

Decision-Making and Confidence Given Uncertain Advice

Michael D. Lee (michael.lee@adelaide.edu.au)

Department of Psychology, University of Adelaide
South Australia, 5005, AUSTRALIA

Abstract

We pick up where the Dry *et al.* paper in this volume left off. We provide an accumulator model account of the four basic regularities in decision-making and confidence observed in that study. The model captures the regularities with interpretable parameter values, and we show its ability to fit the data is not due to excessive model complexity. Indeed, all of the regularities in human performance are elegantly accounted for in terms of the adaptive processes of the accumulator model.

Introduction

We pick up where the Dry *et al.* paper in this volume left off. Figures 1 and 2 show the basic patterns of empirical confidence and decision-making behaviour that we are attempting to model. Figure 1 shows the pattern of change of mean confidence for advice trials and no-advice trials across the six experimental conditions. Figure 2 shows the pattern of change in advice acceptance behavior across the six conditions. For advice trials, the mean proportion of trials for which the advice was accepted is shown. For no advice trials, the mean proportion “go left” decisions is shown.

There are four essential regularities of these data we would like to model. These are:

- Confidence on advice trials shows a slight inverted U-shape across the conditions.
- Confidence on no-advice trials decreases across the conditions.
- Advice is almost always accepted.
- Decisions are consistent with guessing on no-advice trials.

The Model

We model the behavior of subjects in this task using an accumulator sequential sampling process (Vickers, 1979). This class of model involves sampling evidence from the environment until some decision criterion is reached, producing the decision behavior, and then potentially adapting in several ways. We first describe

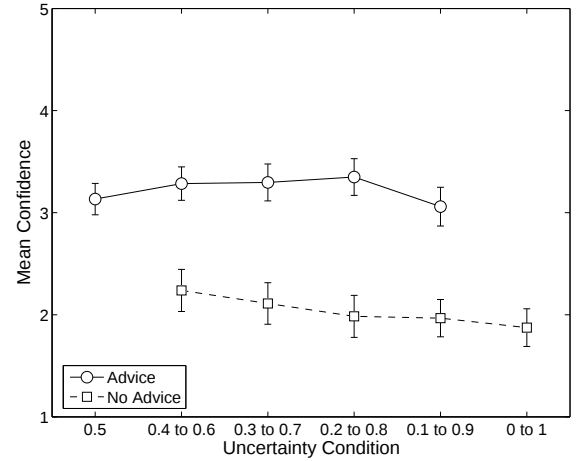


Figure 1: The pattern of change of mean confidence for advice trials and no-advice trials across the six experimental conditions. Error bars show one standard error.

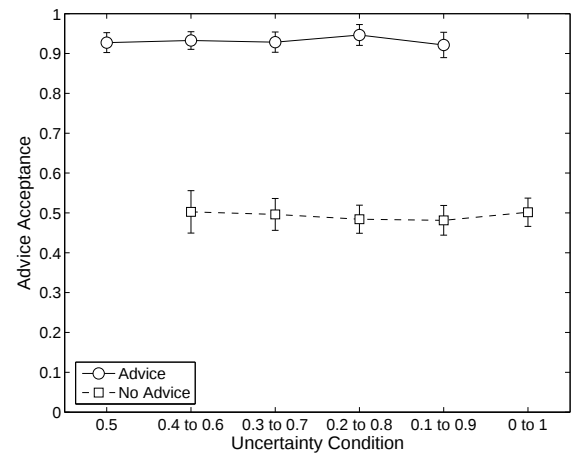


Figure 2: The pattern of change in advice acceptance behavior across the six conditions. Error bars show one standard error.

the information environment we assume to be sampled, then the sequential sampling process itself, and then the ways in which the model learns from feedback, and self-regulates its decision-making mechanisms.

