



AFRL-SR-AR-TR-2012-0242

Inference in Dynamic Environments

Mark Steyvers (PI)
Scott Brown (co-PI)
UNIVERSITY OF California

03/24/2008
Final Report

DISTRIBUTION A: Distribution approved for public release.

Air Force Research Laboratory
AF Office Of Scientific Research (AFOSR)/ RTA1
Arlington, Virginia 22203
Air Force Materiel Command

REPORT DOCUMENTATION PAGE

Form Approved
OMB No. 0704-0188

Public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing this collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden to Department of Defense, Washington Headquarters Services, Directorate for Information Operations and Reports (0704-0188), 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to any penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number. PLEASE DO NOT RETURN YOUR FORM TO THE ABOVE ADDRESS.

1. REPORT DATE (DD-MM-YYYY) 3-24-2008		2. REPORT TYPE Final Report		3. DATES COVERED (From - To) 7/1/2004 to 12/31/2007	
4. TITLE AND SUBTITLE Inference in Dynamic Environments				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER FA9550-04-1-0317	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S) Mark Steyvers (PI) Scott Brown (co-PI)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) University of California, Irvine. Department of Cognitive Sciences. 3151 Social Sciences Plaza University of California Irvine, CA 92697-5100				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING / MONITORING AGENCY NAME(S) AND ADDRESS(ES) Air Force, Office Of Scientific Research.				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S) AFRL-SR-AR-TR-12-0242	
12. DISTRIBUTION / AVAILABILITY STATEMENT For Unlimited Distribution					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT Real world decision contexts are continually varying, and good decision makers must continually adjust their behavior to track environmental changes. This research investigated decision-making behavior in dynamically changing decision contexts. New experimental paradigms were developed in which dynamic decision making environments forced observers to change their decision-making strategies. Computational models for the decision process in dynamic environments were developed. One model is an ideal observer system in which statistical evidence for a changed environment is weighed in optimal fashion against evidence for a stable environment. The ideal observer analysis resulted in estimates for the optimal number of trials it takes to adjust to new decision environments. The comparison of the ideal observer model against empirical results revealed that some decision-makers adjust too quickly to changes in the environment, and hence perform sub-optimally due to excess variability. Other individuals adjusted too slowly and failed to notice short-term temporal trends. A psychologically plausible particle filter model was also developed. This model accounted for all data, describing both optimal and sub-optimal performance. By varying the number of particles, and the prior belief about the probability of a change occurring in the environment, most of the observed individual differences could be modeled.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT	18. NUMBER OF PAGES	19a. NAME OF RESPONSIBLE PERSON
a. REPORT	b. ABSTRACT	c. THIS PAGE			19b. TELEPHONE NUMBER (include area code)

FINAL REPORT FOR AFOSR GRANT

1. Basic Information

Principal Investigator: Mark Steyvers (msteyver@uci.edu)

Co-Principal Investigator: Scott Brown (scottb@uci.edu)

Title: Inference in Dynamic Environments

Administrative Department: Cognitive Sciences

Institution: University of California, Irvine
3151 Social Sciences Plaza
University of California
Irvine, CA 92697-5100

Award Number: FA9550-04-1-0317

Start date: 7/1/2004

Reporting activity between: 7/1/2005 and 12/31/2007

2. Summary of Research Findings

2.1 Dynamic decision-making

Understanding how people make decisions, from very simple perceptual to complex cognitive decisions, is an important area of research in psychology. We examined decision-making behavior in dynamically changing decision contexts. Real world decision contexts are continually varying, and good decision makers must continually adjust their behavior to track environmental changes.

Consider the case of a military observer making decisions about the identity (friend vs. enemy) of noisy stimuli from reconnaissance pictures. The difficulty of these decisions will change throughout the task, as more or less clear pictures are used, or more or less uniform terrain is observed. An ideal observer must dynamically adjust their decision making process to reflect changes in the environment. For example, if it becomes easier to identify friendly stimuli in new terrain, observers should relax their criterion for identifying enemy stimuli.

Previous research has often assumed static models for decision making that ignore sequential dependencies between environments and the effect of history on current decision making. Some more recent research has focused on dynamic models of decision making, in which certain parameters of the decision process are allowed to vary from decision to decision. However, this research typically makes another static assumption: namely that the environment is stationary.

In this part of the AFOSR funded research (reported in Brown & Steyvers, 2005, *JEP:LMC*, and Brown, Steyvers, & Hemmer, 2007, *Psychological Science*), we developed new experimental paradigms in which dynamic decision making environments forced participants to change their decision making processes in order to remain (approximately) ideal. This paradigm allowed us to observe decision makers tracking changes in the environment. One of the experimental paradigms involved an aircraft flying through a canyon environment (see screenshots below). During the flight, the aircraft is attacked by incoming missiles. There are two types of incoming missiles and the participant in the experiment has to make a quick decision about the correct type of missile in order to choose the appropriate counter-measures. The goal was to measure decision speed in natural environments and also to measure how well participants adapt to changes in the decision making environment (e.g., by making the two types of missiles more or less similar during the course of the experiment).

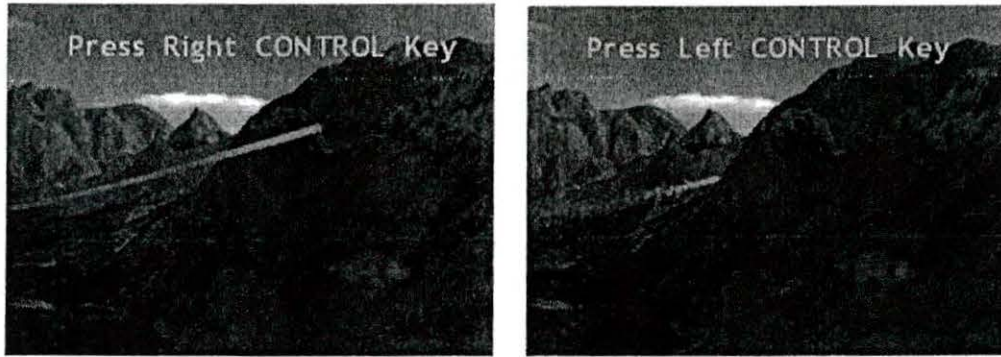


Figure 1. Screenshots of the experiment to track decision-making in dynamic environments

We also developed two models for the decision process in dynamic environments. One model is an ideal observer system in which statistical evidence for a changed environment is weighed in optimal fashion against evidence for a stable environment. The ideal observer analysis results in estimates for the (optimal) number of trials it takes to detect and adjust to new decision environments (see Figure 2). Our other model is a dynamic SDT model that estimates how long it actually takes for individual decision makers to adapt to novel decision environments. By comparing predictions from the ideal observer model to the parameter estimates from the decision model (from individual decision makers), we can quantify the degree of mismatch between ideal and actual observer.

