
Burn-in, bias, and the rationality of anchoring

Falk Lieder

Translational Neuromodeling Unit
ETH and University of Zurich
lieder@biomed.ee.ethz.ch

Thomas L. Griffiths

Psychology Department
University of California, Berkeley
tom.griffiths@berkeley.edu

Noah D. Goodman

Department of Psychology
Stanford University
ngoodman@stanford.edu

Abstract

Bayesian inference provides a unifying framework for learning, reasoning, and decision making. Unfortunately, exact Bayesian inference is intractable in all but the simplest models. Therefore minds and machines have to approximate Bayesian inference. Approximate inference algorithms can achieve a wide range of time-accuracy tradeoffs, but what is the optimal tradeoff? We investigate time-accuracy tradeoffs using the Metropolis-Hastings algorithm as a metaphor for the mind’s inference algorithm(s). We characterize the *optimal time-accuracy trade-off* mathematically in terms of the number of iterations and the resulting bias as functions of time cost, error cost, and the difficulty of the inference problem. We find that reasonably accurate decisions are possible long before the Markov chain has converged to the posterior distribution, i.e. during the period known as “burn-in”. Therefore the strategy that is optimal subject to the mind’s bounded processing speed and opportunity costs may perform so few iterations that the resulting samples are biased towards the initial value. The resulting cognitive process model provides a rational basis for the anchoring-and-adjustment heuristic. The model’s quantitative predictions match published data on anchoring in numerical estimation tasks. In conclusion, *resource-rationality*—the optimal use of finite computational resources—naturally leads to a biased mind.

1 Introduction

What are the algorithms and representations used in human cognition? This central question of cognitive science has been attacked from different directions. The *heuristics and biases* program started by Kahneman and Tversky [1] seeks to constrain the space of possible algorithms by the systematic deviations of human behavior from normative standards. This approach has demonstrated a large number of cognitive biases, i.e. systematic “errors” in human judgment and decision-making that are thought to result from people’s use of heuristics [1]. Heuristics are problem-solving strategies that are simple and efficient but not guaranteed to find optimal solutions. By contrast, *Bayesian models of cognition* seek to identify and understand the problems solved by the human mind, and to constrain the space of possible algorithms by considering the optimal solutions to those problems [2].

Bayesian models and heuristics characterize cognition at two different levels of analysis [3], namely the computational and the algorithmic level respectively. Explanations formulated at the computational level do not specify psychological mechanisms, but they can be used to derive or constrain psychological process models [3]. Conversely, heuristics can be thought of as algorithms that approximate the solutions to a Bayesian inference problems. Thus approximate inference algorithms,

such as sampling algorithms and variational Bayes, can inspire models of the psychological mechanisms underlying cognitive performance. In fact, models developed in this way (e.g. [4–7]) provide a rational basis for established psychological process models such as the Win-Stay-Lose-Shift algorithm [6] and exemplar models [5], as well as cognitive biases such as probability matching [6, 7]. Such rational process models are informed by the teleological constraint that they should approximate Bayesian inference as well as by processing constraints on working memory (e.g. [6]) or computation time (e.g. [7]). In addition to the success of these models—which are based on sampling algorithms—there is some direct experimental evidence that speaks to sampling as a psychological mechanism: people appear to sample from posterior probability distributions when making predictions [8, 9], categorizing objects [10], or perceiving ambiguous visual stimuli [11, 12]. Furthermore, sampling algorithms can be implemented in biologically plausible networks of spiking neurons [13], and individual sampling algorithms are applicable to a wide range of problems including inference on complex structured representations. In what follows, we will therefore use sampling algorithms as our starting point for exploring how the human mind might perform probabilistic inference.

