

Optimal Training Systems

Phase II STTR

Contract # FA9550-05-C-0168



Final Report

April 30, 2008

Submitted to:

Air Force Office of Scientific Research (AFOSR)

Submitted by:

Michael Matessa (mmatessa@alionscience.com)
Alion Science and Technology - Micro Analysis and Design Operation
1789 S. Braddock Ave, Suite 400
Pittsburgh, PA 15218

Marsha Lovett (lovett+@cmu.edu)
Eberly Center for Teaching Excellence
Psychology Department
Carnegie Mellon University
5000 Forbes Ave
Pittsburgh PA 15213

REPORT DOCUMENTATION PAGE				<i>Form Approved</i> 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) 04-30-2008		2. REPORT TYPE Final		3. DATES COVERED (From - To)	
4. TITLE AND SUBTITLE OPTIMAL TRAINING SYSTEMS				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER FA9550-05-C-0168	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S) Dr. Michael Matessa Dr. Marsha Lovett				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) ALION SCIENCE AND TECHNOLOGY 1789 S. BRADDOCK AVE, SUITE 400 PITTSBURGH, PA 15218 PSYCHOLOGY DEPARTMENT CARNEGIE MELLON UNIVERSITY 5000 FORBES AVE PITTSBURGH, PA 15213				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING / MONITORING AGENCY NAME(S) AND ADDRESS(ES) AFOSR 875 N Randolph St Arlington, VA 22203				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S) AFRL-SR-AR-TR-08-0258	
12. DISTRIBUTION / AVAILABILITY STATEMENT Distribution A: Approved for public release. Distribution is unlimited.					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT This report investigates the training implications provided by models built for two domains: learning biology in an online course and learning basic flight maneuvers in an unmanned aerial vehicle simulator. The biology model uses rules to decide on study behavior. The flight maneuver model uses instances of expert behavior to decide on correct flight control actions. Both models are implemented in the ACT-R cognitive architecture. Essentially, the path to optimal training in both these cases involves finding the key domain feature to which learning progress is very sensitive. Based on our results, we would posit that explicitly training on these key features would promote more efficient learning.					
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)

Summary

Successful training in complex environments is normally accomplished through the interaction of a trainee and a skilled expert, but due to resource constraints, experts' use in training can be problematic. Computational models that learn task performance subject to human constraints may be useful in understanding the details of training and offer training suggestions that can be implemented in computerized tutoring systems.

This report investigates the training implications provided by models built for two domains: learning biology in an online course and learning basic flight maneuvers in an unmanned aerial vehicle simulator. The biology model uses rules to decide on study behavior. The flight maneuver model uses instances of expert behavior to decide on correct flight control actions. Both models are implemented in the ACT-R cognitive architecture.

In modeling both the biology and flight maneuver domains, it was found that information needed for good performance is at some level available to the trainee but might not be used. In the biology domain, the option to re-visit a previous topic is implicitly available. In the flight maneuver domain the rate of change information is indirectly available. The key insight for training is to make explicit to the student these aspects of the environment/representation so that the natural learning mechanisms can unfold in more productive ways. This relates to the idea of optimal training because our goal is to take best advantage of the human learning system. Essentially, the path to optimal training in both these cases involves finding the key domain feature to which learning progress is very sensitive. Based on our results, we would posit that explicitly training on these key features would promote more efficient learning. This position is in line with results such as Klahr and Nigam (2004), which show that direct instruction is more effective than discovery learning.

1 Introduction

Successful training in complex environments is normally accomplished through the interaction of a trainee and a skilled expert, but due to resource constraints, experts' use in training can be problematic. Computational models that learn task performance subject to human constraints may be useful in understanding the details of training and offer training suggestions that can be implemented in computerized tutoring systems.

This report investigates the training implications provided by models built for two domains: learning biology in an online course and learning basic flight maneuvers in an unmanned aerial vehicle simulator. The biology model uses rules to decide on study behavior. The flight maneuver model uses instances of expert behavior to decide on correct flight control actions. Both models are implemented in the ACT-R cognitive architecture.

ACT-R (Anderson et al., 2004) is a production system theory that tries to explain human cognition by developing a model of the knowledge structures that underlie cognition. There are two types of knowledge representation in ACT-R -- declarative knowledge and procedural knowledge. Declarative knowledge corresponds to things we are aware we know and can usually describe to others. Examples of declarative knowledge include "George Washington was the first president of the United States" and "An atom is like the solar system". Procedural knowledge is knowledge which we display in our behavior but which we are not conscious of. For instance, no one can describe the rules by which we speak a language and yet we do. In ACT-R declarative knowledge is represented in structures called chunks and held in the Declarative module, whereas procedural knowledge is represented as rules called productions and held in the Procedural module. A production rule is a statement of a particular contingency that controls behavior. An example might be

IF the goal is to add two digits d_1 and d_2 in a column
and $d_1 + d_2 = d_3$ is retrieved
THEN set as a subgoal to write d_3 in the column

The condition of a production rule (the IF part) consists of a specification of the chunks in various modules. The action of a production rule (the THEN part) consists of modifications of the chunks in modules, requests for other chunks to be placed into the modules, or requests for other actions to be taken.

1.1 Subsymbolic attributes of ACT-R

At a subsymbolic level, facts have an activation attribute which influences their probability of retrieval and the time it takes to retrieve them. Rules have a utility attribute which influences their probability of being used. The activation A_i of a chunk i is computed from three components – the base-level, a context component and a noise component. The base-level activation B_i reflects the recency and frequency of practice of the chunk. The equation describing learning of base-level activation for a chunk i is

$$B_i = \ln\left(\sum_{j=1}^n t_j^{-d}\right) \quad (\text{Equation 1.1})$$

where n is the number of presentations for chunk i , t_j is the time since the j th presentation, and d is the decay parameter.

The equation for the activation A_i of a chunk i including context is defined as:

$$A_i = B_i + \sum_k \sum_j W_{kj} S_{ji} \quad (\text{Equation 1.2})$$

Measures of Prior Learning, B_i : The base-level activation reflects the recency and frequency of practice of the chunk as described above.

Across all modules: The elements k being summed over are the modules.

Sources of Activation: The elements j being summed over are the chunks which are in the slots of the chunk in module k .

Weighting: W_{kj} is the amount of activation from source j in module k .

Strengths of Association: S_{ji} is the strength of association from source j to chunk i .

The weights, W_{kj} , of the activation spread defaults to an even distribution from each module. The total amount of source activation for a module is called W_k and is settable for each module. The W_{kj} values determined by the following equation:

$$W_{kj} = W_k / n_k \quad (\text{Equation 1.3})$$

where n_k is the number of chunks in the slots of the chunk in module k . The strength of association, S_{ji} , between two chunks is 0 if chunk j is not in a slot of chunk i or is not itself chunk j and is set using this equation when chunk j is in a slot of chunk i or is itself chunk j :

$$S_{ji} = S - \ln(fan_j) \quad (\text{Equation 1.4})$$

Where S is a parameter to be estimated (set with the maximum associative strength parameter)

And fan_j is the number of chunks in which j is the value of a slot plus one for chunk j being associated with itself.

1.2 Partial matching in ACT-R

In some situations a chunk that exactly matches a request cannot be retrieved but it is desirable to retrieve a closely matching chunk. This is what the partial matching mechanism is designed to address. When partial matching is enabled, the similarity between the chunks in the retrieval request and the chunks in the slots of the chunks in declarative memory are taken into consideration. The chunk with the highest activation is still the one retrieved, but with partial matching enabled that chunk might not have the exact slot values as specified in the retrieval request.

The activation A_i of a chunk i is defined fully as:

$$A_i = B_i + \sum_k \sum_j W_{kj} S_{ji} + \sum_l PM_{li} + \varepsilon \quad (\text{Equation 1.5})$$

B_i , W_{kj} , S_{ji} , and ε have been discussed previously. The new term is the partial matching component.

Specification elements l: The matching summation is computed over the slot values of the retrieval specification.

Match Scale, P: This reflects the amount of weighting given to the similarity in slot l . This is a constant across all slots with the value and a typical setting is 1.0.

Match Similarities, M_{li} : The similarity between the value l in the retrieval specification and the value in the corresponding slot of chunk i .

1.3 Recall probability in ACT-R

If we make a retrieval request and there is a matching chunk, that chunk will only be retrieved if it exceeds the retrieval activation threshold, τ . The probability of this happening depends on the expected activation, A_i , and the amount of noise in the system which is controlled by the parameter s :

$$\text{recall probability}_i = \frac{1}{1 + e^{\frac{\tau - A_i}{s}}} \quad (\text{Equation 1.6})$$

Inspection of that formula shows that, as A_i tends higher, the probability of recall approaches 1, whereas, as τ tends higher, the probability decreases. In fact, when $\tau = A_i$, the probability of recall is 50%. The s parameter controls the sensitivity of recall to changes in activation. If s is close to 0, the transition from near 0% recall to near 100% will be abrupt, whereas when s is larger, the transition will be a slow sigmoidal curve.

