borders make different predictions about the probability of observing one of two possible scene images. The predictive strength θ_{mem} of each cue is defined by the frequency with which it is followed by one of the two scene images. (B) After learning, participants are presented with stochastic visual evidence inside of the colored borders. Their task is to report which image is dominant (i.e., presented more frequently) in a 60Hz stream of visual evidence, followed by a confidence rating in that perceptual decision.

3.2 Methods: experimental design

In our task, observers first learn expectations by observing an interleaved series of probabilistic cue-image pairings (Figure 3.1A). After making subjective reports for all learned

cues, observers perform a cued perceptual decision task that stochastically manipulates the reliability of sensory evidence in a trial-by-trial manner (Figure 3.1B).

3.2.1 Stimuli

Visual evidence consisted of two grayscale scene images. There were two sets of possible scene images (i.e., four images total), with pairs of images and their mappings to keyboard responses randomized across participants. The probability of a given image being the ‘dominant’ image in a stream of visual evidence displayed to the observer on each trial varied across three possible conditions, which were communicated to the observer via a colored border that circumscribed the visual evidence during the stream. There were two sets of possible borders, with each set comprised of triadic colors (set1 = red, blue, yellow; set2 = orange, green, purple). The set of border colors, along with the borders’ assignments to dominance probabilities for particular images in the visual evidence stream, were randomized across participants. Half of the subjects learned a set of borders with true values of $\Pi = \{0.8, 0.8, 0.5\}$ and the other half learned a set of borders with true values of $\Pi = \{0.65, 0.65, 0.5\}$.

3.2.2 Participants

A sample of 44 undergraduate students (age $M=20.6$ years; $SD=1.69$ years; 35 female, 7 male, 3 non-binary) were recruited from the University of California, Irvine. Four subjects were excluded from the analysis due to chance-level responding during decision-making, resulting in a final sample size of $n=40$ ($n=21$ in the 0.8/0.5 condition; $n=19$ in the 0.65/0.5 condition). Participants were compensated for their time either with course credit or a pro-rated cash payment of \$15/hour. This study was approved by the Institutional Review Board at the authors’ University and all subjects provided written informed consent.

3.2.3 Task design

Data were collected in a single experimental session ranging between 60-90 minutes. Each session began by instructing participants on the mapping between the two scene images and the 1 and 2 number keys on a US keyboard.

3.2.3.1 Calibration

We used two interleaved QUEST staircases (Watson and Pelli, 1983) to identify values of visual evidence coherence that, for each image, resulted in 70% accuracy on the perceptual decision task (described in more detail below). Evidence coherence θ_{vis} was defined as the proportion of signal frames in the visual evidence stream that contained the target stimulus. The calibration procedure ensured that decision difficulty in the Cued Inference phase would be identical for both target stimuli, effectively controlling for low-level visual differences between the images that might systematically bias choices toward one of the options. Participants completed 80 trials of the decision task (40 trials per staircase) during the calibration procedure and received feedback on their choices.

3.2.3.2 Cue Learning

Next, participants learned the predictive probability θ_{mem} of each cue (i.e., each colored border) by observing a series of cue-image pairings in which the cue was presented prior to the image it was probabilistically paired with (Figure 3.1, top). To ensure active engagement—and to build up associative motor memories—participants were instructed to respond on each trial indicating which image appeared on screen after the cue using the previously-learned image-key mappings. Participants were told that there was a predictive relationship between the cues and scene images, and that their broader goal for this phase of the experiment was to learn that relationship. Finally, participants were also told that they were permitted to respond in the inter-stimulus interval (ISI) between the onset of the cue and scene image

if they desired. Regardless of when participants responded, they received feedback on their response accuracy on each trial.

In order to maximally align learning and decision environments, we permitted the ISI between cue and image onset to vary across trials according to a truncated exponential distribution. This approach ensures a fixed hazard rate across learning trials, such that participants are maximally uncertain about the temporal onset of scene images across learning trials (Maniscalco and Peters, 2024). ISIs ranged from 750ms to 1500ms (mean = 892ms, mode = 783ms). After a scene image appeared on screen, participants had up to 1500ms to make their response. Post-trial feedback was displayed for 1500ms, and participants were told that they should respond faster on the next trial if the feedback screen appeared before they made a response.

Each cue was presented a total of 30 times and the order of cues was fully randomized, providing participants 90 total observations of cue-image pairings. The objective probability of each cue was defined as the frequency with which it preceded one of the two scene images: each 80% cue thus preceded its dominant image 24/30 times it was presented and the 50% cue preceded each image 15 times. Colors were randomly assigned to cue probabilities.

3.2.3.3 Cued Inference

In the final phase of the experiment, participants observed a rapidly-alternating (60Hz) “stream” of the two scene images interleaved with pure noise frames (phase-scrambled superpositions of the images) (Figure 3.1, bottom). Observers’ task was to report which of the two scene images was presented more frequently (i.e., was the “target”) on each trial. The proportion of target frames on each trial was defined by the calibrated coherence value for that target’s trial, as estimated during the preceding Calibration phase. Crucially, this visual evidence was presented *inside* the colored borders, ensuring that information about the prior probability was always accessible to the observer. Participants were told that the

predictive relationships they just learned between the colored borders and scene images also applied in this phase of the experiment (i.e., “the correct answer is usually the one predicted by the cue”). They were also instructed to respond as quickly and accurately as possible. However, because we also elicited decision confidence ratings on each trial, participants did not get feedback about their choice accuracy.

We incorporated two periods of stochastic visual noise into each Cued Inference trial. The durations of these noise periods, as well as the brief signal period in between them, were all drawn from separate truncated exponential distributions in order to guarantee a fixed hazard rate across trials (Maniscalco and Peters, 2024). The maximum duration of any trial was 4100ms, and any remaining time after the second noise period consisted of threshold-level visual evidence. Immediately after making a decision, participants had up to 3000ms to report their confidence in that decision’s accuracy on a scale of 1-4 (not confidence-quite confident). Trials were separated by 1000ms intertrial interval.

Each subject completed 150 trials for each predictive cue (80% or 65%) and 75 trials for the 50% cue, thus completing 375 trials in total. The assignment of cue and target was fully randomized across trials, with the probability of a scene image being a target for a particular cue being defined by that cue’s true probability. This means that, for 20% of trials with an 80% cue, the cue was *incongruent* with respect to the true trial target (i.e., its effective prediction was 0.2).

3.3 Methods: computational modeling

We used two classes of computational models to interpret the experimental data. One class of models was used to simulate data under the parallel and serial integration processes, and the other class was used to perform parameter estimation (or “fitting”), on both human subjects and simulated data. Simulation models used custom scripts run using MATLAB

2025a,b and parameter estimation models were built using the `pyddm` package in Python (Shinn et al., 2020).

3.3.1 Model specifications: simulation models

We simulated data from two models that correspond to unique theories about how expectations modulate perceptual decisions.

3.3.1.1 Simulation model: parallel process

The model in this section is introduced in Khoudary et al. (2026) and reproduced here for convenience. Simulated agents sequentially sample evidence from two Bernoulli-distributed evidence sources $s = \{memory, vision\}$. Each evidence source has a unique generating parameter θ_s , which we assume to be unknown by the observer. We assume that observers have varying degrees of evidence encoding noise ϵ , and that stochastic periods of sensory noise are treated as samples from a perfectly uncertain Bernoulli distribution, such that vision evidence is generated according to:

$$o_{vis_t} \sim \begin{cases} Bernoulli(0.5) & \text{if } t < t_{s1} \vee t_{n2} \leq t < t_{s2} \\ Bernoulli(\theta_{vis}(1 - \epsilon) + (1 - \theta_{vis})\epsilon), & \text{if } t \geq t_{s1} \wedge t < t_{n2} \vee t \geq t_{s2} \end{cases} \quad (3.1)$$

where t_{s1} is the onset of sensory evidence, t_{n2} is the onset of the intermediate noise period, and t_{s2} is the onset of the sensory evidence after being interrupted by the sensory noise.

A core idea of the parallel model, however, is that memory functions as its own independent source of evidence. Accordingly, we assume that memory evidence is comprised of a stream of samples that are uninterrupted by sensory noise, but still corrupted by evidence encoding noise ϵ . Crucially, we assume that memory samples are generated at a fraction of the rate of visual evidence generation, such that:

$$o_{mem_t} \sim \text{Bernoulli}(\theta_{mem}(1 - \epsilon) + \epsilon(1 - \theta_{mem})) \quad \text{if } (\gamma \mid t) \quad (3.2)$$

where γ is a positive-valued integer. For all simulations in this paper, we assume $\gamma \in \{5, 6, 7, 8, 10, 12, 15, 20, 30\}$. Relative to a visual evidence sampling rate of 60Hz, this generates memory samples 2-12Hz, which spans the low and high range of theta-oscillatory activity that has been empirically linked to task-related associative memory retrieval (Vivekananda et al., 2021; Seger et al., 2023).

Next, we posit that observers use samples from both evidence sources to form and continuously-update a Beta-distributed belief $g_s(t)$ about the value of θ_s for each source (memory and vision). Each new sample from each source updates the belief $g_s(t)$ according to:

$$g_s(t) \equiv \begin{cases} \text{Beta}(\alpha_{s_{t-1}} + 1, \beta_{s_{t-1}}), & \text{if } o_{s_t} > 0 \\ \text{Beta}(\alpha_{s_{t-1}}, \beta_{s_{t-1}} + 1), & \text{if } o_{s_t} < 0 \end{cases} \quad (3.3)$$

We define internal evidence reliability λ_{s_t} as a term that captures both the central tendency and variance of $g_s(t)$:

$$\lambda_{s_t} = \sum (g_s(t) \cdot H(X))^{-1} \quad (3.4)$$

where X is a vector of probability values ranging from 0 to 1 and H is the information entropy function. Each source's dynamically-estimated reliability value thus corresponds to the dot product of the observer's belief about that source's signal strength θ_s with the information entropy function $H(X)$. This creates a reliability estimate that quantifies how informative an ideal observer ought to believe a particular evidence source is, based on how informative it has been up to this point in the trial. We then define the weights on each evidence source w_{s_t} using the standard ideal-observer model of cue integration:

$$w_{memory_t} = \frac{\lambda_{memory_t}}{\lambda_{vision_t} + \lambda_{memory_t}}; \quad w_{vision_t} = \frac{\lambda_{vision_t}}{\lambda_{vision_t} + \lambda_{memory_t}} \quad (3.5)$$

We posit that observers weight each observation from each evidence source by these source-specific weights, creating a decision variable x_t that is a time-varying reliability-weighted sum of samples from independent streams of memory and visual evidence:

$$x_t = x_{t-1} + o_{memory_t} w_{memory_t} + o_{vision_t} w_{vision_t} \quad (3.6)$$

The decision process terminates according to

$$|x_t| > z \quad (3.7)$$

where z is a scalar defining the boundary separation (threshold) value.

3.3.1.2 Simulation model: serial process

To test whether reliability-weighted integration is restricted only to pre-evidence anticipation—rather than continuously-evolving in parallel with sensory evidence accumulation—we specified a model where reliability-weighted integration stops once sensory evidence becomes available. Visual evidence was modeled identically as in Equation 3.1. Memory evidence was again generated using Equation 3.2, but terminated once $t > t_{s1}$. We assumed that the initial value of the decision variable is given by the terminal value of the memory accumulator (Bornstein et al., 2023), such that:

$$x_{t=s1} = x_{memory_{s1-t}} \quad (3.8)$$

We then assume that the decision variable is a cumulative sum of samples from Equation 3.1, such that:

$$x_{t>s1} = x_{t-1} + o_{vis_t} \quad (3.9)$$

again with the process terminating once $|x_t| > z$.

3.3.2 Model specifications: fitting models

We used two different models with a piecewise-constant approximation of the continuously-varying drift rate specified by the parallel process. One model (“simple drift”) enforces an identical effect of cue across all trials, whereas the other (“congruence-modulated drift”) allows the effect of cue on drift to depend on the congruence of the sensory evidence with the prediction from memory. Both models assume that observers use a fixed decision threshold that is identical across levels of θ_{mem} , that the starting point of the process x_0 is fixed to 0, and that the “non-decision” time is 0. These settings allow for the most conservative statistical test of the core ideas in our dynamic integration model: that memory is sampled during the first noise (anticipation) period, that this sampling sets the starting point for visual evidence accumulation (i.e., it is not set immediately upon the start of the trial), and that time-varying effects of expectations are driven solely by modulating the drift rate (rather than, e.g., the rate of threshold collapse). Internal evidence noise/diffusion was fixed to $c^2 = 1$ for all fits.

3.3.2.1 Fitting model: simple drift

The simple drift model specifies one drift rate per epoch (Figure 3.1C) on a trial-by-trial basis, such that:

$$\nu_t = \begin{cases} \nu_{n1} & \text{if } t < t_{s1} \\ \nu_{s1} & \text{if } t \geq t_{s1} \wedge t < t_{n2} \\ \nu_{n2} & \text{if } t \geq t_{n2} \wedge t < t_{s2} \\ \nu_{s2} & \text{if } t \geq t_{s2} \end{cases} \quad (3.10)$$

where t_{s1} , t_{n2} , and t_{s2} are the onsets for the first signal period, second noise period, and second signal period respectively. Coupled with estimating the value of the boundary separation parameter B , this model has $n = 5$ free parameters.

3.3.2.2 Fitting model: congruence-modulated drift

The congruence-modulated drift model further allows epoch-level drift rates to vary as a function of the congruence between sensory and memory evidence, such that:

$$\begin{aligned} \nu_{n1} &= \begin{cases} \nu_{n1:cong} & \text{if congruent} \\ -\nu_{n1:cong} & \text{if incongruent} \\ \nu_{n1:neut} & \text{if neutral} \end{cases} & \nu_{s1} &= \begin{cases} \nu_{s1:cong} & \text{if congruent} \\ \nu_{s1:incong} & \text{if incongruent} \\ \nu_{s1:neut} & \text{if neutral} \end{cases} \\ \nu_{n2} &= \begin{cases} \nu_{n2:cong} & \text{if congruent} \\ \nu_{n2:incong} & \text{if incongruent} \\ \nu_{n2:neut} & \text{if neutral} \end{cases} & \nu_{s2} &= \begin{cases} \nu_{s2:cong} & \text{if congruent} \\ \nu_{s2:incong} & \text{if incongruent} \\ \nu_{s2:neut} & \text{if neutral} \end{cases} \end{aligned} \quad (3.11)$$

where the conditionals on ν_t are identical to Equation 3.10. Note that in the first trial epoch—the anticipation period—we constrain the value of ν_{n1} to be identical between congruent and incongruent trials. This constraint reflects the fact that observers have no information about cue-stimulus congruence during the initial anticipation period, and thus would not be able

to modulate their drift rates by this factor at that point in time. Coupled with boundary separation B , this model has $n = 12$ free parameters.

3.3.3 Model fitting

When fitting these models to data, we used the default method provided by `pyddm` which has been shown to reliably recover time-varying drift rates (Shinn et al., 2020). The method uses differential evolution to perform maximum likelihood estimation on the negative log likelihood of simulated choice-RT distributions. Both models constrained the range of possible drift rate values to $\nu = [0, 12]$ and boundary separation $B = [0.5, 12]$. Each subject was fit individually, using their specific trial-level evidence changepoints for t_{s1} , t_{n2} , and t_{s2} . Trials where no response was made were omitted prior to fitting.

3.3.4 Data simulation procedure

We used a two-part procedure to simulate the serial and parallel models on the same stochastically-varying evidence timeseries used in the experiment (see Section 3.2.3.3). First, we identified combinations of sensory evidence strength θ_{vis} and decision threshold z that resulted in 70% choice accuracy for a vision-only model when presented with noisy, but uninterrupted, sensory evidence. This mimics the calibration procedure used in human subjects to fix the time-constant reliability of sensory evidence across trials. Then, we ran simulations on the stochastically-varying sensory evidence using the paired θ_{vis} and z values to identify steady-state model behavior across a range of plausible values for γ and ϵ . The following subsections detail the procedure used in each step.

3.3.4.1 Threshold-coherence calibration

To identify pairs of sensory evidence strength θ_{vis} and threshold z that fixed vision-only choice performance to 70% accuracy, we ran 5,000 simulations of 100 trials for each combination of θ_{vis} , z , and ϵ . Visual evidence presentation rate was set to 60Hz which, for a stimulus

duration of 3000ms, resulted in $n = 180$ frame-level timesteps on each simulated trial. We simulated the full-factorial combination of $\theta_{vis} \in [0.52, 0.65]$ with a step size of 0.01 and $z \in [3, 30]$ with a step size of 0.5. and $\epsilon \in \{0.01, 0.0575, 0.1050, 0.1525, 0.2\}$. Within each $\theta_{vis} * z * \epsilon$ combination, we computed the running average probability of making a correct response across each run of 100 trials ($n = 150,000$ trials per cell). A response was classified as correct if the decision variable exceeded the positive threshold value z , or, in the case of trials where the boundary was not reached, if it had a positive sign at the last timestep of the simulation. With these steady-state choice probabilities in hand, we then identified four unique values of z that lead to 70% choice accuracy at varying levels of θ_{vis} :

sensory evidence θ_{vis}	threshold z	p(correct)
0.58	5	0.699
0.6	3.5	0.705
0.62	4	0.714
0.54	2.5	0.699

Table 3.1: Calibrated threshold-coherence tuples used to generate data from both the serial and parallel models. Proportion correct responses averaged over 150,000 simulated trials.

3.3.4.2 Simulating behavior on the experimental task

After obtaining the calibrated threshold-coherence tuples, we then simulated both the serial and parallel model on our experimental task. For each model, we simulated behavior with the threshold-coherence tuples in Table 3.1, with $\theta_{mem} \in \{0.5, 0.65, 0.8\}$, with $\gamma \in \{5, 6, 7, 8, 10, 12, 15, 20, 30\}$ and with $\epsilon \in \{0.01, 0.082, 0.155, 0.227, 0.3\}$. To capture differences in evidence encoding rate between the simulation and human subjects, we added a variable κ that slows down the rate of sensory evidence encoding identically to how γ slows memory evidence generation, and simulated both processes with $\kappa = [2, 4]$. For each of these combinations, we ran 1000 simulations of 150 trials.

3.4 Results

We developed a perceptual decision task where the relative reliability of expectations and sensory evidence varied stochastically *within* the course of a single trial. Trial-level expectation strength was set using the predictive probabilities θ_{mem} associated with each colored border (Figure 3.1A), and trial-level sensory evidence strength was identical across all signal epochs in the experiment. Thus, stochastically interrupting the visual evidence with epochs of pure sensory noise ($\theta_{vis} = 0.5$) modulates the *relative* reliability of sensory and memory evidence in a time-varying manner on each trial (Figure 3.1C). Because both starting point and serial models predict that the decision variable should plateau during this period of time, the key test of our reliability-weighted integration theory is whether the drift rate during these intermittent noise periods (ν_{n2}) is greater than 0. To test this, we fit two complementary time-varying drift rate models to the data and compared their performance against serial and start point alternatives. We also report exploratory analyses investigating whether true and estimated values of θ_{mem} (Figure 3.1B) exert dissociable effects on behavior.

3.4.1 Learned cues modulate model-free choice behavior

An initial series of mixed-effects linear regression models confirmed that the learned cues exerted canonical effects on choice-RT behavior. When modeling $\theta_{mem} = [0, 1]$, such that values less than 0.5 correspond to incongruent trials, a logistic regression with a fixed effect of θ_{mem} and random intercepts for subjects confirmed that choice accuracy increased with the predictive match between expectation and sensory evidence ($\beta = 0.838$, $z_{15204} = 4.28$, $p < .001$). Further, log-transformed and z-scored reaction times (RT) were significantly speeded as a function of θ_{mem} ($\beta = -0.20$, $t_{14991} = -4.5$, $p < .001$) as estimated by a linear model with fixed effect for cue only. Random effects for subjects were omitted due to the within-subject transformation of RTs prior to analysis. Finally, a linear model with random slopes for subjects revealed that decision confidence also increased as a function of

θ_{mem} ($\beta = 0.23$, $t_{15145} = 5.96$, $p < .001$), such that observers were generally more confident as a function of cue-evidence congruence. Together, these findings confirm that observers’ learned expectations modulated behavior on the decision task.

Because choice, RT, and confidence are known to interact with each other in evidence accumulation tasks (e.g., Kiani and Shadlen, 2009), we additionally investigated how their interactions were modulated by expectations (θ_{mem}). A linear model of RTs using θ_{mem} and choice accuracy revealed a significant interaction, such that the effect of θ_{mem} on RTs was significantly modulated by choice accuracy ($\beta_{correct} = -0.76$, $t_{14989} = -7.95$, $p < .001$). Pairwise contrasts at each level of θ_{mem} confirmed that incorrect responses were speeded by incongruent cues and slowed by congruent cues (Table 3.2). Critically, the neutral cue did not generate any differences in RT as a function of choice accuracy (Table 3.2, middle row), further demonstrating behavioral sensitivity to the strength of cue predictiveness. The full regression summary table can be found in Table B.1.

θ_{mem}	Δ_{acc}	SE	t	p
0.20	-0.23	0.04	-5.81	< .001
0.35	-0.12	0.02	-4.18	< .001
0.50	-0.00	0.02	-0.11	0.92
0.65	0.11	0.02	5.55	< .001
0.80	0.23	0.03	7.78	< .001

Table 3.2: Contrasts (error - correct) on estimated marginal mean RT (z-scored and log-transformed within subject) at each value of θ_{mem} .

We then fit a linear regression with a three-way interaction among θ_{mem} , choice accuracy, and RT as fixed effects with random intercepts for subjects. The model returned a significant interaction between θ_{mem} and accuracy ($\beta = 0.89$, $t_{14400} = 11.02$, $p < .001$), indicating that observers became increasingly confident in correct responses as a function of cue-stimulus congruence (Table 3.3). Interestingly, there was no significant interaction between θ_{mem} and RT nor a significant three-way interaction, suggesting that RT did little to modulate confidence in this task. Full regression summary is available in Table B.2.

θ_{mem}	Δ_{acc}	SE	t	p
0.20	-0.25	0.03	-7.43	< .001
0.35	-0.38	0.02	-16.25	< .001
0.50	-0.52	0.02	-30.84	< .001
0.65	-0.65	0.02	-37.46	< .001
0.80	-0.78	0.03	-31.56	< .001

Table 3.3: Contrasts (error - correct) on estimated marginal mean confidence at each value of θ_{mem} .

3.4.2 Predictive cues induce expectation-driven responding during intermittent sensory noise

Having confirmed that behavior reflects canonical effects of expectations, we then conducted the first statistical test of our reliability-weighted integration theory by fitting a response-coded version of the “simple drift” model (Equation 3.10) to human subjects data. The upper boundary corresponded to participants making a response that agrees with the cue’s prediction, and the lower boundary corresponds to choosing the alternative not predicted by the cue. Accordingly, positive drift rates reflect a choice bias in the direction of the cue. In direct support of our theory, second noise drift rate ν_{n2} averaged across participants and cues was significantly greater than 0 ($\bar{\nu}_{n2} = 1.91$, $t_{40} = 4.9$, $p < .001$). Table 3.4 shows that other drift rates during all other periods of the trial were also significantly greater than 0, indicating that observers were persistently biased toward responding according to the cue’s prediction.

epoch	$\bar{\nu}$	SEM	t	df	p
noise1	0.854	0.23	3.76	40	.001
signal1	2.21	0.40	5.53	40	< .001
noise2	1.91	0.39	4.93	40	< .001
signal2	1.84	0.29	6.44	40	< .001

Table 3.4: Fitted drift rates from the response-coded model in each epoch, averaged across subjects. P-values reflect a one-sample t-test against 0; SEM = standard error of the mean (n=41).

