

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 the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing the 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) 07-07-2006		2. REPORT TYPE Final Report		3. DATES COVERED (From – To) 1 December 2004 - 01-Dec-05	
4. TITLE AND SUBTITLE Modeling Random Walk Processes In Human Concept Learning			5a. CONTRACT NUMBER FA8655-04-1-3052		
			5b. GRANT NUMBER		
			5c. PROGRAM ELEMENT NUMBER		
6. AUTHOR(S) Professor Richard M Young			5d. PROJECT NUMBER		
			5d. TASK NUMBER		
			5e. WORK UNIT NUMBER		
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) University College London Remax House 31-32 Alfred Place London WC1E 7DP United Kingdom			8. PERFORMING ORGANIZATION REPORT NUMBER N/A		
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES) EOARD PSC 821 BOX 14 FPO 09421-0014			10. SPONSOR/MONITOR'S ACRONYM(S)		
			11. SPONSOR/MONITOR'S REPORT NUMBER(S) Grant 04-3052		
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution is unlimited.					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT This report results from a contract tasking University College London as follows: The primary objective was to investigate the scope and implications of random-walk-like processes in human concept learning, and the construction of cognitive models of concept learning valid at the level of individual subjects. The project distinguished two levels of cognition deployed in the task: passive (automatic) learning of stimuli and their classification, as against deliberative processes such as hypothesis testing. Although the focus of research was intended to be on deliberative processing, realistic baseline assumptions about passive processing needed to be tested which proved to be unexpectedly difficult. Eventually, the source of random walk effects were found, in the time subjects take to learn to distinguish between confusingly similar stimuli. It is predicted that allowing a longer interval between trials will allow subjects more time to carry out deliberative processes, leading to better task performance.					
15. SUBJECT TERMS EOARD, Modeling & Simulation, Training, Human Factors, Psychology, Cognition					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT UL	18, NUMBER OF PAGES 17	19a. NAME OF RESPONSIBLE PERSON ROBERT N. KANG, Lt Col, USAF
a. REPORT UNCLAS	b. ABSTRACT UNCLAS	c. THIS PAGE UNCLAS			19b. TELEPHONE NUMBER (Include area code) +44 (0)20 7514 4437

Final Report

Grant FA8655-04-1-3052

Modelling random walk processes in human concept learning

May 2006

Principal Investigator:

Richard M Young
UCL Interaction Centre
University College London
Remax House, 31-32 Alfred Place
London WC1E 7DP
UK

email: r.m.young@acm.org

telephone: 0207 679 5248

Report submitted to

Lt Col Robert Kang, USAF
Program Manager, Life Sciences and Human Effectiveness
European Office of Aerospace Research and Development
223/231 Old Marylebone Rd
London NW1 5TH

Effort sponsored by the Air Force Office of Scientific Research, Air Force Material Command, USAF, under grant number FA8655-04-1-3052. The U.S. Government is authorized to reproduce and distribute reprints for Government purpose notwithstanding any copyright notation thereon.

The views and conclusions contained herein are those of the author and should not be interpreted as necessarily representing the official policies or endorsements, either expressed or implied, of the Air Force Office of Scientific Research or the U.S. Government.

I certify that there were no subject inventions to declare during the performance of this grant.

Final Report

Grant FA8655-04-1-3052

Modelling random walk processes in human concept learning

Summary

Following work funded by a previous grant from EOARD, this project aimed to investigate the role that random walk processes within human cognition play in explaining the high degree of inter-subject variability observed in concept learning. We distinguish two levels of cognition deployed in the task: *passive* (or *automatic*) learning of stimuli and their classification, as against *deliberative* processes such as hypothesis testing. Although the focus of the research was intended to be on deliberative processing, we needed first to make some realistic baseline assumptions about passive processing, and that proved unexpectedly difficult. Eventually (and after the project had ended) we believe we found a source of random walk effects, in the time subjects take to learn to distinguish between confusingly similar stimuli. We also predict that allowing a longer interval between trials will allow subjects more time to carry out deliberative processes, leading to better task performance.

Introduction

This project concerns the role that random walk processes within human cognition may play in explaining the high degree of variability observed in data on concept learning. Previous work by the PI (Young, Cox & Greaves, 2002) — partly supported by EOARD (Young & Cox, 2002) — had shown how, in a model of a simple concept learning task, random walks lead to a wide range of different performances between individual runs of the model, because of the implicit learning required to discover which dimension(s) of the stimuli are relevant to the task. Paying attention to the wrong dimension causes performance to hover at a random 50% level. Because the solution is correct half the time anyway, the cognitive system shifts only slowly (by a random walk) to paying attention to the right dimension, thereby generating a heavily skewed distribution of learning times.

Empirical data

Detailed experimental data for this project were supplied to me in February 2005 by Dr Kevin Gluck of AFRL in Mesa, Arizona, the technical advisor for the project. The data derive from an experiment by Dr Gluck and others, of which only a summary has so far been published (Gluck, Staszewski, Richman, Simon & Delahanty, 2001), closely based on a classic study by Medin & Smith (1981) of what is known as the “5-4 concept learning task”. The stimuli for the task are schematic faces which differ along four binary dimensions: eye height (EH: high or low), eye separation (ES: wide or narrow), nose length (NL: long or short), and mouth height (MH: high or low). Nine of the sixteen faces are used in training, grouped into two categories, five in A and four in B. The remaining seven faces are available for transfer testing.