1.4 Choice probability in ACT-R

If there are a number of productions competing with expected utility values U_j the probability of choosing production i is described by the formula

$$\text{Probability}(i) = \frac{e^{U_i / \sqrt{2}s}}{\sum_j e^{U_j / \sqrt{2}s}} \quad (\text{Equation 1.7})$$

where the summation is over all the productions which are currently able to fire (their conditions were satisfied during the matching). Note however that that equation only serves to describe the production selection process. It is not actually computed by the system. The production with the highest utility (after noise is added) will be the one chosen to fire. The utilities of productions can be adjusted according to the rewards they receive. If $U_i(n-1)$ is the utility of a production i after its $n-1$ st application and $R_i(n)$ is the reward the production receives for its n th application, then its utility $U_i(n)$ after its n th application will be

$$U_i(n) = U_i(n-1) + \alpha[R_i(n) - U_i(n-1)] \quad (\text{Equation 1.8})$$

where α is the learning rate. This is also basically the Rescorla-Wagner learning rule (Wagner & Rescorla, 1972). According to this equation the utility of a production will be gradually adjusted until it matches the average reward that the production receives.

2 Natural learning interactions in an online course

2.1 OLI platform

We have conducted this part of the work in the context of Carnegie Mellon's Open Learning Initiative (OLI; Smith & Thille, 2004), a collection of freely available online courses, funded by the Hewlett Foundation. OLI courses are developed collaboratively by teams that consist of content experts, cognitive scientists, and instructional technologists. The courses are designed to incorporate principles from the learning sciences and then to be continually refined through user testing and formative assessment.

OLI courses do not follow the model of simply putting a textbook online for students to (passively) read. Rather, OLI courses are fully online, interactive courses that *enact instruction*. In other words, just like regular courses, OLI courses have many components, including exercises, reflective activities, problems, interactive animations and simulations as well as both low- and high-stakes assessments sprinkled throughout each module.

Because of this high level of interactivity, OLI courses provide a productive platform for studying teaching and learning. In particular they include real students taking real courses within a highly instrumented system. Each student interaction in OLI is automatically logged, producing a rich database (e.g., individual students' answers to specific questions, precise times to complete particular activities, and patterns of feedback/hint use). This produces "laboratory quality" data, akin to what are collected in learning science experiments, but with the duration (i.e., a semester) and authenticity of real course-based learning contexts – a combination rarely achieved in laboratory experiments.

These automatically collected data have already been useful for informing various improvements to the OLI courses, and they have offered additional information on which to compare students' learning in OLI vs. more traditional courses. In contrast, the current work aims to use these OLI-log data to investigate the nature of students' learning within the OLI course and, in particular, the learning consequences of the various choices students make within OLI courses.

2.2 Context for current study

The current study involves two particular OLI courses, OLI-Biology (primarily) and OLI-Statistics (secondarily). Both courses were taught in a *blended* mode in which students completed particular portions of the online course (by specific dates) and then attended lectures (three times a week for OLI-Biology and two times a week for OLI-Statistics). Besides the OLI materials in each course being substantially different in format from a conventional textbook, the lectures were conducted in a somewhat non-traditional way as well. For both the Biology and Statistics lectures, the instructors were able to track (at least in a loose way) students' progress through the OLI material, and they used this information to adjust the content of their lectures to more directly address the kinds of difficulties students were facing.

2.2.1 Biology Course Data

From the enrollment rosters, we had access to information on students' "home college" (e.g., Carnegie Institute of Technology, Mellon College of Science, or Humanities and Social Sciences) and their year in college (i.e., first-year, sophomore, junior, or senior). We also administered a beginning-of-semester survey to gather additional data about (a) students' expectations for the course, (b) their beliefs about learning, and (c) their reported use of effective study strategies. Students' expectations were simply measured by asking what they expected as their final course grade. The beliefs about learning questions were taken from Schommer's (1990) epistemological beliefs questionnaire, specifically those questions that related to students' belief that learning is easy and fast. The questions about effective study strategies were taken from Pintrich's metacognitive strategies learning questionnaire (cf. Garcia & Pintrich, 1996).

The primary measures of students' learning outcomes were two paper-and-pencil, in-class exams (given after weeks 4 and 7). There were also five high-stakes quizzes, administered online but outside of the OLI system. The first quiz served as a baseline of sorts because it occurred fairly early in week 1.

In addition to these assessments, a rich data stream capturing each students' interaction with the OLI-Biology system offered another source of data. From the OLI log files, we first culled data to measure how much students were using different parts of the OLI system. Specifically, we calculated the following measures: total time spent in OLI, number of OLI sessions initiated, number of OLI pages viewed, number of self-assessments (i.e., low-stakes assessments within OLI) submitted, and number of times the objectives list was viewed. In addition, we focused on students' interactions with a particular learning tool that occurred during weeks 3 and 4. The number of times students accessed the tool, total "steps" taken within the tool, and total amount of time spent using the tool were also culled from the OLI log data. All of these data describing students' interactions with the OLI-Biology course – both in general and with respect to the particular learning tool – were used to identify patterns of what students chose to use within the course as well as to test for correlations between use and various learning measures.

Finally, the general measures of students' number of sessions connected to the OLI-Biology course were also broken down into specific time windows within each three-week phase of the study. These time windows were set as ever-narrowing spans of time leading up to the subsequent exam, namely, first two weeks of the three-week phase, next 6 days (i.e., a week before the subsequent exam up to a day before the exam), and the day before the exam. These data serve as a launching point for considering how students distribute their study time when using the OLI course.

2.2.2 Statistics Course Data

Data from the Statistics course were fairly similar in that we had both paper-and-pencil assessments (e.g., quizzes and tests from the course) and the OLI log data that tracked their ongoing interactions (e.g., viewing pages, completing activities, and taking low-stakes assessments). We did not have students' epistemological beliefs and course

expectations in the case of Statistics, but instead we had a baseline measure of their incoming statistics knowledge by administering a Statistics knowledge assessment developed by statistics education researchers (delMas, Ooms, Garfield, & Chance, 2006). This test is named the Comprehensive Assessment of Outcomes in a first Statistics course (CAOS), and it is a 40-item multiple choice test designed to measure students' basic statistical reasoning. We also administered the CAOS test at the end of the course to measure students' *learning gain* on those items. Note that the CAOS test was designed to emphasize the basic skills of statistical reasoning.

2.3 Empirical results on learning

2.3.1 Biology: Basic learning and performance results and individual differences

Table 2.1 shows the descriptive statistics for each of the background measures on students' demographics (major, year in school). These background variables show that our study sample consisted of mainly science and engineering majors and that the majority of the students were in their first year of college. This is consistent with many introductory science courses' enrollments. It is also worth noting that the average final course grade expected by students in this course was 3.6 on a 4.0 scale (i.e., a low "A"). And, the two key indices from the epistemological beliefs and learning strategies survey showed that (a) students tended *not* to believe that "learning is easy/fast" (i.e., average of 2.1 on a 1 to 7-point scale, where 1 = strongly disagree and 7 = strongly agree) and (b) students tended to use effective study strategies but not overwhelmingly so (i.e., average of 4.6 on a 1 to 7-point scale). These two results are rather encouraging with respect to the students being in a good position to learn science effectively.

Table 2.1: Number of students in various major areas and in different years of college

Major Area	% of Students	Year in College	% of Students
Engineering	25	First year	66
Science	41	Sophomore	20
Humanities	16	Junior	9
Other	18	Senior	5

Table 2.2 shows the results of the primary learning assessments mentioned above, namely the two exams and five quizzes, all of which were administered outside of the OLI system. One methodological difficulty that these scores reveal is a potential ceiling effect in the quiz scores. This would suggest that there may be a restricted range in students' scores, making it difficult to show correlations between students' learning behaviors and their quiz/exam performance.

Table 2.2: Mean scores on primary assessments

Assessment	Score (out of 100)
Exam 1	73
Exam 2	66
Quiz 1	81
Quiz 2	92
Quiz 3	85
Quiz 4	86
Quiz 5	89

2.3.2 Biology: Learning/study-strategy results and relationships

The next set of measures we collected come from the automatically logged OLI data. These measures describe students' use of the OLI materials. Table 2.3 shows means of the five "OLI usage" measures discussed above. Looking at these descriptive statistics, it is noteworthy that students did not make much use of the lists of objectives for each module. In fact, the low number of times objectives were viewed on average is actually the result of the majority of students *never* accessing the objectives lists, and a very small minority of students viewing them multiple times for each module. This issue warrants some further consideration in terms of the role these objectives should play. It is also worth noting that the average number of self-assessments submitted by students is far lower than the number of self-assessments made available to students in the course of this study. In fact, on average, these numbers indicate that students were working through and submitting only about half of the possible self-assessments available in the OLI-Biology materials. Again, this is an issue worth investigating further.