The problem of having to trade off accuracy for speed is ubiquitous not only in machine learning and artificial intelligence, but also in human cognition. The mind’s bounded computational resources necessitate tradeoffs. Ideally these tradeoffs should be chosen such as to maximize expected utility net the cost of computation. In AI this problem is known as rational meta-reasoning [14], and it has been formalized in terms of maximizing the value of computation (computational utility) [15–17]. One may assume that evolution and cognitive development have endowed the mind with the meta-cognitive ability to stop thinking when the expected improvement in accuracy falls below the cost of thought. Therefore, understanding optimal time-accuracy tradeoffs may benefit not just AI but also cognitive modelling. In cognitive science, a recent analysis concluded that time costs make it rational for agents with bounded processing speed to decide based on very few samples from the posterior distribution [7]. However, generating even a single perfect sample from the posterior distribution may be very costly, for instance requiring thousands of iterations of a Markov chain Monte Carlo (MCMC) algorithm. The values generated during this burn-in phase are typically discarded, because they are biased towards the Markov chain’s initial value. This paper analyses under which conditions a bounded rational agent should decide based on the values generated in the burn-in phase of its MCMC algorithm and tolerate the resulting bias. The results may further our understanding of *resource-rationality*—the optimal use of finite computational resources—which may become a framework for developing algorithmic models of human cognition as well as AI systems. We demonstrate the utility of this resource-rational framework by presenting an algorithmic model of probabilistic inference that can explain a cognitive bias that might otherwise appear irrational.

2 The Time-Bias Tradeoff of MCMC Algorithms

Exact Bayesian inference is either impossible or intractable in most real-world problems. Thus natural and artificial information processing systems have to approximate Bayesian inference. There is a wide range of approximation algorithms that could be used. Which algorithm is best suited for a particular problem depends on two criteria: time cost and error cost. Iterative inference algorithms such as MCMC gradually refine their estimate over time allowing the system to smoothly adjust to time pressure. This section formalizes the problem of optimally approximating probabilistic inference for real-time decision-making under time costs for a specific MCMC algorithm, the Metropolis-Hastings algorithm [18].

2.1 Problem Definition: Probabilistic Inference in Real-Time

The problem of decision-making under uncertainty in real-time can be formalized as the optimization of a utility function that incorporates decision time [19]. Here, we analyze the special case that we will use to model a psychology experiment in Section 3. In this special case, actions are point estimates of the latent variable of a Normal-Normal model,

$$P(X) = \mathcal{N}(\mu_p, \sigma_p^2), P(Y|X = x) = \mathcal{N}(x, \sigma_l^2), \quad (1)$$

where X is unknown and Y is observed. The utility of predicting state a when the true state is x after time t is given by the function

$$u(a, x, t) = -\text{cost}_{\text{error}}(a, x) - c \cdot t. \quad (2)$$

This function has two terms—error cost and time cost—whose relative importance is determined by the time cost c (opportunity cost per second).

In certain situations the mind may be faced with the problem of maximizing the utility function given in Equation 2 using an iterative inference algorithm. Under this assumption the reaction time in Equation 2 can be modeled as a linear function of the number of iterations i used to compute the decision, with

$$t = i/v + t_0, [v] = \text{iterations/sec.} \quad (3)$$

The slope of this function is the system’s processing speed v measured in iterations per second, and the offset t_0 is assumed to be constant. The next section solves this bounded optimality problem assuming that the mind’s inference algorithm is based on MCMC.

2.2 A Rational Process Model of Probabilistic Inference

We assume that the brain’s inference process iteratively refines an approximation Q to the posterior distribution. The first iteration computes approximation Q_1 from an initial distribution Q_0 , the t^{th} iteration computes approximation Q_t from the previous approximation Q_{t-1} , and so on. Under the assumption that the agent updates its belief according to the Metropolis-Hastings algorithm [18], the temporal evolution of its belief distribution can be simulated by repeated multiplication with a transition matrix T : $Q_{t+1} = T \cdot Q_t$. Each element of T is determined by two factors: the probability that a transition is proposed, i.e. $P_{\text{propose}}(x_i, x_j) = \mathcal{N}(x_i; \mu = x_j, \sigma^2)$, and the probability that this proposal is accepted, i.e. $\alpha(x_i, x_j) = \min \left\{ 1, \frac{P(x_i|y) \cdot P_{\text{propose}}(x_j, x_i)}{P(x_j|y) \cdot P_{\text{propose}}(x_i, x_j)} \right\}$. In brief, $T_{i,j} = P_{\text{propose}}(x_i, x_j) \cdot \alpha(x_i, x_j)$.