The data from Dr Gluck consist of detailed records of individual subjects on a trial-by-trial basis in three conditions of the experiment. Some of the subjects provided a verbal protocol of their problem-solving after each stimulus presentation. The data files include the stimulus presented, the subject's latency and response, along with detailed analysis of error patterns and preliminary qualitative analysis of the verbal protocols (on a subject-by-subject basis). Considerable effort was devoted in February and March 2005 to becoming familiar with the data files and interpreting the preliminary analyses. Dr Gluck was helpful at explaining the format of the data files and answering queries about details of the experimental procedure.

Preliminary analysis

Our reading of the research literature (e.g. Murphy, 2002), and also our informal analysis of the cognitive requirements of the task, suggest that subjects can perform the task at either of two levels of cognitive processing. Expressed informally, these are

- 1) *Passive* (or *automatic*) processing: paying attention to the feedback given, and letting repeated exposure to the stimuli and their classifications gradually build up effective knowledge of the correct answers.
- 2) *Deliberative* processing, of which the most obvious form is *hypothesis testing*: where subjects construct and pursue conscious hypotheses such as "I think Eyes High and Mouth Low means it's class B".

Our previous work (Young *et al.*, 2002) led us to believe that the high variability between subjects would derive from differences in the deliberative processing. For one thing, there is obviously considerable room for individual differences in the strategy followed for hypothesis testing, as is well documented in the literature; whereas subjects' passive memory learning presumably all take the same form, even if the individual learning happens at somewhat different rates. Another reason for focusing on deliberative processing as the primary source of individual variation is that in our earlier models, the random walk effect occurs between productions, not in declarative memory. Productions are involved in deliberative processing, but not in declarative learning. Put informally, subjects can take a long time to get the right idea of what to do in their deliberate processing.

Trajectory of the project

Our prior belief that the high degree of variability in concept learning derives from the deliberative aspect of subjects' processing (rather than the passive aspect) motivated our intention to focus on investigating and modelling the deliberative processing. However, there was a slight problem in addressing the deliberative aspect from the outset of the project. According to Act-R, the theoretical framework we are working within (Anderson & Lebiere, 1998), even if subjects are doing deliberative processing, the passive, declarative learning still occurs and makes a contribution to the task performance. We therefore have to make some assumptions, albeit minimal ones, about the passive processing before we are in a position to build models that focus on the deliberative processing. The initial plan for the project was to begin with a short investigation of minimal passive processing (and one or two other necessary preliminaries) before moving to the main topic of deliberative processing.

In the event, and as described in more detail below, dealing with the passive processing proved quite problematic. By the time of the interim report, written around 9 months into the project, we still did not have a satisfactory story about the passive aspect to provide a secure foundation for investigating the deliberative aspect. In consequence, as stated in the interim report, the topic of random walk effects had “dropped below the radar”.

This situation continued to the end of the project as funded (November 2005). During the months since, however, and in preparation for the writing of this present final report, we realised that there *is* probably a random walk effect in the passive learning itself, which could well contribute significantly to the variability in subjects’ learning. The next section of the report explains this discovery. Unfortunately, however, by the time the discovery was made the project was already finished, so it remains theoretical, although ripe for future modelling.

Passive (declarative) learning

As just outlined, subjects’ deliberative processing, such as hypothesis testing, which is the intended focus of this project, takes places within the context of (and makes use of) a declarative memory system which itself contributes to the learning of the task. There are plenty of ad-hoc memory models in the literature for performing this kind of task, some of them very successful at fitting the experimental data. There seems little point, however, and no theoretical interest in simply picking one of these standard models and using brute force to “program” the model in Act-R.

Instead, if one “listens to the architecture” (Newell, 1990) and asks how the Act-R cognitive architecture would “naturally” perform on this concept learning task, a clear answer comes back in the form of the simple strategy of attending to each stimulus, and to its correct classification when it is given, and attempting to recall the classification associated earlier with the given stimulus. This provides a kind of default or baseline strategy for performing the task. For example, a deliberative, hypothesis testing strategy could be compared for its performance against this baseline passive learning strategy.

The problem that arises is that when the baseline strategy is implemented in Act-R in the most straightforward way, using standard Act-R methods and assumptions, the resulting model learns much too fast. We found that such a model consistently learns the nine stimuli presented in the 5-4 experiment after just 2 runs through the stimuli. Such performance is way in excess of what we observe subjects to do, or that we would expect them to do. This observation led us to analyse more carefully the nature of the processing for passive learning.

Confusions, and the lack of them

Consider the following task, with materials that are formally equivalent to those in our concept learning task but easier to write about and to develop intuitions about. Instead of Bruswick faces, we use the letters ‘abcd’ (with order irrelevant), with each letter in either upper or lower case: *bcad*, *BADC*, *aBCd*, etc. Suppose we present a simple old/new recognition task: we show the subject a series of such items, and each time ask whether the item has been seen before. So, first we show *caBD*; it’s new. Then we show *bcad*. The subject hasn’t yet seen an all-lower-case item, so identifies it as new. Next is *CaDb*. Has the subject seen that before? She is

unsure, thinking that there was a previous item with two upper and two lower case letters, but unsure which of the possibilities it was. And so on. After the first few items, the subject is behaving at essentially chance level because of the similarity of the items. However, we assume that given enough exposure to the items, the subject will eventually get to the point of recognising each of them as individuals (instead of as clusters of features) and the confusion is likely to disappear.