Table 2.3: Averages for the five OLI usage measures

Measures of OLI use	Average Use
Number of self-assessments	7.2
Number of pages viewed	26.9
Number of times objectives viewed	3.2
Number of sessions initiated	13.2
Total time viewing pages	2158.1

Relating the OLI usage and performance metrics, the number of objectives viewed was a marginally significant predictor of exam performance. This suggests that there may be a study strategy of self-monitoring that differentiates students and that, not surprisingly, predicts better learning outcomes. Of the remaining measures of OLI system usage, total time was the only measure that showed a correlation with exam performance.¹ In particular, total time was significantly correlated with exam performance, but only when

¹ Note that for these analyses correlating total time with exam performance, we excluded students who did not use the OLI-Biology course at all (9 students altogether) as well as students whose total time was so small as to be nearly equivalent to not having used the OLI-Biology course (an additional 22 students). The students with such low measures for total time would not contribute meaningfully to an analysis of how much engagement with the OLI course predicts learning outcomes. Note that when these students are included in the analyses, the results are similar but simply weaker.

the period of usage was measured separately for the time spent before each exam. That is, students' OLI time-on-task up to week 4 predicted their exam scores for exam 1 (but not exam 2), and their time-on-task from weeks 5, 6, and 7 predicted their exam scores for exam 2 (but not exam 1). Still, this linear relationship is a rather a weak one as Figure 2.1 shows. Further investigation supported the relationship's nonlinear trend in that a log-linear regression of exam performance on total time (with quiz 1 as a covariate) showed a significant relationship for the corresponding exam to time-on-task period (i.e., exam 1 for weeks 2-4 and exam 2 for weeks 5-7) but not vice versa. This suggests that the more students engaged with the OLI materials, the greater their corresponding exam score. The lack of a relationship between OLI time-on-task and the non-corresponding exam scores suggests that this former positive relationship is not simply the result of better students showing greater engagement and higher scores, but rather it shows a content-specific relationship between what students spend time studying in the OLI-biology course and how well they perform on tests of that OLI material.

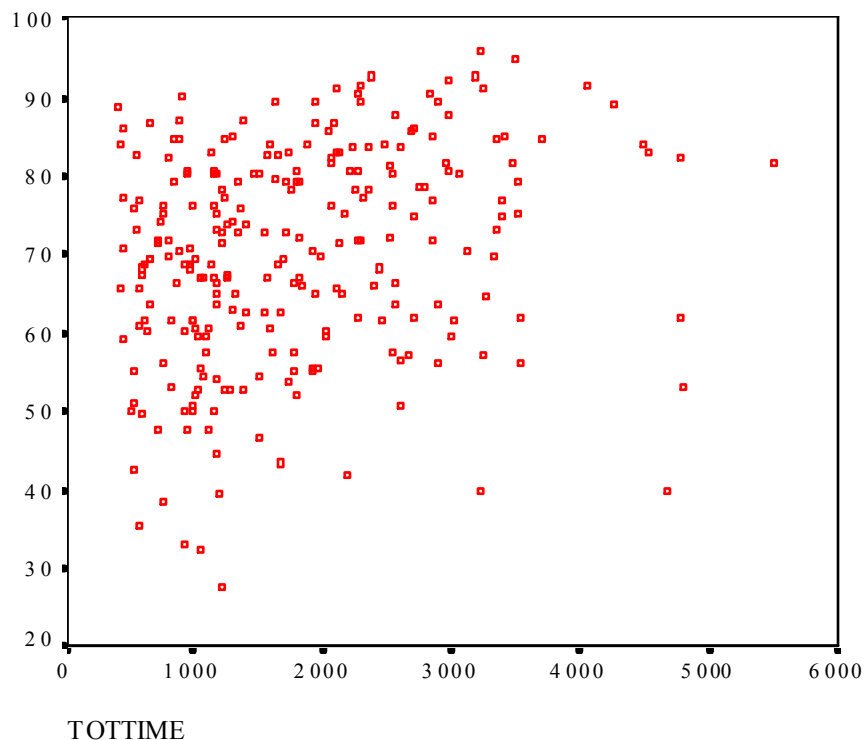


Figure 2.1: Total OLI time (TOTTIME spent before a given exam) as a predictor of the corresponding exam score

A more detailed analysis of students' interaction with the OLI-Biology materials was performed in the case of a particular interactive learning tool that appeared in weeks 3 and 4 of the course. This tool offered students practice working with functional groups, and it provided detailed logging of (a) the time students spent working with the tool, (b) the number of times students revisited it, and (c) the total number of exercises (also called "steps") that they worked through with it. What is striking about the second and third of these measures is that they showed a fairly wide distribution. So, it was possible to look for correlations with students' learning outcomes. As in the case of students' total time

predicting only the corresponding OLI-based exam scores, we would predict that our measure of students' engagement with this functional groups tool would be correlated with students' performance on any quiz related to functional groups but *not* correlated with their performance on other quizzes. Because we have students' scores on five high-stakes quizzes and only quiz 3 related to functional groups, we have an opportunity to test this prediction. Indeed, the correlations between students' quiz 3 performance and both "number of visits to the function groups tool" and "total number of exercises (steps) completed with the functional groups tool" were significant and positive ($r=.30$). At the same time, students' quiz 1, 2, 4, and 5 scores were not significantly or not as strongly correlated with these two measures of engagement with the functional groups tool (all r 's $< .16$). Specifically, the strength of the engagement-performance relationship was significantly stronger for quiz 3 than for the other quizzes. This set of results supports the notion that the more students worked with the functional groups tool, the better they performed on a functional groups quiz. The fact that their degree of engagement with the functional groups tool did not predict performance on other topic quizzes rules out the alternative explanation that such a correlation is simply the result of better students doing better overall.

The last category of quantitative results involves specific time-based measures of students' use of the OLI-Biology course. The total number of OLI sessions were broken down into three time windows relative to the exam date: the day before the exam, the preceding six days (i.e., a week before up to a day before the upcoming exam), and the preceding two weeks (i.e., the first two weeks of a new exam phase). If students' distribution of study sessions working with the OLI-Biology course were evenly distributed across time, we would expect these three time windows to present data in the ratios of 1:6:14 (for the number of days in each time window). As Table 2.4 shows, however, students' usage of the OLI-Biology course was not distributed in this way. In fact, students made almost as many visits to the OLI-Biology course in the day before an exam as compared to the preceding six days put together.

Table 2.4. Number of OLI sessions in different windows, relative to upcoming exam

Time window before exam	Average # of Sessions
Day before exam	1.6
"Week" before exam (not incl. day before)	2.3
All other times (preceding week before)	7.3

2.3.3 Also Statistics

For a similar analysis of students' learning and the effectiveness of the OLI course for Statistics, see Lovett, Meyer, and Thille (in press).

2.4 Model of study behaviors and learning

2.4.1 Model (qualitative and mathematical)

As a first step in modeling students' study behaviors and ultimate learning gains, we reviewed the log data to collect various measures of students' choices regarding what

material to study. For example, which activities do students choose to do (or not to do) and how many times do they choose to repeat a given activity? Similarly, do students have different profiles of behavior when it comes to choosing to review a previously viewed topic (e.g., re-visiting a given page, interactive activity, or assessment)? It is worth noting that, at some level, all of the OLI components are optional activities for students to complete. Students' progress through the OLI material was not being directly graded.

Table 2.3 from above shows that there was a wide range in students' usage of the different OLI components. Perhaps more importantly for current purposes, there was a range of usage *across students* for a given OLI component. For example, even though the average use of the "learning objectives" pages was low, this number reflects large differences in whether and how students used these pages: many students did not view these pages at all, several students viewed them infrequently, and a few students viewed them often. This was one metric that showed a relationship with learning outcomes in the course: the more students viewed the objectives pages, the better their learning. This is suggestive of a meta-cognitive advantage in which students who are more interested in or focused on the objectives of their own learning (i.e., what they are supposed to be able to do at the end of each OLI unit) are more effective at learning.

The other OLI usage measure that was shown above to relate to performance was the total time students spent using the OLI course. This predictive capacity of the time-on-task variable, however, appeared to be specific in that the relationship only held between study time before a given exam and performance on that exam. This result suggests that the more time students spent studying a particular set of OLI materials, the better they learned those particular materials.

But besides investigating students' usage of different components of the OLI course and their overall time using the course, we sought to investigate students' *patterns of use* and the *adaptivity of their use* of the OLI course. In particular, we were interested in any differences in the degree to which students chose to review material they had already covered, and if so, the degree to which students' choices to review (or not to review) were sensitive to their performance/progress at that point in time. One can image several possibilities (with differences in apparent adaptivity):

- Students who rarely or never review previous viewed pages and rather forge ahead through the material regardless of their situation.
- Students who go back to review all or most of the material, and do so in a way that does not depend on their learning situation.
- Students who are more likely to review a topic (or re-do an activity) when their current progress suggests they are performing below expectations, and who are more likely to forge ahead when their current progress looks strong.

Note that, according to these descriptions, only the third case mentioned involves students who are sensitive to current conditions and thus *adapting* to their own learning needs. This would be an example of self-regulated learning in that the students are

monitoring their own situation enough to be aware of a need to review and then actually go back to review relevant material.

Based on an analysis of a sample of students, the rough breakdown into these three cases was as follows:

- Approximately 20% of students simply do not review previous material (regardless of their current performance). Interestingly, this profile goes together with a pattern of skipping optional, low-stakes assessments.
- Approximately 20% of students go back to review much of the material, either doing so in a way that does not depend on their current situation or that follows the opposite pattern relative to their apparent needs (i.e., going to back to review material that was already well understood).
- Approximately 60% of students are more likely to review when they have not fully understood a topic and are less likely to review when they have understood a topic.

These three profiles need not be mapped onto individual students per se, but rather they may reflect different study strategies that students may have in their repertoire. In other words, in a basic ACT-R model of students' study strategy choices, we posit three different strategies (written below in English form):

ALWAYS-MOVE-ON:

When you have come to the end of a current topic, then move on to the next topic in the sequence.

ALWAYS-REVIEW:

When you have come to the end of a current topic, then go back to repeat an activity.
[Variations on this strategy might involve repeating a particular kind of activity or even repeating a particularly high-scoring topic.]