Next, we fit the congruence-modulated drift model (Equation 3.11) to the same choice data: whether participants made a response according to the cue’s prediction. This more flexible model allows us to test whether the probability of responding according to cue varies between congruent and incongruent trials, and addresses potential concerns that the previous drift rate was driven by the necessarily unequal proportion of congruent trials relative to incongruent. Critically, two-sample t-tests returned no significant difference in fitted drift rates between congruent and incongruent trials, indicating that observers were equally likely to respond according to the cue’s prediction regardless of whether that prediction accorded with the correct answer as defined by sensory evidence (Table 3.5). Further, the model specifying equal effects across trial types was preferred both by AIC ($\Delta_{same-cong} = -111.77$) and BIC ($\Delta_{same-cong} = -457.67$), suggesting that the additional complexity introduced by allowing drifts to vary by congruence was not warranted given the data.

epoch	$\Delta\bar{\nu}_{cong-incong}$	t	95% CI	df	p
signal1	1.02	1.72	[-0.15, 2.19]	77	.089
noise2	0.39	0.65	[-0.80, 1.58]	77	.515
signal2	0.23	0.47	[-0.74, 1.20]	77	.6421

Table 3.5: No difference in fitted drift rates for the probability of making a cue-congruent response between congruent and incongruent trials.

3.4.2.1 Threshold values reveal distinct response strategies

Our theory further makes the prediction that the drift rate during periods of sensory noise should increase with the predictiveness of the expectation θ_{mem} . Testing this prediction, however, requires that drift rates are on approximately the same numeric scale. Inspection of the fits revealed that (1) fitted boundary values exhibit a strong positive skew (Shapiro-Wilk $W = 0.79$, $p < .001$; Figure B.1B) and (2) fitted drift rates across all epochs are positively correlated with fitted boundary (all $r \geq 0.6$, all $p < .001$; Table B.3). Accordingly, we median-split our sample into subjects with relatively low vs. relatively high fitted boundary values in order to investigate effects of θ_{mem} on drift rates ν . Figure 3.2A shows that participants with low vs. high thresholds appeared to take qualitatively different approaches

to completing the task, and that these trends are preserved across both experimental conditions. Observers with higher thresholds generated RT distributions that are clearly bimodal, with peaks during the second noise period (shaded gray region between $x=1$ and $x=2$, left panel of Figure 3.2A) and the second signal period (white area after shaded region). Hartigan’s dip test on RTs in the high threshold group only rejected the null hypothesis of unimodality, providing statistical evidence for a bimodal response strategy ($D_{highThresh} = 0.24$, $p < .001$). The same test returned a null result when applied to the low threshold group ($D_{lowThresh} = 0.003$, $p = 0.98$), further corroborating a different response strategy between the two groups. Observers with a low threshold appeared to treat the sensory evidence as one continuous “stream,” with responses also peaking during the second noise period (Figure 3.2A, right panel). Interestingly, approximately one-fourth of responses were made during the modal duration of the first noise period (shaded gray region from $x = 0$ to $x = 0.75$), corroborating a similar finding reported by Bornstein et al. (2023).

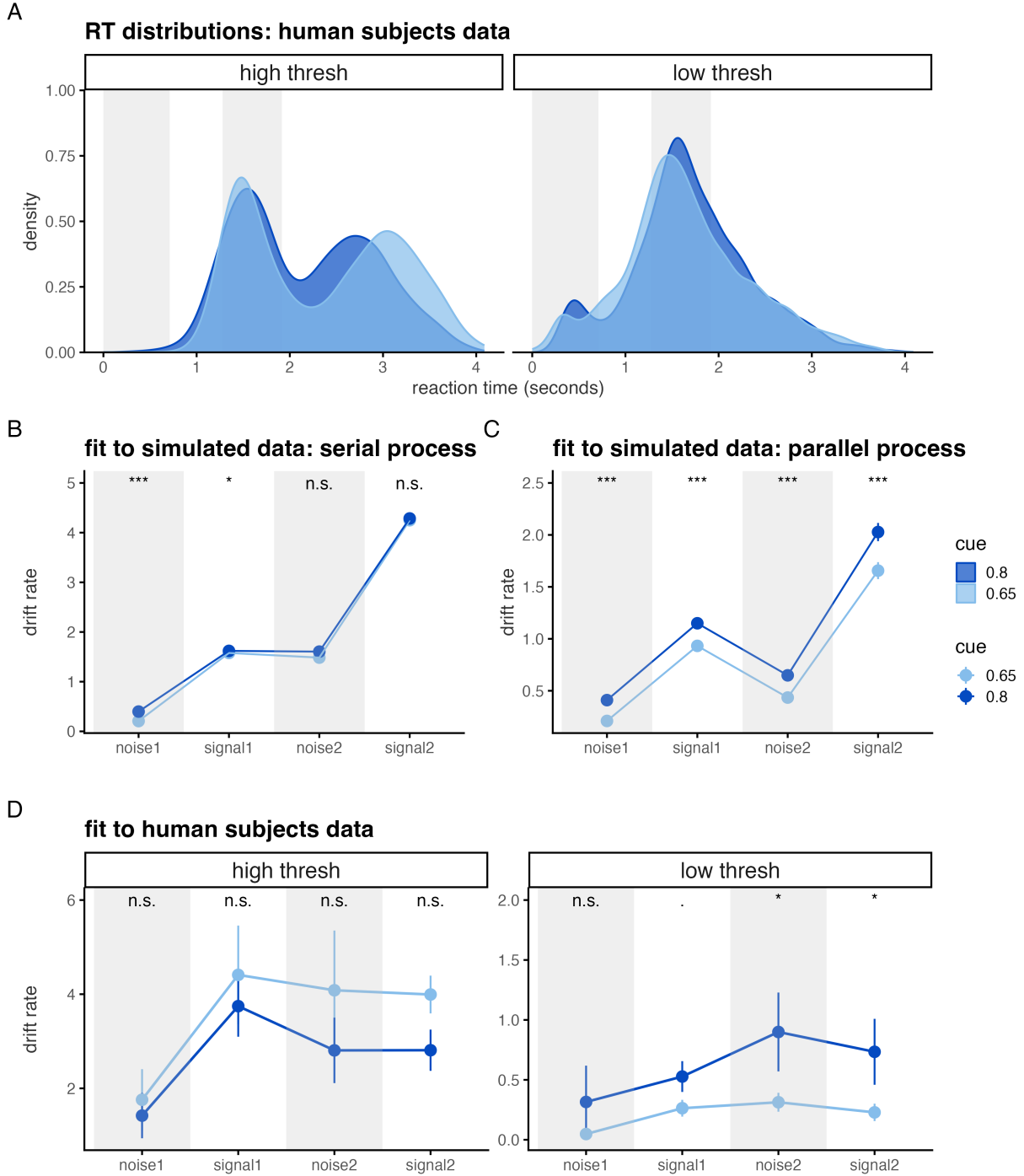


Figure 3.2: Response-coded model fits to simulated and human-subjects data. Shaded gray areas depict modal onsets and durations of the sensory noise periods, which followed a constant hazard rate across trials. Note the difference in y-axis values across Figures B-D. (A) RT distributions split between low and high threshold groups. (B) Fit of the response-coded model to data generated by a serial process. (C) Fit of the response-coded model to data generated by our parallel reliability-weighted process. (D) Fit of the response-coded model to human subjects data, split by fitted threshold value. *** = $p < .001$, ** = $p < .01$, * = $p < .05$, · = $p < .1$.

3.4.2.2 Threshold values are robustly recoverable for both serial and parallel processes

Because our theory is agnostic about how threshold values should be set, we first assessed the recoverability of the threshold parameter using data simulated according to the procedure described in Section 3.3.4. For each combination $\epsilon * \gamma * \kappa * \theta_{vis}$, we applied the response-coded simple drift model to the binary outcome of whether a cue-congruent response was made on that trial to 10 unique runs of 150 trials. Across all runs and combinations, the correlation between generating and fitted threshold was strong for both the serial ($r = 0.912$, $t_{10346} = 225.91$, $p < .001$) and parallel ($r = 0.894$, $t_{10797} = 207.28$, $p < .001$) processes. Thus, regardless of whether observers in our task deployed the parallel process under investigation or the serial process previously reported by Bornstein et al. (2023), our fitting model should reliably estimate group-level differences in the threshold/boundary separation parameter.

3.4.2.3 Cue strength modulates drift rates in observers with a low threshold

Now that we have shown that our model reliably discriminates high vs. low boundary values across data-generating sources, we can investigate cue-driven differences in drift rates within the low and high threshold groups. Figures 3.2B and C display fits of the response-coded drift rate model to data generated from the serial and parallel processes, respectively, and Figure 3.2D displays fits to human subjects data. To assess the effect of θ_{mem} on drift rate in the simulated data, we first averaged the fits in each epoch across each of $n = 10$ independent fits per $\gamma * \epsilon * \kappa * \theta_{vis}$ cell, and then used linear regression with fixed effects of each binning variable as control regressors. For the human subjects data, we ran separate linear models with a simple effect of θ_{mem} only on data pooled within each threshold group (low vs. high).

The key test of interest for our theory is whether human observers' fitted drift rates during the second noise period are significantly modulated by θ_{mem} (or "cue"). As shown in Figure 3.2D, this effect was observed in the low threshold human subjects group (right; $\beta = 3.91$,

$t_{18} = 2.26$, $p = .037$) but not the high threshold human subjects group (left; $\beta = -8.51$, $t_{19} = -0.96$, $p = 0.35$). Interestingly, this trend carried over into the second signal period, again with the low threshold group showing a significant effect of cue ($\beta = 3.37$, $t_{18} = 2.29$, $p = .035$) that was not present in the high threshold group ($\beta = -7.89$, $t_{19} = -1.73$, $p = .101$). Taken together, these findings indicate that observers with lower threshold values were not only increasingly more likely to choose the cue-predicted option during the second noise period, but also *after* the signal became available again post-noise. This finding lends support both to our reliability-weighted integration theory, and conceptually replicates the original Hanks et al. (2011) finding in a substantially different perceptual decision task.

3.4.3 Threshold modulates model-free effects of learned expectations on choice behavior

Given that neither our own nor Hanks et al.’s (2011) theory offer guidance on interpreting qualitatively different choice dynamics as a function of response threshold, we conducted another series of linear regressions to investigate one possible interpretation of the group-level differences in threshold: that observers with higher thresholds for making a cue-based response relied primarily on sensory information to make their choice. If this is the case, then we ought to observe that the model-free signatures of expectation integration reported in Section 3.4.1 differ between the low and high threshold groups. As shown in Figure 3.3, this is precisely what we found. A linear regression on raw RT returned a significant interaction among between threshold group, accuracy, and θ_{mem} ($\beta = -0.31$, $t_{14985} = -2.18$, $p = .029$). This interaction can be interpreted visually by contrasting the effect of cue on error RT density in Figure 3.3A. Whereas both threshold groups demonstrated some amount of slowing as a function of cue-evidence congruence ($\theta_{cue} \rightarrow 1$), this difference only reached statistical significance in the low threshold group, consistent with an interpretation that these observers more strongly incorporated expectations into their decision process.

Further evidence comes from the regressions on choice accuracy and confidence presented in Figures 3.3B and C. The low threshold group (light green) exhibited a significantly steeper increase in choice accuracy as a function of θ_{mem} ($\beta = 0.66$, $z_{14989} = 3.4$, $p = .001$). This trend is particularly notable because the sensory evidence was calibrated to a level that fixes choice accuracy to 70%, which is the value the high threshold group clustered around for all values of θ_{mem} . It follows, then, that the low threshold group exhibited the drop in choice accuracy because they were more strongly weighting the evidence from their highly incongruent expectations when making the perceptual decision.

One final piece of evidence comes from a significant three-way interaction among θ_{mem} , choice accuracy, and threshold in predicting choice confidence. Whereas both the high and low threshold groups demonstrated an increase in choice confidence as a function of cue-evidence congruence (Figure 3.3C, right panel), the low threshold group became increasingly *less* confident in their *incorrect* decisions, the high threshold group was uniformly unconfident in their incorrect decisions ($\beta = 0.75$, $t_{14943} = 4.18$, $p < .001$). This interpretation was confirmed on post-hoc contrasts on the interaction term, which returned a significant difference in the slope over θ_{mem} between the groups in *both* incorrect ($\Delta_{high-low} = 0.5$, $t_{14943} = 3.482$, $p < .001$) and correct ($\Delta_{high-low} = -0.22$, $t_{14943} = -2.31$, $p < .021$) decisions. This finding lends yet further support to the idea that the low threshold group more strongly weights expectations in their choice behavior relative to the high threshold group.

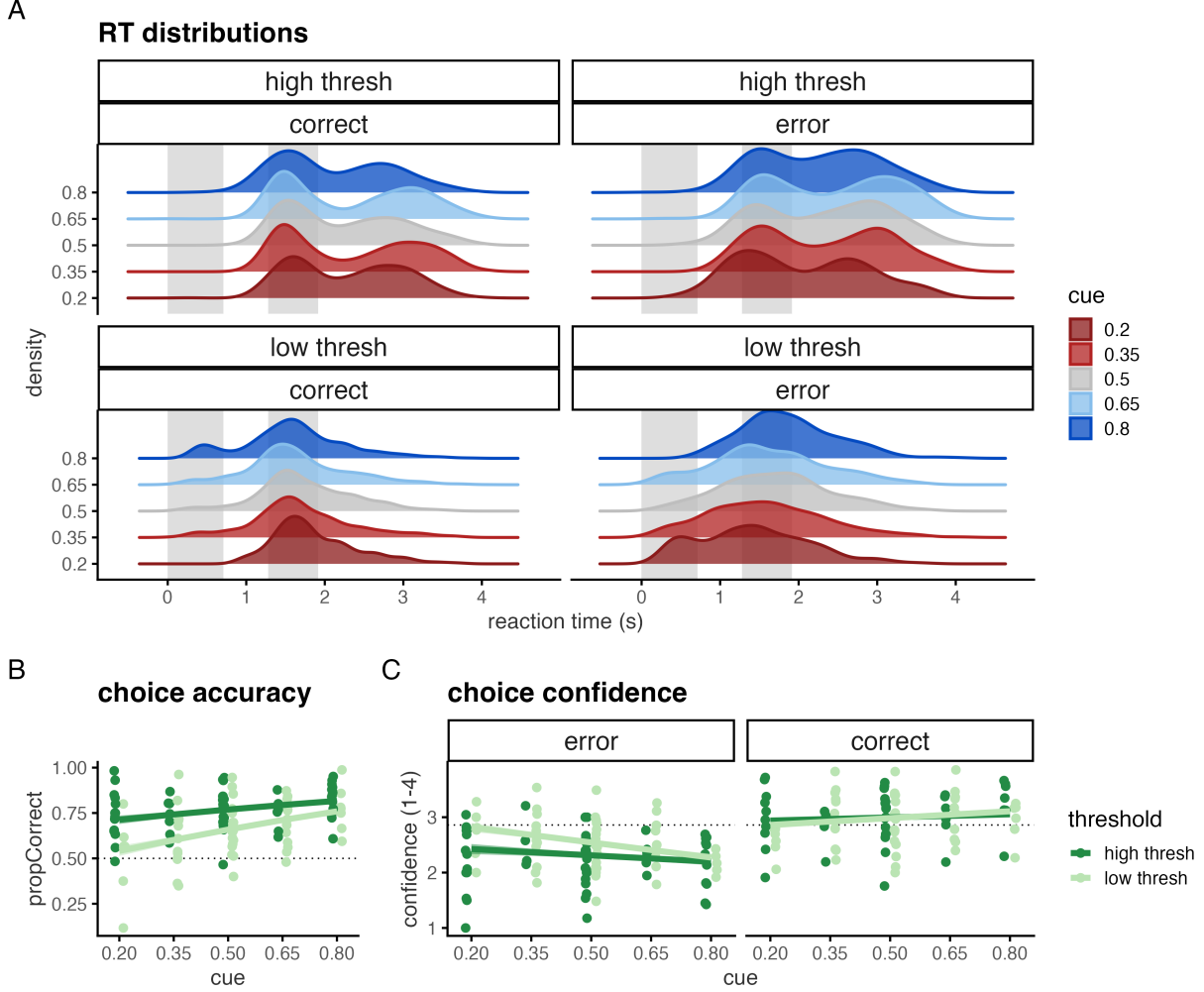


Figure 3.3: Fitted boundary values modulate model-free signatures of expectation integration. (A) RT distributions demonstrating that group-level differences in choice dynamics are conserved across correct and error response. Gray regions depict modal onsets and durations of sensory noise periods. (B) Observers with lower thresholds demonstrated greater accuracy gains as a function of θ_{cue} . (C) Confidence in incorrect decisions is more strongly modulated by θ_{cue} for observers with lower thresholds.

3.5 Discussion

We developed a cue-guided perceptual decision task that manipulated the *relative* reliability of expectations and sensory evidence within a single decision by stochastically interrupting the sensory evidence with periods of pure sensory noise. Critically, the onsets and durations

of these noise epochs followed a fixed hazard rate across trials, meaning that observers could not use elapsed time as a proxy for when sensory evidence would cease being informative. This design allowed us to conduct a behavioral test of our time-varying reliability-weighted alternative to Hanks et al.'s (2011) explanation for dynamic effects of expectations, which posits that observers use elapsed time to infer the reliability of sensory evidence on a trial-by-trial basis.

In direct support of our theory, we found that observers exhibited significant cue-based responding during these periods of intermittent sensory noise (Figure 3.2). Further, this increased probability of cue-based responding persisted through the end of the trial, conceptually replicating the original finding of Hanks et al. (2011) under conditions of substantially greater uncertainty. Interestingly, these trends were significantly modulated by the threshold that observers adopted when performing the decision task. Only observers with thresholds below the median fitted value demonstrated behavioral signatures of dynamic effects of expectations. Observers with higher thresholds had higher drift rates overall, but they were not modulated by the predictiveness of the cue. Further, their reaction time distributions exhibited significant bimodality that was not observed in the lower threshold group, suggesting that the boundary parameter distinguished between two classes of strategies that observers used to complete our task.

What might those strategies be, and how do they relate to the psychological questions of interest? One possibility is that observers with higher boundary values opted to rely more on vision than memory for performing this task. We conducted a series of additional regression analyses to investigate this possibility, and found that threshold group significantly modulated the effect of expectations on all three behavioral measures collected in the experiment (choice accuracy, RT, and confidence). The form of these interactions across all three measurements lend converging support to an interpretation that higher thresholds for making cue-based responses were associated with dampened behavioral signatures of expectation-

guided choice (Figure 3.3) or—by the same token—increased time-constant weighting of sensory information in the decision process.

These findings raise a number of interesting questions to be addressed by future research. Do these group-level differences in evidence weighting generalize to other tasks, in the laboratory or beyond, such that they reflect stable individual differences in dispositions to attend to internal versus external information? Or do they reflect rational, context-sensitive adjustments based on metacognitive monitoring of learning, memory, and decisional processes? Although the specific model we test in this article is agnostic with respect to these questions, the broader theoretical framework to which it belongs—experience-dependent uncertainty as a normative constraint in and of itself—suggests that these two possibilities are not mutually exclusive. Individual differences in weighting of internal and external information can—and *should*—reflect the expected value of acquiring information from that source (Landy et al., 2011; Wang et al., 2021; Sutton and Barto, 2018; Kobayashi et al., 2021). These “unimodal” expectations about how reliable internal evidence is relative to the information available from the sensory environment are learned and updated across multiple timescales, from early-life experiences that scaffold our development over time (Harhen and Bornstein, 2024) to highly localized, task-specific beliefs formed in laboratory settings (Rouault et al., 2019). Elucidating the neural and computational principles driving—and constraining—the learning and integration of these expectations can thus advance basic scientific understanding of how learning, memory, and perceptual systems work in concert to facilitate flexible, robust, and adaptive behavior.

3.6 Conclusion

We have shown that learned expectations exert time-varying effects on perceptual decisions. Contra early starting point models, which posit a time-constant, pre-decisional locus of expectation-integration, we provided experimental evidence supporting a theory whereby

expectations function as their own independent source of internal evidence that observers sample from *in parallel* to sensory information. Critically, our theory posits that expectations and sensory evidence are combined according to a time-varying estimate of the *relative* reliability of the two evidence sources, where reliability is defined as how *informative* that source is with respect to the ongoing decision.

While these findings constitute preliminary empirical support for the theory, they are also limited in their scope. First, although we demonstrated statistical support for the key signature of our theory, we did not compare its quantitative performance against other plausible fitting models (e.g., enforcing that expectations have their effects only at the starting point or during the anticipation period). Ongoing work specifying appropriately-parameterized alternative models, and determining the relevant methods of model comparisons, will address this gap. A second limitation of the current findings is that effects of cue strength were manipulated between-subjects, potentially confounding effects of expectation strength with the overall uncertainty in the decision environment. We plan to run a follow-up study manipulating expectation strength within-subjects, and also including some proportion of trials where the second noise period continues on until the decision times out. This latter manipulation ought to minimize the prevalence of “waiting” strategies observed in the high-threshold group, thus allowing for a more statistically powerful test of our hypothesis.

Chapter 4

Reasoning goals and representational decisions in computational cognitive neuroscience: lessons from the drift diffusion model

4.1 Introduction

Formal theorizing via computational modeling is often lauded as a solution to meta-scientific challenges in the brain and behavioral sciences (Forstmann et al., 2011; Grahek et al., 2021; Guest and Martin, 2021; Muthukrishna and Henrich, 2019; Press et al., 2022; Robinaugh et al., 2021; Turner et al., 2019b). While formal computational models indeed have much to offer toward this end, we also recognize them as incredibly powerful and flexible tools that—when misused—can create even worse problems that they were initially applied to solve. Further, subfields of psychology that rest upon decades of formal modeling work face their own meta-scientific challenges that—by their natures—cannot be resolved by use of a model alone.

In this article, we aim to demonstrate the utility of philosophical research for mitigating and/or resolving such meta-scientific challenges. To do this, we have compiled insights from foundational and recent work in philosophy of modeling to develop a “philosophical toolkit” for computational cognitive modeling. We take as an applied example sequential sampling models of choice-reaction time, with a focus on the drift diffusion or diffusion decision model (DDM).

A number of factors motivated our goal to focus on sequential sampling models. First, this family of models, and the DDM in particular, is highly prominent in computational cognitive (neuro)science. This fact, together with the aforementioned goal of integrating formal models into psychological research, has created demand for a number of different software packages that increase the accessibility of these formal modeling tools (e.g., the “EZ-DDM” by Wagenmakers et al. (2007) and the “PyDDM” by Shinn et al. (2020)). A philosophical and conceptual analysis of how these models work and what factors need to be considered when using them both supports newer researchers’ entry to this research area and helps mitigate confused applications or erroneous conclusions that might arise in the course of learning how to use these tools.¹ We further hope that our toolkit offers experienced researchers an opportunity to reflect on their own modeling practice and the goals that shape it. Finally, we aim to contribute to the philosophical literature a rich case study about a family of models that have been almost entirely overlooked by the discipline. Whereas the merits and limitations of, e.g., connectionist and Bayesian models have received substantive philosophical attention, the DDM has only featured marginally in recent work, primarily as an example of cognitive models more broadly (Figdor, 2018; Drayson, 2020; Gamboa, 2024). Thus, we hope this introduction to the DDM and its broader model family can serve as a starting point for future philosophical research about the role(s) of these models in computational cognitive (neuro)science.

¹This article is not intended as a guide for software usage. Conceptual and practical considerations related to implementing computational cognitive models can be found in Cooper and Guest (2014), Lee and Wagenmakers (2014), Lin and Strickland (2020), and Wilson and Collins (2019).