In other words, our intuition is that the concept learning task is difficult at least in part because the different items are so similar to each other and therefore so confusable.

Unfortunately, it is easy to program Act-R to perform perfectly on this task from the beginning, quite unlike human subjects. This happens because when a declarative chunk — defined as a cluster of attributes and values — is stored in declarative memory, it is *exactly* that chunk which is stored, no matter how complex it is. And a later retrieval using the retrieval buffer mechanism will retrieve *exactly* that chunk, and no other, regardless of how similar it is and confusable with other chunks. (It is true that Act-R includes various bells & whistles that can be adjusted to provide some of the effects of confusability. However, they have no principled basis, are not amenable to learning, and are of little theoretical interest.)

Role and limitations of spreading activation

In order to model confusability, we eschew the retrieval buffer and instead rely on spreading activation. Because the items densely share features (such as upper-case ‘A’), activation spreads to many of the “wrong” items as well as the “right” item, and the items are therefore confusable.

One limitation of this mechanism is that the items remain confusable, even after long experience. Suppose we have the 16 items just described, consisting of the four features a, b, c, d , each either upper or lower case. Each value of a feature occurs in exactly 8 of the items, so according to Act-R theory, the activation strength S_{ji} from any feature to any item is

$$S_{ji} = S - \ln(8) \approx S - 2$$

The quantity S is nominally equal to $\ln(m)$, where m is the number of chunks in memory. It is difficult to give a meaningful figure to that quantity. If, more plausibly, m were taken to be the number of chunks of a given type, then we 16 chunks of the relevant type, and we would have $S \approx 2.8$. If m were taken to be several hundred, say 400, then we would have $S \approx 6$. We will carry both values through the calculations, writing the larger value in square brackets to give an idea of the range. Thus we have

$$S_{ji} = S - \ln(8) \approx S - 2 = 0.8 [4.0]$$

Consider now what happens when Act-R tries to retrieve the chunk corresponding to a given item from declarative memory. In a balanced experiment, all the items will be retrieved equally often, so their base-level activations will be similar and we can disregard them. Suppose the given item is $abcd$, and the corresponding target is the chunk $*abcd*$. If the given item $abcd$ is in the goal buffer and providing source activation, then since all four features are feeding spreading activation to the target, and assuming there is one other symbol in the goal buffer, the total spreading activation reaching the target item is $4 * 1/5 * S_{ji} = 0.65 [3.2]$.

But various other chunks are highly confusable with the target. Consider the distractor item *abcD* corresponding to the chunk **abcD**. Three of its four features are shared with the target, so the total spreading activation reaching it is $3 * 1/5 * S_{ji} = 0.48$ [2.4]. The same is true for the chunks corresponding to items *abCd*, *aBcd*, and *Abcd*.

So the difference between the activation of the target chunk and any of its close distractors is only around 0.17 [1.2] of an activation unit. Given that activation noise is usually set to around 0.3, this indicates that retrieval of the target chunk can never become reliable. This is especially so considering that there are four such distractors, any of which might be retrieved instead of the target.

Associative learning in Act-R

The analysis above implies that, for the model ever to be able to “recognise” a particular item and distinguish it from its near neighbours in feature space, an Act-R model will have to make use of associative learning, i.e. learning of the S_{ji} associative strengths based upon experience.

Associative learning is something of a Cinderella topic within Act-R. It is rarely used, the existing theory is unsatisfactory, and the implementation even more so. The most extensive exploration of associative learning has been by Lebiere (1998) in his PhD thesis, which simulates the “lifetime learning” of someone learning the basic arithmetic facts of addition and multiplication. Some of those facts are highly confusable (e.g. seven eights are 54?), but with enough exposure, Lebiere’s simulation shows that the facts do eventually become distinct. Our present analysis draws heavily on Lebiere’s work.

The way associative learning comes to make initially confusable items distinct, is by acquiring, in addition to the positive S_{ji} links from features to chunks in which they appears, also negative S_{ji} links from features to chunks in which they do *not* appear. Once these negative links are in place, they serve to push the distractors farther away from the target chunks. In the example used above, it will still be the case that the item *abcd* shares three of its positive links with the distractor item *abcD*. But there will also be a negative S_{ji} from the *d* feature to the chunk **abcD** in which it does not appear, which will prevent that chunk from being retrieved instead of the target **abcd**.

Random walks, again

Consider now the process by which Act-R learns those positive and negative S_{ji} s to make the items more distinct, i.e. by making the retrievals more reliable. Lebiere (1998) emphasises the observation that, whereas production parameter learning in Act-R is a form of supervised learning, depending upon signalled success and failure, associative learning by contrast is a form of *unsupervised* learning, dependent upon experience rather than correctness, in which the internal dynamics of the learning process play the key role.