Concretely, the psychological process underlying probabilistic inference can be thought of as a sequence of adjustments of an initial guess, the anchor x_0 . In each iteration a potential adjustment is drawn from a Gaussian distribution with mean zero ($\delta \sim \mathcal{N}(0, \sigma^2)$). The adjustment will either be accepted, i.e. $x_{t+1} = x_t + \delta$, or rejected, i.e. $x_{t+1} = x_t$. If a proposed adjustment makes the estimate more probable ($p(x_t + \delta) > p(x_t)$), then it will always be accepted. Otherwise the adjustment will be made with probability $\alpha = \frac{p(x_t + \delta)}{p(x_t)}$, i.e. according to the posterior probability of the adjusted relative to the unadjusted estimate. This simple scheme implements the Metropolis-Hastings algorithm [18] and thereby renders the sequence of estimates $(x_0, x_1, x_2, \dots)$ a Markov chain converging to the posterior distribution. Beyond this instantiation of MCMC we assume that the mind stops after the optimal number of (potential) adjustments that will be derived in the following subsections. This algorithm can be seen as a formalization of anchoring-and-adjustment [20] and therefore provides a rational basis for this heuristic.

2.3 Bias decays exponentially with the number of MCMC Iterations

The bias of a probabilistic belief Q approximating a normative posterior P can be measured by the absolute value of the deviation of its expected value from the posterior mean,

$$\text{Bias}[Q; P] = \|\mathbb{E}_Q[X] - \mathbb{E}_P[X]\|. \quad (4)$$

If the proposal distribution is positive, then the distribution Q_i that the Metropolis-Hastings algorithm samples from in iteration i converges to the posterior distribution geometrically fast under mild regularity conditions on the posterior [21]. Here, the convergence of the sampling distribution Q_i to the posterior P was defined as the convergence of the supremum of the absolute difference between the expected values $\mathbb{E}_{Q_i}[f(X)]$ and $\mathbb{E}_P[f(X)]$ over all functions f with $\forall x : \|f(x)\| < V(x)$ for some function V that depends on the posterior. Using this theorem, one can prove the geometric convergence of the bias of the distribution Q_i to zero and the geometric convergence of the expected utility of actions sampled from Q_i to the expected utility of actions sampled from P :

$$\begin{aligned} \exists M, r \in \mathbb{R}^+ : \text{Bias}[Q_i; P] &\leq M \cdot r^i, \\ \text{and } |\mathbb{E}_{Q_i(A)}[\text{EE}(a)] - \mathbb{E}_{P(A)}[\text{EE}(a)]| &\leq M \cdot r^i, \end{aligned} \quad (5)$$

where $\text{EE}(a)$ is $\mathbb{E}_{P(X)}[\text{cost}_{\text{error}}(a, x)]$, and the tightness of the bound M , the initial bias, and the convergence rate r are determined by the chosen proposal distribution, the posterior, and the initial value.

We simulated the decay of bias in the Normal-Normal case that we are focussing on (Equation 1) and $\text{cost}_{\text{error}}(a, x) = \|a - x\|$. All Markov chains were initialized with the prior mean. Our results show that the mean of the sampling distribution converges geometrically as well (see Figure 1). Thus the reduction in bias tends to be largest in the first few iterations and decreases quickly from there on, suggesting a situation of diminishing returns for further iterations: an agent under time pressure may do well stopping after the initial reduction in bias.

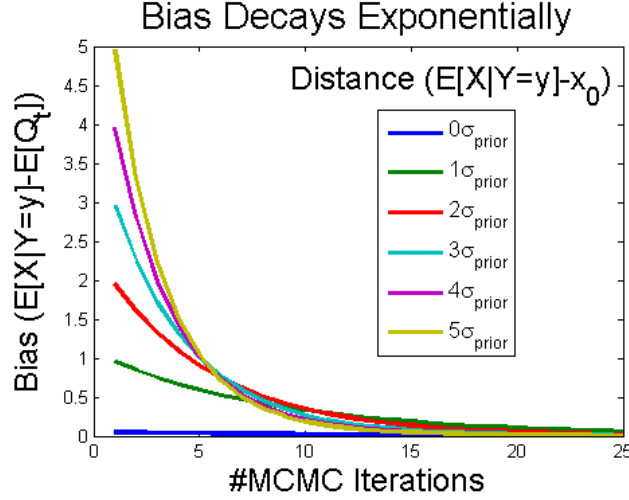


Figure 1: Bias of the mean of the approximation Q_t , i.e. $|\mathbb{E}[\tilde{X}_t] - \mathbb{E}[X|Y=y]|$ where $X_t \sim Q_t$, as a function of the number of iterations t of our Metropolis-Hastings algorithm. The five lines show this relationship for different posterior distributions whose means are located $1\sigma_p, \dots, 5\sigma_p$ away from the prior mean (σ_p is the standard deviation of the prior). As the plot shows, the bias decays geometrically with the number of iterations in all five cases.