Evidence Distributions

The model assumes subjects make decisions by sampling evidence for the left and right alternatives, based on their current knowledge of the task environment. Separate probability distributions, π_l , π_r and π_u are maintained for representing the evidence provided by “go left”, “go right” and “uncertain” advice, respectively. We assume that initially subjects are maximally ignorant about the rate that “go left” or “go right” advice will be correct, but that the “uncertain” advice conveys the information that either choice could possibly be correct.

As Jaynes (2003) demonstrates, these two different states of knowledge must be represented using different probability distributions. In the “uncertain” advice case, where it is at least known both choices are possibly correct, the uniform prior captures the partial ignorance of the subject. For the “go left” and “go right” advice, where it is not even known that both choices could possibly be correct (i.e., it could be the advice is always correct), the complete ignorance of subjects is represented by the Haldane distribution. Conveniently, both of these distributions can be represented as Beta distributions, so that we have the priors

$$\begin{aligned}\pi_l | \mathcal{H} &\sim \text{Beta}(0, 0), \\ \pi_r | \mathcal{H} &\sim \text{Beta}(0, 0), \\ \pi_u | \mathcal{H} &\sim \text{Beta}(1, 1),\end{aligned}$$

where $\mathcal{H}$ explicitly recognizes the background assumptions we are making.

We assume the probability distributions associated with “go left” and “go right” advice are modified as this advice proves to be ‘good’ (i.e., correct) or ‘bad’ (i.e., incorrect). This information is summarized by four counts: g_l , the number of times “go left” advice has proven to be good; b_l , the number of times “go left” advice has proven to be bad; g_r , the number of times “go right” advice has proven to be good; and b_r , the number of times “go right” advice has proven to be bad. Using these counts, π_l and π_r are updated using Bayes theorem, giving

$$\begin{aligned}\pi_l | g_l, b_l, \mathcal{H} &\sim \text{Beta}(g_l, b_l), \\ \pi_r | g_r, b_r, \mathcal{H} &\sim \text{Beta}(g_r, b_r)\end{aligned}$$

For the “uncertain” advice, we assume that no updating takes place, and that π_u continues to be the uniform distribution.

Sequential Sampling Process

Given these assumptions about evidence distributions, the accumulator sequential sampling process proceeds

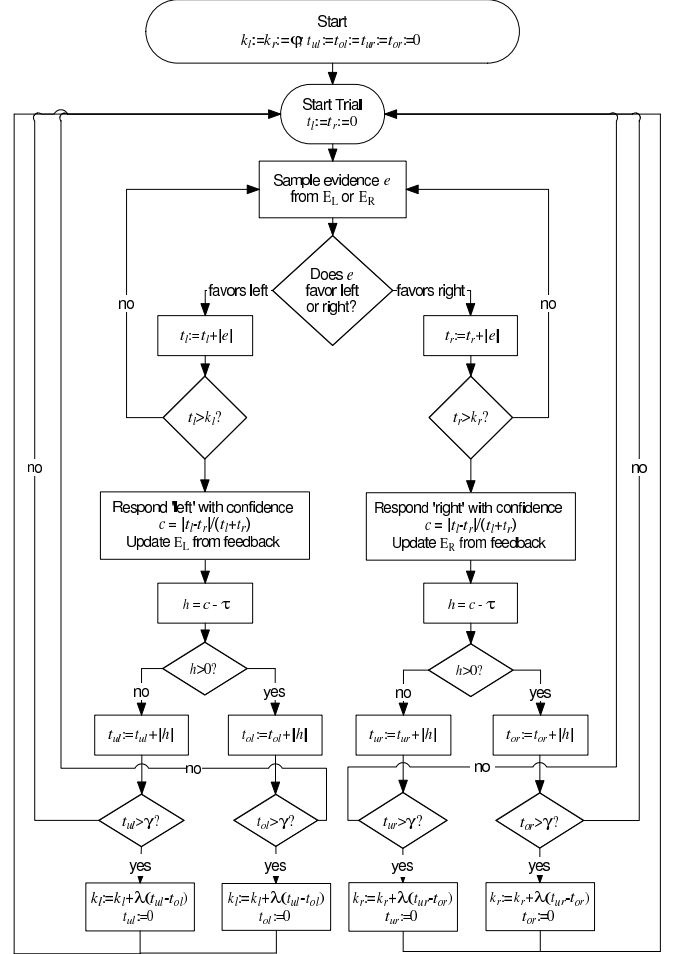


Figure 3: A flowchart describing the self-regulation accumulator model.

according to the flow diagram shown in Figure 3. At the beginning of a series of trials (i.e., at the beginning of a condition), criterion evidence totals, k_l and k_r , for making left and right decisions respectively, are both set to a value κ .