We first provide readers with a collection of philosophical tools for thinking about the practice and products of computational cognitive neuroscience research. We then demonstrate the utility of such a “toolkit” in two ways: philosophically introducing the DDM, and then presenting a novel conceptual analysis of a long-standing debate about the form of decision boundaries in the DDM (the collapsing bounds debate; Figure 4.1). This analysis is also one of the first contributions to the debate that does not aim to argue in favor of one form versus the other. Rather, by explicating the reasoning goals motivating each “side” in the debate, we argue (1) that the forms offer complementary rather than competing explanations and (2) that the difference in those explanations reflect the different aims of two groups of users (cognitive psychometricians and theoretical decision/neuroscientists). We conclude our analysis with a principled heuristic for choosing between forms of the model implicated in the debate.

4.2 Philosophical tools for thinking about computational cognitive models

The term “model” is used to refer to a number of related but functionally distinct objects involved in scientific research. One thing common to most models is that they are *representations* of a **target**—some phenomenon in the natural world—that make it easier for a human (or group of humans) to reason about that target (e.g., van Rooij, 2022). This article focuses on models that are abstract representations of the latent entities and processes generating an organism’s behavior, often called *cognitive* or *process* models. Scientists use these models both to **predict** how an organism (or agent) will behave in response to experimental manipulations, and to **explain** why a particular manipulation induced the behavioral response that it did. In some cases, these models are also tasked with predicting and explaining changes in neural activity related to changes in behavior. Finally, scientists sometimes use these models simply as measurement devices for quantifying individual differences in latent properties or

processes that are relevant for a broader type of explanation. Thus, the *same* model can be used in the service of different goals (prediction, explanation, measurement) directed toward targets at different conceptual/spatiotemporal scales (cognitive processes and/or neural activity). This multiplicity is both what makes models so useful for scientific practice ***and*** is a primary factor contributing to confusion and debates about how best to use models in scientific research. Reviewing some philosophical research concerned with the questions of what models are and how they work can help mitigate these confusions while also offering a salve against the “existemic”² concerns that model-based cognitive (neuro)science research can so frequently incur.

Based on the uses described above, the types of models we consider in this article are simultaneously (i) abstract representations of latent entities and processes that generate an organism’s behavior, (ii) theories that predict and explain changes in the organism’s behavior, and (iii) devices for measuring the latent entities and processes represented in the model. Based on the topics discussed in our case study, we focus our attention on properties (i) and (ii), with a particular focus on how these properties relate to model-based explanations. Following Weisberg (2013), we consider these models *computational* if their theoretical content appeals to transition rules or an algorithm (i.e., a series of steps specifying how an initial state is transformed into an output)³. Importantly, these properties do not apply to all of the models used in the behavioral and brain sciences, and certainly not to all of the objects reasonably considered scientific models.⁴ But they do aptly characterize

²We coin “existemic” to refer to feelings of existential dread brought about by reflecting on the epistemic limitations of particular ways of knowing about the world.

³Weisberg (2013) considers computational models a special case of mathematical models because of how scientists use them to explain. On his account, mathematical models use sequences of states or an equilibrium as the explanation, whereas computational models use the procedure that specifies how an input is transformed into an output. See also Richmond (2024) for a complementary perspective.

⁴For example, some models are concrete representations, like scale models used in architectural engineering and rodent models used in translational neuroscience research. Other models are abstract representations that compactly summarize many observations about a target process, or describe how a target responds when it is intervened upon, rather than comprising theories that can be used to explain why the target responded how it did. These types of models are often called “descriptive”, “phenomenological”, or “data” models.

the use of sequential sampling models in contemporary computational cognitive neuroscience (see also Forstmann et al. (2016), and Turner et al. (2019a) for additional examples). The rest of this section introduces the philosophical research that has clarified our own thinking about these models. Readers primarily interested in learning about the DDM can skip ahead to Section 4.3.

4.2.1 Tools for thinking about models as representations

Our toolkit begins with a survey of some philosophical literature concerned with the representational nature of formal models, with a focus on concepts most relevant to our case study in the next chapter.

A large body of research spanning philosophy of science, philosophy of mind, and philosophy of language is devoted to explicating the concept of representation; Frigg and Nguyen (2021) offer a helpful overview for readers interested in the role of representations in science. For our purposes, it suffices to say that a model becomes a representation of a target when a **model user**—a human who builds or interprets the model—*decides* that they are going to map components and processes of the target onto components and processes of the model, ultimately with the goal of using properties of the model to reason about corresponding properties of the target (Winsberg and Harvard, 2024). On this view, a model represents a target by “standing in” for the target in a way that permits the model user to reason about their target on the basis of interacting with the model. This process has been called surrogate reasoning in the philosophical literature. When coining this term, Swoyer (1991) offers a helpful example:

By using numbers to represent the lengths of physical objects, we can represent facts about the objects numerically, perform calculations of various sorts, then translate the results back into a conclusion about the original objects. In such

cases, we use one sort of thing as a surrogate in our thinking about another (p. 87).

Computational cognitive models use mathematical objects (numbers, distributions, vectors, etc.) and equations to represent the cognitive entities and operations that generate an organism’s behavior. Scientists thus use these mathematical representations as surrogates for thinking about what processes “under the hood” are driving behavior. Recent work in philosophy helpfully distinguishes between a model’s structure and its construal: the first being the mathematical objects and equations comprising the representation, and the second being how those abstract objects are meant to be mapped onto objects and processes in the physical world, respectively (Weisberg, 2013; Andrews, 2021). The simple example above demonstrates how construals are necessary for interpreting the structure of a model. If a user was presented only with the set of numbers that correspond to the length of physical objects, but was not told what properties of the physical world those numbers correspond to, they would (1) have little to no guidelines for constraining the types of mathematical reasoning appropriate to perform with that representation and (2) have no way of translating the results of even simple mathematical operations on the representation into actions in the real world. Construals thus function as a “bridge” between the mathematical domain where we construct and manipulate a representation of the target and the physical domain wherein we intervene upon and collect measurements from it. As such, they are central to the practice of model-based scientific research.

On Weisberg’s (2013) account, a model’s construal consists of three components: assignment, scope, and fidelity criteria. The *assignment* of a model explicitly specifies how parts of the target system are mapped onto parts of the model; in other words, it states what each variable in the model is supposed to correspond to inside an organism’s head. The *scope* of a model specifies which aspects of the target the model aims to represent (and, by extension, which aspects it does *not* aim to represent). Considering a model’s intended scope is essential

for determining which types of measurements from the target are meaningful for assessing the performance of the model. The final components of a construal, fidelity criteria, are the benchmarks that model users reference in order to determine whether they have a “good” model of their target. Our case study in the next chapter demonstrates how differences in fidelity criteria between different subgroups of researchers using the DDM have resulted in current “competing” forms of the model.

A point we wish to emphasize is that both a model’s structure and its construal are the products of *decisions* made by users who build and apply the model to data. And just like in cognitive decision-making, these ***representational decisions***—determining *what* properties of the target to include in the model and *how* to represent those properties mathematically (Harvard and Winsberg, 2022)—are subject to variation across users in different subfields, across users within the same subfield, and even within a single user applying the same model in different contexts. We argue that this variability reflects the multifaceted roles that models play in scientific reasoning, and that understanding the factors contributing to variability in representational decision-making is essential for informed, critical engagement with model-based scientific research. Echoing recent work in philosophy (Danks, 2015; Pottochnik and Sanches de Oliveira, 2020; Weisberg, 2013; Winsberg and Harvard, 2024) and computational neuroscience (Blohm et al., 2020; Kording et al., 2018), we aim to demonstrate how model users’ ***reasoning goals***—which aspect(s) of the target they aim to reason about and how they wish to perform that reasoning via the model—are primary drivers of variability in representational decisions. Because representational decisions determine both a model’s structure and its construal (which itself specifies fidelity criteria), this position implies that reasoning goals shape the model-based research process all the way down to quantitative model comparison procedures.

One might worry that permitting even quantitative model comparison procedures to vary according to a user’s goals might be too flexible of a philosophy of modeling for the prac-

ting scientist. On our view, however, this position is an inevitable consequence of the fact that most—if not all—formal models are *incomplete* representations of their target systems. Because this property of models makes them technically false with respect to their target (Frigg and Hartmann, 2020; Wimsatt, 1987), model users cannot rely simply on the “truth” or falsity of models in order to select the best among them. A prominent line of reasoning in philosophy thus suggests that models be evaluated by their adequacy for purpose rather than verisimilitude (i.e., true or accurate representation of the target) alone (e.g., Parker, 2020)⁵. By positing that the goal of modeling is to identify representations that are useful, rather than strictly-speaking true, the adequacy-for-purpose view further emphasizes the central role of reasoning goals in model-based science: they define the purpose against which a model’s adequacy is evaluated. Because different models are built and/or applied in the service of different reasoning goals, it is desirable to allow fidelity criteria to vary according to the reasoning goal being pursued. The task of model-based science thus becomes identifying which representations are adequate for which purposes, and, at a higher-level, identifying the reasoning goals that are most useful for making progress on particular scientific questions.

Further motivation for the adequacy-for-purpose view comes from the heavy use of idealization in formal models of complex systems (Cartwright, 1983; Potochnik, 2018). When scientists build idealized representations of their target systems, they build representations that intentionally misrepresent features of their target in order to make reasoning about it easier. It can be useful to distinguish between *omissive* and *disortive* idealizations: those that remove certain properties of the target from the representation (e.g., choosing not to represent neural dynamics in a cognitive model) and those that represent known properties in a way that is known to be inaccurate (e.g., assuming that observers have perfect knowledge of the environment), respectively. These complementary forms of idealization—at play in nearly all formal models—permit users to build representations that “selectively attend”

⁵Readers might already be familiar with this position on the basis of the popular aphorism “all models are wrong, but some are useful” (Box and Draper, 1987, p. 424)

to components of the objects that the model user wishes to reason about (Portides, 2021). In doing so, the model reduces the complexity of the target such that reasoning about it becomes more tractable for the model user. The task of modeling, again, becomes identifying the degree and type(s) of idealizations that are most useful for one’s purposes. The rest of this chapter aims to provide tools for thinking about this set of decisions.

4.2.2 Tools for thinking about models as explanations

Equipped with some tools for thinking about models as representations, we next turn to research in computational neuroscience and philosophy to consider the types of reasoning goals enabled by formal models. Kording et al. (2018) helpfully identify twelve different reasoning goals most commonly pursued in computational neuroscience, while noting that “it is impossible to produce an exhaustive list” (p.3). Wimsatt (1987) also proposes “Twelve things to do with false models” (pp. 7-8) based on his work with formal models in engineering. That these two lists minimally overlap both reflects how broad the space of possible reasoning goals is in model-based science, and highlights the utility of integrating insights developed independently in neuroscience and philosophy. Based on the argument we develop in the case study, we will focus our discussion here on reasoning goals related to explanation.⁶ First we discuss how formal models give explanations of empirical targets and then contrast different types of explanations formal models can give.

Foundational work in the philosophy of scientific explanation distinguishes between an observation or phenomenon that scientists aim to explain (i.e., an *explanandum*) and the explanation that scientists give of it (i.e., the *explanans*; Hempel and Oppenheim, 1948).⁷ When a formal model is used as an explanation for some explanandum, the model’s *structure* func-

⁶Explanation is one of the oldest topics in philosophy of science and thus cannot be exhaustively covered here. Helpful introductions to ongoing philosophical work on explanation in the cognitive and brain sciences can be found in Kaplan (2017) and a special issue edited Colombo and Knauff (2020).

⁷Whether an explanandum is different from a model’s target depends on what the user construes as the target of a particular model. In the case when the target is a particular pattern of observations observed in particular experimental contexts, there is no meaningful difference between a target and an explanandum. However, when the target of a model is a general (neuro)cognitive process (e.g., speeded two-alternative deci-

tions as the explanans (Weisberg, 2013). In other words, the target’s behavior is explained by virtue of the structure of its formal/mathematical representation. Weisberg (2013) argues that, in computational models, this structure is comprised of the procedures (or algorithms) that specify how an input state is transformed into an output, a position we adopt here as well. On this account, computational cognitive models explain the behavior of their targets by relating (or “mapping”) changes in observed behavior onto changes in the parameter values and/or configurations of the latent procedures represented in the model’s structure. This mapping can be achieved both by simulating the target’s behavior on different parameter values or configurations of the model’s structure (i.e., “simulation”) and/or by identifying parameter values of components of the structure that maximize the likelihood of observing a particular set of observations from the target (i.e., “model fitting”).

At the heart of both these approaches to model-based explanation is the notion of a model “capturing” properties of its target. Weisberg (2013) proposes that scientists use two different types of *fidelity criteria* to assess whether—and in what ways—a model “captures” features of its target. Dynamical fidelity refers to the quantitative similarity between measurements taken of the target and numerical estimates generated by the model (e.g., its numerical “goodness-of-fit”), whereas representational fidelity refers to how closely a model’s structure matches the causal structure of the real-world phenomenon (Weisberg, 2013). For our purposes, the relevant sense of “causal structure” is synonymous with the “type of explanation” a user wishes to attain by reasoning with the model, a topic we treat in the next paragraph. We will call our notion “explanatory fidelity” to differentiate it from Weisberg’s (2013) causal notion of representational fidelity. Importantly, these notions of fidelity are defined independently of the methods used to evaluate them, meaning that model users have to make representational decisions in order to evaluate the fidelity of their model: decid-

sion making), then particular observations of the target in particular contexts constitute different explananda that the model is tasked with unifying into a single generative formal structure.

ing what kind of criteria are most appropriate for their overall goals and how they wish to evaluate their model with respect to those selected criteria.

Our case study in the next chapter demonstrates the utility of fidelity criteria as tools for thinking about model-based research. In particular, we highlight how the two “competing” forms of the DDM (with and without collapsing decision boundaries) reflect differences in fidelity criteria between communities of scientists using the model, and argue that these differences reflect the different explanatory aims of each community. To do this, we make use of a popular taxonomy that distinguishes what, how, and why approaches to giving explanations (Dayan and Abbott, 2005; Ross and Woodward, 2023). The “what” approach focuses on characterizing how the target behaves under various circumstances, and is commonly thought to be the goal of descriptive models that are not generally thought to be explanatory. The “how” approach focuses on decomposing a target’s behavior into its constituent parts and processes, and can be said to explain the behavior of the target by virtue of those constituent components; this is commonly considered as the goal of mechanistic models in neuroscience (Craver, 2007; Dayan and Abbott, 2005). The “why” approach focuses on identifying properties of the target that require it to behave in particular ways under particular circumstances; this is how we understand the goal of normative models in neuroscience (Anderson, 1990; Dayan and Abbott, 2005).

In computational neuroscience, labels from the above taxonomy are commonly used to refer to the entirety of a model’s structure. Our case study, however, demonstrates that these properties can also be applied to individual components of a model’s structure, such that a particular stage in the input-transformation process can be normative (or mechanistic, or descriptive) even if the structure as a whole is not. Although this perspective might complicate usage of the taxonomy, it reflects growing agreement among philosophers that not all components of a model need to contribute to the user’s ultimate purpose in the same way (Weisberg, 2013). Some model components, for example, are included simply because

the fitting process would be intractable without them; they are thus included for practical reasons. Other components might be included because a user thinks that component is important for their ability to reason about their target with the model, but particular details about the component might be irrelevant for how it ultimately figures into the primary reasoning goal (e.g., the non-decision term in the DDM). Ideally, model users are explicit about how they intend for each component of their model to be construed with respect to their target, and models are constructed in such a way that these “convenience variables” are not central to the explanation a particular model provides. But because this is not always the case, it is crucial that users keep the heterogeneity of justifications for representational decisions in mind when engaging with the products of model-based research.

4.2.3 Tools for thinking about models as normative explanations

This final section of the toolkit provides a brief conceptual analysis of the practice of normative modeling. Our motivations for doing so are (1) to equip readers with the necessary resources for engaging with key topics in the case study, (2) to articulate our own perspective on the role of normative models in computational cognitive neuroscience, and (3) to give readers tools for thinking about a concept that figures heavily in both scientific and meta-scientific debates about perceptual decision making (e.g., Rahnev and Denison, 2018).

Like the term “model”, the term “normative” has a number of different but closely related meanings. In general, a normative statement is one that prescribes a particular course of action in a particular scenario. These statements (or logical attitudes) thus take the form: *if* your goal is X, then you *ought* to Y. Norms thus create a standard or benchmark for evaluating behaviors as good/correct or bad/incorrect. Crucially, as the logical form makes clear, the accuracy or “goodness” of a particular action is determined relative to a particular goal. Normative models in computational cognitive science leverage this property to offer

explanations about why the target behaved as it did. Users achieve this goal by formally specifying

1. an *objective function*: the problem the target is trying to solve (e.g., choosing between two options in the shortest possible amount of time)
2. the *environment* in which it is tasked with solving that problem (e.g., one with or without feedback after each choice)
3. the *procedure* that a target uses to solve that problem (e.g., sequential sampling to a fixed threshold).

Optionally, users can add a fourth component capturing any *constraints* they wish to impose on the procedure (e.g., leakage or loss of accumulated evidence over time)⁸. This comprehensive formal representation allows users to identify the logical or mathematical limit of the specified target’s ability to solve the specified problem in the specified environment, i.e., *optimal* behavior on the experimental task. A common explanatory approach is thus to state that the target exhibited a particular pattern of behavior because it is the optimal solution to the formally-specified problem.

This type of optimality explanation is made possible by *formal frameworks*, which can be defined a set of axioms/postulates (i.e., statements accepted as true), formal/mathematical objects, and rules constraining the relationships among those objects. In other words, formal frameworks offer model users a “grammar” for expressing questions and answers in a format that permits quantitative assessment (Guest and Martin, 2021; Press et al., 2022). Sutton and Barto’s (2018) textbook on reinforcement learning offers a helpful example⁹:

⁸Identifying normative solutions to formal problems that represent different types of constraints a user wishes to incorporate into their model is the premise of the resource-rational approach to cognitive modeling (Lewis et al., 2014; Lieder and Griffiths, 2019)

⁹Some common formal frameworks in psychology and neuroscience include signal detection theory, sequential analysis, information theory, Bayesian inference, and Markov decision processes.

The MDP [Markov decision process] framework is a considerable abstraction [idealization] of the problem of goal-directed learning from interaction. It proposes that whatever the details of the sensory, memory, and control apparatus, and whatever objective one is trying to achieve, any problem of learning goal-directed behavior can be reduced to three signals passing back and forth between an agent and its environment: one signal to represent the choices made by the agent (the actions), one signal to represent the basis on which the choices are made (the states), and one signal to define the agent’s goal (the rewards). (p. 50, bracketed text added)

The above quotation also highlights the central role that idealization plays in normative modeling. Because axioms of formal frameworks primarily function to promote mathematical expressivity, modelers often have to make a number of distortive assumptions about their targets in order to build normative models of its behavior. A common example is the widely-used assumption that agents have perfect knowledge of the statistical properties of their environment (e.g., Wald and Wolfowitz, 1948). This assumption exemplifies a distortive idealization of the target because modelers do not often think that the target actually has this perfect knowledge—either in the real world or in the experiment—but still represent their target in this way because it allows them to compute a normative benchmark for performance on the task. Our case study demonstrates two approaches users of the DDM have taken when the assumptions of existing models are inadequate for their reasoning goals.

The pervasiveness of idealization in normative models complicates the question of how to interpret the *failure* of a normative model to explain its target. Discussions concerning the “Great Rationality Debate” in economic decision making¹⁰ suggest three classes of possible interpretations: (1) that the target is generally suboptimal on this task, (2) that more

¹⁰This is one of the first debates in the behavioral and brain sciences to turn a critical eye to notions of optimality shaping the collective scientific project. The debate concerned whether human economic decisions are “rational” according to formal theories developed in economics. Overviews of the debate can be found in Stanovich and West (2000) and Tetlock and Mellers (2002).

empirical data are needed to understand what factors drive suboptimal performance on this task, or (3) that the formal definition of optimality is inadequate for explanatory purposes. A complementary debate has recently begun in the study of perceptual decision making, with researchers questioning broadly-scoped claims about the optimality of perceptual decisions. Rahnev and Denison (2018) comprehensively review evidence demonstrating that humans behave suboptimally on every type of task where their behavior has been considered optimal. The authors do not suggest that the data license an inference that humans are perceptually suboptimal (option 1 from the rationality debate), but instead that the field drop its emphasis on optimality in favor of building detailed models that capture *all* aspects of the perceptual decision process (a model-focused version of option 2 from the rationality debate).

Rahnev and Denison (2018) identify two conceptual challenges of normative models that motivate their suggestion. The first pertains to the variability of optimality definitions across model specifications, and the second pertains to the utility of optimality claims for predicting and explaining behavior. At a broader level, Rahnev and Denison’s (2018) critique can be understood as a much-needed reminder for the field that optimality is not a property of targets that can be “discovered” using scientific methods because it is not something that exists independently of the mathematical models that are used to define it. This is exemplified by optimality claims being inherently contextual in nature (i.e., dependent on formal specifications of the objective function and the environment) and based on idealized representations of the target. Our suggestion, in line with option 3 from the rationality debate, is that users leverage these properties to build normative models that better align with their goals in reasoning about a target. Exploring normative solutions to formally-specified questions not yet addressed in the literature is a principled and straightforward way to advance theorizing in neuroscience, and often motivates the creation of novel experiments that measure the target’s behavior in these differently-idealized settings (e.g., Harhen and Bornstein, 2023; Khoudary et al., 2022). Our case study, which we turn to in the next chapter, demonstrates

how this approach to normative modeling has been pursued by communities of users of the DDM.

4.3 A philosophical introduction to the DDM

We now turn to the target of our case study: the diffusion decision (or drift diffusion) model of decision making (DDM). The DDM is one member of a family of models that represent decision-making as a process of evidence accumulation, or adding up information over time. These models typically assume that decision-makers are forced to decide between two possible options (i.e., two-alternative forced decision-making), but variants that permit reasoning about more than two options continue to be developed (e.g., Tajima et al., 2019; Villarreal et al., 2024). Models in this family are conceptually united by the sequential sampling framework in psychology and neuroscience (Forstmann et al., 2016; Gold and Shadlen, 2007; Ratcliff et al., 2016; Shadlen and Shohamy, 2016), which itself draws on the framework of sequential analysis in statistics (Barnard, 1946; Wald, 1945).¹¹

The conceptual framework of sequential sampling posits that humans and other animals make decisions by continuously sampling information from an evidence source, extracting and integrating decision-relevant information over time, and committing to one of the options once the accumulated evidence surpasses a threshold value. The formal framework of sequential analysis permits specifying normative solutions to the accumulation process, and thus can be used to build models that optimize the tradeoff between decision accuracy and deliberation time (i.e., the speed-accuracy tradeoff). Importantly, however, not all models in this family are normative. A major strength of the sequential sampling framework is the generality and flexibility of its conceptual entities (e.g., “evidence”), both of which permit users

¹¹An interesting and relevant bit of history is that two different goals led to independent developments of the sequential analysis framework during World War II. One goal was enhancing efficiency of industrial output (Barnard, 1946; Wald, 1945), and another was cryptanalysis to decode enciphered German messages. This latter method was derived by Alan Turing, and the relationship of Turing’s framework to sequential sampling is detailed quite nicely in Gold and Shadlen (2002).

to specify a wide range of models that can be construed at various spatiotemporal scales. In what follows, we use Weisberg’s (2013) notions of structure, scope, and assignment both to offer a philosophical introduction to the DDM. On Weisberg’s account, a construal consists of a scope, assignment, and fidelity criteria; we use this final component to structure the case study in Section 4.4.