2.4 Optimal Time-Bias Tradeoffs

This subsection combines the result reported in the previous subsection with time costs and computes the optimal bias-time tradeoffs depending on the ratio of time cost to error cost and on how large the initial bias is. It suggests that intelligent agents might use these results to choose the number of MCMC iterations according to their estimate of the initial bias. Formally, we define the optimal number of iterations i^* and resulting bias b^* as

$$i^* = \arg \max_i \mathbb{E}[u(a_i, x, t_0 + i/v) | Y = y], \text{ where } a_i \sim Q_i \quad (6)$$

$$b^* = \text{Bias}[Q_{i^*}; P] \quad (7)$$

using the variables defined above. If the upper bound in Equation 5 is tight, then the optimal number of iterations and the resulting bias can be calculated analytically,

$$i^* = \frac{1}{\log(r)} \cdot (\log(c) - M \cdot \log(M \cdot \log(1/r))) \quad (8)$$

$$b^* \leq \frac{c}{\log(1/r)}. \quad (9)$$

Figure 2 shows the optimal number of iterations for various initial distances and ratios of time cost to error cost (c), and Figure 3 shows the resulting biases. Over a large range of time cost to error cost ratios, the optimal number of iterations is about 10 and below. Thus by the standard use of MCMC algorithms, the samples used in the resource-rational decisions would have been discarded as “burn-in”. As one might expect, the optimal number of iterations increases with the distance from initial value to the posterior mean, but the conclusions hold regardless of those differences. As a result, rational agents should tolerate a non-negligible amount of bias in their approximation to the posterior.

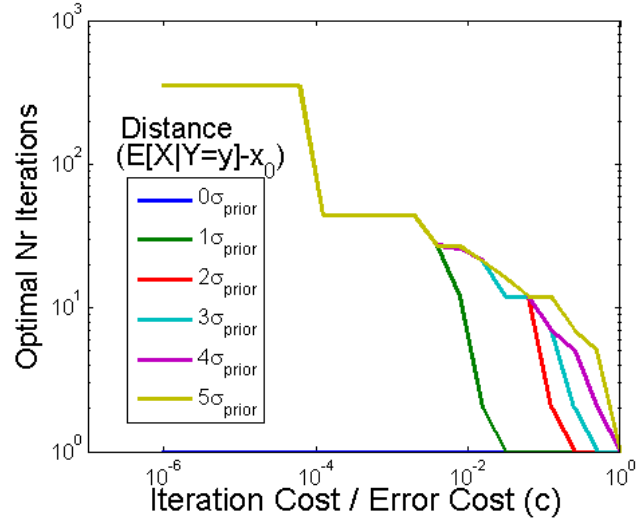


Figure 2: Number of iterations i^* that maximizes the agent’s expected utility as a function of the ratio between the cost per iteration and the cost per unit error. The five lines correspond to the same posterior distributions as in Figure 1.

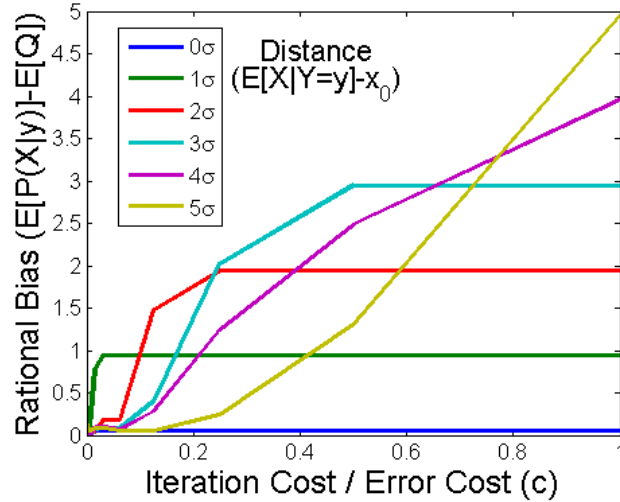


Figure 3: Bias of the approximate posterior (Equation 4) after the optimal number of iterations shown in Figure 2. The five lines correspond to the same posterior distributions as in Figure 1.