REVIEW-AS-NEEDED:

When you have come to the end of a topic and perceived performance is below expectations, then go back to repeat a relevant activity.

Note that this model posits an interesting strategy choice situation whenever the student comes to the end of a unit having performed below expectations. In this situation, all three of the above strategies are potentially applicable. In this situation, the model chooses the one to fire that has the highest estimated utility (where utility is a measure of the time-weighted reward associated with that strategy, see Introduction). The higher a production's utility, the more likely it will be selected. In our model, the ALWAYS-REVIEW strategy begins with a lower utility because it is necessarily more costly (in time) and less "rewarding" than its competitors. A straightforward application of the ACT-R utility-based choice mechanism (cf. Equation 1.7) can produce the distribution of study strategies we observe.

This leads to the question of where those strategies (i.e., productions) and their associated utilities come from. As mentioned in the Introduction, a utility value is updated based on ones' experience in the world after having applied the corresponding production, i.e., utility is increased as a function of the reward received upon applying the production. This implies that the more rewarding a production is, the higher its utility and hence the more likely it will be selected in the future. This leads to two predictions: (1) Across time, students should learn to prefer the REVIEW-AS-NEEDED strategy because it should lead to the best reward, and (2) Students with prior knowledge/experience in effective learning should begin with a bias toward the REVIEW-AS-NEEDED strategy.

Regarding the first prediction, we did find that students increased their tendency to review as needed across time in the OLI course. Specifically, students' tendency to review past material decreased across units in the cases where they were performing well and their tendency to review past material increased across units in the cases where they were performing poorly. These two trends combined led students to show review behaviors that were sensitive to performance: When students' performance was low (less than 2/3 of questions answered correctly), students chose to review 18% of the time, and when their performance was high (greater than 2/3 of questions answered correctly), students chose to review less than 10% of the time. Put another way, students' average performance on topics where they ultimately chose to do some review was 48% whereas the average performance on topics where they ultimately chose *not* to review was 67%.

Regarding the second prediction, none of our baseline or demographic/beliefs measures was a significant predictor of students' tendency to use the ALWAYS-MOVE-ON, ALWAYS-REVIEW, or REVIEW-AS-NEEDED approach. In the case of our measures of epistemological beliefs and study strategies, it may be that our sample had a somewhat restricted range that impaired the chance of finding such a relationship. So, in our modeling, we simply started different models (representing different potential students) with different initial values for the utilities of the three productions and let utility-value learning carry on from there.

These explorations of models for choosing to review or not to review revealed an interesting "absorbing state" in some cases that appeared to mimic a profile shown by some students as well. This is the situation where the model has a high-utility for the ALWAYS-MOVE-ON production such that the other two competitors do not get a chance to fire (and hence to learn to be preferred based on experience). In students, this is the profile in which students do not go back to review *even when they arguably should* (e.g., they have performed poorly/learned little from a given unit in the course). How can the model get absorbed into the state of always moving on? Given that the utility learning mechanism of ACT-R is updating utility values as a function of experienced reward, and given that the experienced reward is a measure of net gain (i.e., goal/value achieved minus time spent achieving it), one can see that the ALWAYS-MOVE-ON production has a time-cost advantage over its competitors. In other words, it's always faster to move on. In the case of students' evaluations of reward for the purposes of strategy choice, it seems that students value time over accuracy (Lovett & Chang, 2007). So, we may have found in this OLI learning situation an example of students selecting a low-cost strategy

even when a more accurate/rewarding one exists. This could occur among students if they come in to the OLI course without any representation of the ALWAYS-REVIEW or REVIEW-AS-NEEDED productions or with such low utilities for these approaches that they never select a review strategy. If this should happen, no matter how effective the system's utility-learning mechanism is in principle, our model cannot discover a review-based strategy. This suggests that a key opportunity for efficient training – a way to tap into the natural adaptive, utility-learning mechanism in ACT-R – is to create an environment that encourages students to at least try a strategy that involves review. Moreover, whenever students do review and then show success in reaching a goal (or when students choose to move on and then show poor performance), the value of the goal (big or small for the amount of time spent) should be highlighted.

2.5 Implications for training

Based on the empirical work conducted in this part of the project, there are several implications for training and feedback that relate to fostering students' effective strategy choices in learning from online courses. While these implications would need to be tested in multiple settings to be established as general, we posit them here based on the fact that (a) they are consistent with the empirical results of the current work and (b) they are based on predictions derived from the ACT-R architecture and, as such, are consistent with a much broader set of research on learning and performance. The first two implications involve strategies for increasing students' tendency to review relevant material (rather than always moving on to the next unit). And the third implication involves giving students ample practice with feedback so that the learning strategies they choose to apply can be refined across time.

First, create the online learning environment in such a way that students are encouraged to try a strategy that involves reviewing previously viewed material appropriately. This guideline comes from the result that a nontrivial proportion of students never (or very rarely) spent any time reviewing past material. And yet, given the complex material involved in the OLI courses, it almost goes without saying that every student could benefit from *some* review of a particular piece of the course. Our modeling results show that, if an effective learning strategy such as “reviewing past material when needed” is never attempted in the first place, this strategy will never get a chance to become more prominent because greater use of a strategy requires an increase in that production's relative utility, which in turn requires some experience at applying the production.

Strategies that might foster students' review of past material include raising students' awareness about the viability of going back to a previous section. In other words, one hypothesis for why students didn't review past material (even when they could have benefited from doing so) is that they did not represent a “review-based” strategy in their repertoire. So making this option explicitly available to students could increase their chances of trying the strategy (and then the natural learning mechanisms for promoting that strategy could unfold). For example, in one of the OLI courses, for a subset of the topics, an explicit choice point was added after the unit's low-stakes assessment. In this way, after students received their feedback on the assessment, they would be prompted to answer the question: “Did I get this yet? If yes, continue on; if not yet, click here to

review this unit.” Although we did not conduct an official experiment to test the effectiveness of this intervention, preliminary analyses suggest that it did increase students’ tendency to review in that they clicked the review button more than half of the times they encountered it (and yet only chose to review past material less than a quarter of the time without the button). Other strategies to encourage students’ review might involve posing a question that (implicitly or explicitly) directs a majority of students back to a previous piece of the course in order to find an answer. This strategy could serve to increase the utility of “going back” in general if students found that doing so was a quick and easy way to find a solution.

The second implication for training involves setting up the utility structure of the learning environment so that applying an effective review strategy will in fact lead students to better outcomes (i.e., faster and/or better learning). This might involve designing the assessments that are administered outside the learning environment in a way that taps students’ deeper conceptual understanding of the material (i.e., a level of learning that would likely require multiple passes through the information or several practice opportunities). In contrast, if students are able to perform well on quizzes and exams after only having skimmed through the learning materials once, then it is a signal that students are actually following what would be predicted according to a rational choice model. In other words, the goal of optimal learning is not to encourage students to review material in *all* contexts but according to their needs and the benefits that can be gained through the effort of extra review. Finally, when students do take action to review past material and show improved performance (or conversely when they do not review past material and show poor performance) these implicit payoffs of their actions can be highlighted in a way that might inform subsequent actions. For example, in an online learning environment, student data are continuously collected and can be “replayed” to students as evidence for effective strategies that actually work.

The third implication for training and feedback involves making sure that students have sufficient practice at applying the strategies you want to promote. In the case of the current work, we identified appropriate review strategies as the apparent gap in students’ skill set. So the strategy here would be to give students ample practice at reviewing relevant past material. In particular, at the end of each unit (or even sub-unit pieces), students could be encouraged to review past material as appropriate. Moreover, students’ actions could be tracked as they work in the online learning environment and then, when they do go back to previous material, an explicit piece of feedback could be offered. Although giving students practice and feedback on metacognitive skills is difficult – because the actions of applying effective strategies usually occur at a level above the content being taught – such instructional interventions could be especially helpful to the degree that they reify what is an abstract piece of effective learning.

3 Instance-based modeling of UAV maneuvers

3.1 UAV task

The goal of this task is to create models of basic aircraft maneuvering using the ACT-R cognitive architecture in order to explore implications for teaching. ACT-R is a computational theory of human performance that incorporates procedural (rule-based) knowledge and declarative (fact-based) knowledge. In this task we use data collected from expert pilots to provide instances of declarative knowledge that indicate an appropriate action to take given a particular circumstance.

3.1.1 Synthetic Task Environment

The Predator UAV Synthetic Task Environment (STE) is a realistic simulation of the flight dynamics of the Predator RQ-1A System 4 UAV with built in tasks and data collection capabilities. The core aerodynamics model of the UAV STE is used in the training system for Air Force Predator pilots at Indian Springs Air Force Auxiliary Field in Nevada. The UAV STE is essentially a scaled down version (hardware wise) of the training system. The three tasks built on top of the core aerodynamics model include: the *Basic Maneuvering Task*, in which a pilot must make very precise, constant-rate changes in airspeed, altitude and/or heading; the *Landing Task* in which the UAV must be guided through a standard approach and landing; and the *Reconnaissance Task* in which the goal is to obtain simulated video of a ground target through a small break in cloud cover. For each task, there are multiple scenarios which manipulate various performance requirements (e.g. turn right, turn left and climb) and external conditions (e.g. wind, no fly zones). During performance of a task, the values of approximately 100 different aircraft and human performance variables are recorded every 200 msec. The design of these synthetic tasks is the result of a unique collaboration between behavioral scientists and expert pilots of the UAV. The aim in developing the tasks was to identify important aspects of the UAV pilot's overall task—aspects that tax the key cognitive and psychomotor skills required of a UAV pilot. They are tasks that lend themselves to laboratory study, yet do not fall prey to oversimplifications. Tests using military and civilian pilots showed that experienced UAV pilots perform better in the STE than pilots who are highly experienced in other aircraft but have no UAV experience, indicating that the STE is realistic enough to tap UAV-specific pilot skill.