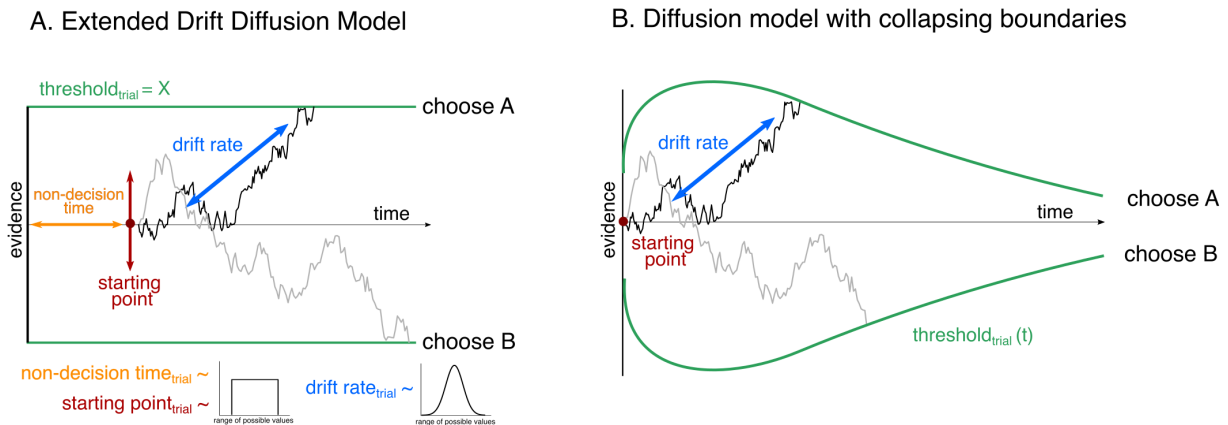


Figure 4.1: Graphical depictions of two standard forms of a diffusion model of speeded two-alternative choice. Standard interpretations of the key variables are provided in the main text. **(A)** The extended drift diffusion model. This form of the model features parameters whose values vary trial-by-trial but remain fixed within the course of a single trial. Starting point (yellow) and non-decision time (red) values are drawn from uniform distributions and drift rates (blue) are drawn from normal distributions. Threshold values (green) are constant within a trial, and can be allowed to vary as a function of experimental condition and/or participant. **(B)** The diffusion model with collapsing decision boundaries. This form of the model features parameters whose values are effectively fixed trial-by-trial and decision boundaries whose values decrease (“collapse”) over the course of a single trial. Drift rate and, in some models, starting point are permitted to vary as a function of experimental condition but are assumed to have fixed effects across trials within an experimental condition.

4.3.1 Formal structure of the DDM

Recall that on Weisberg’s (2013) account, the structure of a model refers to the mathematical objects and relations that comprise the formal representation of a target. In the DDM, this representation consists of components commonly termed starting point, decision variable, drift rate, internal noise, decision threshold, and non-decision time. This formal represen-

tation is presented graphically in Figure 4.1. Mathematically, the evidence accumulation process operating in the DDM is expressed as:

$$dx = A dt + c dW, \quad x(0) = 0 \quad (4.1)$$

where $x(0)$ represents the starting point of the decision variable, dx represents a change in the decision variable x over a unit of time dt , A represents the drift rate, and $c dW$ represents noise/diffusion in the accumulation process which follows a normal distribution with mean 0 and variance $c^2 dt$ (Bogacz et al., 2006). This mathematical structure is equivalent to a random walk in probability theory or Brownian motion in physics. In the DDM, the structure commonly corresponds to a continuously-updating log likelihood ratio quantifying the probability that one of the two possible outcomes is correct, based on the evidence observed thus far. In other words, the decision variable reflects the time-evolving *difference* of evidence in favor of either option. The sampling/accumulation process ends when the value of the decision variable x exceeds a scalar threshold value z , at which point the decision maker commits to the choice corresponding to the threshold value reached ($+z$ or $-z$). The procedure whereby a sequential sampling model specifies how the accumulation process will be terminated and converted into a choice can be called the *decision rule* of the model.

This structure permits the DDM to generate predictions both of *which* option a decision-maker will choose and *how long* it takes them to commit to that choice on a trial-by-trial basis. This is what is meant when it is referred to as a joint model of choices and reaction times. Importantly, because of the noise term in the structure, the DDM assumes that choice behavior is stochastic (i.e., subject to random variation). This structure permits the DDM to jointly predict a decision-maker's overall decision accuracy and the reaction time (RT) *distributions* corresponding to correct and incorrect decisions.

When the DDM is used to fit behavioral data, the starting point, drift rate, threshold, and non-decision time are commonly treated as “free parameters,” meaning that the model fitting process aims to identify the values of these components that maximize the likelihood of the data being fitted. These can be contrasted with “model variables” which are components specified in the model whose values are either fixed by the user or vary probabilistically. Because the DDM formalizes the relationship between choices and RTs, the quantitative model fitting process is often tasked with fitting measurements from both of these elements, such that the model is really tasked with identifying parameter values that best describe how a decision maker traded off speed and accuracy. The resulting parameter values thus reflect the model’s best estimate of the “settings” of the decision process that generated a particular pattern of target behavior. Based on the research question, it is often desirable to permit one or more the free parameters to vary as a function of experimental condition; this is what permits decomposing behavioral effects of experimental interventions into specific changes in the configuration of the latent decision process.

The structure we describe above corresponds to the “original” DDM (Stone, 1960). Presently, there are two variations of this original DDM that are most prominently used. Both of these forms retain the evidence accumulation process described in Equation 4.1, but posit different procedures within which that accumulation process generates choice behavior. The first process—commonly called the “extended DDM”—allows the drift rate, starting point, and non-decision time components to vary probabilistically on a trial-by-trial basis (Figure 4.1A). The second form—commonly called a “collapsing bound” diffusion model—posits that the threshold value for committing to a choice decreases over the course of a single decision (Figure 4.1B). This latter form initiated the “collapsing bound debate” that we conceptually analyze in Section 4.4.

4.3.2 Structure of the random dot motion task

Much of the empirical success of the DDM is due to the development of the random dot motion discrimination task (Britten et al., 1992). In this task, observers are presented with a display of stochastically moving visual elements (usually dots), some proportion of which consistently move from one direction to another (e.g. left to right). On each trial, observers report which direction of motion appeared on the display, a judgment whose difficulty scales with the proportion of consistently moving dots (i.e., the coherence of the stimulus; Palmer et al., 2005a). Importantly, although the overall proportion of coherently-moving elements remains constant within a trial, the actual elements whose position is displaced varies randomly with each refresh of the digital display. This “limited lifetime” property of the stimuli ensures that observers cannot make the decision via smooth visual pursuit: they must continuously sample the stimulus in order to accumulate evidence about the overall direction of coherence motion.

This task is useful for a number of reasons. First, it requires integrating evidence over time, so researchers have good reason to believe that the task utilizes the core mechanism of evidence accumulation (although see Stine et al., 2020 for challenges to this idea). Next, it is simple enough that rodent, primate, and human observers are all capable of performing it (Hanks and Summerfield, 2017), allowing for powerful cross-species comparisons. Finally, it permits a clean one-to-one mapping of stimulus properties onto components of the model, thus giving both experimenters and theorists alike a powerful tool for probing decision making processes under various tightly-controlled circumstances. The next subsection offers an overview of the wide-ranging pieces of empirical support for the DDM, the vast majority of which was generated using behavior in some variant of a motion discrimination task.

4.3.3 Assignment of components in the DDM

Recall that, on Weisberg’s (2013) account, a model’s assignment specifies what each component of a model structure is meant to represent with respect to its target. It’s important to note the standard assignment of components in the DDM that we describe here are specified with respect to the random dot motion task. Different applications of the DDM warrant slightly different assignments of model components, a point that can be overlooked in the meta-scientific writing that touts the utility of models for standardizing interpretations of behavior across task contexts. While it is true that the *structure* of the DDM provides common grounding for different measurements of behavior, how that structure is construed with respect to the target crucially depends on the measurements a user has of that target. For this reason, we highly encourage users to think critically about how components of the DDM are assigned to neural and/or cognitive processes as measured within a particular task setting (Bompas et al., 2023; Jones and Dzhafarov, 2014).

With those caveats in mind, we can now introduce the standard assignment of DDM components in the context of a motion discrimination task. The threshold separation variable defines the quantity of accumulated evidence required for an observer to commit to a choice, and thus is commonly interpreted as reflecting the observer’s response caution or decision policy under different speed-accuracy regimes. The starting point variable ($x(0)$ in Equation 4.1) defines the initial value of the decision variable, and thus is conventionally interpreted as quantifying the bias an observer has toward one of the two choice options. The decision variable represents the decision maker’s internal representation of accumulated evidence, sometimes called their time-evolving belief about the correct answer. The drift rate quantifies the rate of change in the decision variable per unit time, and thus is commonly thought to reflect the “quality” of the internal evidence driving a particular decision.¹² The non-decision time variable is intended to aggregate various kinds of delays in response times

¹²More recently, the drift rate variable has also been construed to represent how quickly an observer internally processes information in general (Schubert et al., 2015).

due to processing occurring before and after evidence accumulation. Processes believed to contribute to the non-decision time include encoding of sensory information and initiating a motor response; in this sense, it can be thought of as a “convenience parameter,” though some work has successfully decomposed this “non-”decision time into decision-relevant components (e.g. Kraemer and Gluth, 2023; Yoo and Bornstein, 2024). Finally, the internal noise variable is thought to reflect imperfections in the encoding and representation of sampled evidence, and also can be thought of as a “convenience variable” because its values are commonly fixed when the DDM is fit to human data (Ratcliff et al., 2016).

4.3.4 Empirical scope of the DDM

On Weisberg’s (2013) account, the scope of a model is the component of its construal that specifies which aspects of the target a user aims to capture in the model. Models thus have an intended empirical and/or theoretical scope and are initially evaluated with respect to those originally intended explananda. One of the reasons why the DDM is considered so successful is because it has continued to display explanatory adequacy for targets well outside its originally intended scope: identifying a formal relationship between choice accuracy and reaction time in human memory retrieval (Link and Heath, 1975; Ratcliff, 1978). Importantly, the DDM was developed as a purely cognitive model—nothing about the original construal mentioned neural activity or measurements as a desired component for the model to capture. It was not until the random dot motion task was developed by Britten et al. (1992) that the DDM was used also to reason about neural processes involved in two-alternative decision making.

Careful early studies conducted on non-human primates offered the first pieces of evidence for the DDM as a model of neural activity (Britten et al., 1996; Gold and Shadlen, 2001; Roitman and Shadlen, 2002; Shadlen and Newsome, 2001). These foundational studies provided evidence suggesting that both single-unit and population level recordings from distinct cor-

tical regions exhibit activity that strongly resembles and correlates with distinct components of the DDM (see Gold and Shadlen, 2007 for a review of this line of work). Since these early findings, the DDM has been used to link behavior with population-level neural responses in rodents (Brunton et al., 2013; Hanks et al., 2015; Khilkevich et al., 2024), as well as intracranial recordings, scalp oscillations, and blood oxygen level dependent signals in humans (Krueger et al., 2017; O’Connell and Kelly, 2021; Polanía et al., 2014; Weber et al., 2024). The early and growing evidence for the DDM as a model of the neural processes generating decision behavior has even led some researchers to posit evidence accumulation as a basic mechanism of decision making that is conserved across species (Hanks and Summerfield, 2017).

While this broad empirical scope lends a great deal of support to the DDM as a theory of speeded decision making, it can also result in confusion, ambiguity, and/or disagreement about precisely what inferences a user can make on the basis of applying the DDM to data. For this reason, among several others, it is desirable for modelers to explicitly discuss a model’s intended construal when communicating their findings from the model.

4.4 Reasoning goals and representational decisions in the collapsing bounds debate

In this section, we further demonstrate the utility of our philosophical tools for reasoning by using them to provide a novel conceptual analysis of the so-called “collapsing bounds” debate in the DDM. As shown in Figure 4.1, this debate concerns whether the threshold in the DDM ought to remain fixed over the course of a single decision (Figure 4.1A) or change as a function of time spent accumulating evidence (Figure 4.1B). Importantly, anyone who uses a DDM to explain or interpret behavioral data necessarily has to “take a stance” in this debate by deciding whether to use fixed or collapsing boundaries in their model. This

ubiquitous representational decision is complicated by two facts: (1) each form of the model makes highly accurate predictions about both behavioral and neural observations and (2) each form has been proven, under different definitions of optimality, to optimize the two-alternative decision problem. Accordingly, much ink has thus been spilled about which form “correctly” represents the decision process (Evans et al., 2017; Hawkins et al., 2015; Miletić and van Maanen, 2019; Palestro et al., 2018; Ratcliff et al., 2016). Our review of the philosophy literature, however, suggests that this question is misguided.

We argue that the two forms of the model implicated in the debate offer complementary—rather competing—explanations of the processes generating speeded two-alternative choices. To do this, we demonstrate how each form was developed with respect to different sets of fidelity criteria that reflect the different explanatory aims (or “styles”, Potochnik and Sanches de Oliveira, 2020) of its users. Although the community of DDM users is broad and heterogeneous, we focus on a distinction between two subgroups implicated in the debate: cognitive psychometricians and theoretical neuroscientists. Cognitive psychometrics, as a field, aims to develop models that offer interpretable and statistically robust decompositions of observed behavior into latent cognitive processes. By extension, the primary explanatory fidelity criterion motivating representational decisions is that of *parsimony*: striking a balance between goodness-of-fit (i.e., dynamical fidelity) and model simplicity (e.g., Vandekerckhove et al., 2015). Theoretical neuroscience, by contrast, aims to explain behavior and brain function using normative models based on formal principles (e.g., Abbott, 2008), the most common of which is optimality (see Section 4.2.3). Models developed toward this end still must exhibit dynamical fidelity with respect to the relevant patterns of behavior, but they might do so by using formal representations that are more mathematically complex than those preferred by psychometricians.

Our case study explicates how reasoning goals have shaped representational decisions along three complementary dimensions of the collapsing bounds debate: definitions of optimality

(Section 4.4.1), structural modifications to the original DDM (Sections 4.4.2-4.4.2.2), and formal model comparisons (Section 4.4.3).

4.4.1 Two types of optimality in the DDM

Recall that the DDM establishes a formal link between decisions and the time it takes to make them; that is, it is a joint model of choices and reaction times. The speed-accuracy tradeoff in the DDM is determined by the threshold, the value of which specifies the “decision rule”: at each point in time, determining whether an agent should keep sampling information or commit to a choice on the basis of already-accumulated evidence. Higher thresholds always result in increased choice accuracy, but often at the cost of taking longer to make a decision. Conversely, lower thresholds allow subjects to make decisions more quickly but often come at the expense of a greater number of errors. Formal definitions of optimality in the DDM thus aim to define a quantitative benchmark for determining whether decision makers use threshold values that maximally balance decision speed and accuracy. Optimality therefore figures prominently in the collapsing bounds debate precisely because it is a debate about the appropriate structure of thresholds in the DDM.

The rest of this section demonstrates how each threshold structure in the debate (fixed/static or collapsing/dynamic) meets different definitions of optimal speed-accuracy tradeoffs. Understanding the sense in which each form is optimal is crucial for understanding when the DDM gives a normative explanation and what notion of normativity (or optimality) the explanation invokes. Recall that normative models use optimality to answer the question of why the target behaved in a particular way: it deployed a process that optimizes the objective function (Section 4.2.3). The explanatory fidelity of a normative model thus corresponds to how well its objective function—and the assumptions required to prove its optimality—align both with properties of the target and how the user wishes to reason about them using the model. We turn next to demonstrating how and why objective functions in normative

DDMs have changed over time, and how that change has resulted in multiple definitions of optimality that are met by different forms of the DDM.

4.4.1.1 SPRT Optimality: definition and idealizing Assumptions

The first formalization of an optimal speed-accuracy tradeoff comes from the sequential probability ratio test (SPRT). The SPRT was developed independently by Barnard (1946) and Wald (1945) originally for purposes of quality control in manufacturing. Shortly after its introduction, the SPRT was mathematically proven to minimize a weighted, linear sum of decision time and error rate (Wald and Wolfowitz, 1948; Bogacz et al., 2006). The optimality of the SPRT as a decision process explicitly motivated the development of the original DDM (Stone, 1960; Laming, 1968), which is the continuous-time equivalent of the SPRT. In both the SPRT and the original DDM, the decision variable is a running estimate of the log-likelihood ratio of one choice option relative to another, which is equivalent to the difference of evidence in favor of each choice. This decision variable continues to accumulate until it surpasses a threshold quantity that is fixed both within a single decision and across different decisions.

Although the SPRT was immensely useful as a starting point for developing the large and successful family of sequential sampling models, its explanatory fidelity has been called into question for a number of reasons. First, the SPRT is only optimal in cases where the signal-to-noise ratio of sensory evidence is identical for every choice (i.e., in a homogeneous environment; Moran, 2015). In the real world and in most common laboratory settings, the strength of sensory signal will vary from decision to decision, and the SPRT does not achieve the optimal speed-accuracy tradeoff in these *heterogeneous* environments. Next, the SPRT is only optimal if observers have an infinite amount of time to sample evidence before committing to a choice (Frazier and Yu, 2007). But again, for *speeded* perceptual decisions made in real life and the lab, there is often a hard limit on how long an observer can keep sampling (e.g., if the real-world stimulus is no longer present in the visual field or

if the laboratory trial times out). Finally, implementing the optimal speed-accuracy tradeoff using the SPRT requires that observers define ahead of time either their desired level of average decision speed or average decision accuracy (Bogacz et al., 2006). Once this is done, employing the SPRT process ensures that an agent will make decisions that either maximize accuracy for that pre-specified decision time, or minimize decision time for that pre-specified level of accuracy. The assumption of observers pre-defining some level of speed or accuracy that they wish to optimize might be plausible in some decision settings, but it also creates a methodological issue of optimality being specific to each observer’s internal criteria. This issue is not insurmountable, but it does create an optimality benchmark that can only be evaluated after estimating the value of a free parameter individually for each observer.

4.4.1.2 Reward Rate Optimality: definition and idealizing assumptions

The above-listed limitations of the SPRT motivated the development of *reward rate* as an alternative metric for assessing the optimality of speed-accuracy tradeoffs in two-alternative decision making. Reward rate is defined as the average expected reward (or proportion of correct responses) divided by the average amount of time that elapses between each decision (Gold and Shadlen, 2001; Bogacz et al., 2006). Accordingly, it corresponds to the change in probability of being rewarded and/or correct per unit of time.

A major advantage of reward rate relative to the SPRT is that it is a *parameter-free* optimality benchmark: i.e., assessing the optimality of behavior on the individual or group levels does not require estimating the weights that different observers place on speed versus accuracy. Reward rate can also be optimized in decision environments that limit sampling time and that are heterogeneous, thus overcoming the conceptual issues faced by the SPRT. Further, the reward rate formalization of speed-accuracy tradeoffs has two properties that make it a more powerful theoretical tool. First, its representation of time permits theoretically and empirically investigating how different environmental dynamics shape choice behavior, which is not possible with the SPRT. Second, by invoking the notion of reward, it

aligns sequential sampling models more closely with other formal models of decision making (e.g., expected utility theory) which enhances the prospects for inter-theoretic model building (more on this in Section 4.4.3). All of these properties enhance the explanatory fidelity of a reward rate formalization of the speed-accuracy tradeoff relative to the SPRT.

Reward rate has thus become the standard formalization of speed-accuracy tradeoffs in the sequential sampling literature. We have two things to note about this. First, when a community of model users adopt a standard formalism for a key concept in their research, they often cease to make explicit both its construal with respect to the target (Weisberg, 2013) *and*, we argue, their justification for that representational decision. This means that when model users communicate findings using the standardized notion, they can write, for example, “We assume that the aim of the decision maker is to maximize the net reward over all trials” (Drugowitsch et al., 2012, p. 3620) without needing to explain why they made this particular assumption and what alternative assumptions might be. A core argument of this article is that these standardized formalisms are hallmarks of “healthy” model-based science, and that users must remain cognizant of their necessarily incomplete nature in order to avoid limiting their scientific imagination.

Second, and relatedly, reward rate optimality rests upon an assumption that all properties of the environment (distribution of signal strengths, timing, reward structure, etc.) are *known* to the observer. As shown above, this distortive idealization gives model users a more powerful formalism for reasoning normatively about the speed-accuracy tradeoff in two-alternative decision making. Another core goal of this article is to demonstrate that there are a variety of ways to engage with idealizing assumptions of formal models. First, as we show in Section 4.4.2, model users can choose not to pursue normative reasoning goals that require strong idealizations about their target, and instead build models that aim to maximize dynamical fidelity with as minimal assumptions as possible. Alternatively, we show in Section 4.4.2.2, model users can posit less-distortive (or “relaxed”) idealizations about the

target and develop new models that offer normative solutions to less-idealized problems. A third, related approach is to investigate what properties of the target cause it to deviate from predictions of normative models (e.g., Balci et al., 2011; Holmes and Cohen, 2014; Drugowitsch et al., 2016). Findings from this last line of research can then be integrated as constraints on future normative models, as in resource-rational approaches to cognitive modeling (Lieder and Griffiths, 2019; Section 4.2.3).

Altogether, this section articulated the role of optimality in the collapsing bounds debate, defined the two notions of optimality at play in the debate—and the broader literature—along with their idealizing assumptions, and articulated different ways that model users can define their reasoning goals with respect to idealizing assumptions required for normative modeling. We turn next to demonstrating how the debate about collapsing bounds can be diffused by recognizing the different reasoning goals motivating the forms of the models implicated in the debate.

4.4.2 Reasoning goals motivating structural changes to the original DDM

As discussed in Section 4.4.1, the original DDM (Stone, 1960; Laming, 1968) is the SPRT-optimal procedure for two-alternative decision making. This section demonstrates two approaches that DDM users have taken to modify the original DDM in order to enhance its fidelity (i.e., ability to support particular reasoning goals) with respect to the target. In doing so, we show how the two models implicated in the collapsing bounds debate can be understood as complementary, rather than competing, explanations of the same target.

4.4.2.1 Parsimony and dynamical fidelity motivated the Extended DDM

Although the original DDM is commonly attributed to Ratcliff (1978), the model structure detailed in Ratcliff (1978) does not correspond to the structure that most authors refer to as

the original DDM. Descriptions of the original DDM—where log-likelihoods are continuously estimated and decision boundaries are symmetric around a starting point of 0—can be found in Stone (1960) and Laming (1968). The structure of Ratcliff’s (1978) model already reflects representational changes made to enhance the model’s dynamical fidelity at the expense of SPRT-optimality: allowing drift rates to vary on a trial-by-trial basis.¹³

Throughout every iteration of the DDM’s development, Ratcliff’s representational decisions were guided by the goal of creating a representation that can reproduce all measured aspects of human behavior on the two-alternative forced choice task (Ratcliff and McKoon, 2008; Ratcliff and Rouder, 1998; Ratcliff and Tuerlinckx, 2002). Of uniquely high importance to Ratcliff is the ability of the DDM to reproduce patterns of RT distributions across task conditions, a point he emphasizes in nearly all papers involving the model (Ratcliff, 1978; Ratcliff and McKoon, 2008; Ratcliff et al., 2016). An empirically robust pattern in this regard is the mixture of slow and fast errors present in a sample of measurements. This mixture varies both across individuals performing a task with the same instructions, and within a single individual when they are instructed to prioritize either the speed or accuracy of their responses. Ratcliff’s solution to capturing this empirical property of his target was to allow the drift rate and the starting point of the evidence accumulation process to vary probabilistically from trial-to-trial according to a uniform and normal distribution, respectively (Ratcliff and Rouder, 1998). A later modification—uniform trial-wise variability in the non-decision term—was added to increase the DDM’s ability to fit RT distributions with a high amount of variance in the 0.1 quantile (Ratcliff and Tuerlinckx, 2002). The extended DDM encompasses all of these changes and functions as the standard form of the model in current research (Ratcliff et al., 2016).