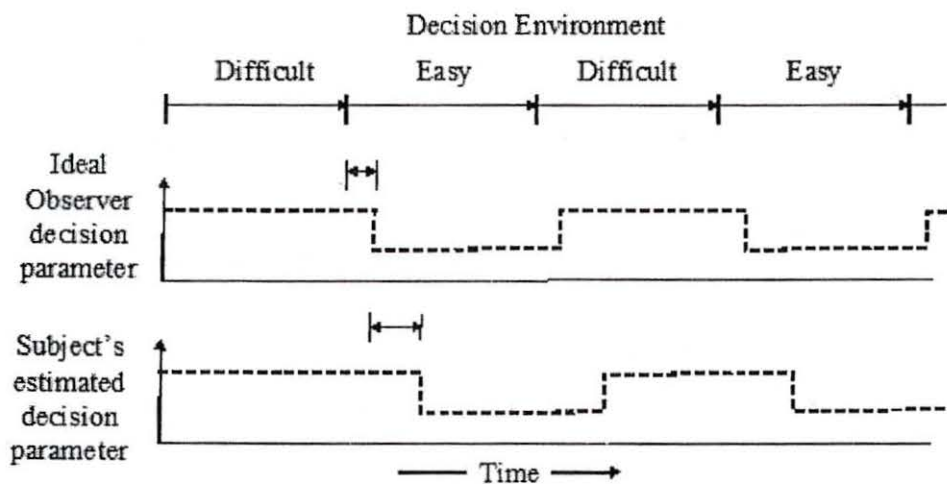


Figure 2.

2.2 The speed of change detection

Many real-world environments involve complex changes over time where behavior that was previously adaptive becomes sub-optimal. These dynamic environments require a rapid response to change. For example, stock analysts need to quickly detect changes in the market in order to adjust investment strategies, and coaches need to track changes in a player's performance in order to adjust team strategy. Reliable change detection requires accurate interpretation of sequential dependencies in the observed data. Research on decision making in probabilistic environments has often called into question our ability to correctly interpret sequential events. When required to predict random events, such as coin tosses, people reliably use inappropriate strategies, such as the famous "Gambler's Fallacy" identified by Kahneman and Tversky (1973, 1979) – if a random coin is tossed several times, people often believe that a tail becomes more likely after a long run of heads. The error of reasoning that underlies the Gambler's Fallacy is the perception of probabilistic regularities in a sequence where no such regularities are present, because the sequence is in fact truly random. Such perceptions might arise because real-world environments rarely produce truly random sequences – often there really is statistical information in the sequence of events. Therefore, the Gambler's Fallacy could simply be the result of people projecting their experience of real-world environments onto laboratory tasks.

Related work in dynamical systems research using response time tasks paints a complementary picture. When the optimal strategy in a task is to provide a series of independent and identically distributed responses, people often perform sub-optimally. Long-range autocorrelations have been observed, where responses depend on earlier responses that occurred quite a long time previously (e.g., Gilden, 2001; Van Orden, Holden, & Turvey, 2003, 2005; Thornton & Gilden, 2005), although not all authors agree on the meaning of the data (e.g., Farrell, Wagenmakers, & Ratcliff, 2006; Wagenmakers, Farrell, & Ratcliff, 2004, 2005). The same criticism applies to dynamical systems research as to the Gambler's Fallacy – tasks requiring long sequences of stationary and conditionally random responses have questionable ecological validity.

Even with real-world environments, people often observe statistical regularities where no such regularities might be present (e.g., Albright, 1993; Gilovich, Vallone & Tversky, 1985). For example when a basketball player makes several successes in a row, observers readily conclude that the player's underlying skill level has temporarily increased; that the player has a "hot hand". Observers make these conclusions even when the data are more consistent with random fluctuations than with underlying changes in skill level. The problem with the hot hand phenomenon is the statistical interpretation of the results. The uncontrolled nature of batting averages and basketball successes make the true state of underlying process impossible to know. Even after detailed statistical analyses of data from many games, statisticians are still unsure whether a "hot hand" phenomenon actually exists in the data (Adams, 1992; Larkey, Smith, & Kadane, 1989; Chatterjee, Yilmaz, Habibullah, & Laudato, 2000). This confusion makes it difficult to draw meaningful conclusions about the optimality of people's judgments.

In this part of the AFOSR funded research, we investigated the ability of human observers to track changes in dynamic environments. In contrast to research on the hot hand phenomenon, we used controlled dynamic environments where we knew exactly how the observations were produced and at what time points the changes occurred. We could therefore assess the degree to which human observers detect changes at the correct times and whether they observe too many or too few changes. When an observer detects too many change points they may behave sub-optimally because they react to perceived changes in the underlying environment that do not exist (e.g., a basketball coach who is prone to seeing hot hands where none are present). Conversely, when an observer detects too few change points, they may fail to adapt to short-lived changes in the environment. This tradeoff between detecting too few and too many change points has often been ignored in previous studies of change detection, which mostly assumed an *offline* experiment where the task is to identify change points in a complete set of data that were observed earlier (see, e.g., Chinnis & Peterson, 1968, 1970; Massey & Wu, 2005; Robinson, 1964). However, real-world examples are invariably *online*: data arrive sequentially, and a detection response is required as soon as possible after a change point passes, before all the data have been observed. Online change detection is also important in clinical settings, particularly for identifying dorsolateral frontal lobe damage. For example, the widely used Wisconsin Card Sorting Task (Berg, 1948) screens patients according to how often they make perseverative errors – that is, how often they fail to detect a change in the task environment, and continue to apply an outdated and sub-optimal strategy. Animal researchers have studied similar behavior in rats (e.g., Gallistel, Mark, King, & Latham, 2001). Rats take some time to detect and adjust to unsignaled changes in reinforcement schedules, but eventually return to optimal probability matching behavior.