Contrary to traditional recommendations to run MCMC algorithms for thousands of iterations (e.g. [22]) our analysis suggests that the resource-rational number of iterations for real-time decision-making can be about two orders of magnitude lower. This implies that real-time decision-making can and should tolerate bias in exchange for speed. In separate analyses we found that the tolerable bias is even higher for decision problems with non-linear error costs. In predicting a binary random variable with 0 – 1 loss, for instance, bias has almost no effect on accuracy unless it pushes the agent’s probability estimate across the 50% threshold (see Figure 4). These surprising insights have important implications for modelling mental algorithms and cognitive biases, as we will illustrate in the next section.

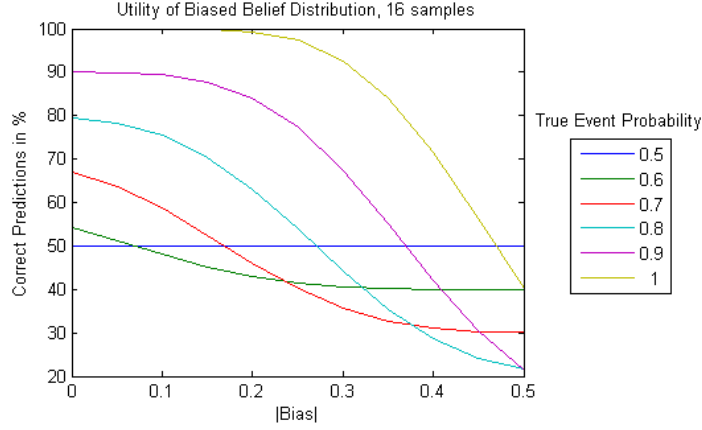


Figure 4: Expected number of correct predictions of a binary event from 16 samples as a function of the bias of the belief distribution from which these samples are drawn. Different lines correspond to different degrees of predictability. As the belief becomes biased, the expected number of correct predictions declines gracefully and stays near its maximum for a wide range of biases; especially if the predictability is high. This illustrates that biased beliefs can support decisions with high utility. Therefore the value of reducing bias can be rather low.

3 A rational explanation of the anchoring bias in numerical estimation

The resource-rationality of biased estimates might be able to explain why people’s estimates of certain unknown quantities (e.g. the duration of Mars’s orbit around the sun) are systematically biased towards other values that come to mind easily (e.g. 365 days). This effect is known as the *anchoring bias*. The anchoring bias has been explained as a consequence of people’s use of the anchoring-and-adjustment heuristic [1]. The anchoring-and-adjustment heuristic generates an initial estimate by recalling a related quantity from memory and adjusts it until a plausible value is reached [20]. Epley and Gilovich have shown that people use an anchoring-and-adjustment for unknown quantities that call to mind a related value by demonstrating that the anchoring bias decreases with subjects’ motivation to be accurate and increases with cognitive load and time pressure [20, 23, 24].

Here we model Experiment 1B from [20]. In this study, one group of 54 subjects estimated six numerical quantities including the duration of Mars’s orbit around the sun and the freezing point of vodka. Each quantity was chosen such that subjects would *not* know its true value but the value of a related quantity (i.e. the intended anchor, e.g. 365 days and 32°F). Another group was asked to provide plausible ranges ($[l_i, u_i]$, $1 \leq i \leq 6$) for the same quantities. If responses were sampled uniformly from the plausible values or a Gaussian centered on their mean, then mean responses should fall into the center of the plausible ranges. By contrast, the first group’s mean responses deviated significantly from the center of the plausible range towards the anchor; the deviations were measured by dividing the difference between the mean response and the range endpoint nearest the intended anchor by the range of plausible values (“mean skew”). Thus subjects’ adjustments were insufficient, but why? Our model explains insufficient adjustments as the optimal bias-time tradeoff of an iterative inference algorithm that formalizes the anchoring-and-adjustment heuristic.

3.1 Modelling Estimation as Probabilistic Inference

When people are asked to estimate numerical quantities they appear to sample from an internal probability distribution [25]. For quantities such as the duration of Mars’s orbit around the sun the process computing those distributions may be reconstruction from memory. Reconstruction from memory has been shown to combine noisy memory traces y with prior knowledge about categories $p(X)$ into a more accurate posterior belief $p(X|y)$ [26, 27]. For other quantities such as the duration of Mars’ orbit there may be no memory trace of the quantity itself, but there will be memory traces of related quantities. In either case the relation between the memory trace and the quantity to be

estimated is probabilistic and can be described by a likelihood function. Therefore estimation always amounts to combining noisy evidence with prior knowledge. The six estimation tasks were thus modelled as probabilistic inference problems.