The learning process is supposed to work like this. Suppose there are two items, X and Y, which are similar in that they share many features, and suppose their corresponding chunks are **X** and **Y**. Suppose there is a single feature *f* on which the items differ. Then initially, because of the positive S_{ji} from *f* to **X**, X is more likely to retrieve **X** than **Y**, even though not reliably so

(as we demonstrated above). Because of the unequal frequencies of retrieval, f will gradually acquire a more negative S_{ji} to $*Y*$ (built up each time $*X*$ is correctly retrieved) than to $*X*$ (which will happen only when $*Y*$ is erroneously retrieved). So, in the end, X will almost always retrieve $*X*$ and not $*Y*$.

But now consider what would happen if, early in the learning, $*Y*$ is retrieved several times by “bad luck” in response to X . This could result in the S_{ji} from f to Y being almost as strong as its S_{ji} to X . The learning process sketched in the last paragraph can kick in only once the S_{ji} s have moved apart. That process of moving apart, depending on the exigencies of experience and random noise, takes the form of a random walk. It can therefore give rise to a long tail in the distribution of learning times, in other words, high inter-subject variability.

(It should also be noted that if, by bad luck, the association of f to $*Y*$ ever became stronger than to $*X*$, the model could end up learning reliably to retrieve the wrong chunk.)

The need for modelling

One of the difficulties of doing research is that ideas do not necessarily present themselves at the right time. The story just sketched clearly cries out for realisation in a running model. However, we discovered this story only recently, while preparing to write the present report, long after the project was over. There therefore has not yet been time to build the appropriate model. We are hopeful this will become possible in the near future.

Other findings

Various other findings were described in the interim report, and are included here for completeness.

Implementing the experiment for a cognitive model

The 5-4 experiment is rather more complicated in design than most experiments modelled in the Act-R architecture we are using, and presented a considerable technical challenge as to how to program the experiment for the model.

The main difficulty concerns the relationship between the control structures of the code for running the experiment and of the model itself. One of the claimed strengths of Act-R modelling is that the “same program” is used to run the experiment with human subjects and with the model. To sustain that claim, it is customary to use the same code to run the experiment with humans or with the model, simply changing a switch to have the program interact with a person or with the Act-R model.

When running the experiment on a human S , it is very clear “who is in charge”. The experiment code determines what happens and presents stimuli to the S , then waits for S ’s response, or times out, or does whatever it has to do next. The S is clearly treated as “subordinate” to the experiment code, being presented with stimuli and then responding (or not). However, this control regime does not transfer well to running the experiment on the model. Having the model be a subroutine of the experiment code is possible only if the model starts only when the

stimulus is presented, and stops once it has responded. But structuring the model that way is cognitively implausible (because a person's cognition "runs" continuously rather than starting and stopping), is regarded as bad form for modelling, and itself leads to further technical problems (for example, the passage of time in the model becomes discontinuous).

The recommended solution is to have the model run continuously, which in terms of control structure clearly makes it the primary procedure, with the experiment code treated as subordinate. In order for the experiment code to maintain its autonomy and its ability to perform actions at times of its own choosing independently of the model, the experiment code has to be "turned inside out" and written in a wholly event-driven manner, with the 'events' deriving from clock time and from the behaviour of the model. Doing this is feasible for very simple experimental designs, but in our case has two major disadvantages:

1. It requires the code for running with the model to be radically different to the code for running a human S, thereby introducing the risk that significant differences will arise between the respective tasks that the human and the model are performing.
2. It requires very awkward, complex, unnatural, and error-prone coding for running with the model, since the event-driven structure means that all control state has to be maintained in variables between successive events.

We did implement the experiment in such a manner, and it worked satisfactorily, although we estimate that it would have been too difficult to do if we had not had the opportunity to debug the basic interaction between experiment and subject (whether human or modelled) first by implementing the code for running human Ss. But it seemed to us that a far better approach is to implement the experiment code and the model as two separate, interacting processes. In that way, the experiment can continue to be coded with a natural control flow as if the subject were indeed subordinate, while the model can also run continuously as if it were the main procedure. The Lisps used to implement Act-R offer facilities for running multiple processes, and indeed make it simple to do so by running the two processes (i.e., experiment and model) in different windows.

So we re-programmed the experiment using a two-process approach. For simplicity, and to respect Dr Gluck's request for portability so that the program can be run in either Allegro Lisp (on a PC) or Mac Common Lisp (MCL, on a Macintosh), we made use of multi-processing primitives provided as part of the Act-R implementation, and minimised our dependence on platform-specific facilities. The resulting experiment code seems to work satisfactorily, is easy and natural to write, and is virtually identical for running the model or a human subject. We ended up using one MCL-specific feature to find the name of the currently running process. We assume this could be translated very easily to run under Allegro; or perhaps with a little more work a way could be found to avoid the platform dependence entirely.

All the code for these implementations is of course available to interested parties.