We theorized that there might be a U-shaped relationship between the speed with which an individual detects changes and their task performance. Figure 3 helps illustrate this theoretical relationship. Imagine there is some task in which optimal performance is only possible when the subject has an accurate idea of the decision environment. Now suppose that environment is dynamic, and changes with time in an unpredictable manner. Then individuals who detect changes very slowly (left side of Figure 3) will perform poorly because they will base their decisions on outdated assessments of the task environment. Individuals who detect changes very rapidly (right hand side) may also be expected to perform poorly. This is because individuals who detect changes very rapidly are the same individuals who are prone to detecting changes when no changes exist – they “see” change in the environment too easily. These people will perform poorly on the task because they sometimes “detect” changes in the environment that do not exist, and adjust their task performance behavior more often than is optimal.

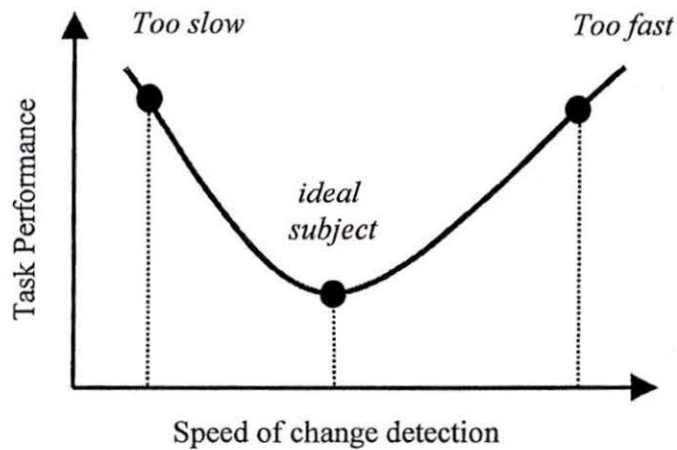


Figure 3: We predicted a U-shaped relationship between performance and speed of change detection. Individuals who detect changes very slowly (left side) will perform poorly, as they base their behavior on old data. Individuals who detect changes very easily (right side) will also perform poorly because they will sometimes “detect” changes that are simply random variability in the environment.

We used sequence prediction experiments to see if we could observe this hypothetical U-shape in real data. In one of the experiments (reported in Steyvers & Brown, 2007; *Neural Information Processing Systems*) presented participants with an 11x11 grid of buttons on a touch-screen computer. We programmed the computer to “light up” buttons using a random sequence, and instructed subjects to try to predict the next button that would light up in the random sequence. We introduced dynamics by making the properties of the random sequence change from time to time. We found that people were very good at predicting the location of the next element in the random sequence. More interestingly, we observed the predicted U-shaped relationship between task error and speed of change detection (see Figure 4).

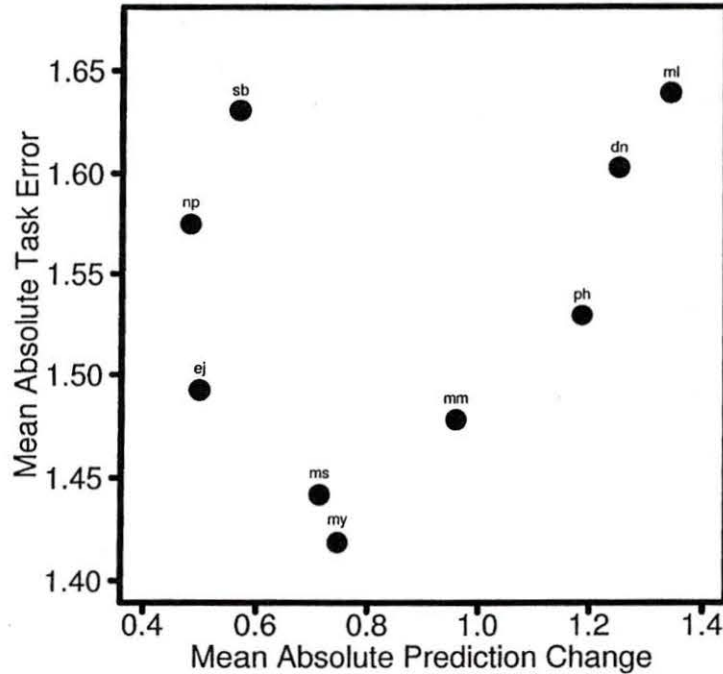


Figure 4. Observed data supporting predicted U-shaped relationship between task error and speed of change detection (“Mean Absolute Prediction Change”).

Given that *all* real-world decision making takes place in a dynamic environment, our results could be very useful. We have shown that some people detect changes in their environment very readily, and some people do not. Both kinds of people perform decision making tasks more poorly than people who detect changes in their environment at an ideal rate. These tools allow us to assess the stability of an individual’s trait of “speed of change detection. This measurement will be useful to know, not just for recruitment, but also for personnel-task matching. In environments where changes are frequent, the ideal individual will be one who usually detects changes too rapidly. Conversely, in a situation where the environment is quite stable, the best performance will be given by people who are slow to detect changes.

2.3 Change Detection: prediction versus inference

We followed up our previous experiments on change detection to get better data for the experimental phenomena and also develop psychologically plausible *particle filters* for our data. In our new experiments (reported in Brown & Steyvers, under review), random numbers were presented to an observer, one number at a time. After each new value was presented, the observer was required to respond, either with an inference about the *mean* value of the process that is currently generating data, or with a *prediction* for the next value. Figure 5 illustrates the particular set of distributions we used in our experiments, along with some example stimuli and two sets of example responses. Each stimulus was sampled from one of four normal distributions, all with the same standard deviation but with different means, shown by the four curves labeled “A” through to “D” in the upper

right corner of Figure 5. The 16 crosses below these distributions show 16 example stimuli, and the labels of the distributions from which they arose are shown just to the right of the stimuli (we call these the “generating distributions”). The five uppermost stimuli were generated from distribution A. Most of these fall close to the mean of distribution A, but there are random fluctuations – for example, the third stimulus is close to the mean of distribution B. After the first five stimuli, the generating distribution switches, and three new stimuli are produced from distribution B. The process continues with six stimuli then produced from distribution D and finally two from distribution A again.