This series of representational changes that made the extended DDM satisfy Ratcliff’s fidelity criteria also break the SPRT-optimality of its initial form (the original DDM). The

¹³We thank Barbara Doshier for directing our attention to this detail.

mathematical form of the decision variable in the extended DDM, however, still corresponds to the optimal procedure for accumulating evidence for two-alternative forced decisions in any environment (Moran, 2015). It is thus the details about how that process is converted into decisions (i.e., with trial-wise variability in key parameters and fixed decision thresholds) that break the optimality of the extended DDM as a whole. Interestingly, a formal analysis undertaken by Jones and Dzhafarov (2014) showed that if the form of trial-wise variability in the extended DDM (i.e., the probability distributions from which starting point, drift rate, and non-decision time are drawn on each trial) is left unconstrained, the model can fit any pattern of speed-accuracy data. The authors used these findings to argue that “the explanatory or predictive content of these models is determined not by their structural assumptions, but, rather, by distributional assumptions that are traditionally regarded as implementation details” (Jones and Dzhafarov, 2014, p.1). In response, developers of the extended DDM argued that the form of the distributions governing trial-wise variability were not chosen ad hoc, but rather on the basis of their ability to robustly and reliably fit human data (Smith et al., 2014).

Taken together, the history of the development of the extended DDM—and how its form is defended against critics—reveals a strong commitment to dynamical fidelity on the part of its developers. This has led to a formal structure that, if used to explain why choice behavior exhibits particular patterns, can only do so by appealing to intrinsic stochasticity of the system whose form was derived from best-fit to large amounts of data. We argue that this structure reflects the aims of cognitive psychometrics as a whole: developing principled models of cognitive processes that permit robust and reliable decomposition of behavior into meaningfully different component parts. Researchers in this subcommunity of DDM users often use statistical parsimony to guide their formal theory building, a position that naturally lends itself to structures with components that prioritize compact description over normative guarantees (Palminteri et al., 2017; Vandekerckhove et al., 2015). In this sense, statistical parsimony is a primary explanatory fidelity criterion for users with these goals.

Because the extended DDM satisfies both the high standards of dynamical fidelity and a guiding explanatory fidelity criterion, it can be considered to have been “optimized” for the explanatory goals of cognitive psychometrics.

4.4.2.2 Explanatory fidelity of optimality motivated collapsing bounds

The structure of the collapsing bound diffusion model, we argue, has likewise been developed to meet the explanatory goals of users we call theoretical decision (neuro)scientists. This community is considerably more heterogeneous than the one discussed above, comprised both of users whose focus is on behavior (e.g., Frazier and Yu, 2007) and those who use the DDM to link behavior with neural activity (e.g., Drugowitsch et al., 2012). The feature uniting these users is their commitment to optimality as an explanatory fidelity criterion. This section thus discusses how the collapsing bound diffusion model is the result of representational decisions that aim to preserve the normative status of the model under assumptions that are less restrictive than those required for SPRT-optimality.

The first specification of a diffusion model with collapsing decision boundaries was proposed by Frazier and Yu (2007). These authors targeted the SPRT’s idealizing assumption that observers have an infinite amount of time to sample evidence before committing to a choice (formally, the assumption of infinite horizon). Frazier and Yu (2007) developed a normative model that incorporates stochastic (i.e., randomly varying) time limits and found that the optimal procedure in this specification involves decision thresholds that decrease over the course of a single choice (i.e., collapsing bounds). Drugowitsch, Moreno-Bote et al. (2012) then targeted the SPRT’s assumption of environmental homogeneity (i.e., same signal-to-noise ratio on each trial). Their solution also relaxes the assumption of infinite horizon as in Frazier and Yu (2007), but did so by adding an assumption that each sample of evidence incurs some cost to the observer. Drugowitsch, Moreno-Bote et al. (2012) found that, on this specification of the formal problem, the optimal procedure also involves thresholds that decrease over the course of a decision.

In order to identify normative solutions to these differently-formalized problems, Frazier and Yu (2007) and Drugowitsch, Moreno-Bote et al (2012) both made use of a notion of optimality known as Bellman optimality. This formalism posits that an optimal procedure is one that maximizes an agent’s expected return, i.e., the amount of reward earned over the course of an experiment, which is conceptually equivalent to maximizing reward rate. Whereas the SPRT’s optimality can be proven using standard mathematical techniques (Bogacz et al., 2006; Edwards, 1965; Wald and Wolfowitz, 1948), identifying Bellman-optimal solutions requires using a computational procedure known as dynamic programming. Details of the procedure vary across applications, but the general idea is that agents recursively update an estimate of their expected return with each new piece of information they get from the environment (Sutton and Barto, 2018). Crucially, this means that the precise form of the Bellman-optimal procedure will depend on details of the environment (e.g., strength of evidence on each trial, reward scheme, timing, etc.), and thus will vary across tasks. But if the mathematical form is determined using a Bellman equation, users have a guarantee that the form maximizes expected return in that environment. In the case of the DDM, collapsing bounds thus maximize expected return—and thus are both Bellman and reward-rate optimal—for evidence accumulation decisions made both within a fixed amount of time and in heterogeneous environments.¹⁴

A procedure whereby evidence accumulates to a threshold value that decreases over time thus emerges as a generally normative solution to slightly less idealized formalizations of the speeded two-alternative decision problem. This structure reflects the explanatory aims of DDM users who desire normative explanations for choice behavior. When the assumptions required for the original DDM’s normativity (i.e., SPRT optimality) were too restrictive, users in this community opted to identify optimal structures on weaker assumptions (i.e., finite amount of time to decide and different stimulus difficulty on each trial) and with more

¹⁴Frazier and Yu (2007) do not explicitly frame their Bellman-optimal solution as one that optimizes reward rate, opting instead to discuss optimality in the Bayesian sense of appropriately updating belief about a latent property of the environment.

a more flexible notion of optimality (i.e., reward rate). This approach preserves the model’s normative status while also enhancing its fidelity with respect to explanatory aims of its users.

4.4.3 Reasoning goals motivating approaches to representational decisions in model comparison

In the previous two sections, we demonstrated how consideration of reasoning goals and fidelity criteria can provide insight into variability in representational decisions made by different users modeling the same target. We first discussed this in the context of optimal speed-accuracy tradeoffs in the DDM (Section 4.4.1), and then contrasted approaches that users can take when they deem the assumptions of normative/optimal models too idealized for their reasoning goals (Sections 4.4.2-4.4.2.2). The present section extends the discussion to demonstrate how reasoning goals shape fidelity criteria and approaches to representational decision-making employed in *model comparison*.

Recognition of sources of variability in representational decisions involved in model comparison is essential because (1) the explanatory success of a particular model is almost always defined *relative* to alternative possible models, and (2) the space of possible alternative models is—in theory—infinite. Accordingly, the representational decisions that users make when constructing alternative models defines the space of possible alternatives against which the theory encoded in a model is tested. In much the same way that experimental scientists must carefully consider how to construct their control conditions, model users must carefully reflect on the specifications of alternative models when interpreting the results of model-based research.

4.4.3.1 Parsimonious dynamical fidelity motivates data-driven approaches

In response to the growing popularity of collapsing-bound accumulator models, Hawkins et al. (2015) undertook a large-scale quantitative comparison of model performance on a variety of perceptual decision making datasets. The authors took great care to ensure the statistical robustness of their results, utilizing data from different research groups and species, specifying multiple forms of collapsing bounds models, and running several computationally intensive sensitivity analyses. Further, they used well-established metrics for assessing how models in the comparison trade off goodness-of-fit with complexity: Akaike Information Criterion (AIC), Bayesian Information Criterion (BIC), and nested likelihood ratio tests. Taken together, the authors’ decisions reflect a “bottom-up” approach to model comparison: tiling the space of possible models, fitting them on large and heterogeneous datasets, and assessing each model’s fidelity on the basis of how parsimoniously it accounts for the wide range of data. Again, we argue, this approach reflects the aims of cognitive psychometrics as a whole: maximizing dynamical fidelity and relying on parsimony as an explanatory fidelity criterion.

One challenge with this approach to model comparison pertains to how the notion of *model complexity* is formalized. Two of the three quantitative metrics used by Hawkins et al. (2015)—AIC and BIC—define complexity in terms of the number of parameters in a model, and linearly penalize models according to how many free parameters they use to explain the data. While this is a standard notion of model complexity, particularly in the field of mathematical psychology, it is not the only formal way to reason about complexity. Villareal et al. (2023), for example, propose that complexity can be formalized in terms “the predictions a model makes and the ability of empirical evidence to falsify those predictions” (p.1), and present a metric that uses Kullback–Leibler divergence to quantify complexity in these terms. The authors go on to show that this approach challenges traditional intuitions about how complexity relations among nested models (i.e., those where one model is a specific case

of another), which is the logical basis of the third metric used by Hawkins et al. (2015) in their large-scale model comparison endeavor. The point we wish to emphasize here is that model complexity—much like optimality—can be formally defined in a number of different ways, some of which are better suited to particular reasoning goals than others. Paying attention to the assumptions embedded into formal definitions of key explanatory fidelity criteria, like complexity and optimality, is essential for engaging critically with model-based research findings.

To demonstrate how the findings of model-based research are jointly sensitive to representational decisions about alternative models and assumptions of performance metrics, we highlight how the pattern of results observed by Hawkins et al. (2015) changed drastically between two sets of analyses reported in the paper. In the first analysis, all of the collapsing bounds models tested by Hawkins et al. (2015) were specified in a manner that made them between 1.3 to 2 times as complex as the extended DDM (which functioned as the null model). One of sources of this heightened complexity was the incorporation of trial-level variability in drift rate, starting point, and non-decision time into the collapsing bounds models that did not originally incorporate this variability (Drugowitsch et al., 2012; Frazier and Yu, 2007). Hawkins et al.’s (2015) logic behind this decision was to create nested model structures, such that the extended DDM represents the simplest possible decision process in the comparison. This type of comparison set has been called a “stacked deck,” since it quantitatively favors the null model (O’Connell et al., 2018). However, even in this loaded inferential context, the collapsing bounds models appeared favored by BIC for datasets acquired from non-human primates.

In a second analysis, Hawkins et al. (2015) performed the same model fitting procedures but removed trial-variability in drift rate, starting point, and non-decision time for the collapsing bounds models. These changes better aligned their specifications of collapsing bounds with the fidelity criteria motivating this form of the model, since neither Drugowitsch, Moreno-

Bote et al. (2012) nor Frazier and Yu (2007) needed to incorporate any trial-varying parameters to achieve their fidelity criteria. Readers are encouraged to compare the “Posterior Model Probability” visualizations in Figures 5 and 6 of Hawkins et al. (2015) to appreciate how drastically these specification changes altered the pattern of results. Primate data that previously strongly supported collapsing bounds now seemed either mixed or to prefer fixed bounds, and human data appeared to favor each form roughly equally (Hawkins et al., 2015, Figure 6).

4.4.3.2 Explanatory fidelity of optimality motivates theoretical approaches

We can contrast Hawkins et al.’s (2015) “bottom-up” approach with the “top-down” approach displayed in Moran (2015). In a thought-provoking series of analyses, Moran (2015) combines mathematical proofs and model simulations to investigate optimal decision procedures in heterogeneous and *biased* decision environments; i.e., environments where signal strength varies trial-by-trial and where one outcome occurs more frequently than the other. This model development strategy is similar to those surveyed in Section 4.4.1, but differs in one important regard: *assuming* that the decisions are made according to the original DDM, which is known to be suboptimal in heterogeneous environments (Section 4.4.1). Moran (2015) then investigated which structural modification(s) to the suboptimal process led to the highest possible reward rate in that environment, effectively testing the optimality of differently suboptimal model structures. We consider this a “top-down” approach to model comparison because it focuses exclusively on the formal/logical relationships among different models and their suitability for supporting particular types of reasoning goals.

The utility of Moran’s (2015) approach can be demonstrated by contrasting his findings with those of van Ravenzwaaij et al. (2012). On the basis of several simulations, van Ravenzwaaij et al. (2012) reported that maximizing reward rate in biased and heterogeneous environments requires adjusting only the starting point of the decision process. This finding contradicted previous theoretical work indicating that both the starting point and drift rate must be

biased to maximize reward rate (Bogacz et al., 2006), as well as empirical evidence supporting the existence of a biased drift in human and non-human primate decision making (Hanks et al., 2011). By asking the same question with a different approach to representational decision making, Moran (2015) identified a key oversight in van Ravenzwaaij et al.’s (2012) simulations: all of their models used the same threshold value, which was chosen arbitrarily. This representational decision simplified the space of possible solutions in a manner that aligned with van Ravenzwaaij et al.’s (2012) specific goal (i.e., challenging the model reported in Hanks et al., 2011). However, it also omitted the fact that optimal solutions also depend on how/where the threshold value is set. Moran (2015) then showed that when the same simulations reported by van Ravenzwaaij et al. (2012) included threshold values in the search space (along with negative values of drift rate), the reward-rate maximizing process is one that imparts a bias both on the starting point and drift rate. Altogether, this contrastive example highlights the utility of formal approaches to model comparison, particularly for purposes of reasoning about formal notions like optimality.

4.4.4 A principled heuristic for deciding whether to use fixed or collapsing bounds

This section provided a novel conceptual analysis of the collapsing bounds debate, highlighting how different reasoning goals motivated diverging approaches to developing and testing models implicated in the debate. In contrast to some of the rhetoric of previous model comparison research surveyed above, we argue that the two models offer complementary explanations about the mechanisms of decision making. The extended DDM, which uses fixed decision boundaries, gives a *parsimonious* explanation, whereas the collapsing bounds model gives a *normative* explanation.

Our suggestion is that the appropriate form to use depends on your explanatory goals. If you wish to reason about the optimality of two-alternative choice behavior in heterogeneous

and/or time-limited environments, then your model must incorporate collapsing decision boundaries. But if the normativity of the process is irrelevant for your explanatory or reasoning goals, then the extended DDM is completely sufficient for that purpose (see also Boehm et al., 2020 for data suggesting that fixed bounds might offer a robust “default” setting that approximates reward rate optimality). Of course, it is always possible to fit both types of models to target data and use careful model comparison methods to determine which structure offers the better explanation (e.g., Palestro et al., 2018).

4.5 Summary and Conclusion

In this article, we summarized some core ideas in philosophy of modeling to develop a “philosophical toolkit” for computational cognitive modeling (Section 4.2). We then demonstrated the utility of such a resource by using it to give a philosophical introduction to an extremely prominent model in the brain and behavioral sciences (Section 4.3) and then offer a novel conceptual analysis of a long-standing debate regarding the form of that model (Section 4.4). Throughout, we emphasized the central role that reasoning goals play in shaping every step of model-based research, echoing recent calls from philosophy (Danks, 2015; Potochnik and Sanches de Oliveira, 2020), computational neuroscience (Kording et al., 2018), and cognitive modeling of human behavior (Wilson and Collins, 2019). Our goal in doing so was to highlight the user-dependence of the insights afforded by these formal tools, an aspect that can be overlooked by both new and seasoned researchers alike. This goal was motivated by two core contributions that philosophy can offer practicing scientists: (1) highlighting implicit beliefs or assumptions they have about how their tools work, and (2) identifying the logical and epistemic limitations of those tools with respect to particular goals. Our toolkit and case study, though focused on topics in the DDM, were constructed with the goal of conveying insights that can generalize to most, if not all, of the formal models that are becoming increasingly common in cognitive neuroscience.

Chapter 5

Conclusion

In reinforcement learning, as in other kinds of ~~learning~~ *modeling*, such representational choices are at present more art than science.

—Sutton & Barto, *Reinforcement Learning*, 2018 (striketrough and “*modeling*” added)

This dissertation has proposed that reliability estimation is a normative principle unifying dynamic models of decision-making. For both the decision threshold and the drift rate of accumulation, we showed how assuming *uncertainty*—rather than perfect knowledge—about an evidence source’s informativeness generates normative solutions that involve the quantity changing in magnitude across the course of a decision. In both cases, the temporal fluctuations in the model parameter are governed by a time-evolving belief about the true informativeness of the evidence bearing on choice. This belief, and its translation into dynamically-changing model components, is what allows observers to optimally trade-off speed and accuracy when the informativeness of evidence is unknown to them in advance of deliberation.

5.1 Summary

Chapter 2 demonstrates the utility of reliability estimation as a normative principle by using it to develop a novel sequential sampling model of expectations in decision-making. We built this model to address the systematic omission of uncertainty about expectations in normative models of perceptual decision-making. We also built it to provide a formalism for theorizing about time-varying interactions between associative memory retrieval and the sampling of sensory information. To do this, we modeled expectations as their own dynamic source of evidence that is integrated into the decision process across to a time-evolving estimate of the *relative* reliability of the two evidence sources (memory and vision). We showed that this simple model can reproduce four unique time-varying effects of expectations, each of which was observed in a different decision task and neuroimaging modality. The ability of our model to synthesize these previously-disparate observations into a common theoretical framework suggests that it offers a unifying explanation, further demonstrating the utility of normative principles for advancing scientific understanding.

Chapter 3 then directly tests the key prediction of our model: that expectations should be weighted most strongly at the moments in time when sensory evidence is maximally uncertain. To do this, we designed a novel perceptual decision-making task that uses intermittent epochs of pure sensory noise to dynamically manipulate the *relative* reliability of the two evidence sources (memory and vision). Standard models predict that the average drift rate during this period of time should be about 0, whereas our model predicts that the drift rate should scale with the predictiveness of the expectation. This is precisely what we observed in human subjects behavior, albeit only in a subset of observers who adopted a low decision threshold for the task. Because lower thresholds reflect a greater emphasis on decision speed rather than accuracy, and because expectations ought to have stronger effects when sensory sampling time is limited, these findings are in line with existing theory about expectations in perceptual decision-making. The higher-threshold group, by contrast, appeared to prioritize

decision accuracy by adopting a strategy to wait until the second noise period has passed to make their response. A follow-up study where the second noise period extends throughout the duration of sensory evidence on several trials will help us better investigate effects of expectations for observers who adopt this more conservative decision strategy.

Chapter 4 then presents an interdisciplinary review of the “collapsing bounds” debate in perceptual decision-making, along with a philosophical toolkit articulating the conceptual foundations for this whole body of work. We discuss how all mathematical models will necessarily be incomplete representations of their targets, and how this means that model performance has to be evaluated with respect to the *purpose* of the model. Then, we analyze the history of representational decisions leading to the current form of each model implicated in the “debate,” showing how time-varying (or “collapsing”) decision boundaries reflect a theoretical commitment to optimality whereas the alternative model (the “extended DDM”) is theoretically committed to statistical parsimony. Specifically, we discuss how a time-varying decision boundary is necessary for optimizing the speed-accuracy tradeoff when relaxing the assumption that observers have full knowledge of the informativeness of sensory evidence on each decision.

5.2 Future Directions

This work raises several possible directions for future research. First among these is augmenting the body of empirical evidence testing our model, both by incorporating more temporally-sensitive measurements of latent processing and by expanding the types of decisions tested. A particularly exciting possibility toward both of these ends is the measurement of eye movements during cue learning and subsequent decision-making. Mounting evidence implicates eye movements in successful memory retrieval, particularly in the context of guiding visual behavior (Wynn et al., 2016; Ryan and Shen, 2020; Zubair et al., 2026). Additional research has linked these eye movements directly with memory sampling during value-based

deliberation (Nicholas et al., 2026), naturally raising the question of how mnemonic content guides—and potentially interferes with—the eye movements made when sampling sensory evidence (Ladyka-Wojcik et al., 2026).

Another exciting direction for future work is synthesizing the philosophical research reviewed in Chapter 4 with the question of how humans flexibly retrieve information from memory to guide behavior in complex, uncertain environments. The concepts of representational decisions and adequacy-for-purpose offer complementary tools for theorizing about information compression during memory encoding and consolidation (e.g., Zhou et al., 2025; Tarder-Stoll et al., 2026), as well as its subsequent expansion during memory retrieval (Kerrén et al., 2026). Further, the notion of fidelity criteria—benchmarks for inferring the adequacy of a model relative to a particular goal—can offer insights into how observers estimate the reliability of information retrieved from memory, especially in cases where these estimates diverge from objectively-definable metrics of that memory’s reliability (e.g., Talarico and Rubin, 2007).

As discussed in Chapter 4, Weisberg’s (2013) account posits that modelers use both “dynamical” and “representational” fidelity criteria to evaluate model performance. Dynamical fidelity refers to the similarity between the quantitative behavior of the model and that of the target, such that models with a better goodness-of-fit have higher dynamical fidelity. Representational fidelity, by contrast, refers to the similarity between the *causal structure* of the model and that of the target. Dynamical fidelity thus asks “does this model capture the data?”, whereas representational fidelity further asks “does it do so for the right reasons?” One possibility, then, is that dynamical fidelity corresponds to a time-evolving estimate of the *perceptual* similarity of incoming information relative to what exists in memory, serving to guide both pattern completion and separation in situations of high input complexity and/or uncertainty (e.g., Zhou et al., 2025). Representational fidelity, by contrast, could correspond to evaluations of the *accuracy* or *relevance* of retrieved content with respect to

decisional context and/or inferred latent cause of perceptual inputs (e.g., Shi and Garvert, 2026), allowing agents greater flexibility and control in how they use contents retrieved from memory to guide their perception and action.

These operationalizations open the door for future work formalizing how these fidelity criteria are applied to information retrieved from memory, and—conversely—how they might be used to guide efficient learning and information-seeking. Are the criteria applied sequentially, such that dynamical fidelity guides the pattern completion process which serves as input to the representational fidelity evaluation? Or do the two work in parallel, such that assessments of similarity and relevance jointly guide memory search? With these possibilities formalized, we can then ask how observers *weigh* each criterion during memory-guided behavior, and whether—as in the original development of the concept—these criteria are flexibly applied with respect to the agents’ goals and/or constraints. A particularly powerful project toward this end would be identifying the optimal balance of dynamical and representational fidelity under a variety of limiting conditions to serve as a principled benchmark against which real-world human strategies can be compared (e.g., Kinney and Lombrozo, 2024). As this dissertation has shown, treating empirical deviations from optimality as information about a mismatch between the normative model and real-world constraints, rather than a description of the suboptimality of real-world solutions, is a powerful approach reverse-engineering the *actual* problems that cognitive systems are tasked with solving.

Bibliography

- Abbott, L. F. (2008). Theoretical neuroscience rising. *Neuron*, *60*, 489–495.
- Afacan-Seref, K., Steinemann, N. A., Blangero, A., & Kelly, S. P. (2018). Dynamic interplay of value and sensory information in high-speed decision making. *Curr. Biol.*, *28*, 795–802.e6.
- Aitken, F., & Kok, P. (2022). Hippocampal representations switch from errors to predictions during acquisition of predictive associations. *Nat. Commun.*, *13*, 3294.
- Alister, M., & Evans, N. J. (2026). A diffusion-based framework for modeling systematic, time-varying cognitive processes. *Psychol. Rev.*
- Aly, M., Barense, M., & Kok, P. (2026). The role of hippocampal predictions in cognition: Bridging perception and memory. *Philos. Trans. R. Soc. Lond. B Biol. Sci.*, *381*.
- Anderson, J. R. (1990). *The adaptive character of thought*. Lawrence Erlbaum Associates.
- Andrews, M. (2021). The math is not the territory: Navigating the free energy principle. *Biol. Philos.*, *36*, 30.
- Angelaki, D. E., Gu, Y., & DeAngelis, G. C. (2009). Multisensory integration: Psychophysics, neurophysiology, and computation. *Curr. Opin. Neurobiol.*, *19*, 452–458.
- Aronowitz, S. (2019). Memory is a modeling system. *Mind Lang.*, *34*, 483–502.
- Bakkour, A., Palombo, D. J., Zylberberg, A., Kang, Y. H., Reid, A., Verfaellie, M., Shadlen, M. N., & Shohamy, D. (2019). The hippocampus supports deliberation during value-based decisions. *Elife*, *8*.