Advice is then provided, corresponding to a probability distribution π . On each iteration, an independent sample is taken from this evidence distribution on the log-odds scale, since the model adds successive evidence values. Thus,

$$e \sim \log \frac{\pi}{1 - \pi}.$$

On each iteration, if e is positive, it is added to a right accumulator, t_r . If e is negative, its absolute value is added to a left accumulator, t_l . Sampling continues until either the total t_l exceeds the criterion total k_l , or the total t_r exceeds the criterion total k_r .

At this point, the model makes a left or right decision accordingly, with a response time corresponding

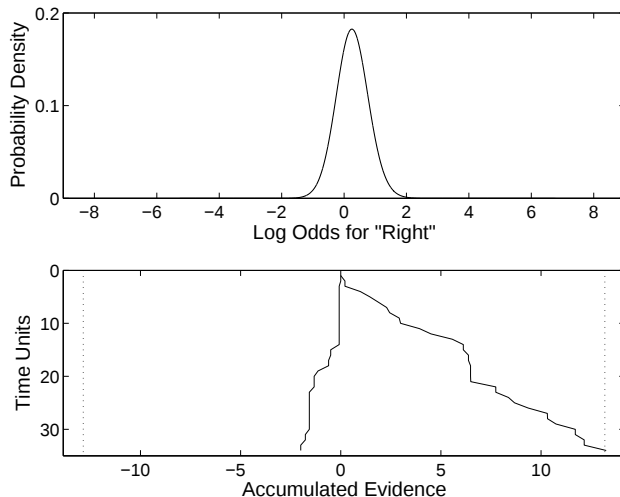


Figure 4: An accumulator sequential sampling process account of decision-making on a single trial.

to the number of iterations. The balance-of-evidence measure of confidence, c , is provided by the difference between the totals, expressed as a proportion of the total evidence accumulated, so that

$$c = |t_l - t_r| / (t_l + t_r).$$

A graphical representation of the model making a decision on a single trial is shown in Figure 4. The top panel shows the evidence distribution in memory corresponding to “go right” advice. The operation of the decision-making process is shown in the bottom panel, with the two accumulators shown as evolving solid lines, and the criterion threshold levels of evidence as dotted lines. In this example, it takes a little over 30 time units for a right decision to be made, with relatively low confidence, because there is similar evidence accumulated for both choices at the time the decision is made.

Adaptation

Having made a decision, the model adapts to its environment in two ways. First, if advice was given, the feedback about the correctness of the advice is used to update the relevant memory counts. In other words, feedback is provided about whether or not the decision was correct, allowing one of the g_l , b_l , g_r and b_r counts to be updated, as appropriate.

Secondly, the decisional confidence is used to self-regulate the criterion threshold levels. The difference between the confidence with which the decision was made, c and a target level of confidence τ is calculated as $h = c - \tau$. If h is positive, it is added to an over-confidence accumulator t_{ol} for left decisions and t_{or} for right decisions. If h is negative, its absolute value is

added to an under-confidence accumulator t_{ul} for left decisions and t_{ur} for right decisions.

If any of these over- or under-confidence accumulators exceeds a critical amount γ , the model undertakes a self-regulating adjustment of a decision threshold. This is done by increasing a decision thresholds making under-confident decisions, or decreasing decision thresholds leading to over-confident decisions, according to a learning rate $0 \leq \lambda \leq 1$ and the difference between accumulated over- and under-confidence totals. For example, if the under-confidence accumulator for left decisions t_{ul} exceeds the critical amount, then the threshold for making left decisions, t_l is increased by $\lambda(t_{ul} - t_{ol})$. The other possibilities for adjustment are formulated similarly, and are detailed in Figure 3.