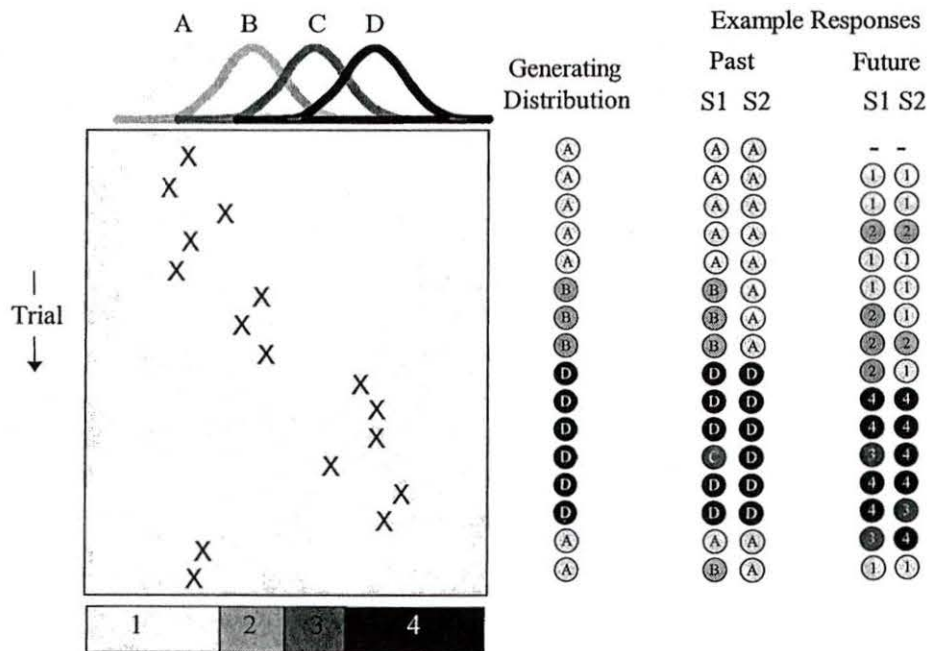


Figure 5. An illustration of the data generating process for the change detection experiments. Each stimulus is sampled from one of four normal distributions, but the selection of the distribution is not observed by the participant. At each timestep, there is a small probability that the generating distribution is switched to a new distribution. Example data for two participants are shown on the right for two different tasks; an inference task where the goal is to identify the generating distribution for the last stimulus, and a prediction task where the goal is to predict the most likely region where the next stimulus will appear. See text for additional details.

There were two different kinds of questions that we asked participants. In the *prediction* condition, we asked about the *future*, requiring participants to predict where the next stimulus would fall. Participants made this response by clicking the mouse in one of four regions, defined by the crossover points of the four distributions and illustrated in Figure 5 by the shaded rectangles in the lower left corner. Some example predictions for two participants (S1 and S2) are shown on the far right of Figure 5. These data illustrate some

general trends that will become important later – for example, the third stimulus is quite close to distribution B, and this causes both participants to predict that the next stimulus will fall in the corresponding region. The other kind of question we asked participants was about the *past*, which we call an *inference* response. In this condition, after each stimulus we asked the participants which distribution they think had generated it. In other words, with the inference task, we require participants to identify the *hidden* state of the generating process. Two example sets of responses are also shown for the inference task on the right side of Figure 5.

```

Sample  $z_0$  randomly from  $\{1, 2, \dots, K\}$ .
Repeat for trial  $i=1, 2, 3, \dots$ 
{
  i. Sample  $x$  uniformly from  $[0, 1]$ 
  ii. If  $(x < \alpha)$  sample  $z_i \sim \{1, 2, \dots, K\} \setminus \{z_{i-1}\}$ 
      else keep  $z_i = z_{i-1}$ 
  iii. Sample observation  $y_i \sim N(\mu_{z_i}, \sigma)$ 
}

```

Figure 6: The stimulus generation algorithm.

Figure 6 provides the details of the process used to generate stimuli. The variable K represents the number of generating distributions, which we set to four. Each generating distribution $j=1, \dots, K$ is normally distributed with mean μ_j and standard deviation σ . We chose mean values μ_j of $\{.2, .4, .6, .8\}$ and $\sigma=.1$ in order to have some overlap in the generating distributions and to insure that almost observations fall on the unit interval. On each trial, an observation was sampled randomly from one of the four normal distributions. After every trial there was a probability α that the distribution would change to one of the other three locations, selected at random. The exact value of α was varied depending on the experimental condition. With a low α , most observations are sampled from the same distribution, requiring few change detections from the participant. With a high α , rapid fluctuations in the underlying generating distributions are possible, requiring participants to carefully interpret the changes in observed stimulus values.

A Particle Filter Model. Our task involved tracking the hidden state of an underlying process whose visible outputs are perturbed by noise. This is an important problem in many statistical and computational applications in which particle filters have had great success. Particle filter models are a type of sequential Monte Carlo integration, and are the real-time (or online) analogues of standard Markov chain Monte Carlo integration methods (see, e.g., Doucet, de Freitas, & Gordon, 2001, or Pitt & Shepard, 1999). For statistical analysis, the central advantages of particle filter models are their simple and efficient computational properties, and that they can be made to approximate the optimal Bayesian solution for some problems, without the intractable integration problems that usually arise in Markov chain Monte Carlo methods. Particle filters also make good candidates for psychological models because they can perform almost optimally, but the

required computations are simple enough to be plausible as a psychological process. The first application of particle filters to psychological data was by Sanborn et al. (2006), who showed that a particle filter model for categorization was able to capture both ideal observer as well as suboptimal behavior simply by varying the number of particles. Particle filters also have similarities to existing, but less formal, models of cognition that have enjoyed considerable support, such as Kahneman and Tversky's (1982) "simulation heuristic". There are many methods for sampling from particle filter models (see, e.g., Doucet, Andrieu, & Godsill, 2000). The most efficient algorithms are the importance weighting methods, but these are too complicated to be plausible as models of human cognition. Instead, we develop a model for the change detection task based on the method of direct simulation, which is much simpler and does not make unreasonable demands on the observer. The mathematical details of this algorithm are set out in the Appendix, but we provide a more intuitive description in the following text, and an illustrative example in Figure 7.