- Balci, F., Simen, P., Niyogi, R., Saxe, A., Hughes, J. A., Holmes, P., & Cohen, J. D. (2011). Acquisition of decision making criteria: Reward rate ultimately beats accuracy. *Atten. Percept. Psychophys.*, *73*, 640–657.
- Barendregt, N. W., Gold, J. I., Josić, K., & Kilpatrick, Z. P. (2022). Normative decision rules in changing environments. *Elife*, *11*.
- Barker, R. M., Armson, M. J., Diamond, N. B., Liu, Z.-X., Wang, Y., Ryan, J. D., & Levine, B. (2025). Remembrance with gazes passed: Eye movements precede continuous recall of episodic details of real-life events. *Cognition*, *268*, 106380.
- Barnard, G. A. (1946). Sequential tests in industrial statistics. *Suppl. J. R. Stat. Soc.*, *8*, 1.
- Batchelder, W. H. (2010). Mathematical psychology. *Wiley Interdiscip. Rev. Cogn. Sci.*, *1*, 759–765.
- Biderman, N., Bakkour, A., & Shohamy, D. (2020). What are memories for? the hippocampus bridges past experience with future decisions. *Trends Cogn. Sci.*, *24*, 542–556.
- Blohm, G., Kording, K. P., & Schrater, P. R. (2020). A how-to-model guide for neuroscience. *eNeuro*, *7*.
- Boehm, U., van Maanen, L., Evans, N. J., Brown, S. D., & Wagenmakers, E.-J. (2020). A theoretical analysis of the reward rate optimality of collapsing decision criteria. *Atten. Percept. Psychophys.*, *82*, 1520–1534.
- 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. *Psychol. Rev.*, *113*, 700–765.
- Bompas, A., Sumner, P., & Hedge, C. (2023). Non-decision time: The higg’s boson of decision. *bioRxiv*, 2023.02.20.529290.
- Bondy, A. G., Charlton, J. A., Luo, T. Z., Kopec, C. D., Stagnaro, W. M., Venditto, S. J. C., Lynch, L., Janarthanan, S., Oline, S. N., Harris, T. D., & Brody, C. D. (2025). Brain-wide coordination of internal signals during decision-making. *bioRxiv*, 2024.08.21.609044.

- Bornstein, A. M., Aly, M., Feng, S. F., Turk-Browne, N. B., Norman, K. A., & Cohen, J. D. (2023). Associative memory retrieval modulates upcoming perceptual decisions. *Cogn. Affect. Behav. Neurosci.*
- Bornstein, A. M., & Daw, N. D. (2012). Dissociating hippocampal and striatal contributions to sequential prediction learning. *Eur. J. Neurosci.*, *35*, 1011–1023.
- Bornstein, A. M., & Daw, N. D. (2013). Cortical and hippocampal correlates of deliberation during model-based decisions for rewards in humans. *PLoS Comput. Biol.*, *9*, e1003387.
- Bornstein, A. M., & Pickard, H. (2020). “chasing the first high”: Memory sampling in drug choice. *Neuropsychopharmacology*, *45*(6), 907–915.
- Box, G. E. P., & Draper, N. R. (1987). Empirical model-building and response surfaces. *Wiley series in probability and mathematical statistics.*, 669.
- Britten, K. H., Newsome, W. T., Shadlen, M. N., Celebrini, S., & Movshon, J. A. (1996). A relationship between behavioral choice and the visual responses of neurons in macaque MT. *Vis. Neurosci.*, *13*, 87–100.
- 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. *J. Neurosci.*, *12*, 4745–4765.
- Brody, C. D., & Hanks, T. D. (2016). Neural underpinnings of the evidence accumulator. *Curr. Opin. Neurobiol.*, *37*, 149–157.
- Brunton, B. W., Botvinick, M. M., & Brody, C. D. (2013). Rats and humans can optimally accumulate evidence for decision-making. *Science*, *340*, 95–98.
- Callaway, F., Griffiths, T. L., Norman, K. A., & Zhang, Q. (2024). Optimal metacognitive control of memory recall. *Psychol. Rev.*, *131*, 781–811.
- Carpenter, R. H., & Williams, M. L. (1995). Neural computation of log likelihood in control of saccadic eye movements. *Nature*, *377*, 59–62.
- Cartwright, N. (1983, June 9). *How the laws of physics lie*. OUP Oxford.

- Chen, J., & Bornstein, A. M. (2024). The causal structure and computational value of narratives. *Trends Cogn. Sci.*, 28, 769–781.
- Colombo, M., & Knauff, M. (2020). Editors’ review and introduction: Levels of explanation in cognitive science: From molecules to culture. *Top. Cogn. Sci.*, 12, 1224–1240.
- Cooper, R. P., & Guest, O. (2014). Implementations are not specifications: Specification, replication and experimentation in computational cognitive modeling. *Cogn. Syst. Res.*, 27, 42–49.
- Craver, C. F. (2007, June 7). *Explaining the brain: Mechanisms and the mosaic unity of neuroscience*. Clarendon Press.
- Danks, D. (2015). Goal-dependence in (scientific) ontology. *Synthese*, 192, 3601–3616.
- Dayan, P., & Abbott, L. F. (2005, August 12). *Theoretical neuroscience: Computational and mathematical modeling of neural systems*. MIT Press.
- de Lange, F. P., Heilbron, M., & Kok, P. (2018). How do expectations shape perception? *Trends Cogn. Sci.*, 22, 764–779.
- de Lange, F. P., Rahnev, D. A., Donner, T. H., & Lau, H. (2013). Prestimulus oscillatory activity over motor cortex reflects perceptual expectations. *J. Neurosci.*, 33, 1400–1410.
- Deneve, S. (2012). Making decisions with unknown sensory reliability. *Front. Neurosci.*, 6, 75.
- Diaz, J. A., Pisauro, M. A., Delis, I., & Philiastides, M. G. (2024). Prior probability biases perceptual choices by modulating the accumulation rate, rather than the baseline, of decision evidence. *Imaging Neurosci. (Camb.)*, 2, imag-2–00338.
- Diederich, A., & Busemeyer, J. R. (2006). Modeling the effects of payoff on response bias in a perceptual discrimination task: Bound-change, drift-rate-change, or two-stage-processing hypothesis. *Percept. Psychophys.*, 68, 194–207.

- Diederich, A., & Oswald, P. (2014). Sequential sampling model for multiattribute choice alternatives with random attention time and processing order. *Front. Hum. Neurosci.*, *8*, 697.
- Diederich, A., & Oswald, P. (2016). Multi-stage sequential sampling models with finite or infinite time horizon and variable boundaries. *J. Math. Psychol.*, *74*, 128–145.
- Drayson, Z. (2020). Why I am not a literalist. *Mind Lang.*, *35*, 661–670.
- Drugowitsch, J., Moreno-Bote, R., Churchland, A. K., Shadlen, M. N., & Pouget, A. (2012). The cost of accumulating evidence in perceptual decision making. *J. Neurosci.*, *32*, 3612–3628.
- Drugowitsch, J., Wyart, V., Devauchelle, A.-D., & Koechlin, E. (2016). Computational precision of mental inference as critical source of human choice suboptimality. *Neuron*, *92*, 1398–1411.
- Dunovan, K., & Wheeler, M. E. (2018). Computational and neural signatures of pre and post-sensory expectation bias in inferior temporal cortex. *Sci. Rep.*, *8*, 13256.
- Dunovan, K. E., Tremel, J. J., & Wheeler, M. E. (2014). Prior probability and feature predictability interactively bias perceptual decisions. *Neuropsychologia*, *61*, 210–221.
- Edwards, W. (1965). Optimal strategies for seeking information: Models for statistics, choice reaction times, and human information processing. *J. Math. Psychol.*, *2*, 312–329.
- Evans, N. J., Hawkins, G. E., Boehm, U., Wagenmakers, E.-J., & Brown, S. D. (2017). The computations that support simple decision-making: A comparison between the diffusion and urgency-gating models. *Sci. Rep.*, *7*, 16433.
- Fang, M., Mao, J., Donner, T. H., & Stocker, A. A. (2026). The resource-rational dynamics of evidence accumulation. *bioRxiv*, 2026.04.15.718716.
- Fetsch, C. R., Kiani, R., Newsome, W. T., & Shadlen, M. N. (2014). Effects of cortical microstimulation on confidence in a perceptual decision. *Neuron*, *84*, 239.

- Fetsch, C. R., Pouget, A., DeAngelis, G. C., & Angelaki, D. E. (2011). Neural correlates of reliability-based cue weighting during multisensory integration. *Nat. Neurosci.*, *15*, 146–154.
- Figdor, C. (2018, April 26). *Pieces of mind: The proper domain of psychological predicates*. Oxford University Press.
- Findling, C., Hubert, F., International Brain Laboratory, Acerbi, L., Benson, B., Benson, J., Birman, D., Bonacchi, N., Buchanan, E. K., Bruijns, S., Carandini, M., Catarino, J. A., Chapuis, G. A., Churchland, A. K., Dan, Y., Davatolhagh, F., DeWitt, E. E. J., Engel, T. A., Fabbri, M., ... Pouget, A. (2025). Brain-wide representations of prior information in mouse decision-making. *Nature*, *645*, 192–200.
- Finn, E. S., Corlett, P. R., Chen, G., Bandettini, P. A., & Constable, R. T. (2018). Trait paranoia shapes inter-subject synchrony in brain activity during an ambiguous social narrative. *Nat. Commun.*, *9*, 2043.
- Forstmann, B. U., Ratcliff, R., & Wagenmakers, E.-J. (2016). Sequential sampling models in cognitive neuroscience: Advantages, applications, and extensions. *Annu. Rev. Psychol.*, *67*, 641–666.
- Forstmann, B. U., Wagenmakers, E.-J., Eichele, T., Brown, S., & Serences, J. T. (2011). Reciprocal relations between cognitive neuroscience and formal cognitive models: Opposites attract? *Trends Cogn. Sci.*, *15*, 272–279.
- Frazier, P., & Yu, A. J. (2007). Sequential hypothesis testing under stochastic deadlines. *Adv. Neural Inf. Process. Syst.*, 465–472.
- Frigg, R., & Hartmann, S. (2020). Models in science. In E. N. Zalta (Ed.), *The stanford encyclopedia of philosophy* (Spring 2020). Metaphysics Research Lab, Stanford University.
- Frigg, R., & Nguyen, J. (2021). Scientific representation (E. N. Zalta, Ed.).
- Gamboa, J. P. (2024). On cognitive modeling and other minds. *Philos. Sci.*, *91*, 615–633.

- Gläscher, J., Daw, N., Dayan, P., & O’Doherty, J. P. (2010). States versus rewards: Dissociable neural prediction error signals underlying model-based and model-free reinforcement learning. *Neuron*, *66*, 585–595.
- Glaze, C. M., Kable, J. W., & Gold, J. I. (2015). Normative evidence accumulation in unpredictable environments. *Elife*, *4*.
- Gold, J. I., & Shadlen, M. N. (2001). Neural computations that underlie decisions about sensory stimuli. *Trends Cogn. Sci.*, *5*, 10–16.
- Gold, J. I., & Shadlen, M. N. (2002). Banburismus and the brain: Decoding the relationship between sensory stimuli, decisions, and reward. *Neuron*, *36*, 299–308.
- Gold, J. I., & Shadlen, M. N. (2007). The neural basis of decision making. *Annu. Rev. Neurosci.*, *30*, 535–574.
- Grahek, I., Schaller, M., & Tackett, J. L. (2021). Anatomy of a psychological theory: Integrating construct-validation and computational-modeling methods to advance theorizing. *Perspect. Psychol. Sci.*, *16*, 803–815.
- Guest, O., & Martin, A. E. (2021). How computational modeling can force theory building in psychological science. *Perspect. Psychol. Sci.*, *16*, 789–802.
- Hachisuka, A., Shor, J. D., Liu, X. C., Friedman, D., Dugan, P., Saez, I., Panov, F. E., Wang, Y., Doyle, W., Devinsky, O., Oermann, E. K., & He, B. J. (2026). Neural and computational mechanisms underlying one-shot perceptual learning in humans. *Nat. Commun.*, *17*, 1204.
- Hanks, T. D., Kopec, C. D., Brunton, B. W., Duan, C. A., Erlich, J. C., & Brody, C. D. (2015). Distinct relationships of parietal and prefrontal cortices to evidence accumulation. *Nature*, *520*, 220–223.
- Hanks, T. D., Mazurek, M. E., Kiani, R., Hopp, E., & Shadlen, M. N. (2011). Elapsed decision time affects the weighting of prior probability in a perceptual decision task. *J. Neurosci.*, *31*, 6339–6352.

- Hanks, T. D., & Summerfield, C. (2017). Perceptual decision making in rodents, monkeys, and humans. *Neuron*, *93*, 15–31.
- Hannula, D. E., Althoff, R. R., Warren, D. E., Riggs, L., Cohen, N. J., & Ryan, J. D. (2010). Worth a glance: Using eye movements to investigate the cognitive neuroscience of memory. *Front. Hum. Neurosci.*, *4*, 166.
- Harhen, N. C., & Bornstein, A. M. (2023). Overharvesting in human patch foraging reflects rational structure learning and adaptive planning. *Proc. Natl. Acad. Sci. U.S.A.*, *120*.
- Harhen, N. C., & Bornstein, A. M. (2024). Interval timing as a computational pathway from early life adversity to affective disorders. *Top. Cogn. Sci.*, *16*, 92–112.
- Harvard, S., & Winsberg, E. (2022). The epistemic risk in representation. *Kennedy Inst. Ethics J.*, *32*, 1–31.
- Hawkins, G. E., Forstmann, B. U., Wagenmakers, E.-J., Ratcliff, R., & Brown, S. D. (2015). Revisiting the evidence for collapsing boundaries and urgency signals in perceptual decision-making. *J. Neurosci.*, *35*, 2476–2484.
- Heilbronner, S. R., & Hayden, B. Y. (2016). The description-experience gap in risky choice in nonhuman primates. *Psychon. Bull. Rev.*, *23*, 593–600.
- Hempel, C. G., & Oppenheim, P. (1948). Studies in the logic of explanation. *Philos. Sci.*, *15*, 135–175.
- Hertwig, R., & Erev, I. (2009). The description-experience gap in risky choice. *Trends Cogn. Sci.*, *13*, 517–523.
- Herweg, N. A., Solomon, E. A., & Kahana, M. J. (2020). Theta oscillations in human memory. *Trends Cogn. Sci.*, *24*, 208–227.
- Holmes, P., & Cohen, J. D. (2014). Optimality and some of its discontents: Successes and shortcomings of existing models for binary decisions. *Top. Cogn. Sci.*, *6*, 258–278.
- Huang, Y., Friesen, A. L., Hanks, T. D., Shadlen, M. N., & Rao, R. P. N. (2012). How prior probability influences decision making: A unifying probabilistic model. *Adv. Neural Inf. Process. Syst.*, *25*.

- Jacobs, J. (2014). Hippocampal theta oscillations are slower in humans than in rodents: Implications for models of spatial navigation and memory. *Philos. Trans. R. Soc. Lond. B Biol. Sci.*, *369*, 20130304.
- Jones, M., & Dzhafarov, E. N. (2014). Analyzability, ad hoc restrictions, and excessive flexibility of evidence-accumulation models: Reply to two critical commentaries. *Psychol. Rev.*, *121*, 689–695.
- Kaelbling, L. P., Littman, M. L., & Cassandra, A. R. (1998). Planning and acting in partially observable stochastic domains. *Artif. Intell.*, *101*, 99–134.
- Kaper, R., & Peters, M. (2026). Stimulus familiarity shapes hierarchical structure learning and metacognitive dynamics. *PsyArXiv*.
- Kaplan, D. M. (2017, December 1). *Explanation and integration in mind and brain science*. Oxford University Press.
- Kelly, S. P., Corbett, E. A., & O’Connell, R. G. (2021). Neurocomputational mechanisms of prior-informed perceptual decision-making in humans. *Nat Hum Behav*, *5*, 467–481.
- Kerrén, C., Linde-Domingo, J., Hanslmayr, S., & Wimber, M. (2018). An optimal oscillatory phase for pattern reactivation during memory retrieval. *Curr. Biol.*, *28*, 3383–3392.e6.
- Kerrén, C., Michelmann, S., & Doeller, C. F. (2026). Hippocampal ripples initiate cortical dimensionality expansion for memory retrieval. *Nat. Commun.*, *17*, 6677.
- Khilkevich, A., Lohse, M., Low, R., Orsolic, I., Bozic, T., Windmill, P., & Masic-Flogel, T. D. (2024). Brain-wide dynamics linking sensation to action during decision-making. *Nature*, 1–11.
- Khodadadi, A., Fakhari, P., & Busemeyer, J. R. (2014). Learning to maximize reward rate: A model based on semi-markov decision processes. *Front. Neurosci.*, *8*, 101.
- Khoudary, A., Bornstein, A., & Peters, M. (2025a). Investigating implicit and explicit expectations in perceptual decision making. *Proceedings of the 47th Annual Conference of the Cognitive Sciences Society*, 6110–6116.

- Khoudary, A., Peters, M. A. K., & Bornstein, A. M. (2022). Precision-weighted evidence integration predicts time-varying influence of memory on perceptual decisions. *Cognitive Computational Neuroscience*.
- Khoudary, A., Peters, M. A. K., & Bornstein, A. M. (2025b). Reasoning goals and representational decisions in computational cognitive neuroscience: Lessons from the drift diffusion model. *Eur. J. Neurosci.*, *61*, e70098.
- Khoudary, A., Peters, M. A. K., & Bornstein, A. M. (2026). Reliability-weighted integration explains time-varying effects of expectations on perceptual decisions. *psyArxiv*.
- Kiani, R., Churchland, A. K., & Shadlen, M. N. (2013). Integration of direction cues is invariant to the temporal gap between them. *J. Neurosci.*, *33*, 16483–16489.
- Kiani, R., Corthell, L., & Shadlen, M. N. (2014). Choice certainty is informed by both evidence and decision time. *Neuron*, *84*, 1329–1342.
- Kiani, R., & Shadlen, M. N. (2009). Representation of confidence associated with a decision by neurons in the parietal cortex [this is the paper where they introduce post decision wagering (with great success)]. *Science*, *324*, 759–764.
- Kilpatrick, Z. P., Holmes, W. R., Eissa, T. L., & Josić, K. (2019). Optimal models of decision-making in dynamic environments. *Curr. Opin. Neurobiol.*, *58*, 54–60.
- Kinney, D., & Lombrozo, T. (2024). Building compressed causal models of the world. *Cogn. Psychol.*, *155*, 101682.
- Kobayashi, K., Lee, S., Filipowicz, A. L. S., McGaughey, K. D., Kable, J. W., & Nassar, M. R. (2021). Dynamic representation of the subjective value of information. *J. Neurosci.*, *41*, 8220–8232.
- Kok, P., Brouwer, G. J., van Gerven, M. A. J., & de Lange, F. P. (2013). Prior expectations bias sensory representations in visual cortex. *J. Neurosci.*, *33*, 16275–16284.
- Kok, P., Jehee, J. F. M., & de Lange, F. P. (2012). Less is more: Expectation sharpens representations in the primary visual cortex. *Neuron*, *75*, 265–270.

- Kording, K., Blohm, G., Schrater, P., & Kay, K. (2018). Appreciating diversity of goals in computational neuroscience.
- Kraemer, P. M., & Gluth, S. (2023). Episodic memory retrieval affects the onset and dynamics of evidence accumulation during value-based decisions. *J. Cogn. Neurosci.*, *35*, 692–714.
- Kragel, J. E., VanHaerents, S., Templer, J. W., Schuele, S., Rosenow, J. M., Nilakantan, A. S., & Bridge, D. J. (2020). Hippocampal theta coordinates memory processing during visual exploration. *Elife*, *9*.
- Krajbich, I., Armel, C., & Rangel, A. (2010). Visual fixations and the computation and comparison of value in simple choice. *Nat. Neurosci.*, *13*, 1292–1298.
- Krueger, P. M., van Vugt, M. K., Simen, P., Nystrom, L., Holmes, P., & Cohen, J. D. (2017). Evidence accumulation detected in BOLD signal using slow perceptual decision making. *J. Neurosci. Methods*, *281*, 21–32.
- Kumle, L., Nobre, A. C., & Draschkow, D. (2025). Sensorimnemonic decisions: Choosing memories versus sensory information. *Trends Cogn. Sci.*
- Ladyka-Wojcik, N., Flick, G., Liu, Z.-X., & Ryan, J. D. (2026). Where we look, what we know: How hippocampal predictions shape, and are shaped by, eye movements. *Philos. Trans. R. Soc. Lond. B Biol. Sci.*, *381*, 20250248.
- Ladyka-Wojcik, N., Schmidt, H., Cooper, R. A., & Ritchey, M. (2025). Neural signatures of recollection are sensitive to memory quality and specific event features. *J. Cogn. Neurosci.*, 1–17.
- Laming, D. R. J. (1968). *Information theory of choice-reaction times*. Wiley.
- Landy, M. S., Banks, M. S., & Knill, D. C. (2011, September 21). Ideal-observer models of cue integration. In J. Trommershauser, K. Kording, & M. S. Landy (Eds.), *Sensory cue integration*. Oxford University Press.
- Lee, M. D., & Wagenmakers, E.-J. (2014, April 3). *Bayesian cognitive modeling: A practical course*. Cambridge University Press.

- Lewis, R. L., Howes, A., & Singh, S. (2014). Computational rationality: Linking mechanism and behavior through bounded utility maximization. *Top. Cogn. Sci.*, *6*, 279–311.
- Lieder, F., & Griffiths, T. L. (2019). Resource-rational analysis: Understanding human cognition as the optimal use of limited computational resources. *Behav. Brain Sci.*, *43*, e1.
- Lin, Y.-S., & Strickland, L. (2020). Evidence accumulation models with R: A practical guide to hierarchical bayesian methods. *Quant. Methods Psychol.*, *16*, 133–153.
- Link, S. W., & Heath, R. A. (1975). A sequential theory of psychological discrimination. *Psychometrika*, *40*, 77–105.
- Link, S. W. (1975). The relative judgment theory of two choice response time. *J. Math. Psychol.*, *12*, 114–135.
- Liu, A., Schartner, M., International Brain Laboratory, & Fiete, I. (2025). How learned expectations shape brain-wide responses. *bioRxiv*, 2025.12.15.694430.
- Malhotra, G., Leslie, D. S., Ludwig, C. J. H., & Bogacz, R. (2017). Overcoming indecision by changing the decision boundary. *J. Exp. Psychol. Gen.*, *146*, 776–805.
- Malhotra, G., Leslie, D. S., Ludwig, C. J. H., & Bogacz, R. (2018). Time-varying decision boundaries: Insights from optimality analysis. *Psychon. Bull. Rev.*, *25*, 971–996.
- Maniscalco, B., & Peters, M. A. K. (2024). Achieving constant hazard rate in psychophysics experiments. *PsyArXiv*.
- Martin, C. B., & Barense, M. D. (2023). Perception and memory in the ventral visual stream and medial temporal lobe. *Annu. Rev. Vis. Sci.*, *9*, 409–434.
- Mazor, M., Moran, R., & Press, C. (2026). Beliefs about perception shape perceptual inference: An ideal observer model of detection. *Psychol. Rev.*, *133*, 271–295.
- Miletić, S., & van Maanen, L. (2019). Caution in decision-making under time pressure is mediated by timing ability. *Cogn. Psychol.*, *110*, 16–29.
- Moran, R. (2015). Optimal decision making in heterogeneous and biased environments. *Psychon. Bull. Rev.*, *22*, 38–53.