Figure 3.1: The UAV Synthetic Task Environment (STE)

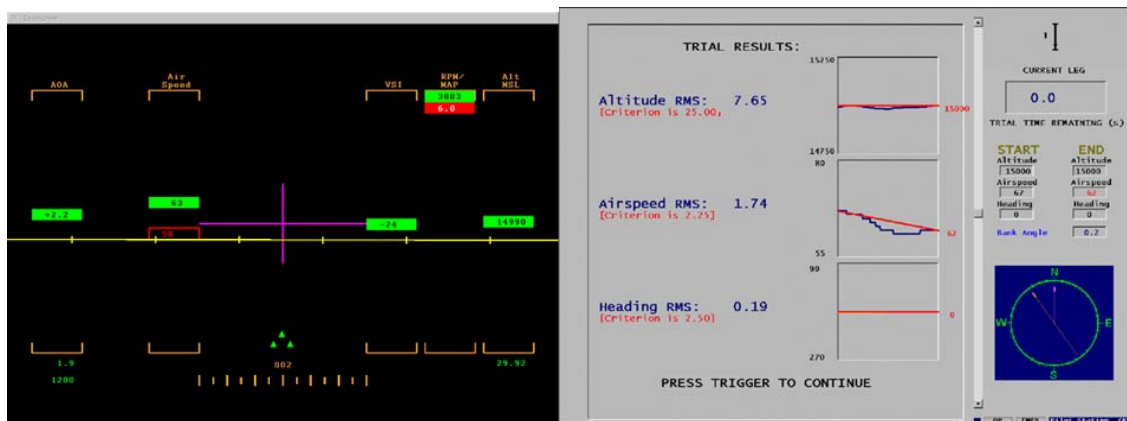


Figure 3.2: Detailed view of the STE screens.

3.1.2 ACT-R representation of environment

Computational cognitive models “see” their visual environment by moving visual attention around within a digital representation of that environment. This is fairly trivial with simple, static tasks that are implemented in the same software language as the cognitive model, but it is more complicated when the architecture must interface with an external simulation. We took advantage of the work done by Gluck et al. (2003) to create an ACT-R 5.0 model of basic aircraft maneuvering that could interface with the Predator STE. Their approach in interfacing models to the STE was to re-implement the visual displays of the STE in Lisp, the programming language in which ACT-R is written. The focus of the reimplementation was on matching the information provided by the visual display without necessarily reverse engineering the full graphics display of the STE. This was facilitated by the use of digital readouts for the flight instruments (other than the horizon line and reticle) in the STE, such that the model was not required to process an analog device in order to determine the value of the flight instrument. In the case of the horizon line and reticle, ACT-R returns a digital value for pitch and bank to the model (as reflected in the orientation of the horizon line with respect to the reticle), even though a graphic depiction of the horizon line and reticle is displayed. Other than the visual

displays, the Predator STE provides a Variable Information Table (VIT) data structure that contains data on most of the flight parameters of the UAV.

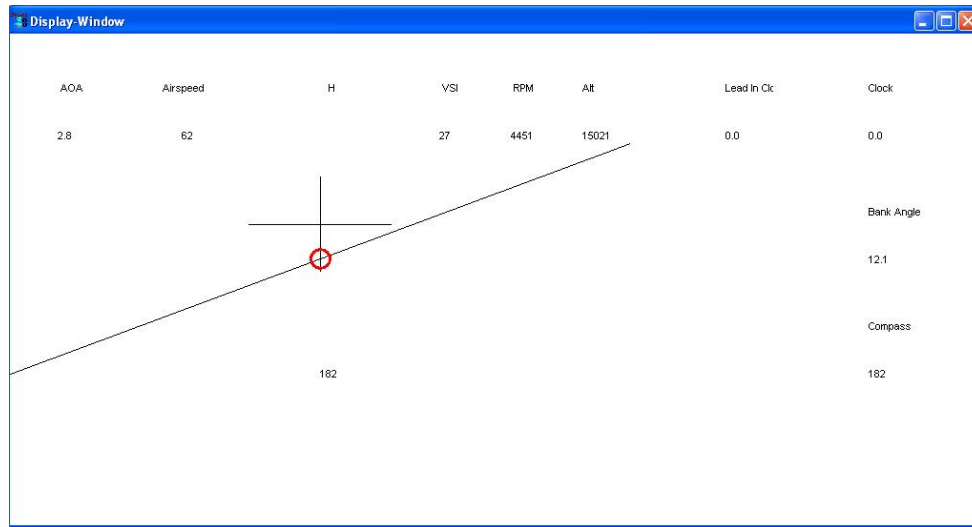


Figure 3.3: Lisp-based visual display of STE used by the ACT-R model.

3.1.3 Basic maneuvers

For a Predator pilot, the knowledge and skills necessary to effectively maneuver are essential to success. A natural place to begin a research program aimed at developing a fine-grained cognitive process model of a Predator pilot/teammate is the basic maneuvering task. This task was inspired by an instrument flight task originally designed by Wickens and colleagues at the University of Illinois at Urbana-Champaign (Bellenkes, Wickens, & Kramer, 1997). The task requires the pilot to fly a number of distinct instrument flight maneuvers. Preceding each maneuver is a 10 second lead-in during which time the pilot is asked to stabilize the aircraft in straight and level flight. Following the lead-in is a timed maneuver of 60 seconds during which time the pilot maneuvers the aircraft by making constant rate changes to altitude, airspeed, and/or heading, depending on the maneuver, as specified in Table 3.1.

Table 3.1: Goals of the basic flight maneuvers.

Maneuver	Airspeed	Heading	Altitude
1	Decrease 67–62 knots	maintain 0°	maintain 15,000 feet
2	maintain 62 knots	Turn Right 0-180°	maintain 15,000 feet
3	maintain 62 knots	maintain 180°	Increase 15,000-15,200 feet

3.2 Instance data of expert performance

3.2.1 Variables

The STE collects data five times per second for 60 variables, including heading, altitude, airspeed, bank angle, pitch, and RPM. Participants performed basic maneuvers over a number of trials, allowing some aspects of performance to stabilize. For example, Figure 3.4 shows that one participant converged on a bank angle of 13 degrees for maneuver 2 after 19 trials.

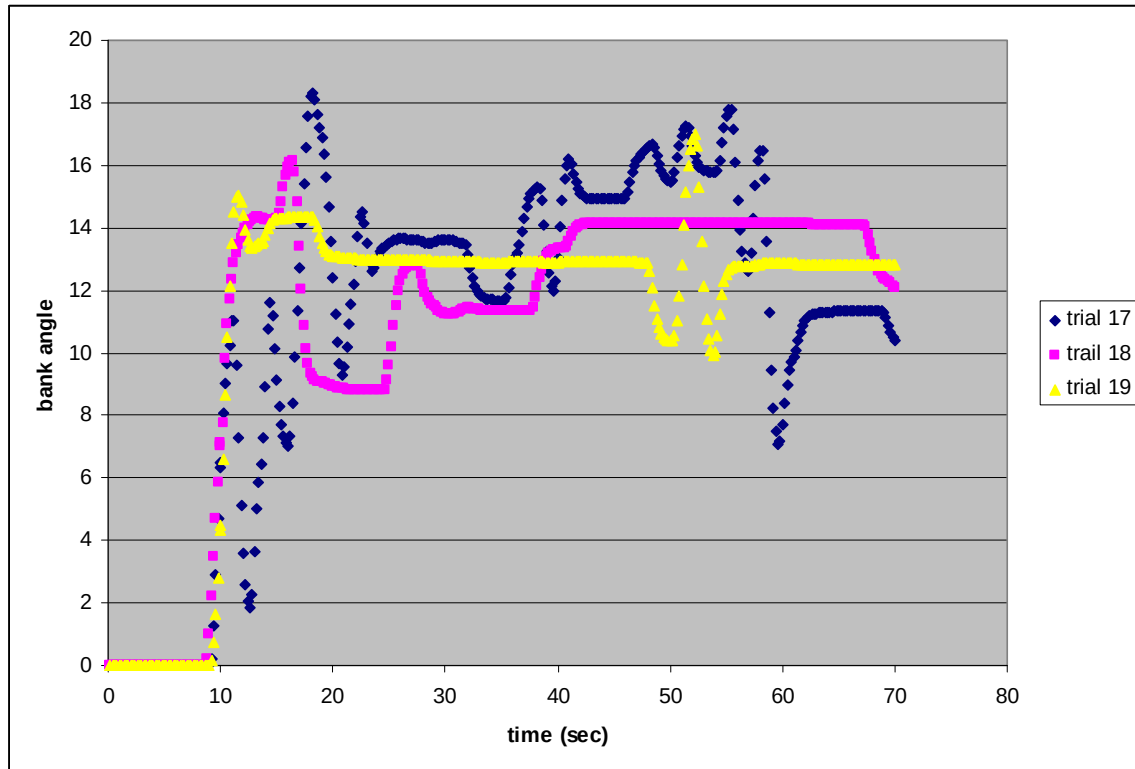


Figure 3.4: Three trials of bank angle as a function of time.