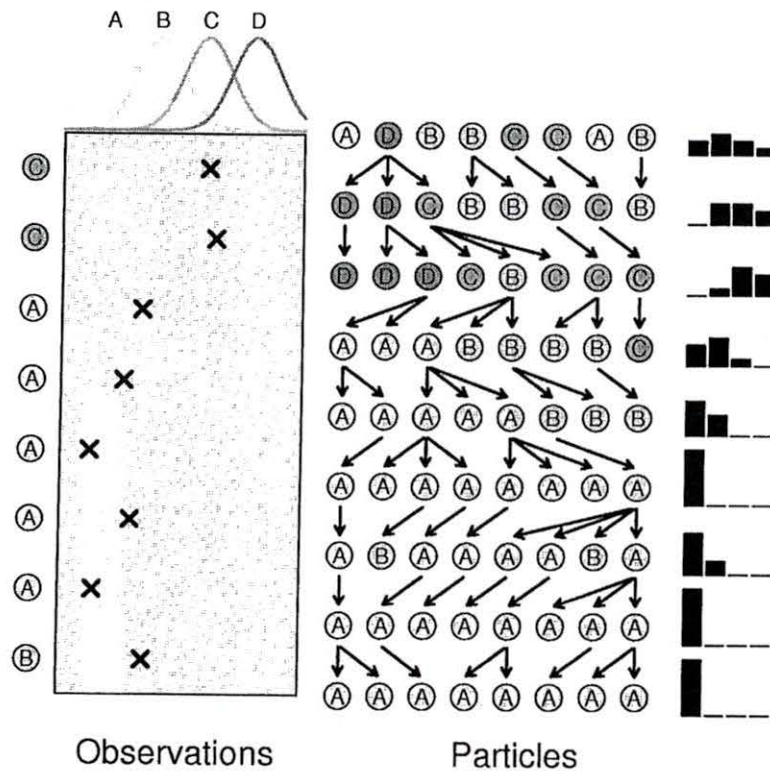


Figure 7. Illustrative example of the particle filter algorithm.

The observer begins with some number of *particles*, say P . Each particle represents a single hypothesis about the mean of the distribution that is currently generating stimuli. An initial set of particles for a $P=8$ system is shown on the top row of the right hand side of Figure 7. These initial particles have randomly distributed guesses – for example, three particles guess that the generating distribution is B , one particle that the distribution is D and so on. The first observation is generated from distribution C , and is illustrated by the uppermost cross on the shaded rectangle. After this observation, the particles are updated

such that particles that are consistent with the observation are kept (or even multiplied) and those that are inconsistent are rejected. The first observation is most consistent with distribution *C*, but also somewhat consistent with *B* and *D*. This causes the two initial particles that hypothesized distribution *A* to be rejected – no arrows show these particles continuing to the next trial. After this filtering, only particles consistent with distributions *B*, *C* and *D* remain. The second observation (also drawn from distribution *C*) is most consistent with distributions *C* and *D*. This time, the filtering operation rejects the three particles from distribution *B* because they are inconsistent with the data. This process continues and the ensemble of *P* particles evolves over trials, and tracks the location of the generating distribution as it moves, illustrated by the histograms on the right side of Figure 7. These show that the distribution of particles tracks the true generating distribution (shown on the far left of the figure).

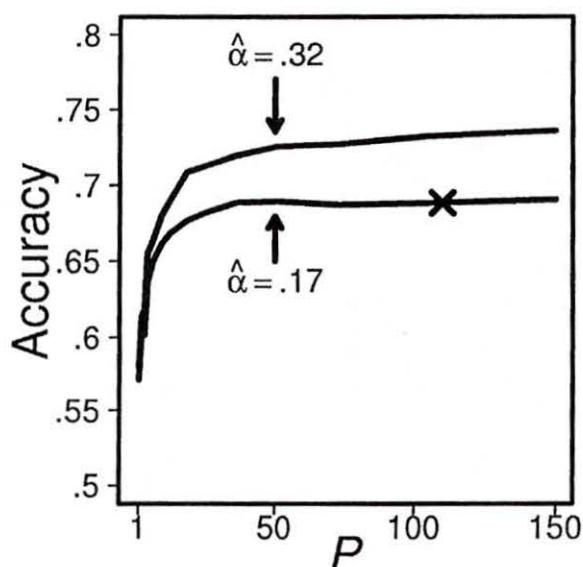


Figure 8. Predicted task accuracy vs. the number of particles (*P*) for the particle filter model, using two different estimates of change probability: $\hat{\alpha}=.32$ (the true value) and $\hat{\alpha}=.17$ (the average participant value).

The particle filter model has two parameters that affect its performance, the number of particles (*P*) and an estimate of the frequency of changes in the generating process ($\hat{\alpha}$). The number of particles imposes a processing limitation, and affects the overall level of task performance. The estimate $\hat{\alpha}$ represents the observer's belief about how often the generating distribution changes – it is the observer's estimate of the parameter α that was used to generate the observations. Figure 8 shows how these two parameters affect task accuracy, using parameter values obtained from our fits to data. We propose that the number of particles might vary from observer to observer, and this will set an overall limit on the accuracy of each observer. Consider for a moment just the upper line ($\hat{\alpha}=.32$) in Figure 8. When there are very few particles (e.g. $P=1$ on the very left side of the graph), accuracy is low, but still well above the chance level of 25%. Such poor performance results from various suboptimal sequential strategies that might mimic human data. Increasing the number of particles increases the overall accuracy of the model. As $P \rightarrow \infty$ the model mimics an ideal observer with statistically optimal decisions,

because the distribution of particles over the hypotheses approaches the posterior distribution conditional on all previous observations. The lower line in Figure 8 shows how overall performance is affected by an inaccurate estimate of change frequency, $\hat{\alpha}$. This line illustrates what happens when the observer believes there are fewer changes in the generating process (17%) than actually occur (32%). Overall performance is lowered, and even with a large number of particles, performance does not grow to the optimal levels. An inaccurate $\hat{\alpha}$ estimate causes decreased performance because it governs how new particles are generated from old ones. Just as in the data, particles remain the same from trial to trial, except for a probability $\hat{\alpha}$ of changing to one of the other three locations. When the estimate accurately matches the environment ($\hat{\alpha}=\alpha$) the model performs most accurately. When the estimate is too small or too large, the model performs inaccurately because it is either too slow to change or too quick to label outlying observations as changes in the generating distribution.

In Experiment 1, we tested participants on the inference task only. That is, after displaying each stimulus, we asked participants which of the four distributions was most likely to have generated that observation. We assessed performance in this task under three conditions, in which there was either a low, medium or high change probability (i.e., α value). The question was whether participants were able to track changes in the hidden state of the data generating process and whether they were sensitive to the changes in the rate at which those changes occurred. Also, this experiment allowed us to investigate individual differences in the accuracy of change detection as well as the number of changes detected. We hypothesized there would be substantial differences in individual ability, with some individuals detecting too few or too many changes, leading to suboptimal performance.