Model Evaluation

We evaluate the model in two stages. In the first, we focus on descriptive adequacy, by considering the ability of the model to fit the data at an appropriate parameterization. Having demonstrated this, we then focus on the important issue of model complexity (e.g., Myung, Forster, & Browne, 2000; Roberts & Pashler, 2000; Pitt, Myung, & Zhang, 2002), and show that the model is highly constrained in the behavior it can produce at reasonable parameterizations.

Model Fitting

The model has four free parameters. These are: κ , the initial evidence level required to make a decision; τ , the target level of confidence; γ , the critical level of over- or under-confidence needed for self-regulation, and λ , the learning rate.

It is possible to give some interpretation of the scale for each of these parameters, and so constrain in meaningful ways the values they can take. The target level of confidence τ is a value between 0 and 1, which ought to be set at a value above (and likely well above) 0.5 in a two-choice task. The learning rate λ also lies between 0 and 1, with the usual tradeoff between speed and stability of adaptation, and so a range of values are worth considering.

The level of evidence needed to make a decision κ , lies on a log-odds scale, and so can be treated in the same way as Bayes Factors and other likelihood ratios, which are commonly given interpreted on this scale. Kass and Raftery (1995), for example, suggest that a value of 2 provides evidence “not worth more than a bare mention”, while a value of 6 is “positive” evidence, and a value of 10 is “strong” evidence. The critical level of over- or under-confidence γ simply accumulates differences on the confidence scale, so that, for example, a value of 2 would correspond to two completely miscalibrated decisions.

Using these interpretations as a guide, we examined the behavior of the model for every possible combination of the following parameter values: $\kappa =$

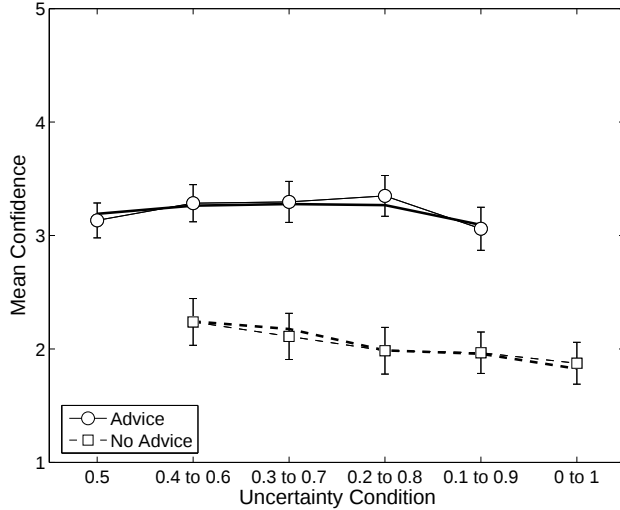


Figure 5: The best fit of the model (heavy lines), over the set of parameterizations considered, to the empirical data measuring mean confidence.

$\{2, 6, 10\}$, $\tau = \{0.6, 0.7, 0.8, 0.9\}$, $\gamma = \{2, 6, 10\}$, and $\lambda = \{0.1, 0.3, 0.5, 0.7\}$. At each of these parameterizations, we measured the fit of the advice acceptance behavior and mean confidence of the model to the empirical data.

Fit was measured on a log-likelihood scale using Gaussian likelihood functions with the means and standard errors shown in Figures 1 and 2, giving equal weight to both the confidence and decision behavioral measures. Because of the different scales on which empirical confidence and the confidence of the model are measured, the single scalar multiple that best mapped the $[0, 1]$ range of model confidence values onto the $[1, 5]$ range for the empirical data was found in each case. That is, the likelihood of a model prediction was assessed once every value was multiplied by the same number that aligned it as closely as possible with the data.