3.2.2 Preprocessing of training data

Every trial contains $5 \times 70 = 350$ instances of performance parameters and control settings. Using all of these instances would not be a realistic training task for humans, so an automated procedure was developed to preprocess the data and limit the number of training instances.

In each maneuver the goal is to change one performance parameter (airspeed, heading, or altitude) while keeping the others constant. The pilots take actions that change control settings (RPM, bank angle, pitch) which then change performance parameters. For performance parameters that were intended to change, the change usually occurred in a continuous manner. For performance parameters that were intended to stay constant, they usually deviated from the initial value then returned as pilots made corrections. Instances where performance parameters started to return to their original value after a

deviation were chosen as examples where pilot saw a deviation and effected a desired change. Control settings for changing performance parameters to goals usually had deviations that returned to stabilized constant values. Instances where control setting started to return to stable values after a deviation where chosen as examples where pilots attempted to maintain a desired control setting. Control settings for maintaining constant performance parameters usually had a high variability. These features can be seen in Figure 3.5.

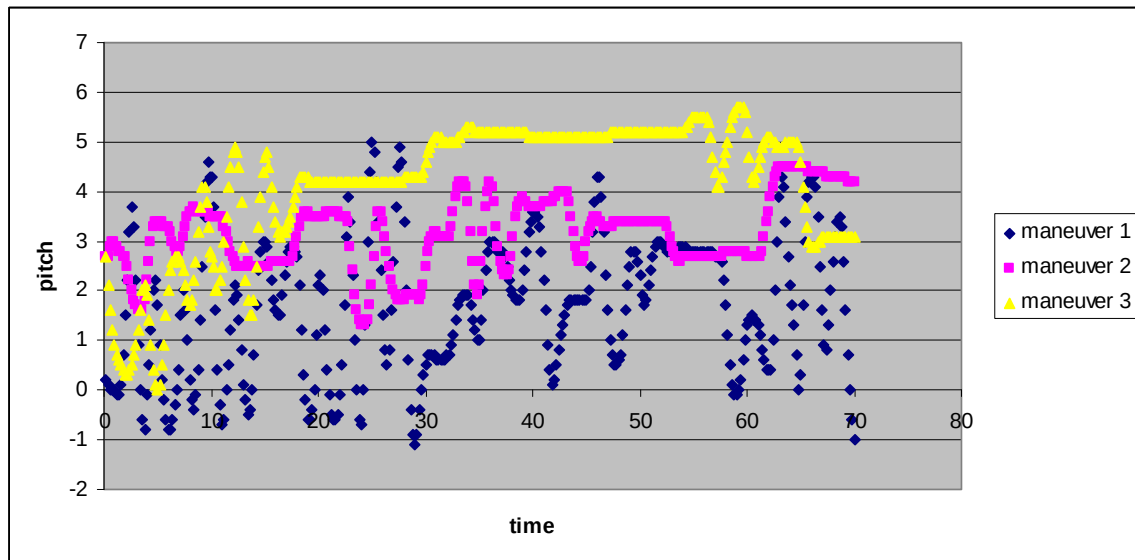


Figure 3.5: Pitch values over time for different maneuvers.

3.3 Model creation

The model of basic aircraft maneuvering used in this study is based on an instrument flight strategy called the “Control and Performance Concept” (Air Force Manual on Instrument Flight, 2000). This aircraft control process involves first establishing appropriate control settings (pitch, bank, power) for the desired aircraft performance, and then crosschecking the instruments to determine whether the desired performance is actually being achieved.

According to the Air Force Manual on Instrument Flight, a key to expert flight performance is knowledge of the appropriate control settings needed to obtain desired flight performance. For example, a pitch of 3 degrees and an engine RPM of 4300 will maintain straight and level flight of the UAV at 67 knots over a range of altitudes and external conditions. The expert pilot need only set the appropriate pitch and engine RPM to obtain the desired performance, subject to monitoring and adjustment based on variable flight conditions like wind and air pressure.

The focus of this project is to create models that use appropriate control settings based on data from expert performance. This allows the easy creation of models by supplying expert data instead of using that data to write ACT-R code. In theory, to obtain level

Table 3.2: Control setting and action used to change performance parameter.

Performance Parameter	Control Setting	Control Action
Airspeed	RPM	Throttle (front, back)
Heading	Bank Angle	Stick (left, right)
Altitude	Pitch	Stick (front, back)

flight the pilot should use power to get the desired airspeed, and use pitch used to maintain altitude (Jeppesen Instrument/Commercial Manual, 1998). The model perceives a current performance parameter (airspeed, heading, or altitude), retrieves a desired control setting (RPM, bank angle, or pitch) based on that parameter, perceives the current control setting, and takes appropriate action to obtain the desired control setting. Table 3.2 shows the control setting and control action used to change performance parameters. Figure 3.6 shows the information used to determine control response.

3.3.1 Perceiving the environment

The model uses procedural knowledge developed by Gluck et al. (2003) to focus the visual buffer of ACT-R on instrument display values in the Lisp representation of the STE environment. These display values include the lead-in clock time, performance parameters (airspeed, heading, altitude), and control settings (RPM, bank angle, pitch). Once the lead-in clock time reaches zero, it is no longer attended. As mentioned in section 3.1.2, perception was facilitated by the use of digital readouts for the flight instruments (other than the horizon line and reticle) in the STE, such that the model was not required to process an analog device in order to determine the value of the flight instrument. In the case of the horizon line and reticle, ACT-R returns a digital value for pitch and bank to the model (as reflected in the orientation of the horizon line with respect to the reticle), even though a graphic depiction of the horizon line and reticle is displayed.

3.3.2 Instance-based decisions

Instances from expert data are retrieved by matching the current performance parameter deviation to the deviation stored in the instance. In a dynamic environment such as flight, there may not be an instance that exactly matches current conditions. The ACT-R theory provides a way to retrieve the nearest instance with partial matching. As was mentioned earlier, with partial matching, the instance with the highest similarity has the highest activation and is more likely to be retrieved (cf. Equation 1.5). The model uses the same parameters (activation noise=0.25, mismatch penalty=1.5) and ratio similarity measure for partial matching that were used by Lebiere (1998) for instance-based learning of arithmetic facts.

Current Param +	Instance +	Current Setting +	Deviation =	Response
Curr-airspeed: 63	Stim-param: Airspeed	Curr-rpm: 4435	Parameter: Rpm	Isa: Move-throttle
Desired-airspeed: 62	Stim-deviation: 2	Desired-rpm: 4362	Deviation: 70	Direction: High
Airspeed-deviation: 1	Resp-param: Rpm		Direction: High	Magnitude: Small
	Resp-value: 4362		Magnitude: Small	R: 0.5
				Theta: 3.14

Figure 3.6: Information used to determine control response.

3.3.3 Vehicle control

The model uses procedural and declarative knowledge developed by Gluck et al. (2003) to map control setting value goals onto control actions for the stick and throttle that are sent to the STE. The declarative knowledge maps deviations in current control settings from desired control settings to response direction and magnitude. The procedural knowledge maps the direction and magnitude to r and theta values for the throttle and stick. Figure 3.7 shows how perception, decision, and control come together in a partial trace of model performance. In the trace, the model is executing maneuver 2 and using the current heading to determine the appropriate control action.

```

Time 12.460: Sync Fired
Time 12.510: Retrieve-Crosscheck-Intent Fired
Time 12.510: Ci-Hlr-Pitch-To-Heading Retrieved
Time 12.545: Module :VISION running command ENCODING-COMplete
Time 12.545: Vision sees HLR710
Time 12.560: Crosscheck-Intent-Set-Context Fired
Time 12.610: Crosscheck-Retrieve-Instrument Fired
Time 12.610: Heading-Indicator Retrieved
Time 12.660: Crosscheck-Find-Instrument Fired
Time 12.660: Module :VISION running command FIND-LOCATION
Time 12.660: Vision found LOC614
Time 12.710: Crosscheck-Attend-Instrument Fired
Time 12.710: Module :VISION running command MOVE-ATTENTION
Time 12.795: Module :VISION running command ENCODING-COMplete
Time 12.795: Vision sees TEXT713
Time 12.845: Crosscheck-Encode-Instrument-Not-Hlr-Stim Fired
Time 12.845: Instance-Heading-Bank-3 Retrieved
Time 12.895: Retrieve-Param-From-Instance Fired
Time 12.945: Crosscheck-Retrieve-Instrument Fired
Time 12.945: Bank-Angle-Indicator Retrieved
Time 12.995: Crosscheck-Find-Instrument Fired
Time 12.995: Module :VISION running command FIND-LOCATION
Time 12.995: Vision found LOC649
Time 13.045: Crosscheck-Attend-Instrument Fired
Time 13.045: Module :VISION running command MOVE-ATTENTION
Time 13.130: Module :VISION running command ENCODING-COMplete
Time 13.130: Vision sees TEXT712
Time 13.180: Crosscheck-Encode-Instrument-Not-Hlr-Resp Fired
Time 13.230: Desired-Bank-Angle-From-Instance Fired
Time 13.280: Crosscheck-Set-Deviation-Bank-Angle Fired
Time 13.280: Bank-Angle-Dev--2 Retrieved
Time 13.330: Crosscheck-Retrieve-Deviation-Bank-Angle Fired
Time 13.380: Assess-Bank-Angle-Left-Small Fired
Time 13.380: Module :MOTOR running command MOVE-STICK
Time 13.430: Module :MOTOR running command PREPARATION-COMplete
Time 13.430: Done-Assess-Bank-Angle Fired

```

Figure 3.7: Partial model trace for determining control action based on current heading.