- Mulder, M. J., Wagenmakers, E.-J., Ratcliff, R., Boekel, W., & Forstmann, B. U. (2012). Bias in the brain: A diffusion model analysis of prior probability and potential payoff. *J. Neurosci.*, *32*, 2335–2343.
- Muthukrishna, M., & Henrich, J. (2019). A problem in theory. *Nat Hum Behav*, *3*, 221–229.
- Newsome, W. T., Britten, K. H., & Movshon, J. A. (1989). Neuronal correlates of a perceptual decision. *Nature*, *341*, 52–54.
- Nicholas, J., Chen, S., & Mattar. (2026). Looking at nothing during deliberation reflects optimal sampling from episodic memory. *9th Conference on Cognitive Computational Neuroscience*.
- Nicholas, J., & Mattar, M. G. (2026). Episodic memory facilitates flexible decision-making via access to detailed events. *Nat. Hum. Behav.*, 1–17.
- Noh, S. M., Singla, U. K., Bennett, I. J., & Bornstein, A. M. (2023). Memory precision and age differentially predict the use of decision-making strategies across the lifespan. *Sci. Rep.*, *13*, 17014.
- O’Connell, R. G., Dockree, P. M., & Kelly, S. P. (2012). A supramodal accumulation-to-bound signal that determines perceptual decisions in humans. *Nat. Neurosci.*, *15*, 1729–1735.
- O’Connell, R. G., & Kelly, S. P. (2021). Neurophysiology of human perceptual decision-making. *Annu. Rev. Neurosci.*, *44*, 495–516.
- O’Connell, R. G., Shadlen, M. N., Wong-Lin, K., & Kelly, S. P. (2018). Bridging neural and computational viewpoints on perceptual decision-making. *Trends Neurosci.*, *41*, 838–852.
- Oliva, A., & Torralba, A. (2007). The role of context in object recognition. *Trends Cogn. Sci.*, *11*, 520–527.
- Pagan, M., Tang, V. D., Aoi, M. C., Pillow, J. W., Mante, V., Sussillo, D., & Brody, C. D. (2025). Individual variability of neural computations underlying flexible decisions. *Nature*, *639*, 421–429.

- Palestro, J. J., Weichart, E., Sederberg, P. B., & Turner, B. M. (2018). Some task demands induce collapsing bounds: Evidence from a behavioral analysis. *Psychon. Bull. Rev.*, *25*, 1225–1248.
- Palmer, J., Huk, A. C., & Shadlen, M. N. (2005a). The effect of stimulus strength on the speed and accuracy of a perceptual decision. *J. Vis.*, *5*, 376–404.
- Palmer, J., McKinley, M. K., Mazurek, M., & Shadlen, M. N. (2005b). Effect of prior probability on choice and response time in a motion discrimination task. *J. Vis.*, *5*, 235–235.
- Parker, W. S. (2020). Model evaluation: An adequacy-for-purpose view. *Philos. Sci.*, *87*, 457–477.
- Pinto, L., Tank, D. W., & Brody, C. D. (2022). Multiple timescales of sensory-evidence accumulation across the dorsal cortex. *Elife*, *11*.
- Polanía, R., Krajbich, I., Grueschow, M., & Ruff, C. C. (2014). Neural oscillations and synchronization differentially support evidence accumulation in perceptual and value-based decision making. *Neuron*, *82*, 709–720.
- Portides, D. (2021). Idealization and abstraction in scientific modeling. *Synthese*, *198*, 5873–5895.
- Poskanzer, C., & Aly, M. (2023). Switching between external and internal attention in hippocampal networks. *J. Neurosci.*, *43*, 6538–6552.
- Potochnik, A. (2018, January 12). *Idealization and the aims of science*. University of Chicago Press.
- Potochnik, A., & Sanches de Oliveira, G. (2020). Patterns in cognitive phenomena and pluralism of explanatory styles. *Top. Cogn. Sci.*, *12*, 1306–1320.
- Press, C., Kok, P., & Yon, D. (2020). The perceptual prediction paradox. *Trends Cogn. Sci.*, *24*, 13–24.
- Press, C., Yon, D., & Heyes, C. (2022). Building better theories. *Curr. Biol.*, *32*, R13–R17.

- Rahnev, D., & Denison, R. N. (2018). Suboptimality in perceptual decision making. *Behav. Brain Sci.*, *41*, e223.
- Rao, R. P. N. (2010). Decision making under uncertainty: A neural model based on partially observable markov decision processes. *Front. Comput. Neurosci.*, *4*, 146.
- Rao, V., DeAngelis, G. C., & Snyder, L. H. (2012). Neural correlates of prior expectations of motion in the lateral intraparietal and middle temporal areas. *J. Neurosci.*, *32*, 10063–10074.
- Ratcliff, R. (1978). A theory of memory retrieval. *Psychol. Rev.*, *85*, 59–108.
- Ratcliff, R. (1980). A note on modeling accumulation of information when the rate of accumulation changes over time. *J. Math. Psychol.*, *21*, 178–184.
- Ratcliff, R., & McKoon, G. (2008). The diffusion decision model: Theory and data for two-choice decision tasks. *Neural Comput.*, *20*, 873–922.
- Ratcliff, R., & Rouder, J. N. (1998). Modeling response times for two-choice decisions. *Psychol. Sci.*, *9*, 347–356.
- Ratcliff, R., Smith, P. L., Brown, S. D., & McKoon, G. (2016). Diffusion decision model: Current issues and history. *Trends Cogn. Sci.*, *20*, 260–281.
- Ratcliff, R., & Tuerlinckx, F. (2002). Estimating parameters of the diffusion model: Approaches to dealing with contaminant reaction times and parameter variability. *Psychon. Bull. Rev.*, *9*, 438–481.
- Richmond, A. (2024). How computation explains. *Mind Lang.*
- Rittershofer, K., Wang, Y., Eimer, M., Kok, P., Yon, D., & Clare Press. (2026). Paradoxical influences of prediction are resolved across time. *bioRxiv*, 2026.03.05.709935.
- Robinaugh, D. J., Haslbeck, J. M. B., Ryan, O., Fried, E. I., & Waldorp, L. J. (2021). Invisible hands and fine calipers: A call to use formal theory as a toolkit for theory construction. *Perspect. Psychol. Sci.*, *16*, 725–743.

- Roitman, J. D., & Shadlen, M. N. (2002). Response of neurons in the lateral intraparietal area during a combined visual discrimination reaction time task. *J. Neurosci.*, *22*, 9475–9489.
- Ross, L., & Woodward, J. (2023, March 17). Causal approaches to scientific explanation. In E. N. Zalta & U. Nodelman (Eds.), *The stanford encyclopedia of philosophy*. <https://plato.stanford.edu/archives/spr2023/entries/causal-explanation-science/>.
- Rouault, M., Dayan, P., & Fleming, S. M. (2019). Forming global estimates of self-performance from local confidence. *Nat. Commun.*, *10*, 1141.
- Roweis, S., & Ghahramani, Z. (1998). A unifying review of linear gaussian models. *Neural Comput.*, *11*, 34.
- Rungratsameetaweemana, N., Itthipuripat, S., Salazar, A., & Serences, J. T. (2018). Expectations do not alter early sensory processing during perceptual decision-making. *J. Neurosci.*, *38*, 5632–5648.
- Rungratsameetaweemana, N., & Serences, J. T. (2019). Dissociating the impact of attention and expectation on early sensory processing. *Curr. Opin. Psychol.*, *29*, 181–186.
- Ryan, J. D., & Shen, K. (2020). The eyes are a window into memory. *Current Opinion in Behavioral Sciences*, *32*, 1–6.
- Schacter, D. L., & Addis, D. R. (2007). The cognitive neuroscience of constructive memory: Remembering the past and imagining the future. *Philos. Trans. R. Soc. Lond. B Biol. Sci.*, *362*, 773–786.
- Schubert, A.-L., Hagemann, D., Voss, A., Schankin, A., & Bergmann, K. (2015). Decomposing the relationship between mental speed and mental abilities. *Intelligence*, *51*, 28–46.
- Seger, S. E., Kriegel, J. L. S., Lega, B. C., & Ekstrom, A. D. (2023). Memory-related processing is the primary driver of human hippocampal theta oscillations. *Neuron*, *111*, 3119–3130.e4.

- Senoussi, M., Galas, L., Busch, N. A., & Dugué, L. (2025). Theta-rhythmic attentional exploration of space. *bioRxiv*, 2025.08.16.670674.
- Seriès, P., & Seitz, A. R. (2013). Learning what to expect (in visual perception). *Front. Hum. Neurosci.*, 7, 668.
- Shadlen, M. N., & Newsome, W. T. (2001). Neural basis of a perceptual decision in the parietal cortex (area LIP) of the rhesus monkey. *J. Neurophysiol.*, 86, 1916–1936.
- Shadlen, M. N., & Shohamy, D. (2016). Decision making and sequential sampling from memory. *Neuron*, 90, 927–939.
- Shi, Y.-L., Zeraati, R., International Brain Laboratory, Levina, A., & Engel, T. A. (2025). Brain-wide organization of intrinsic timescales at single-neuron resolution. *bioRxiv.org*.
- Shi, Y., & Garvert, M. M. (2026). Forming, selecting and updating cognitive maps to support flexible behaviour: Relevance for anxiety. *Philos. Trans. R. Soc. Lond. B Biol. Sci.*, 381, 20250246.
- Shinn, M., Lam, N. H., & Murray, J. D. (2020). A flexible framework for simulating and fitting generalized drift-diffusion models. *Elife*, 9.
- Shushruth, S., Zylberberg, A., & Shadlen, M. N. (2022). Sequential sampling from memory underlies action selection during abstract decision-making. *Curr. Biol.*, 32, 1949–1960.e5.
- Simen, P., Contreras, D., Buck, C., Hu, P., Holmes, P., & Cohen, J. D. (2009). Reward rate optimization in two-alternative decision making: Empirical tests of theoretical predictions. *J. Exp. Psychol. Hum. Percept. Perform.*, 35, 1865–1897.
- Smith, P. L., Ratcliff, R., & McKoon, G. (2014). The diffusion model is not a deterministic growth model: Comment on jones and dzhafarov (2014). *Psychol. Rev.*, 121, 679–688.
- Srivastava, V., Feng, S. F., Cohen, J. D., Leonard, N. E., & Shenhav, A. (2017). A martingale analysis of first passage times of time-dependent wiener diffusion models. *J. Math. Psychol.*, 77, 94–110.

- Stanovich, K. E., & West, R. F. (2000). Individual differences in reasoning: Implications for the rationality debate? *Behav. Brain Sci.*, *23*, 645–65, discussion 665–726.
- Stine, G. M., Zylberberg, A., Ditterich, J., & Shadlen, M. N. (2020). Differentiating between integration and non-integration strategies in perceptual decision making. *Elife*, *9*.
- Stone, M. (1960). Models for choice-reaction time. *Psychometrika*, *25*, 251–260.
- Summerfield, C., & de Lange, F. P. (2014). Expectation in perceptual decision making: Neural and computational mechanisms. *Nat. Rev. Neurosci.*, *15*, 745–756.
- Sutton, R. S., & Barto, A. G. (2018, November 13). *Reinforcement learning: An introduction, 2nd edition*. MIT Press.
- Swoyer, C. (1991). Structural representation and surrogate reasoning. *Synthese*, *87*, 449–508.
- Tajima, S., Drugowitsch, J., Patel, N., & Pouget, A. (2019). Optimal policy for multi-alternative decisions. *Nat. Neurosci.*, *22*, 1503–1511.
- Talarico, J. M., & Rubin, D. C. (2007). Flashbulb memories are special after all; in phenomenology, not accuracy. *Appl. Cogn. Psychol.*, *21*, 557–578.
- Tarder-Stoll, H., Sekeres, M. J., Levine, B., & Moscovitch, M. (2026). Adaptive episodic memory: How multiple memory representations drive behavior in humans and non-humans. *Physiol. Rev.*, *106*, 841–889.
- Tavares, G., Perona, P., & Rangel, A. (2017). The attentional drift diffusion model of simple perceptual decision-making. *Front. Neurosci.*, *11*, 468.
- Ter Wal, M., Linde-Domingo, J., Lifanov, J., Roux, F., Kolibius, L. D., Gollwitzer, S., Lang, J., Hamer, H., Rollings, D., Sawlani, V., Chelvarajah, R., Staresina, B., Hanslmayr, S., & Wimber, M. (2021). Theta rhythmicity governs human behavior and hippocampal signals during memory-dependent tasks. *Nat. Commun.*, *12*, 7048.
- Tetlock, P. E., & Mellers, B. A. (2002). The great rationality debate. *Psychol. Sci.*, *13*, 94–99.
- Thura, D., Beauregard-Racine, J., Fradet, C.-W., & Cisek, P. (2012). Decision making by urgency gating: Theory and experimental support. *J. Neurophysiol.*, *108*, 2912–2930.

- Tsvinev, A., Pilipenko, A., & Samaha, J. (2025). A common neural signal of evidence accumulation for perceptual and mnemonic decisions. *bioRxiv*, 2025.11. 13.688140.
- Turner, B. M., Forstmann, B. U., & Steyvers, M. (2019a). *Joint models of neural and behavioral data* (1st ed.). Springer Nature.
- Turner, B. L., Caruso, E. M., Dilich, M. A., & Roese, N. J. (2019b). Body camera footage leads to lower judgments of intent than dash camera footage. *Proc. Natl. Acad. Sci. U. S. A.*, *116*, 1201–1206.
- van Ravenzwaaij, D., Mulder, M. J., Tuerlinckx, F., & Wagenmakers, E.-J. (2012). Do the dynamics of prior information depend on task context? an analysis of optimal performance and an empirical test. *Front. Psychol.*, *3*, 132.
- van Rooij, I. (2022). Psychological models and their distractors. *Nature Reviews Psychology*, *1*, 127–128.
- van Vugt, M. K., Beulen, M. A., & Taatgen, N. A. (2019). Relation between centro-parietal positivity and diffusion model parameters in both perceptual and memory-based decision making. *Brain Res.*, *1715*, 1–12.
- Vandekerckhove, J., Matzke, D., & Wagenmakers, E. (2015, April 30). Model comparison and the principle of parsimony. In J. R. Busemeyer, Z. Wang, J. T. Townsend, & A. Eidels (Eds.), *The oxford handbook of computational and mathematical psychology* (pp. 300–319). Oxford University Press.
- Verschooren, S., Dahl, M. J., Aly, M., & Mittner, M. (2026). Transition dynamics of external and internal attention across on-task and off-task states. *Nat. Rev. Psychol.*, 1–17.
- Vickers, D. (1970). Evidence for an accumulator model of psychophysical discrimination. *Ergonomics*, *13*, 37–58.
- Villarreal, M., Chávez De la Peña, A. F., Mistry, P. K., Menon, V., Vandekerckhove, J., & Lee, M. D. (2024). Bayesian graphical modeling with the circular drift diffusion model. *Comput. Brain Behav.*, *7*, 181–194.

- Vivekananda, U., Bush, D., Bisby, J. A., Baxendale, S., Rodionov, R., Diehl, B., Chowdhury, F. A., McEvoy, A. W., Miserocchi, A., Walker, M. C., & Burgess, N. (2021). Theta power and theta-gamma coupling support long-term spatial memory retrieval. *Hippocampus*, *31*, 213–220.
- Vul, E., Goodman, N., Griffiths, T. L., & Tenenbaum, J. B. (2014). One and done? optimal decisions from very few samples. *Cogn. Sci.*, *38*, 599–637.
- Wagenmakers, E.-J., van der Maas, H. L. J., & Grasman, R. P. P. P. (2007). An EZ-diffusion model for response time and accuracy. *Psychon. Bull. Rev.*, *14*, 3–22.
- Wald, A., & Wolfowitz, J. (1948). Optimum character of the sequential probability ratio test. *Ann. Math. Stat.*, *19*, 326–339.
- Wald, A. (1945). Sequential tests of statistical hypotheses. *Ann. Math. Stat.*, *16*, 117–186.
- Wald, A. (1947). Sequential analysis. *212*.
- Wang, S., Feng, S. F., & Bornstein, A. M. (2021). Mixing memory and desire: How memory reactivation supports deliberative decision-making. *Wiley Interdiscip. Rev. Cogn. Sci.*, *13*, e1581.
- Watson, A. B., & Pelli, D. G. (1983). QUEST: A bayesian adaptive psychometric method. *Percept. Psychophys.*, *33*, 113–120.
- Weber, J., Solbakk, A.-K., Blenkmann, A. O., Llorens, A., Funderud, I., Leske, S., Larsson, P. G., Ivanovic, J., Knight, R. T., Endestad, T., & Helfrich, R. F. (2024). Ramping dynamics and theta oscillations reflect dissociable signatures during rule-guided human behavior. *Nat. Commun.*, *15*, 637.
- Weisberg, M. (2013). *Simulation and similarity: Using models to understand the world*. Oxford University Press.
- Wert, S., Seidle, A., Rissman, J., & Knowlton, B. J. (2025). The temporal evolution of implicit bias in perceptual decision-making. *CogSci*, *47*.
- Wilson, R. C., & Collins, A. G. (2019). Ten simple rules for the computational modeling of behavioral data. *Elife*, *8*.

- Wimsatt, W. C. (1987). False models as means to truer theories. In M. H. Nitecki & A. Hoffman (Eds.), *Neutral models in biology* (pp. 23–55). Oxford University Press.
- Winsberg, E., & Harvard, S. (2024, January). Scientific models and decision making. In *Elements in the philosophy of science*. Cambridge University Press.
- Wynn, J. S., Bone, M. B., Dragan, M. C., Hoffman, K. L., Buchsbaum, B. R., & Ryan, J. D. (2016). Selective scanpath repetition during memory-guided visual search. *Vis. Cogn.*, *24*, 15–37.
- Yang, X., & Krajbich, I. (2023). A dynamic computational model of gaze and choice in multi-attribute decisions. *Psychol. Rev.*, *130*, 52–70.
- Yoo, J., & Bornstein, A. (2024). Temporal dynamics of model-based control reveal arbitration between multiple task representations. *PsyArXiv*.
- Zeithamova, D., Schlichting, M. L., & Preston, A. R. (2012). The hippocampus and inferential reasoning: Building memories to navigate future decisions. *Front. Hum. Neurosci.*, *6*, 70.
- Zhang, H., Watrous, A. J., Patel, A., & Jacobs, J. (2018). Theta and alpha oscillations are traveling waves in the human neocortex. *Neuron*, *98*, 1269–1281.e4.
- Zhou, D., Noh, S. M., Harhen, N. C., Banavar, N. V., Kirwan, C. B., Yassa, M. A., & Bornstein, A. M. (2025). A compressed code for memory discrimination. *bioRxiv.org*, 2025.10.12.681901.
- Zhou, Z., & Geng, J. J. (2024). Learned associations serve as target proxies during difficult but not easy visual search. *Cognition*, *242*, 105648.
- Zubair, H. N., Stangl, M., Topalovic, U., Inman, C., Seeber, M., Hiller, S., Rao, V. R., Halpern, C. H., Eliashiv, D., Fried, I., & Suthana, N. (2026). Eye movements reflect memory-related theta activity in the human brain. *PLoS Biol.*, *24*, e3003695.
- Zylberberg, A., Bakkour, A., Shohamy, D., & Shadlen, M. N. (2024). Value construction through sequential sampling explains serial dependencies in decision making. *Elife*, *13*, RP96997.

Zylberberg, A., Fetsch, C. R., & Shadlen, M. N. (2016). The influence of evidence volatility on choice, reaction time and confidence in a perceptual decision. *Elife*, 5.

Appendix A

Supplementary Materials for Chapter 2

A.1 Qualitative comparison of normative starting point models

Visualizing prescriptions made by the starting point models of Edwards (1965) and Link (1975) demonstrates the exponential scaling of optimal starting points by the signal-to-noise ratio of sensory evidence in the environment.

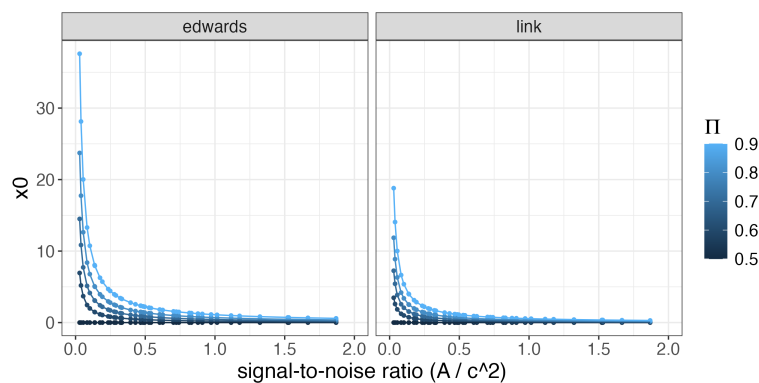


Figure A.1: Optimal starting points on the models of Edwards (1965) and Link (1975).

A.2 Sequential effects analysis

This section contains details about the two studies used to demonstrate sequential effects of expectations on RTs. For each study, we provide a brief summary, followed by the regression model specification and a full summary table of fitted coefficients. Reaction times were log-transformed and z-scored within subject prior to being entered into the regression. All predictors were defined categorically to maximize the descriptive ability of each model (i.e., we did not force reaction times to scale numerically with cue predictiveness), except for sensory evidence strength and time on task, which were treated as continuous. Models were fit using the `lm()` function in R 4.4.2, and summary tables were produced using the `tbl_regression()` function from the `gtsummary` package.

A.2.1 Diaz, Pisauro et al. (2024): instructed expectations

The data from Diaz, Pisauro et al. (2024) were generated by human observers (n=16) performing a face/car discrimination task in a fully heterogeneous environment with explicitly instructed Π_{cue} . Each trial began with a brief (750ms) presentation of one of three possible cues—30F, 50F, or 70F—overtly instructing observers on the trial-level probability that the correct answer will be face. A noise-degraded stimulus was then briefly (50ms) presented, and observers had up to 1250ms to indicate whether the stimulus contained an image of a car or a face. Stimuli were presented at moderate levels of signal strength (coherence = 32.5% and 37.5%), and observers received no feedback on the accuracy of their choice. Each cue was presented 360 times for a total of 1080 trials per subject.

We modeled trial-by-trial transformed reaction times (RT_t) using the following formula:

$$RT_t \sim Accuracy_t + Validity_t + PreviousCue_t + PreviousTarget_t + PreviousResponse_t \quad (A.1)$$

The results of this model are displayed in Table A.1 below. All factors canonically known to modify RTs—accuracy, cue validity, and the accuracy of the previous response—exhibited significant main effects ($\beta_{accuracy} = -0.23$, $\beta_{valid} = -0.12$, $\beta_{invalid} = 0.05$, $\beta_{prev_correct} = -0.03$). Crucially, switching cues across trials slowed RTs ($\beta_{prev_cue} = 0.03$) but switching *targets* across trials had no effect on RTs ($\beta_{prev_target} = 0.00$).

Predictor	Beta	SE	t-statistic	95% CI	p-value
accuracy					
correct	-0.20	0.011	-17.4	-0.22, -0.17	< 0.001
incorrect	—	—	—	—	
validity					
valid	-0.12	0.010	-12.0	-0.14, -0.10	< 0.001
invalid	0.05	0.010	5.23	0.03, 0.07	< 0.001
neutral	—	—	—	—	
prev_cue					
different	0.03	0.008	3.86	0.01, 0.04	< 0.001
same	—	—	—	—	
prev_target					
different	0.00	0.007	0.481	-0.01, 0.02	0.6
same	—	—	—	—	
prev_response					
correct	0.00	0.011	0.290	-0.02, 0.03	0.8
error	—	—	—	—	
logTrial	-0.23	0.007	-31.2	-0.25, -0.22	< 0.001

Table A.1: Full regression table for Diaz, Pisauro et al. (2024) analysis. No. Obs. = 16,696; Sigma = 0.915; Residual df = 16,688.