Selective and sequential encoding of dimensions

Selective encoding

A simple assumption often made by models of concept learning is that the subject encodes the stimulus ‘fully’, i.e. encodes it on at least all the relevant dimensions. However, Anderson & Betz (2001) cite various studies claiming to show that Ss attend to and encode the different dimensions separately and sequentially. Earlier models by the PI have included both sequential and selective encoding of dimensions, in the sense that not all the dimensions are necessarily encoded on each presentation of each stimulus. One of the motivations for such selective encoding is the hope that, with time, Act-R’s parameter learning mechanisms will come into play in such a way that the model automatically learns to encode only the dimensions relevant to the categorisation task.

Such a hope has not been fulfilled. One of the reasons why selective encoding does not spontaneously occur in the model is a phenomenon we describe as “piggy-backing”. Suppose in some concept learning experiment that stimuli have three relevant dimension R1, R2, R3, and one irrelevant dimension I4, and suppose that provided the Ri’s are encoded then the model has learned to make the correct categorisation. Now consider a run of the model in which first R1, R2, and R3 are encoded, and then I4, followed by a correct categorisation. The parameter learning mechanisms for the rule(s) that encode I4 find themselves in a situation where:

- the rule(s) for encoding I4 share credit for the successful outcome with the rules for encoding the Ri;
- the estimated cost parameters for the I4 rule(s) are fractionally lower than the costs for the Ri rules, given that I4 is encoded last and is therefore closest in time to the eventual success.

In consequence, the parameter learning mechanism reinforces the rule(s) for encoding I4. The encoding of I4 is piggy-backing on the success of the rules for the Ri. In fact, the reinforcement for I4 will be slightly greater than for the Ri, given the marginally lower cost. We have made an approximate mathematical analysis of the situation, and although the details become rather messy, a clear outcome is that the hoped-for kind of automatic focusing on the relevant dimensions will not occur.

Sequential encoding

The simplest mechanism in a model for sequential encoding of dimensions is to have independent productions for encoding each of the stimuli, which therefore compete against each other. Predictions of such a mechanism are that (a) dimensions are encoded in a random (though possibly biased) order; and (b) the order of encodings is independent, in the sense that whichever dimension happens to be encoded first does not affect the probabilities for which of the remaining dimensions is encoded next. We were interested in the extent to which these properties hold true in the empirical data.

Details of the analysis and model fitting are presented in the Appendix to this report. In summary, we examined the first few passes through the training set for the Ss who gave protocols, looking for protocol evidence of which dimension was attended to first, which second,

which third, and which fourth. These data turned out to have a distinctive shape suggesting that Ss have two favourite dimensions for initial encoding, with one being preferred for first encoding and the other for second, in a way which indicates non-independence of the encoding. We tuned various models to the observed frequencies of choosing each of the dimensions for first encoding, and examined the resulting patterns for second and later encodings. Their relations to the first encoding differ systematically to those in the empirical data.

We draw conclusions about the subjects' encoding preferences, and believe that we have a simple regime for this "front end" of the task that is adequate to support modelling of the further cognitive processing. Again, details are in the Appendix.

Basic considerations of match between task and Act-R

Rather than imposing on Act-R a preconceived notion of what a model of concept learning should be like, our approach to modelling is to follow Newell's (1990) advice to "listen to the architecture", to try to understand how Act-R best and most naturally lends itself to the performance of the task. In other words, and very approximately: we try to approach the ideal of simply "giving" Act-R the task and seeing what it does with it. To that end, we focus on "local tactics". For example, given a stimulus, what does it make sense for S (or a model) to do? Well, to start encoding the stimulus, for one thing. For another, once the stimulus is at least partly encoded, to see if it brings to mind anything about that or a similar stimulus. And if we have a specific hypothesis, to apply it to the stimulus to see what classification it predicts. And so forth. The idea is that the overall strategy — or better, a wide range of strategies — should emerge from the interactions among these local tactics.

As a start, we report here briefly on three issues that arise from this exercise. The first focuses on the short time available to the S for processing feedback during training. The other two reflect aspects of memory in Act-R that do not seem cognitively realistic.

Short time available for processing feedback

In the original Medin & Smith (1981) experiment, during training, after S makes a classification, feedback is provided and remains visible for 2 seconds. In the experiment by Gluck *et al.*, feedback time may also have been 2 seconds, or it may have been as short as 1 second. Basic task analysis suggests that 2 seconds is a rather short time for S to do what ideally would seem to be required:

- read the feedback;
- recall one's own response;
- compare the two to determine whether was right or wrong; or alternatively, if feedback is given as to whether right or wrong, then if necessary reverse one's remembered response to deduce the correct classification;
- focus on and rehearse the stimulus item together with its correct classification;
- draw conclusions about current hypothesis, update it, and focus on it and rehearse it.

Obviously, a person will not have time to do all that, and even an approximate cognitive model should not either. Two conclusions follow from this simple analysis:

1. We predict that giving Ss a longer inter-trial interval (ITI) to provide more time to process feedback should result in better learning and performance.
2. Such fine details of experimental design are not usually of interest to experimental psychologists. They would normally be unconcerned whether the ITI is 1, 2, 5, 10, or more seconds, and often do not report the duration clearly in published descriptions. It requires a modelling approach to reveal how important such details can be for Ss' performance.