3.4 Model performance

3.4.1 Comparison to optimal performance

Optimal performance was considered to be a constant change in performance parameter from initial conditions to goal conditions from the end of the lead-in clock at 10 seconds to the end of the trial at 70 seconds. Formulae for optimal values are found in Table 3.3. RMS deviation from optimal performance was calculated in the following manner:

1. For each sample, loop through the UAV states array and take the difference between `uav_state_actual` and `uav_state_desired` and square it.
 - o For heading, the values are adjusted so that they will always be between 180 and -180 degrees.
2. Sum all the samples collected in step 1.
3. Total RMS = square-root (total_sigma_of_all_samples / total_number_of_samples)

Participants were considered to pass the maneuver if the RMS deviation of performance parameters was less than the criterion value found in Table 3.3.

Table 3.3: Optimal value formulae and RMS criterion for basic flight maneuvers.

Maneuver	Changed Parameter	Initial Value	Goal Value	Optimal Value	RMS Criterion
1	Airspeed	67	62	$-1/12 \cdot \text{time} + 67 + 5/6$	1.75
2	Heading	0	180	$3 \cdot \text{time} - 30$	7.50
3	Altitude	15000	15200	$10/3 \cdot \text{time} + 14966 + 2/3$	25.00

An initial model that retrieved instances of RPM values based on current airspeed was able to pass maneuvers two and three but not one. This is because the main control settings for maneuvers two and three (bank and pitch) indicate the change of desired direction (bank right and positive pitch) and therefore performance parameters (heading and altitude) are changed correctly. The main control setting for maneuver one (RPM) does not indicate a change in desired direction and therefore the performance parameter (airspeed) does not change.

An improved model was created that retrieved change in RPM instead of RPM. The model then calculated a desired RPM based on this change. This model was able to pass all three maneuvers. Figure 3.8 shows the RMS deviation for the model for maneuver 1 compared to the last three trials of the best pilot and the criterion value for passing. Figures 3.9 and 3.10 show the same comparisons for maneuvers 2 and 3. Note that due to the scheduling of maneuvers for pilots, the last three trials are not the same for the different maneuvers.

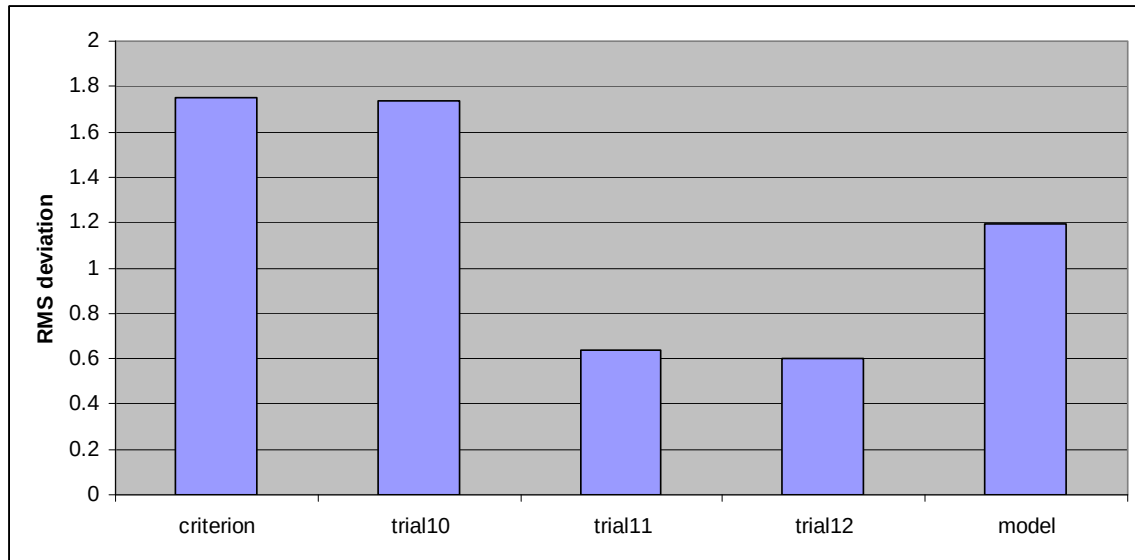


Figure 3.8: Airspeed deviation in maneuver 1.

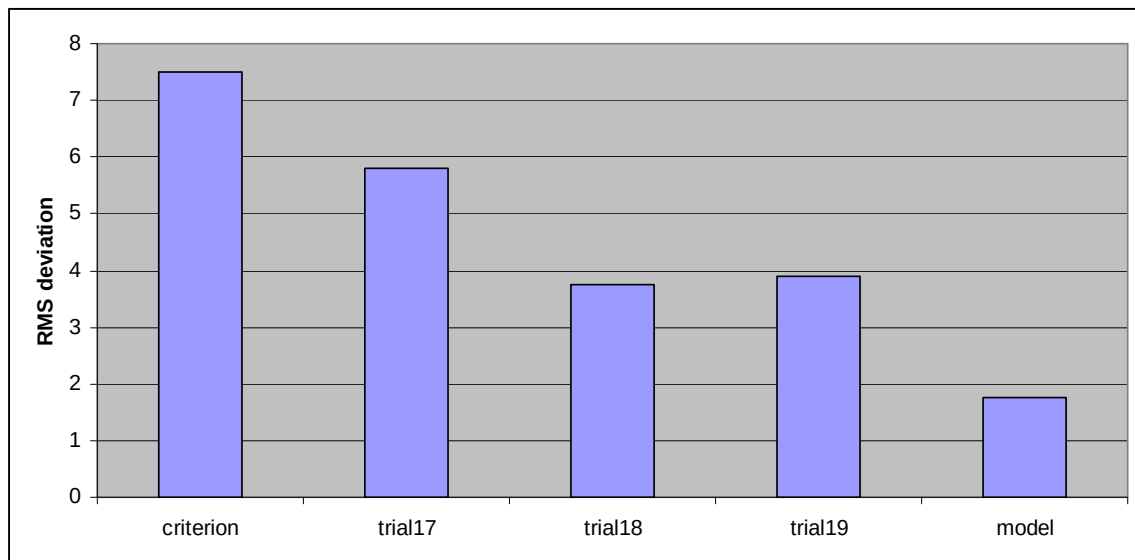


Figure 3.9: Heading deviation in maneuver 2.

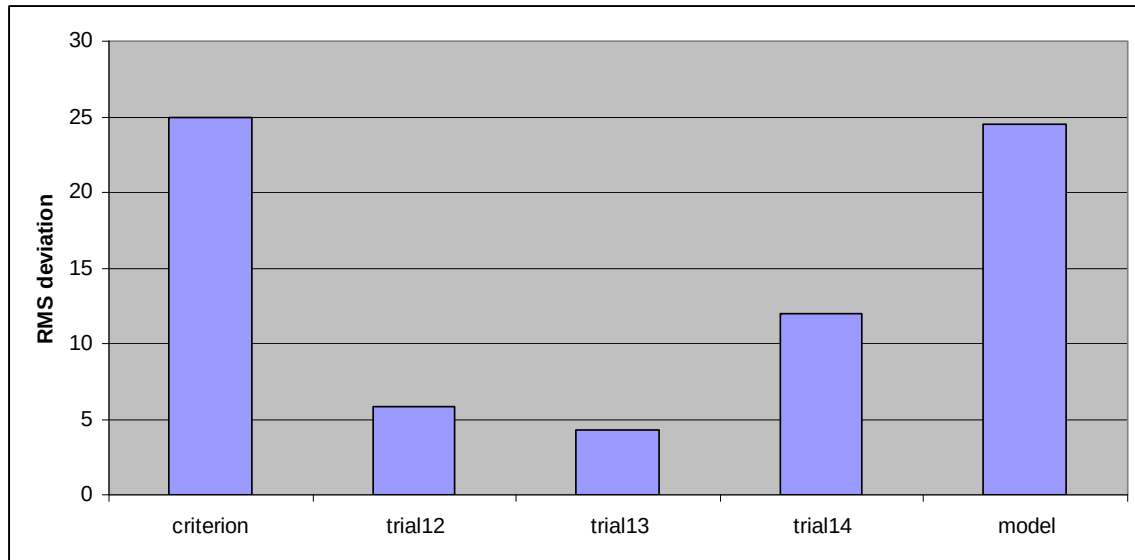


Figure 3.10: Altitude deviation in maneuver 3.

3.4.2 Comparison to expert performance

In addition to being compared to optimal performance, the model performance can also be compared to expert performance. Looking at the last trial of the best pilot, results from maneuver 2 demonstrate how the model can produce similar performance while drawing on a subset of the training instances. Figure 3.11 shows the performance heading deviation as a function of time for the model and pilot. Figure 3.12 shows bank angle as a function of heading deviation for the model and pilot, with instances plotted in yellow. One reason that the model produces performance with a smaller RMS deviation than the expert is that feedback delays result in an averaging of the instances of extreme bank angle values.