Methods. One hundred and three undergraduates from the University of California, Irvine, participated in Experiment 1. We generated stimuli using the algorithm from Figure 2, and illustrated these using a “tomato processing factory” in which cans of tomatoes were produced from one of four horizontally-separated “production pipes” at the top of the screen (to view a working example of the experiment, visit <http://psiexp.ss.uci.edu/>). The locations of these four pipes correspond to the mean values of four distributions that generate the stimuli. Using simple animations, tomato cans were produced one at a time from one of the four pipes. The standard deviation of each distribution was half of the horizontal separation between pipes. After each new stimulus appeared, all other stimuli on screen were moved downwards to make room, simulating a conveyor belt. No more than 15 stimuli were ever visible on screen at any time, with the lowest one falling off screen. There were four response buttons at the top of the screen, one above each pipe. Participants used these buttons to indicate their responses to the inference questions (“Which pipe generated the most recent stimulus?”). The experiment began with the participant reading through instructions that described the generating process and the random movement of the cans. After this, the first cans rolled out of the machine with all elements of the task visible. This “familiarization phase” lasted for 10 stimuli at the beginning of each block. The participant’s task at this time was trivial, as they could directly observe which pipe had generated the most recent stimulus. The familiarization phase allowed us to identify and exclude participants who failed to pay

attention. It also made the participants more familiar with the idea of inferring which pipe generated the stimuli, and illustrated that the stimuli did not always fall nearest to the pipe that had generated them. After 10 familiarization trials, a curtain covered the machinery that generated the stimuli, beginning the 40 trials of the decision phase. During this phase, the participant's task remained unchanged, but was difficult since the only information available to the participant was the final location of the stimuli. Participants completed one short practice block followed by 12 full blocks divided into four blocks in each of three conditions. The three conditions were defined by the frequency of changes in the underlying generating process: $\alpha=8\%$, $\alpha=16\%$ and $\alpha=32\%$. The pipe used to generate each stimulus was either a repeat of the pipe used for the previous trial (with probability $1-\alpha$) or a new pipe drawn randomly from the other three possibilities (with probability α). All four blocks of each condition occurred consecutively, but the order of the three conditions was counterbalanced across participants. We constrained the pseudo-random stimulus sequences to ensure that there was at least one change in the generating distribution during each familiarization phase. Importantly, we used the same stimulus sequence for corresponding blocks for all participants to reduce variability in comparisons across participants.

Results. The two central attributes of participants' responses were their accuracy (i.e., how often they correctly inferred which distribution had produced the data), and their variability; these measures are summarized in the left and right columns of Figure 9, respectively. Participants averaged about 70% correct responses, with a tendency for accuracy to decrease with increases in the proportion of changes in the data generating process (α). We summarized response variability by calculating the proportion of trials on which the participant provided a different response than they had provided for the previous trial. This measures how often the participant changes their belief about the mean of the underlying data generation process. The histograms for response variability (right column of Figure 9) show that variability increased as the proportion of changes in the data generating distribution increased. On average, participants made 11 response changes per block during the low frequency ($\alpha=8\%$) condition, rising to 13 changes per block in the medium frequency ($\alpha=16\%$) condition and 16 changes per block in the high frequency ($\alpha=32\%$) condition.

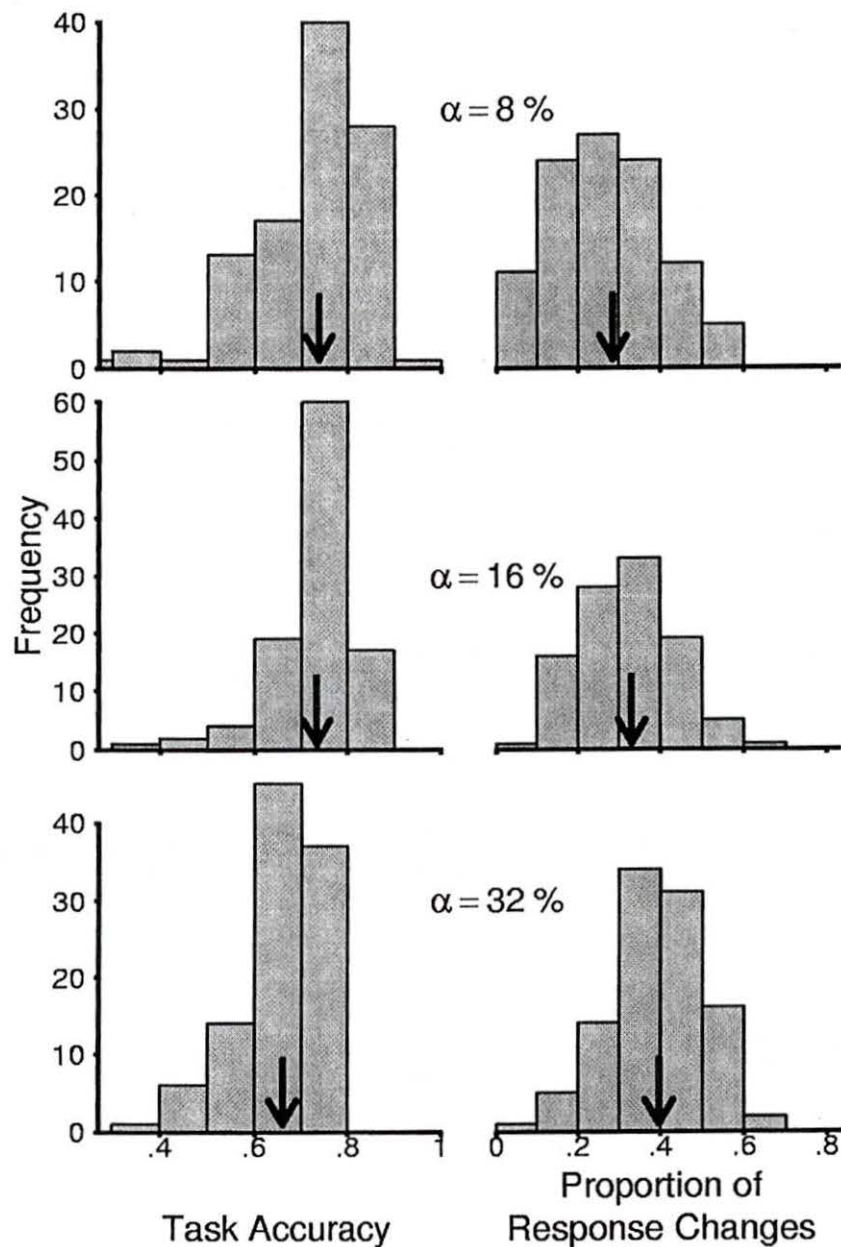


Figure 9. Histograms of task accuracy and response variability across participants in Experiment 1. The arrows show the mean of each distribution.