A.2.2 Bornstein et al. (2023): learned expectations

Data from Bornstein et al. (2023) were generated by human observers ($n=33$) performing a category-level discrimination task of noise-degraded, perceptually-similar grayscale images (i.e., discriminating Face A from Face B). Π_{cue} was held constant for within-category stimuli, such that one category was always more predictable than the other, but each stimulus within the category had a unique cue (colored fractal). The predictability of each category (faces and scenes) was manipulated across blocks, such that each observer learned distinct cues on the range $\Pi_{cue} = \{0.5, 0.6, 0.7, 0.8\}$. Within each block, observers had 30 trials per cue (120 trials total) to learn the predictive relationship of each unique colored fractal and its corresponding grayscale image; this learning procedure ensures that observers have a source-related uncertainty in their expectations. Sensory evidence was presented at “low” and “high” signal strength, corresponding to 65% and 85% discrimination accuracy in the absence of cues. Coherence was yoked to stimulus category, such that stimuli with more predictive cues were always presented with low coherence whereas stimuli with less predictive cues were always presented with high coherence. Each trial began with a brief (750ms) presentation of one of the four fractal cues, followed by a variable delay of 4000, 6000, or 8000ms. Observers were then shown a 60Hz “stream” of within-category images, and they had up to 3000ms to report which image was presented more frequently. Immediately after responding, the correct image appeared on screen surrounded by a colored box indicating whether the response was correct (green) or incorrect (red). There was no time penalty for incorrect responses. Each block contained 20 decision trials per cue, such that each observer contributed 160 trials total.

We modeled trial-by-trial transformed reaction times (RT_t) using the following formula:

$$RT_t \sim \text{cueLevel}_t * \text{validity}_t + \text{Accuracy}_t + \text{PreviousCue}_t + \text{PreviousTarget}_t + \text{PreviousResponse}_t \quad (\text{A.2})$$

This model is slightly more complex than the others due to the nature of the task. Bornstein et al. (2023) measured effects of four different cue levels for each subject, whereas the other studies only had one level of cue. Accordingly, it was necessary for cue validity to interact with the strength of the cue for this dataset, whereas the other datasets could use cue validity alone as the predictor. We treated `cueLevel` as a categorical variable to ensure that any variability due to errors in learning would be captured by the model (as opposed to enforcing a strictly linear effect of the cue’s true predictive probability).

The results of this model are displayed in Table A.2 below. Choice accuracy, cue level, and previous cue all exhibited significant main effects on reaction times ($\beta_{\text{accuracy}} = 0.16$, $\beta_{\text{cueLevel}_{0.5}} = 0.21$, $\beta_{\text{cueLevel}_{0.7}} = -0.06$, $\beta_{\text{prev_cue}} = 0.06$). The positive relationship between correct responses and reaction times is due to including both “early” and “late” responses in the same regression; the canonical speed-accuracy relationship is observed when β values are estimated independently for “early” and “late” responses (Bornstein et al., 2023).

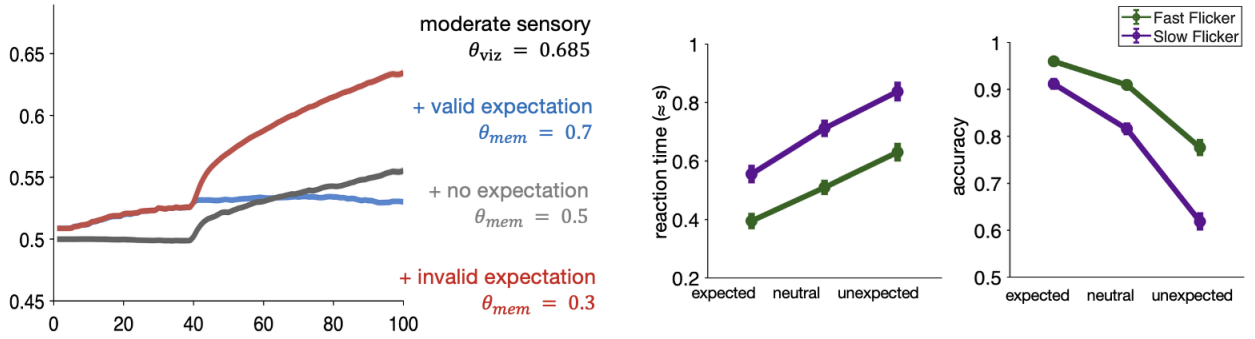
Predictor	Beta	SE	t-statistic	95% CI	p-value
Accuracy					
correct	0.16	0.016	9.95	0.13, 0.19	< 0.001
incorrect	—	—	—	—	
cueLevel					
0.5	0.21	0.025	8.16	0.16, 0.26	< 0.001
0.6	0.05	0.026	1.79	0.00, 0.10	0.074
0.7	-0.05	0.027	-2.03	-0.11, 0.00	0.043
0.8	—	—	—	—	
validity					
invalid	0.02	0.016	1.36	-0.01, 0.05	0.2
valid	—	—	—	—	
prev_cue					
different	0.06	0.017	3.15	0.02, 0.09	0.002
same	—	—	—	—	
prev_target					
different	-0.01	0.020	-0.343	-0.05, 0.03	0.7
same	—	—	—	—	
prev_response					
correct	-0.01	0.015	-0.926	-0.04, 0.02	0.4
error	—	—	—	—	
log(Trial)	-0.24	0.085	-2.81	-0.41, -0.07	0.005
cueLevel * validity					
0.5 * invalid	-0.03	0.026	-1.24	-0.08, 0.02	0.2
0.6 * invalid	0.00	0.026	-0.051	-0.05, 0.05	> 0.9
0.7 * invalid	0.02	0.027	0.571	-0.04, 0.07	0.6

Table A.2: Full regression table for Bornstein et al. (2023). No. Obs. = 4,686; Sigma = 0.973; Statistic = 18.9; Residual df = 4,673.

A.2.3 Rungratsameetaweemana, Itthipuripat et al. (2018) effect

This section shows that the main-text results reported in Figure 2.5 are preserved when averaging across all levels of evidence encoding noise ϵ (Figure A.2A) and when splitting up the time-varying weights by low and high flicker rate (Figure A.2B).

(A) Weights on memory and choice behavior averaged across all levels of ϵ



(B) Weights on memory in low and high flicker conditions averaged across all ϵ

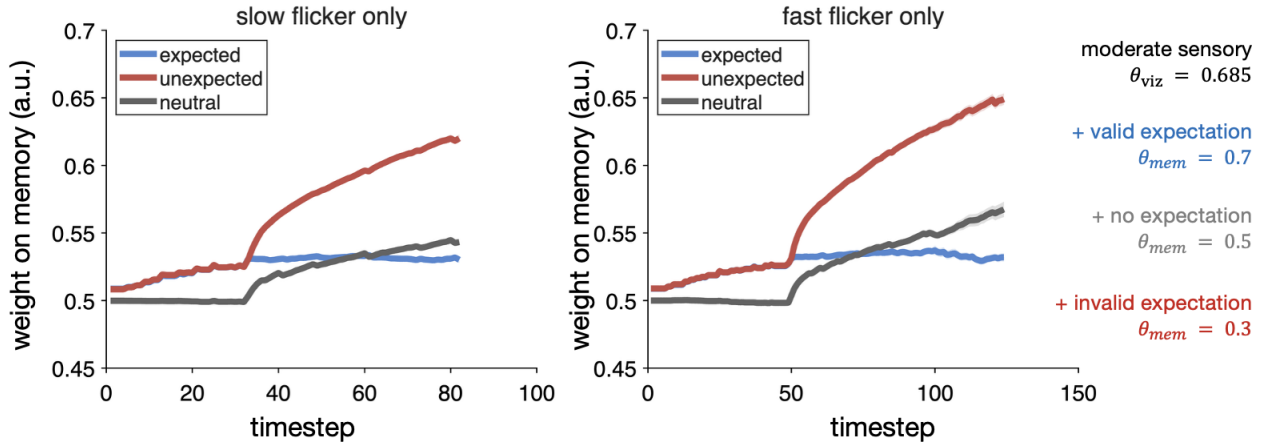


Figure A.2: (A) The qualitative effects produced by our model are preserved when averaging over evidence encoding noise ϵ . (B) Dynamic weights on memory generated in the low (33Hz) and high (50Hz) flicker rate simulations of Rungratsameetaweemana, Itthipuripat et al. (2018). Simulations with low flicker rates used $\gamma \in \{4, 6, 8, 10, 12, 20\}$ corresponding to activity at 8.25, 5.5, 4.125, 3.3, 2.5, and 1.5Hz, and simulations with high flicker rates used $\gamma \in \{4, 6, 8, 10, 12, 20\}$ to generate memory samples at 8.3, 6.25, 4.16, 3.5, 2.5, and 1.5Hz. Threshold values z were fixed across conditions within a participant, but varied across participants; data were generated using $z \in \{3, 4, 5\}$. Weights in each condition cease being estimated once the decision variable hits threshold, leading to an increased weighting of invalid expectations at the group level.

A.3 Effect of cue-stimulus interval duration on Bornstein et al. (2023) effect

Figure A.3 shows that the uncertainty-adaptive anticipatory sampling effect displayed in Figure 2.7 is preserved for both the medium (top) and long (bottom) cue-stimulus interval durations. Visual evidence onset is indicated by the solid vertical line; note difference in x-axes across plots.

Steady-state evidence weights for middle and long cue-stimulus intervals in Bornstein et al. (2023) averaged across all levels of ϵ

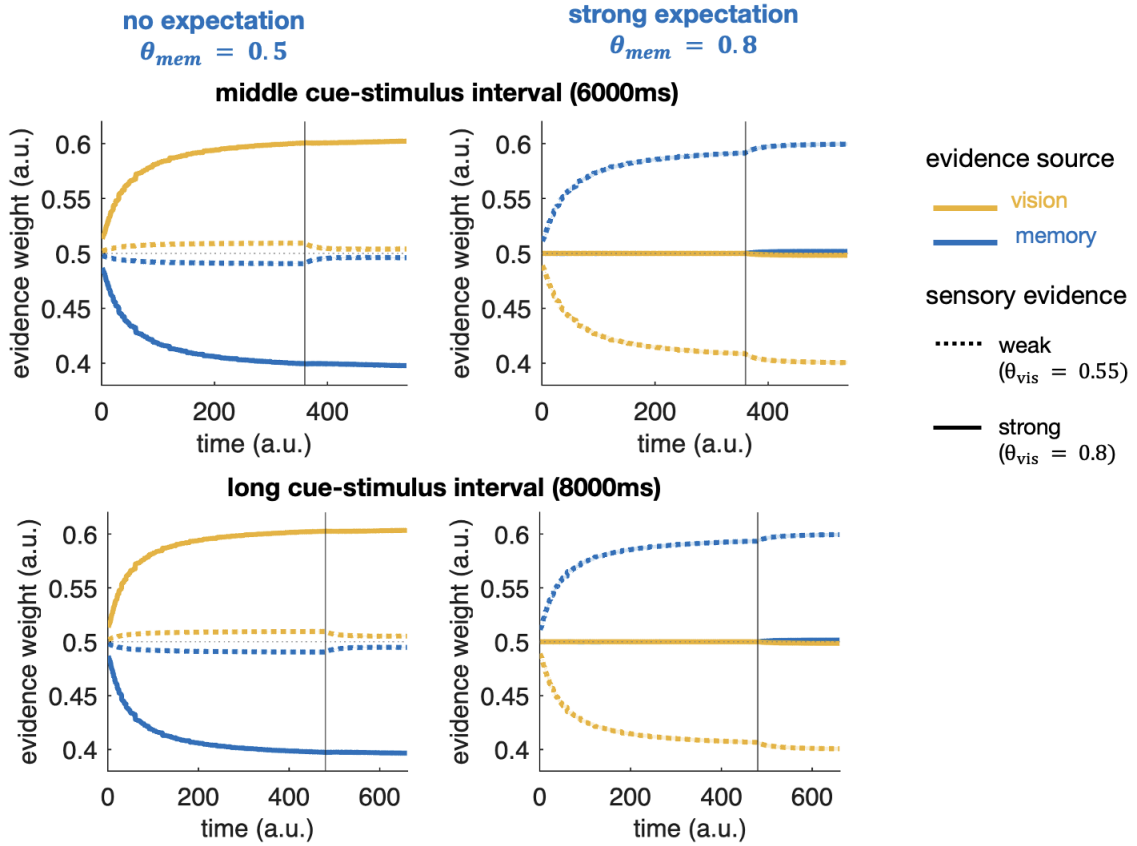


Figure A.3: Steady-state evidence weights displaying the uncertainty-adaptive sampling effect observed by Bornstein et al. (2023) at intermediate (top) and long (bottom) cue-stimulus interval durations. Standard error across $n=50$ steady-state timeseries is plotted but not visible due to its magnitude.

A.4 Hachisuka, Shor, Liu et al. (2026) effect

This section shows that the main-text results in Figure 2.8 effect are preserved when averaging across all levels of evidence encoding noise ϵ .

Hachisuka, Shor, Liu et al. (2026) “late” effect averaged across all levels of ϵ

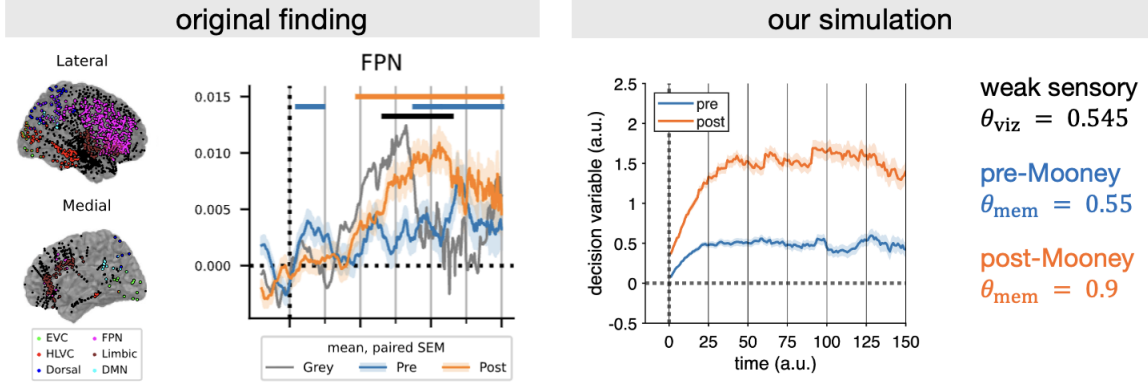


Figure A.4: Timeseries from run-averaged simulations of Hachisuka, Shor, Liu et al. (2026). Solid lines are cue-level decision variables averaged across the mean decision variable for each combination of γ , z , and ϵ ; shaded areas reflect standard error of the mean ($n=100$).

Appendix B

Supplementary Materials for Chapter 3

B.1 Full regression tables from Section 3.4.1

B.1.1 Effect of accuracy and cue on RT

Regressor	Beta	SE	t-statistic	95% CI	p-value
congCue	0.38	0.082	4.61	0.22, 0.54	< 0.001
accuracy	0.38	0.057	6.68	0.27, 0.49	< 0.001
congCue * accuracy	-0.76	0.096	-7.95	-0.95, -0.57	< 0.001

Table B.1: Full regression table: effect of choice accuracy and cue on RT. Abbreviations: CI = Confidence Interval, SE = Standard Error

No. Obs. = 14,993; Sigma = 0.995; Residual df = 14,989

B.1.2 Effect of accuracy, RT, and cue on confidence

Regressor	Beta	SE	95% CI	p-value
congCue	-0.55	0.069	-0.68, -0.41	< 0.001
accuracy	0.07	0.048	-0.02, 0.16	0.15
zlogRT	-0.15	0.037	-0.22, -0.07	< 0.001
congCue * accuracy	0.89	0.081	0.73, 1.0	< 0.001
congCue * zlogRT	0.12	0.063	-0.01, 0.24	0.064
accuracy * zlogRT	0.05	0.047	-0.05, 0.14	0.3
congCue * accuracy * zlogRT	-0.09	0.077	-0.24, 0.06	0.2

Table B.2: Full regression table: effect of choice accuracy, RT, and cue on choice confidence. Abbreviations: CI = Confidence Interval, SE = Standard Error. No. Obs. = 14,951; Sigma = 0.822; Residual df = 14,941.

B.2 Fitted boundary values in response-coded same drift model

B.2.1 No effect of cue condition on fitted boundary

Fitted boundary values from the response-coded same drift model in Section 3.4.2 do not differ as a function of cue condition ($\beta = 3.382$, $t_{39} = 0.54$, $p = 0.595$).

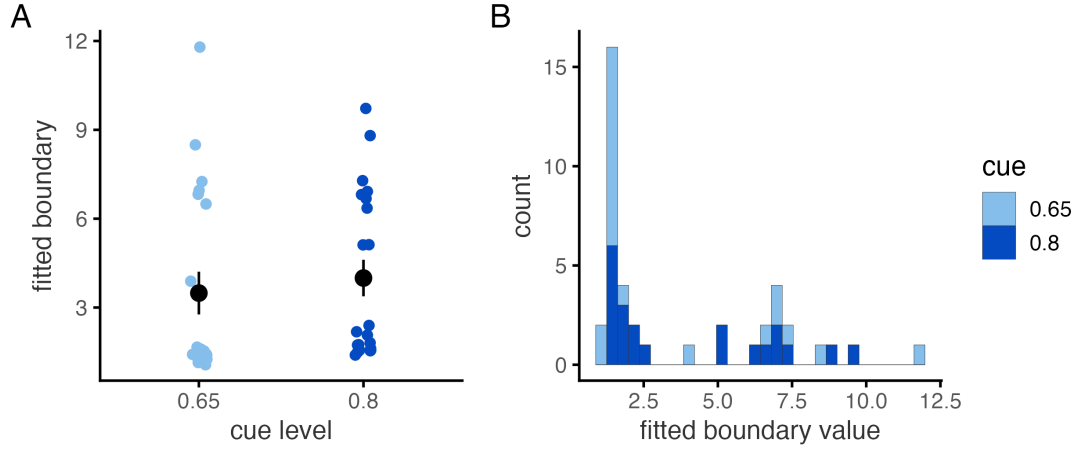


Figure B.1: Fitted boundary values as a function of cue condition.

B.2.2 Boundary is correlated with drift in all trial epochs

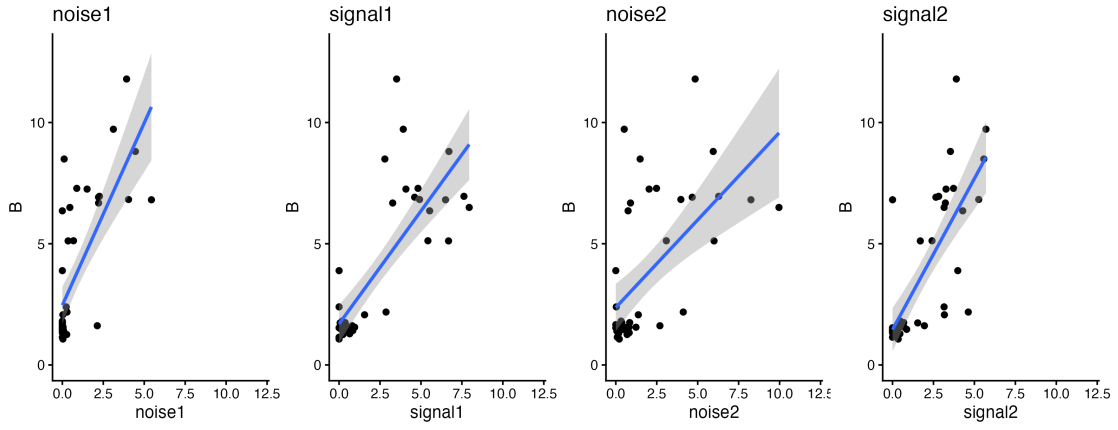


Figure B.2: Correlation between fitted drift rate and boundary value in each trial epoch

epoch	Pearson's r	t	df	p
noise1	0.73	6.64	39	< .001
signal1	0.79	8.09	39	< .001
noise2	0.60	4.64	39	< .001
signal2	0.76	7.28	39	< .001

Table B.3: Correlation coefficients between fitted boundary and drift rate.

B.2.3 Boundary recoverability

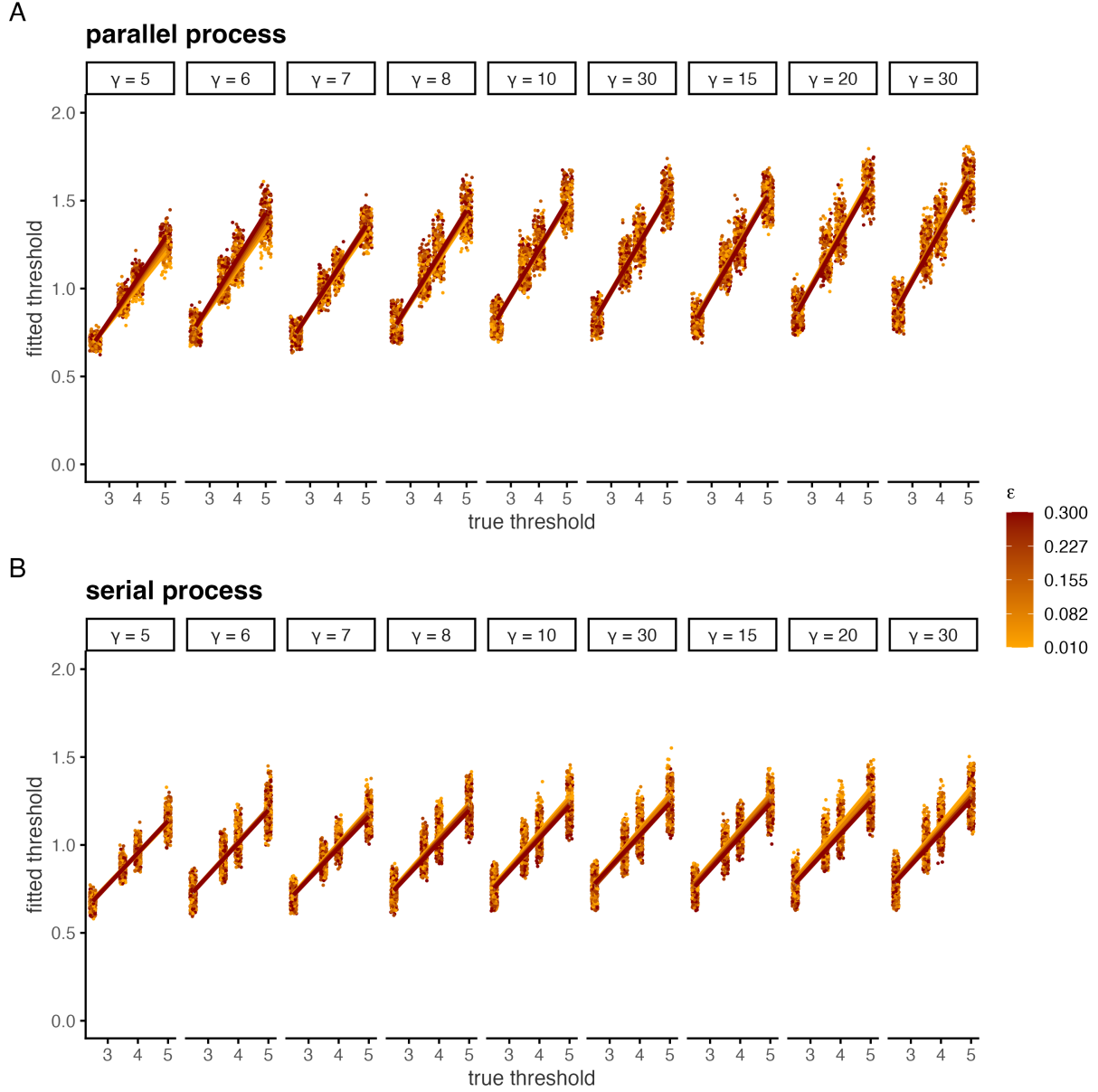


Figure B.3: Fitted thresholds as a function of generating thresholds for the parallel (A) and serial (B) processes. γ is the parameter scaling the relative “sampling rate” of memory to vision, and ϵ is the proportion of evidence samples incorrectly encoded from both evidence sources. Details on how the data were generated can be found in Section 3.3.4.