The unknown posterior distributions were approximated by Gaussians with a mean equal to the actual value and a variance such that the 95% posterior credible intervals are as wide as the plausible ranges reported by the subjects ($\sigma_i = \frac{u_i - l_i}{\Phi^{-1}(0.975) - \Phi^{-1}(0.025)}$). Epley and Gilovich [20] assumed that the anchoring-and-adjustment process starts from self-generated anchors, such as the duration of a year in the case of estimating the duration of Mars' orbit. The initial distributions were thus modelled by delta-distributions on the subjects' self-generated anchors. Considered adjustments were assumed to be sampled from $\mathcal{N}(\mu = 0, \sigma = 10)$. Thus numeric estimation can be modelled as probabilistic inference according to our rational model. The model's optimal bias-time tradeoff depends only on the ratio of the cost per iteration to the cost per unit error. This ratio ($\tilde{c} = c/v$) was estimated from subjects' mean responses reported in [20] by weighted least squares. The weights were the precisions (inverse variances) of the assumed posterior distributions. The resulting point estimate was used to predict subjects' mean responses.

3.2 Results

Figure 5 illustrates the model-fit in terms of the adjustment scores. The predicted adjustment scores are highly correlated with those computed from subjects' responses ($r = 0.95$). Importantly, whenever the experiment found an insufficient adjustment ($\text{skew} < 0.5$) this is also predicted by our model. Conversely, in the one condition where the experiment did not find an insufficient adjustment that is also what our model predicts. Thus the anchoring bias reported in [20] may result from a resource-rational probabilistic inference algorithm. In another study (Experiment 2C in [20]) Epley

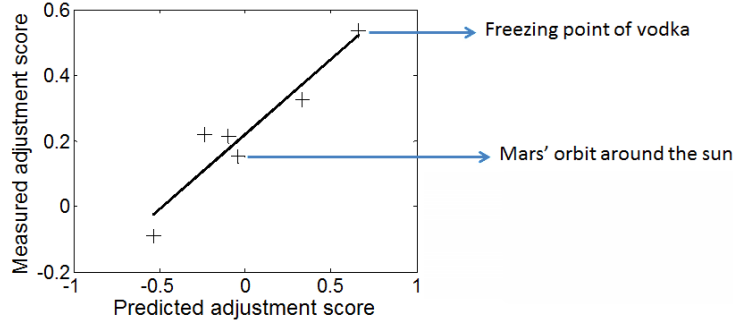


Figure 5: Comparison of people's mean adjustment scores in the six estimation tasks of [20] with those predicted by our model. The model achieves a good fit with just a single parameter. The correlation between predicted and measured adjustment scores is about 0.95.

and Gilovich found that the anchoring bias was stronger when while estimating numerical quantities subjects had to remember an eight-letter string (high cognitive load) than when they did not (low cognitive load). We used this data set to assess the model's validity. If the model's parameter $\tilde{c}$ measures the relative cost per iteration, then its estimate should be higher when the simultaneous performance of an additional task restricts subjects' computational resources. This is indeed what we found. For subjects under high cognitive load the estimated relative cost per iteration ($\tilde{c}_{\text{busy}} = 0.31$) was higher than for subjects under low cognitive load ($\tilde{c}_{\text{not busy}} = 0.18$). This is consistent with the parameter's intended interpretation, and the resulting model predictions captured the increases of the anchoring bias under cognitive load.

4 Discussion

In this article we have performed a resource-rational analysis of the Metropolis-Hastings algorithm. Across a wide range of time costs the resource-rational solution performs so few iterations that the resulting estimates are biased towards the initial value. This provides a rational basis for the anchoring-and-adjustment heuristic [1] as a previous analysis [28] did for the representativeness

heuristic [1]. By deriving the anchoring-and-adjustment heuristic from a general approximate inference algorithm we open the door to understanding how this heuristic should apply in domains more complex than the estimation of a single number. From previous work in psychology, it is not obvious what an adjustment should look like in the space of objects or events, for instance. By reinterpreting adjustment in terms of MCMC, we extend it to almost arbitrary domains of inference. Furthermore our model illustrates that heuristics in general can be formalized by and derived as resource-rational approximations to rational inference. This provides a new perspective on the resulting cognitive biases: the anchoring bias is tolerated because its consequences are less costly than the time that would be required to eliminate it. Thus the anchoring bias can be interpreted as a sign of resource-rationality rather than irrationality. This may be equally true of other cognitive biases, because eliminating bias can be costly and biased beliefs do not necessarily lead to poor decisions (see Figure 4).