Unrealistic working memory load

A further task analysis in much the same spirit reveals that most accounts of how Ss perform the task, whether modelled or just verbally stated, implicitly assume a rather large pool of dynamic information to be carried in working memory (WM). For example, a straightforward account of that stage of processing feedback typically supposes that Ss can carry the following information:

- a full description of the current stimulus (4 independent items);
- their most recent response (1 item);
- the feedback (1 item);
- the current hypothesis (say 3 items for a 1-dimensional hypothesis: dimension, value, classification; at least 5 items for a 2-dimensional hypothesis);
- possibly an episodic item retrieved from memory (say minimum 3 items as for simple hypothesis);
- and often more.

That gives a count of 10-12 items in WM, minimum. Given that most estimates of dynamic memory capacity are of the order of 3-4 items, there is clearly a discrepancy. Suggesting that, in Act-R terms, much of the information may be carried in declarative memory rather than in the goal chunk, is probably correct, but does not by itself solve the problem, since the standard accounts require a wide range of information to be available in directly testable WM in order to drive the processing by matching the conditions of productions; and information held in declarative memory has to be retrieved and stored in WM before it can influence behaviour.

Act-R does not directly limit the size of the goal chunk, and hence the number of independent items in WM. However, using more than a small number of slots (say around 4) is regarded as poor style, as cognitively unrealistic, and has negative consequences for memory retrieval and learning. As part of an initial exploration of this issue, we implemented a simple, instance-based model in which items are memorised with their correct classification. With unrestricted access to information, its performance is very unrealistic: it spends just a couple of seconds on each stimulus, and learns perfectly after the first pass. We then placed some restrictions on access to information, for example (a) by storing information in declarative memory and having to retrieve it, and (b) by having to look again at the current stimulus in order to know its value on a specific dimension, instead of carrying full details of the encoding in WM. With these changes, the model moved considerably towards more realistic performance, taking 6-7 seconds to process a stimulus, and requiring 4 passes to achieve perfect learning.

References

- Anderson, J. R. & Lebiere, C. (1998) *The Atomic Components of Thought*. Erlbaum.
- Gluck, K. A., Staszewski, J. J., Richman, H., Simon, H. A. & Delahanty, P. (2001). The right tool for the job: Information-processing analysis in categorization. Proceedings of the twenty-third annual conference of the Cognitive Science Society, 330-335. London: Erlbaum.
- Lebiere, C. (1998) *The dynamics of cognition: An ACT-R model of cognitive arithmetic*. PhD dissertation. Technical Report CMU-CS-98-186.
- Medin, D. L. & Smith, E. E. (1981) Strategies and classification learning. *Journal of Experimental Psychology: Human Learning and Memory*, 7 (4), 241-253.
- Murphy, G. L. (2002) *The Big Book of Concepts*. MIT Press.
- Newell, A. (1990) *Unified Theories of Cognition*. Harvard University Press.
- Young, R. M. & Cox, A. L. (2002) Final report on EOARD project: Comparing human concept acquisition to models in a cognitive architecture. Unpublished report, University of Hertfordshire.
- Young, R.M., Cox, A. L. & Greaves, M. (2002) Random walk processes in Act-R mechanisms lead to a wild distribution of learning times. Talk presented at Act-R Workshop, Carnegie Mellon University, Pittsburgh, PA.

Appendix: Analysis and Modelling of Encoding Order

The following is a lightly edited copy of an email sent to Dr Kevin Gluck.

To: Kevin.Gluck@mesa.afmc.af.mil
From: Richard M Young <r.m.young@acm.org>
Subject: Initial looking is not at random
Date: 7 Mar 2005

Hi Kevin,

My "model" of the Brunswick faces at present doesn't do anything with the faces, it just looks at them. But I did a bit of comparative data analysis that you might find interesting.

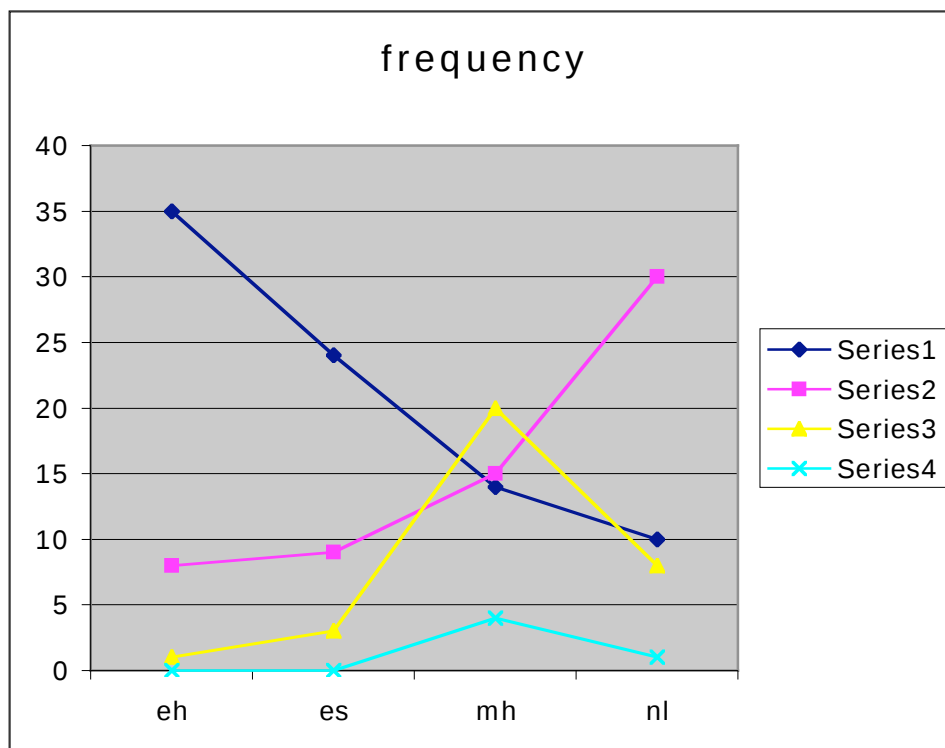
You know that part of the approach I'm taking to modelling variability -- not because I think it's true, but to see how far it will take us -- is to assume that Ss are all "alike", and their differences come partly as the result of randomness and history-dependence (i.e. what they happen to do initially affects what they do next). Well I wondered whether that would apply to the order of encoding the four features of the faces.