Figure 10 compares accuracy and response variability on a within-subjects basis, separately for the three conditions: low, medium and high α . Each participant is represented by a black dot in each of the three panels (the grey areas are predictions from the particle filter model, discussed next). The upper boundary of the accuracy data forms an inverted-U shape for all three conditions. This inverted-U shape illustrates the trade-off between subjects who were too quick to detect changes and those who were too cautious.

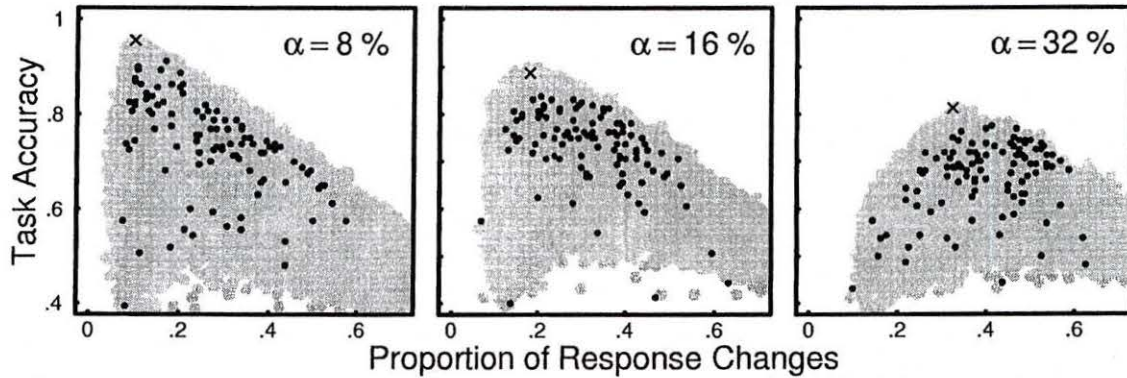


Figure 10. Accuracy vs. response variability (the proportion of trials on which the participant changed their response). From left to right, the panels show data from the three conditions of Experiment 1: low, medium and high probabilities of a change (α) in the data generating process. Black dots show human data, the grey shading shows predictions from the particle filter model, under all possible parameter values. The crosses show “ideal” performance, using the particle filter model with a correct estimate of change probability ($\hat{\alpha}=8\%$, 16% or 32%) and many particles ($P=1000$).

Our participants demonstrated increased response variability and decreased accuracy across the three conditions, as the frequency changes in the generating process (α) increased. This may suggest that participants were appropriately reactive to the experimental manipulation, but the situation may be more complicated. Increasing the frequency of changes in the generating process also increases the variability in the data, and so naturally leads to greater response variability. Put another way, even if participants were completely insensitive to the experimental manipulation of α , they would still exhibit increased response variability with increased α , just by virtue of tracking the (more variable) stimuli. The interesting question is whether participants demonstrated *enough* extra response variability to suggest they were sensitive to the experimental manipulation of α . This question can only be answered by model-based analyses.

Particle filter model analyses. The particle filter model maintains a set of particles that represent a distribution of beliefs about which of the four processes is currently being used to generate data – see, e.g., the histograms on the extreme right of Figure 7. In Experiment 1, participants were forced to give just one response, not a distribution. This constraint is included in the model in the statistically optimal manner, by having the model return the mode of the particle distribution as its inference about the generating distribution (assuming that the goal of the task is to maximize the number of correct decisions). We use two different approaches to assess the model. We first take a global approach, in which we evaluate the model’s predictions for *all* possible parameter settings, and compare the range of behaviors produced by the model with participants’ behavior. Next, we take a focused approach by estimating model parameters for each participant, and investigating the effects of the experimental manipulations on the parameter estimates. The grey circles in Figure 10 illustrate the entire range of behavior

that can be predicted by the particle filter model when given the same stimuli as our participants. We varied the estimate of change probability from $\hat{\alpha}=0$ to $\hat{\alpha}=1$ in small steps, and the number of particles from $P=1$ to $P=1000$ (predicted accuracy reached asymptote around $P=250$). The model captures the observed variability in participant behavior in this task in two ways. Firstly, the model successfully predicts data that were actually observed – almost all data fall inside the range of model predictions. Secondly, the model does not predict data that are wildly different from those that were observed – the grey circles do not generally fall very far from the data. The model also captures the important qualitative trends in Figure 10, including the decreased accuracy and increased variability with increasing change probability (α). The performance of an “ideal observer” is shown by the crosses, which represent the particle filter’s predictions with a large number of particles ($P=1000$) and with perfect estimates of the frequency of underlying changes ($\hat{\alpha}=8\%$, 16% and 32%). Surprisingly, some participants performed close to optimally on this task, and indeed the *average* participant performance was not far below optimal. The upper boundary of the grey shading illustrates the optimal performance level that can be achieved for all different estimates of $\hat{\alpha}$ from zero to one. Even though the accuracy of the participants decreases as they detect too many and too few changes, the accuracy of many participants remains close to the top border of the grey shading. This indicates that many participants could be characterized as operating optimally, except that they used an inaccurate estimate of $\hat{\alpha}$.

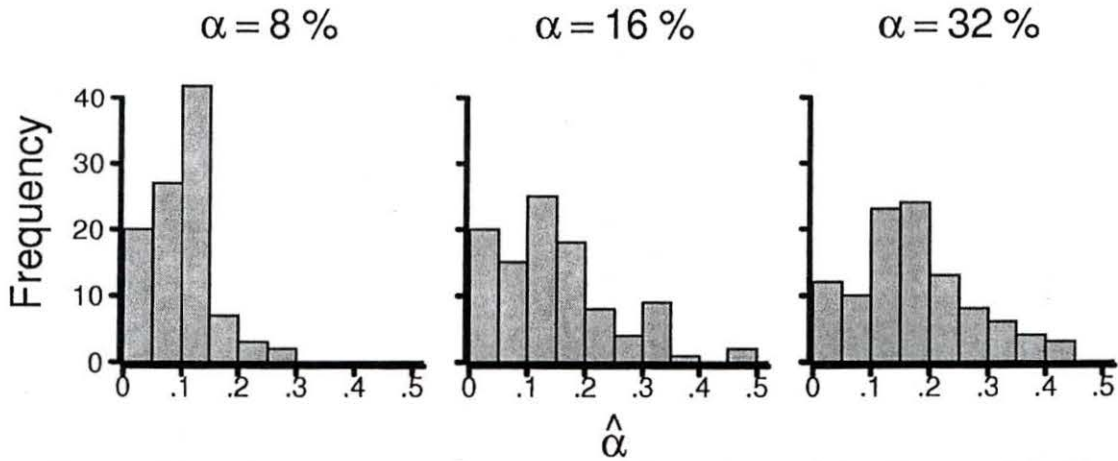


Figure 11. Estimates of the $\hat{\alpha}$ parameter from the particle filter model. The estimates increase as the true value (α) increases from 8% to 32%.