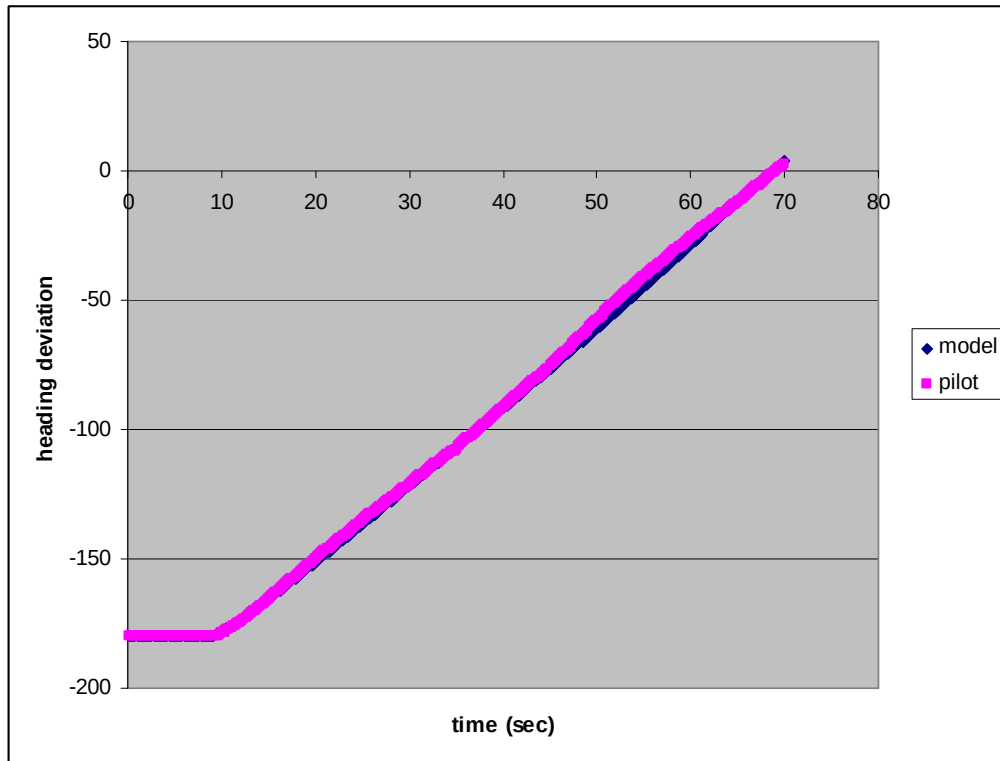


Figure 3.11: Heading deviation as a function of time for the model and pilot.

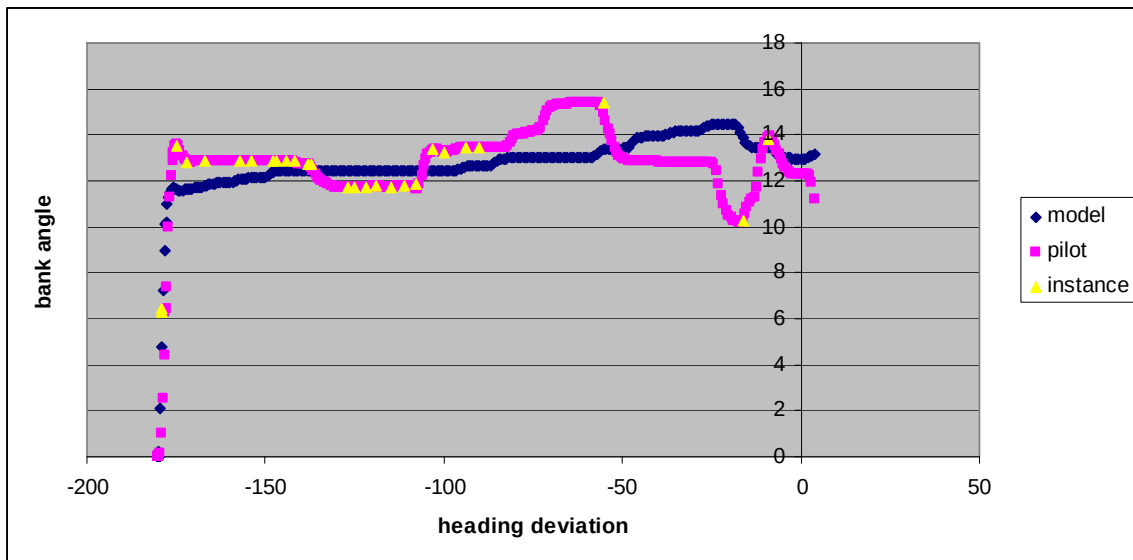


Figure 3.12: Bank angle as a function of heading deviation for the model and pilot.

3.5 Possible Improvements

Although the model passed all three basic maneuvers, the deviation from optimal performance for model performance was sometimes greater than that of pilot performance. This could be due to poor instances chosen by the preprocessing procedure and the model's lack of representation of global time. Since the preprocessing procedure focuses on a return of deviation to the norm, constantly increasing or decreasing control settings may not be noted. Adding a representation for global time could provide a way to add instances of particular control settings for particular times.

3.6 Implications for training

The model that successfully learned basic flight maneuvers uses instance-based examples of expert performance. For a human to acquire the same information, the examples would have to either be learned from unstructured observation or from structured training. In a complex dynamic task such as flight, there are too many parameters to memorize all of them at any particular instant, and it is unlikely that a trainee would notice critical parameter combinations by chance. Therefore it is important to explicitly train what information is needed to make control decisions. We have also found that time is a critical factor in dynamic tasks. Instance-based training that does not include a representation for global time needs to incorporate variables that indicate rate and direction of desired change. Again, these variables need to be explicitly visible to the student. Once training instances are made explicit, the learning systems of the model and human can be used to make informed decisions given a particular context.

4 Conclusions

In modeling both the biology and flight maneuver domains, it was found that information needed for good performance is at some level available to the trainee but might not be used. In the biology domain, the option to re-visit a previous topic is implicitly available. In the flight maneuver domain the rate of change information is indirectly available. The key insight for training is to make explicit to the student these aspects of the environment/representation so that the natural learning mechanisms can unfold in more productive ways. This relates to the idea of optimal training because our goal is to take best advantage of the human learning system. Essentially, the path to optimal training in both these cases involves finding the key domain feature to which learning progress is very sensitive. Based on our results, we would posit that explicitly training on these key features would promote more efficient learning. This position is in line with results such as Klahr and Nigam (2004), which show that direct instruction is more effective than discovery learning.

5 References

- Air Force Manual 11-217, "Instrument Flight Procedures," Vol. 1, 2000
- Anderson, J. R., Bothell, D., Byrne, M. D., Douglass, S., Lebiere, C., & Qin, Y. (2004). An integrated theory of the mind. *Psychological Review* 111, (4). 1036-1060.
- Bellenkes, A. H., Wickens, C. D., and Kramer, A. F., (1997) "Visual scanning and pilot expertise: The role of attentional flexibility and mental model development," *Aviation, Space, and Environmental Medicine*, Vol. 68, No. 7, 1997, pp. 569-579.
- delMas, R., Ooms, A., Garfield, J., & Chance, B. (2006). Assessing students' statistical reasoning. In *Proceedings of the Seventh International Conference on the Teaching of Statistics*. Salvador, Brazil.
- Garcia, T., & Pintrich, P. R. (1996). Assessing students' motivation and learning strategies in the classroom context. In M. Birenbaum, & F. J. R. Dochy (Eds.) *Alternatives in assessments of achievement, learning processes and prior knowledge* (pp. 319-339). New York: Kluwer.
- Gluck, K. A., Ball, J. T., Krusmark, M. A., Rodgers, S. M., & Purtee, M. D. (2003). A computational process model of basic aircraft maneuvering. In F. Detje, D. Doerner, & H. Schaub (Eds.), *In Proceedings of the Fifth International Conference on Cognitive Modeling* (pp. 117-122). Bamberg, Germany.
- Instrument Commercial Manual. (1998). Englewood, CO: Jeppesen Sanderson Inc.
- Klahr, D. & Nigam, M. (2004) The equivalence of learning paths in early science instruction: effects of direct instruction and discovery learning. *Psychological Science*, 15(10).
- Koedinger, K. R., & MacLaren, B. A. (1997). Implicit strategies and errors in an improved model for early algebra problem solving. In *Proceedings of the Nineteenth Annual Conference of the Cognitive Science Society*. Hillsdale, NJ: Erlbaum.
- Lebiere, C. (1998). The dynamics of cognition: An ACT-R model of cognitive arithmetic. Ph.D. Dissertation. CMU Computer Science Dept Technical Report CMU-CS-98-186. Pittsburgh, PA.
- Lovett, M. C., & Chang, N. C. (2007). Data-analysis skills: What and how are students learning? In M. Lovett, & P. Shah (Eds.) *Thinking with Data*. Mahwah, NJ: Erlbaum.
- Lovett, M. C., Meyer, O., & Thille, C. (in press). Open Learning Initiative: Testing an accelerated learning hypothesis in Statistics. *Journal of Interactive Media Education*.
- Schommer, M. 1990, Effects of beliefs about the nature of knowledge on comprehension, *Journal of Educational Psychology*, 82, 498-504.
- Wagner, A. R., & Rescorla, R. A. (1972). Inhibition in Pavlovian conditioning: Application of a theory. In Boakes & Halliday (Eds.) *Inhibition and learning*. London: Academic Press.