On the model side, because at present it doesn't do anything with the faces, it just looks at them, it ends up always encoding all 4 features. But the order isn't fixed. There are 4 separate productions which compete, so the order is random. However, they don't necessarily have equal PG-Cs, so there is a bias to encode some features before others.

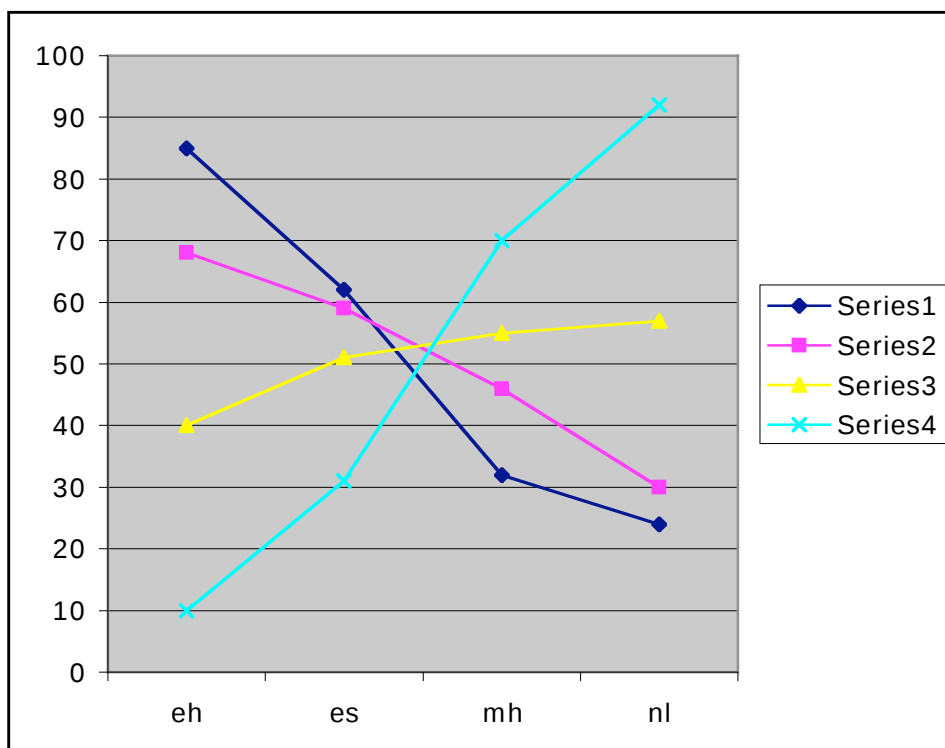
On the human data side, I looked at the detailed tables which show what features each of the protocol Ss mentioned for each stimulus on each pass (what you call a trial, and I'm calling a run). Now each S, or at least many of them, settle down into an idiosyncratic order as the experiment progresses, so I didn't want to use the later passes/runs/trials. Also, of course, not the Rule+X Ss, who are given biasing instructions. So I ended up going through the VP Ss, in the Standard and Prototype conditions, for their first pass through the 9 training faces only, looking at the order in which they mention whichever of the features they say something about. Of course, this is pretty rough-and-ready, but it should give at least an approximate indication of what they look at, at the beginning of the experiment before there's much chance for them to form hypotheses about what's important to look at and what's not.

Here [top, next page] is a diagram of the results. The blue line (Series1) shows the numbers of each feature mentioned 1st, the pink line (Series2) shows the numbers mentioned 2nd, and so on -- but because there are variable number of total features mentioned and not often reaching 4, the curves necessarily flatten off for 3 and 4.

The blue line shows that the eye features EH and ES are mentioned first most often, then MH and least frequently NL. But look at the pink curve for second: it's the exact reverse. That is not at all what we would expect to happen if the features were being chosen randomly but independently in proportions indicated by the blue points. We would expect that in 2nd position, the second-most-probable feature would dominate, i.e. ES. But that's not what happens: the probabilities actually reverse.