Figure 5 shows the best-fitting behavior of the model to the data, achieved using the parameterization $\kappa = 6$, $\tau = 0.7$, $\gamma = 2$ and $\lambda = 0.7000$.

Figure 6 shows the advice acceptance behavior of the model at the best-fitting parameterization. As with the human data, the model nearly always accepts advice, and displays behavior consistent with guessing when no advice is provided.

It is clear that the model is able to emulate closely the empirical regularities in which we are interested.

Complexity Analysis

Of course, one possible explanation for the ability of the model to fit the data is that it is a very compli-

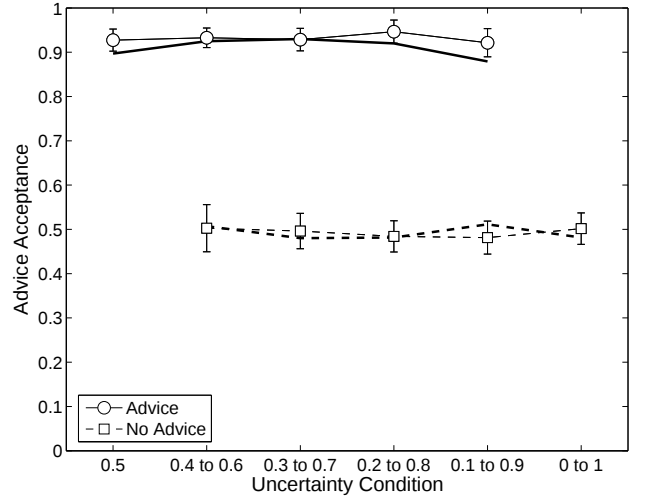


Figure 6: The relationship between the model (heavy lines) and the empirical data measuring advice acceptance behavior, at the best-fitting parameterisation.

cated model, potentially able to fit all sorts of qualitatively different patterns of behaviour by using different parameterizations. To consider this issue, we examined the data patterns generated by all of the $3 \times 4 \times 3 \times 4 = 144$ parameterizations considered.

Figure 7 shows the full range of model behavior for the mean confidence measure across conditions. It is clear that, for advice trials, there are two possible patterns of change in mean confidence. Both have the slight inverted U-shape of the empirical data, but one is more confident than the other, and shows less change in confidence across conditions.

For no advice trials, there are three distinguishable levels of confidence, one of which is too high to agree with the empirical data. Within the lower bands, some of the model behavior shows a linear decrease in mean confidence across conditions that agrees with empirical data, while others shows little or no change in mean confidence.

Figure 7 shows the full range of model behavior for the advice acceptance measure across conditions. It can be seen that, essentially, the model is only able to show one pattern of behavior, which involves accepting advice with high probability, and guessing when no advice is provided.

Taken together, Figures 7 and 8 make it clear that the model is not complicated, in the formal sense explained by Myung, Balasubramanian, and Pitt (2000), since it is able to index a only small number of distinguishable predictions about confidence and decision-making behavior on the task.

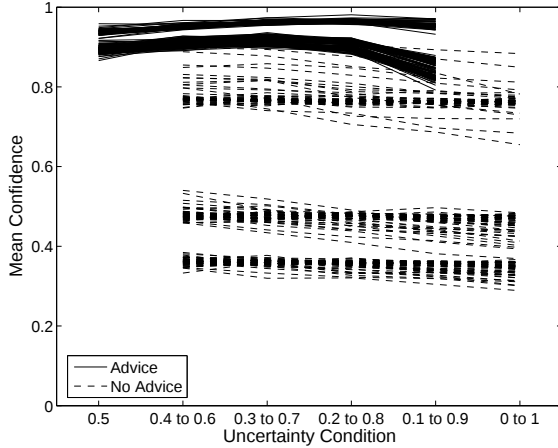


Figure 7: The confidence behaviour of the model, over the set of parameterizations considered, for both advice and no advice trials.

Discussion

Our evaluation shows that the model can account for the data, and does so without benefitting from excessive complexity. What it does not offer is an explanation of *why* the model behaves as it does, and hence why people behave as they do in this task.