This article illustrates the value of resource-rationality as a framework for deriving models of cognitive processes. This approach is a formal synthesis of the function-first approach underlying Bayesian models of cognition [2] with the limitations-first approach that starts from cognitive and perceptual illusions (e.g. [1]). This synthesis can be achieved by augmenting the problems facing the mind with constraints on information processing and solving them formally using optimization. Rather than determining optimal beliefs and actions, this approach seeks to determine resource-rational processes. Conversely, given a process the same framework can be used to determine the processing constraints under which it is resource-rational. This idea is so general that it can be applied to reverse-engineering the whole spectrum of mental algorithms and processing constraints.

Understanding the mind’s probabilistic inference algorithm(s) will require many more steps than we have been able to take so far. We have demonstrated the resource-rationality of very few iterations in a trivial inference problem for which an analytical solution exists. Thus, one may argue that our conclusions could be specific to a simple one-dimensional inference problem. However, our results follow from a very general mathematical property: the geometric convergence of the Metropolis-Hastings algorithm. This property is also true of many complex and high-dimensional inference problems [21]. Furthermore, Equation 8 predicts that if the inference problem becomes more difficult, then the bias b^* tolerated by a resource-rational MCMC algorithm increases, because higher complexity leads to slower convergence and this means $r \rightarrow 1$. This suggests, that resource-rational solutions to the challenging inference problems facing the human mind are likely to be biased, but this remains to be shown. In addition, the solution determined by our resource-rational analysis remains to be evaluated against alternative solutions to real-time decision-making under uncertainty. Furthermore, while our model was consistent with existing data, our formalization of the estimation task can be questioned, and since [20] did not report standard errors, we were unable to perform a proper statistical assessment of our model. It therefore remains to be shown whether or not people’s bias-time tradeoff is indeed near-optimal. This will require dedicated experiments that systematically manipulate probability structure, time pressure, and error costs.

The prevalence of Bayesian computational-level models has made a sampling based view of mental processing attractive—after all, sampling algorithms are flexible and efficient solutions to difficult Bayesian inference problems. Vul et al. [7] explored the rational use of a sampling capacity, arguing that it is often rational to decide based on only one perfect sample. However, perfect samples can be hard to come by; here we have shown that it can be rational to decide based on one *imperfect* sample. A rational sampling-based view of mental processing thus leads naturally to a biased mind.

Acknowledgments. This work was supported by grant number FA-9550-10-1-0232 from the Air Force Office of Scientific Research (T. L. Griffiths) and a John S. McDonnell Foundation Scholar Award (N. D. Goodman).

References

- [1] A. Tversky and D. Kahneman, “Judgment under uncertainty: Heuristics and biases,” *Science*, vol. 185, no. 4157, pp. 1124–1131, 1974.
- [2] T. L. Griffiths, C. Kemp, and J. B. Tenenbaum, “Bayesian models of cognition,” in *The Cambridge handbook of computational psychology* (R. Sun, ed.), ch. 3, pp. 59–100, Cambridge University Press, 2008.
- [3] D. Marr, *Vision: A Computational Investigation into the Human Representation and Processing of Visual Information*. W. H. Freeman, 1983.