I tuned the PG-Cs in the model (4 numbers, but 3 degrees of freedom) so that the first-choice proportions (blue line) are similar to those in your data (percentages in data 42/29/17/12, in model 42/31/16/12), and the choice curves look like this (see below). Because all 4 features are always encoded, curves 3 & 4 don't fall away in



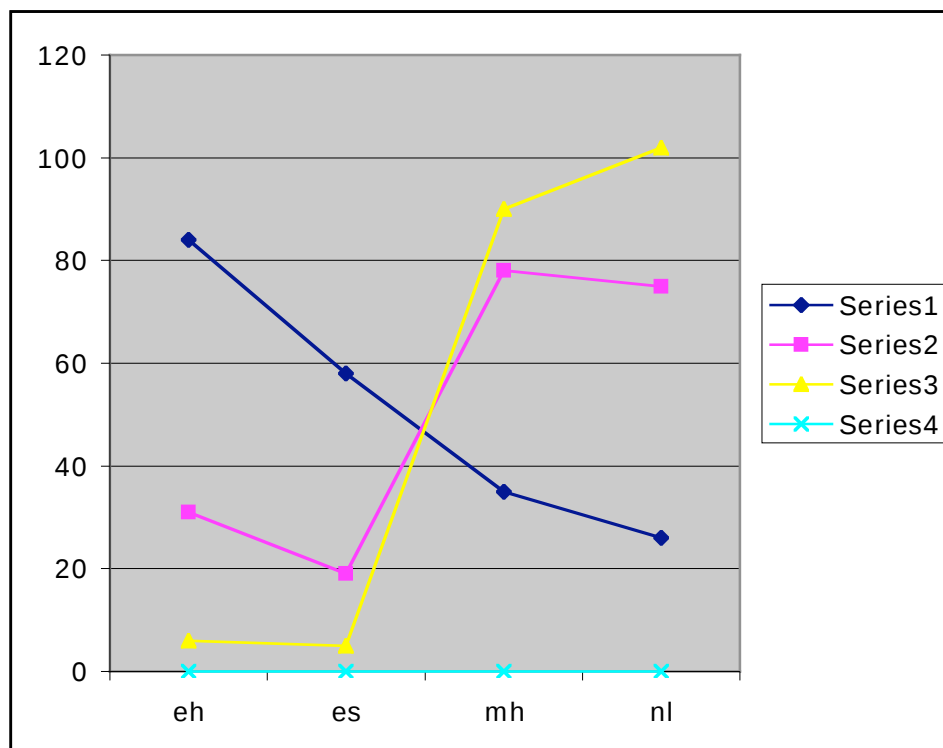
the same way as in the human data, so it's best to concentrate again on the dark blue & pink lines. You can see that the pattern is entirely different. The lines slowly rotate counter-clockwise: the initially most probable features get "used up", and the less probable features dominate the later positions, so that the turquoise curve (4th position) is the mirror image of the dark blue (1st).

So the pattern in the human data has a "signature" that tells us that it *doesn't* come from a random process of the type I've described. It's rather as if, having chosen one of the two eye features first (EH or ES on 71% of occasions), Ss next go mostly for the feature they previously avoided: a pattern of "eyes first, nose next".

This isn't terribly surprising, i.e. that Ss aren't making their encodings at random, and we shouldn't take any of this too seriously, if only because the data are so rough-and-ready. But out of curiosity I changed the model so that the eyes are handled as a single choice (in competition with nose and mouth), with a further choice between EH and ES if E is chosen. This revised model still has 3 degrees of freedom: 2 for the 3-way choice between E, MH, and NL, and 1 for the 2-way choice between EH and ES.

Strictly speaking the revised model can't be right because it means the model can *never* encode both eye features -- but it does have the consequence that only 3 of the 4 features are encoded, and it's close to the human pattern that neither of the eye features is ever encoded in the 4th position. Furthermore it still can't make the pink curve (2nd position) be the mirror of the dark blue (1st), because the MH (second most probable choice) will continue to dominate the NL (least probable). Also there's no way for the EH and ES to swap their relative probabilities, as in the human data. So the pink curve will be a zigzag, but it's conceivable that the overall pattern will look more like the human data.

Once again tuning to the dark blue 1st position, we get the diagram below. The tuning on the 1st pass is again very good (percentages 42/29/17/12 in data, 41/29/17/13 in model). As expected, the pink line for the 2nd pass is zigzag, but it is now showing the reversal, with the eye features low and the mouth & nose features high.



Obviously, one could press on with post-hoc statistical modelling, but given that Ss clearly differ, it's unclear what benefit that would have, or what the resulting model would be "modelling".

TENTATIVE CONCLUSIONS from this:

1. The verbal protocol Ss tend to lump the two eye features together, so that their primary choice is between encoding an eye, mouth, or nose feature, and only secondarily between eye-height and eye-separation.

2. The encoding choices are not made randomly and independently. Rather, for their choice of second feature to encode, Ss appear to be reversing their priorities from the first choice.

(To some extent, 1 & 2 are alternative -- competing? -- explanations for the same data pattern.)

3. For the modelling enterprise, for the feature encoding order at the beginning of the experiment, before it can be guided strongly by hypotheses and feedback, we have available a fairly simple regime that seems to reflect reasonably closely what Ss are doing (at least as measured at this gross level).

4. Ss really do start from different points. Not all the divergence can be attributed to different experiences in the experiment.

Any comments welcomed.

-- Richard