Clear insights along these lines are provided by considering which parameter values are responsible for the observed variation in model behavior. If the initial level of evidence to make a decision, κ is set to the lowest value of 2, there is high and relatively constant confidence for both advice trials and no-advice trials (although, on average, no-advice trials are a little lower). This corresponds to the upper bands for both advice and no-advice trials in in Figure 7. Intuitively, this is because often only one or two samples will be sufficient to trigger a decision, and no evidence for the alternative choice will be accumulated, leading to perfect confidence. If κ is set to the more conservative levels of 6 and 10, corresponding to requiring positive or strong evidence before making a decision, then the change in mean confidence across conditions falls into the band that agrees with the empirical data, and confidence for no-advice trials is low.

For the model to agree with the empirical data on no-advice trials, it must show the further pattern of declining as more “uncertain” advice is provided. This is achieved by having either a low criterion γ on the critical threshold for adapting to over- or under-confidence, or a high learning rate λ . A decrease in confidence on no-advice trials is evident when either of these parameter setting is used, and is most exaggerated when they are both present.

Taken together, these conclusions mean that the

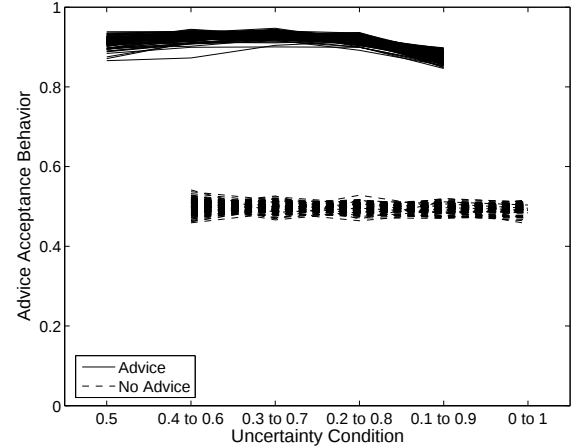


Figure 8: The advice acceptance behaviour of the model, over the set of parameterizations considered.

model captures the empirical data, as long as (a) a conservative initial decision threshold is used, and (b) the learning rate and threshold for adaptivity allow large changes in the self-regulation of the decision threshold. These two conditions correspond neatly with the two ways in which the model adapts: through external adaptation to the environment based on feedback, and through internal self-regulation based on confidence. We discuss each in turn, providing an account of why people behave as they do on this task.

External Adaptation

Because the evidence in memory almost always leads the model to accept advice, what determines confidence is the extent to which evidence for the alternative choice is sampled. For a trial where “go right” advice is given, this is measured by the extent to which the evidence distribution gives probability to negative log-odds. Equivalently, it depends on the extent to which the probability distribution gives density to probabilities less than 0.5.

There are two ways in which the evidence distributions can give density to values less than 0.5. One way is to have a mean near 0.5, and some variance. This is what happens for the experimental conditions where advice is always offered, and so the counts of correct and incorrect advice are relatively close. The other way is to have a large variance, whatever the mean. This is what happens for the experimental conditions where little advice is given, and so the counts of correct and incorrect advice are small.

This state of affairs is represented graphically in Figure 9, which shows the expected probability distributions for the “0.5”, “0.3 to 0.7” and “0.1 to 0.9” experimental conditions at trial 30 out of 50, under the best-fitting parameterization. The “0.3 to 0.7” distri-

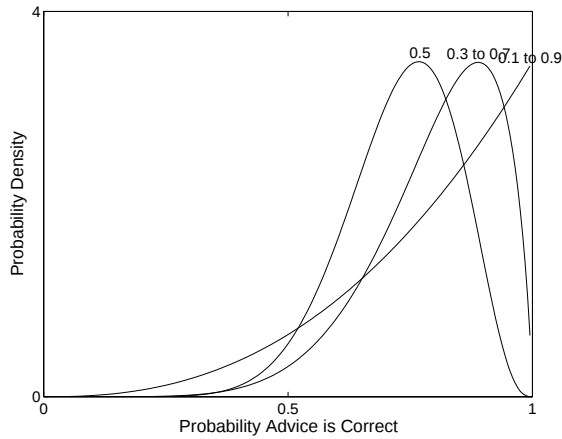


Figure 9: The mean evidence distributions after 30 trials for three of the experimental conditions.

bution gives the least probability to values less than 0.5 because it is centered at a value far away (thanks to accurate advice) and has relatively small variance (thanks to plentiful advice).