- [4] A. N. Sanborn, T. L. Griffiths, and D. J. Navarro, "Rational approximations to rational models: Alternative algorithms for category learning," *Psychol. Rev.*, vol. 117, no. 4, pp. 1144–1167, 2010.
- [5] L. Shi, T. L. Griffiths, N. H. Feldman, and A. N. Sanborn, "Exemplar models as a mechanism for performing Bayesian inference," *Psychon B Rev*, vol. 17, no. 4, pp. 443–464, 2010.
- [6] E. Bonawitz, S. Denison, A. Chen, A. Gopnik, and T. L. Griffiths, "A simple sequential algorithm for approximating bayesian inference," *Proceedings of the 33rd Annual Conference of the Cognitive Science Society, Austin, TX: Cognitive Science Society*, 2011.
- [7] E. Vul, N. D. Goodman, T. L. Griffiths, and J. B. Tenenbaum, "One and done? Optimal decisions from very few samples.," *Proceedings of the 31st Annual Conference of the Cognitive Science Society*, 2009.
- [8] T. L. Griffiths and J. B. Tenenbaum, "Optimal predictions in everyday cognition," *Psych. Sci.*, vol. 17, no. 9, pp. 767–773, 2006.
- [9] T. L. Griffiths and J. B. Tenenbaum, "Predicting the future as bayesian inference: People combine prior knowledge with observations when estimating duration and extent," *J. Exp. Psychol. Gen.*, vol. 140, no. 4, pp. 725–743, 2011.
- [10] N. Goodman, J. B. Tenenbaum, J. Feldman, and T. L. Griffiths, "A rational analysis of rule-based concept learning," *Cognitive Science*, vol. 32, no. 1, pp. 108–154, 2008.
- [11] S. J. Gershman, E. Vul, and J. B. Tenenbaum, "Multistability and perceptual inference," *Neural Comput.*, vol. 24, no. 1, pp. 1–24, 2011.
- [12] R. Moreno-Bote, D. C. Knill, and A. Pouget, "Bayesian sampling in visual perception," *Proc. Natl. Acad. Sci. U.S.A.*, vol. 108, no. 30, pp. 12491–12496, 2011.
- [13] L. Buesing, J. Bill, B. Nessler, and W. Maass, "Neural dynamics as sampling: a model for stochastic computation in recurrent networks of spiking neurons.," *PLoS Comput. Biol.*, vol. 7, no. 11, pp. e1002211+, 2011.
- [14] S. Russell, "Rationality and intelligence," *Artificial Intelligence*, vol. 94, pp. 57–77, July 1997.
- [15] E. J. Horvitz, "Reasoning under varying and uncertain resource constraints," in *Proceedings of the Seventh National Conference on Artificial Intelligence*, pp. 111–116, 1988.
- [16] S. Russell and E. H. Wefald, *Do the Right Thing: Studies in Limited Rationality (Artificial Intelligence)*. The MIT Press, 1991.
- [17] N. Hay, S. Russell, D. Tolpin, and S. E. Shimony, "Selecting computations: Theory and applications," in *Uncertainty in Artificial Intelligence: Proceedings of the Twenty-Eighth Conference*, 2012.
- [18] W. K. Hastings, "Monte Carlo sampling methods using Markov chains and their applications," *Biometrika*, vol. 57, no. 1, pp. 97–109, 1970.
- [19] E. Horvitz, "Reasoning about beliefs and actions under computational resource constraints," in *Uncertainty in Artificial Intelligence 3 Annual Conference on Uncertainty in Artificial Intelligence (UAI-87)*, (Amsterdam, NL), pp. 301–324, Elsevier Science, 1987.
- [20] N. Epley and T. Gilovich, "The anchoring-and-adjustment heuristic," *Psych. Sci.*, vol. 17, no. 4, pp. 311–318, 2006.
- [21] K. L. Mengersen and R. L. Tweedie, "Rates of convergence of the Hastings and Metropolis algorithms," *The Annals of Statistics*, vol. 24, no. 1, pp. 101–121, 1996.
- [22] A. E. Raftery and S. Lewis, "How many iterations in the Gibbs sampler," in *Bayesian Statistics 4*, vol. 4, pp. 763–773, 1992.
- [23] N. Epley and T. Gilovich, "When effortful thinking influences judgmental anchoring: differential effects of forewarning and incentives on self-generated and externally provided anchors," *J Behav. Decis Making*, vol. 18, no. 3, pp. 199–212, 2005.
- [24] N. Epley, B. Keysar, L. Van Boven, and T. Gilovich, "Perspective taking as egocentric anchoring and adjustment," *J. Pers. Soc. Psychol.*, vol. 87, no. 3, pp. 327–339, 2004.
- [25] E. Vul and H. Pashler, "Measuring the crowd within," *Psych. Sci.*, vol. 19, no. 7, pp. 645–647, 2008.
- [26] J. Huttenlocher, L. V. Hedges, and J. L. Vevea, "Why do categories affect stimulus judgment?," *J. Exp. Psychol. Gen.*, vol. 129, no. 2, pp. 220–241, 2000.
- [27] P. Hemmer and M. Steyvers, "A Bayesian account of reconstructive memory," *Topics in Cognitive Science*, vol. 1, no. 1, pp. 189–202, 2009.
- [28] J. B. Tenenbaum and T. L. Griffiths, "The rational basis of representativeness," in *Proceedings of the 23rd Annual Conference of the Cognitive Science Society*, pp. 84–98, 2001.