The model analyses provide insight into the question of whether our participants were sensitive to the experimental manipulation of α . We used the particle filter model to generate maximum-likelihood estimates of the parameters (P and $\hat{\alpha}$), separately for each participant. Various techniques exist to estimate the likelihood of a particle filter model (Hürzeler & Künsch, 2001), but we used a simple brute force method; see Appendix B for details. We constrained the parameter estimates to a single, fixed value of P for all three conditions, but allowed three independent values of $\hat{\alpha}$ for the three experimental conditions. The parameters revealed that participants were indeed sensitive to manipulations of the underlying frequency of changes in the data generating process. Figure 11 shows histograms of the $\hat{\alpha}$ estimates across participants, separately for the

three experimental conditions. The $\hat{\alpha}$ estimates were smallest in the $\alpha=8\%$ condition, where the average value across participants was $\hat{\alpha}=10.9\%$, increasing to an average of $\hat{\alpha}=15.3\%$ in the $\alpha=16\%$ condition, and $\hat{\alpha}=17.2\%$ for the $\alpha=32\%$ condition. These estimates confirm that participants were sensitive to the experimental manipulation, although their estimates of $\hat{\alpha}$ were set sub-optimally. There was some shrinkage towards the mean condition (i.e., $\hat{\alpha}$ was too low in the high- α condition and too high in the low- α condition). The average estimate of the number of particles ranged from $P=12$ to $P=400$. The mean value was $P=110$, which is quite large, given that model performance becomes close to asymptotic at around $P=250$. The high value for P indicates that participants were close to optimal, although most often they had an incorrect estimate of the underlying change frequency (i.e. $\hat{\alpha} \neq \alpha$).

The remarkable finding from Experiment 1 was that some participants performed almost as well as a statistically optimally “ideal observer”. The performance of most other participants was well described by the same (ideal) particle filter model, although with an incorrect estimate of the underlying frequency of changes in the generating process. The reader may wonder why our participants were close to optimal, when data from other paradigms reliably demonstrate the fallibility of human decision making. Our dynamic data generating process may be similar to “real world” environments than static (i.i.d.) processes, resulting in better performance. A second feature that separates our paradigm from others is the question posed to participants. We asked our observers to identify which of the four stimulus distributions generated the most recent stimulus. This question asks about the *past*, requiring participants to make inferences about prior states of the world using existing data. Most other analyses of human decision making ask instead about the *future*, requiring participants to make predictions about the outcomes of gambles, for example (e.g., see Kahneman & Tversky, 1973). It is possible that people treat questions about the past and about the future very differently, even when there is no statistical reason to do so. Similarly, Jones and Pashler (2007) have shown that making predictions about the future is not an especially privileged task. Indeed, they found that when similar questions were asked about the future (predictions) and the past (retrodictions), participants never performed better when making predictions than retrodictions.

Advantage of Particle Filters. Modern Monte Carlo techniques provide a natural way to model inference in probabilistic environments, including many decision making tasks commonly used in psychological research. Standard Markov chain Monte Carlo (MCMC) methods, and the newer sequential Monte Carlo (SMC) methods, are both useful solutions to some of the associated computational problems, and variants of each have been proposed as psychological models (e.g. Brown & Steyvers, 2005; Sanborn et al., 2006). While both MCMC and SMC are equally appropriate analyses from a statistical viewpoint, it seems that SMC methods such as particle filters have several advantages as models of human cognition. Firstly, SMC algorithms place lighter and more plausible computational demands on the observer. SMC methods do not require the observer to keep a memory for a long chain of previous events. A related advantage for SMC methods is their ease of application to online (as opposed to post hoc, or offline) experiments – when observations arrive sequentially, and a response is required after

each observation. SMC methods are naturally designed for such sequential paradigms, and employ incremental update algorithms between observations. In contrast, many of the standard MCMC approaches require extensive re-calculation between each observation, making them computationally inefficient and unlikely as psychological process mechanisms. A final advantage of SMC approaches to cognitive modeling is the ability to model both optimal and sub-optimal behavior within the same framework. Standard Bayesian approaches to statistical inference based on MCMC naturally provide optimal performance, under the assumption of accurate knowledge of the data structure. When endowed with a large number of particles, the particle filter models provide the same optimality as the MCMC techniques, but they can also accommodate suboptimal behavior, when the number of particles is small.

Appendix: Particle filter based on direct simulation

Let U and ϕ be the uniform and normal density functions, respectively. Suppose that there are K generating distributions, P particles, and that the observation on trial i , say y_i is normal with standard deviation σ and mean μ_j , where j is the generating distribution. The algorithm is initialized on trial $i=0$ with particles $u_0=\{u_{0,1}, u_{0,2}, \dots u_{0,P}\}$, where each element is one of the integers $1..K$. On each subsequent trial i set a counter $p=0$ and then repeat the following:

1. Sample z randomly from $\{u_{i,1}, u_{i,2}, \dots u_{i,P}\}$.
2. Sample $x_1 \sim U[0,1]$. If $x_1 < \alpha$ sample $z \sim \{1, 2, \dots K\} \setminus \{z\}$.
3. Sample $x_2 \sim U[0,1]$. If $[\phi(y_i | \mu_z, \sigma) > x_2]$ then $p=p+1$ and $u_{i,p}=z$.
4. If $p=P$ stop, else return to 1.

This method of simulation directly mimics the data generation process, but creates an evolving set of particles that approximates the desired posterior distribution.

3. Publications

Brown, S.D., & Steyvers, M. (2005). The Dynamics of Experimentally Induced Criterion Shifts. *Journal of Experimental Psychology: Learning, Memory & Cognition*, 31(4), 587-599.

Brown, S.D., Steyvers, M., & Hemmer, P. (2007). Modeling Experimentally Induced Strategy Shifts. *Psychological Science*, 18(1), 40-45.

Brown, S.D., & Steyvers, M. (under review). Detecting and predicting changes. *Cognitive Psychology*.

Steyvers, M., & Brown, S. (2006). Prediction and Change Detection. In Y. Weiss, B. Scholkopf, and J. Platt (Eds.) *Advances in Neural Information Processing Systems*, 18, pp. 1281-1288. MIT Press.