Intuitively, when the advice is poor, the model considers the alternative, and loses confidence. When little advice is given, the model is unsure, and loses confidence. When a reasonable amount of reasonably accurate advice is given, the model is most confident. This adaptation to the accuracy and volume of advice is responsible for the inverted U-shape in confidence on advice trials. To be captured, the only requirement is that enough evidence be sampled for the effect to be clear. Hence it is necessary that the initial decision threshold κ not be too lenient.

Internal Adaptation

Figure 10 shows the change in the mean final decision thresholds (i.e., the mean of both t_l and t_r) for each condition, at the best-fitting parameterization. It is clear that, as conditions include more no-advice trials, and more guessing decisions are consequently made with low confidence, the thresholds are adapted to larger values. This self-regulation is responsible for the decline in confidence for no-advice trials. As the criterion thresholds become larger, the many samples taken from the uniform distribution lead to consistent very low balance-of-evidence confidence values. To be evident, the only requirement is that the model self-regulates its decision threshold often enough, and by a large enough amount. Hence it is necessary that criterion γ for adapting to over- or under-confidence not be too lenient, and that the learning rate λ not be too low.

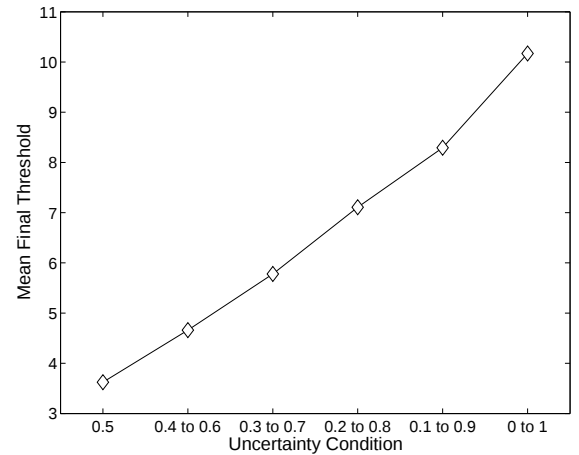


Figure 10: The change in the mean final decision thresholds across the six experimental conditions.

Conclusion

The accumulator model provides a simple, interpretable and elegant account of the four interesting empirical regularities in this task. Future work is intended to test whether the predictions the model makes for different environments—especially those in which accuracy and availability of advice are not perfectly related—are observed in human decision-making.

References

- Jaynes, E. T. (2003). *Probability theory: The logic of science* (G. L. Bretthorst, Ed.). New York: Cambridge University Press.
- Kass, R. E., & Raftery, A. E. (1995). Bayes factors. *Journal of the American Statistical Association*, 90(430), 773-795.
- Myung, I. J., Balasubramanian, V., & Pitt, M. A. (2000). Counting probability distributions: Differential geometry and model selection. *Proceedings of the National Academy of Sciences*, 97, 11170-11175.
- Myung, I. J., Forster, M., & Browne, M. W. (2000). A special issue on model selection. *Journal of Mathematical Psychology*, 44, 1-2.
- Pitt, M. A., Myung, I. J., & Zhang, S. (2002). Toward a method of selecting among computational models of cognition. *Psychological Review*, 109(3), 472-491.
- Roberts, S., & Pashler, H. (2000). How persuasive is a good fit? A comment on theory testing. *Psychological Review*, 107(2), 358-367.
- Vickers, D. (1979). *Decision processes in visual perception*. New York, NY: Academic Press.