

AD-A240 775



2

LEARNING PHYSICS
VIA
EXPLANATION-BASED LEARNING OF CORRECTNESS
AND ANALOGICAL SEARCH CONTROL.

Technical Report AIP - 137

Kurt VanLehn and Randolph M. Jones

**The Artificial Intelligence
and Psychology Project**

Departments of
Computer Science and Psychology
Carnegie Mellon University

Learning Research and Development Center
University of Pittsburgh

9 12 1988

Approved for public release; distribution unlimited.

91-10977



**LEARNING PHYSICS
VIA
EXPLANATION-BASED LEARNING OF CORRECTNESS
AND ANALOGICAL SEARCH CONTROL.**

Technical Report AIP - 137

Kurt VanLehn and Randolph M. Jones

Learning Research and Development Center
and Computer Science Department
University of Pittsburgh
Pittsburgh, PA

September 1991

This research was supported by the Computer Sciences Division, Office of Naval Research, under Contract Number N00014-86-K-0678. Reproduction in whole or in part is permitted for purposes of the United States Government. Approved for public release; distribution unlimited.

In L. Birnbaum and G. Collins (Eds.) Machine Learning: Proceedings of the Eighth International Workshop. San Mateo, CA: Morgan Kaufman. 110-114.

Classification: ☒ CONFIDENTIAL
Distribution: ☐ UNCLASSIFIED
Availability: ☐ UNCLASSIFIED
Availability: ☐ UNCLASSIFIED
A-1

Unclassified

SECURITY CLASSIFICATION OF THIS PAGE

REPORT DOCUMENTATION PAGE

1a. REPORT SECURITY CLASSIFICATION Unclassified		1b. RESTRICTIVE MARKINGS	
2a. SECURITY CLASSIFICATION AUTHORITY		3. DISTRIBUTION / AVAILABILITY OF REPORT Approved for public release; distribution unlimited.	
2b. DECLASSIFICATION / DOWNGRADING SCHEDULE			
4. PERFORMING ORGANIZATION REPORT NUMBER(S) AIP-137		5. MONITORING ORGANIZATION REPORT NUMBER(S)	
6a. NAME OF PERFORMING ORGANIZATION Carnegie Mellon University	6b. OFFICE SYMBOL (If applicable)	7a. NAME OF MONITORING ORGANIZATION Computer Sciences Division Office of Naval Research (Code 1133)	
6c. ADDRESS (City, State, and ZIP Code) Department of Psychology Pittsburgh, PA 15213		7b. ADDRESS (City, State, and ZIP Code) 800 North Quincy Street Arlington, VA 22217-5000	
8a. NAME OF FUNDING / SPONSORING ORGANIZATION Same as Monitoring Organization	8b. OFFICE SYMBOL (If applicable)	9. PROCUREMENT INSTRUMENT IDENTIFICATION NUMBER N00014-86-K-0678	
8c. ADDRESS (City, State, and ZIP Code)		10. SOURCE OF FUNDING NUMBERS	
		PROGRAM ELEMENT NO. N/A	PROJECT NO. N/A
		TASK NO. N/A	WORK UNIT ACCESSION NO. N/A
11. TITLE (Include Security Classification) Learning physics via explanation-based learning of correctness and analogical search			
12. PERSONAL AUTHOR(S) VanLehn, K. A. and Jones, R. control.			
13a. TYPE OF REPORT	13b. TIME COVERED FROM _____ TO _____	14. DATE OF REPORT (Year, Month, Day) August 1991	15. PAGE COUNT 4
16. SUPPLEMENTARY NOTATION In L. Birnbaum and G. Collins(Eds.) Machine Learning: Proceedings of the Eighth International Workshop.			
17. COSATI CODES		18. SUBJECT TERMS (Continue on reverse if necessary and identify by block number)	
FIELD	GROUP	SUB-GROUP	
19. ABSTRACT (Continue on reverse if necessary and identify by block number) Cascade models humans learning college physics by studying examples and solving problems. It simulates the main qualitative phenomena visible in human protocols of learning, including several strategies for analogical and non-analogical problem solving, and two strategies for studying examples. It learns at the knowledge level by acquiring new physics rules, and it learns search control knowledge. Most importantly, it models a recently observed phenomenon, the self-explanation effect, which correlates students' example studying strategies with the amount they learn.			
20. DISTRIBUTION / AVAILABILITY OF ABSTRACT ☐ UNCLASSIFIED/UNLIMITED ☒ SAME AS RPT. ☐ DTIC USERS		21. ABSTRACT SECURITY CLASSIFICATION	
22a. NAME OF RESPONSIBLE INDIVIDUAL Dr. Alan L. Meyrowitz	22b. TELEPHONE (Include Area Code) (202) 696-4302	22c. OFFICE SYMBOL N00014	

Learning Physics Via Explanation-based Learning of Correctness and Analogical Search Control

Kurt VanLehn and Randolph M. Jones
Learning Research and Development Center
University of Pittsburgh
Pittsburgh, PA 15260
VanLehn@cs.pitt.edu, Jones@cs.pitt.edu

Abstract

Cascade models humans learning college physics by studying examples and solving problems. It simulates the main qualitative phenomena visible in human protocols of learning, including several strategies for analogical and non-analogical problem solving, and two strategies for studying examples. It learns at the knowledge level by acquiring new physics rules, and it learns search control knowledge. Most importantly, it models a recently observed phenomenon, the self-explanation effect, which correlates students' example studying strategies with the amount they learn.

INTRODUCTION

The long term goal of the Cascade project is to develop a model of human cognitive skill acquisition in scientific, mathematical and engineering task domains. The goal for the present version of the system, Cascade 3, is to model the data from a study by Chi, Bassok, Lewis, Reimann and Glaser (1989), who took protocols as students studied a chapter on Newton's laws from a standard college physics textbook. The chapter consisted of three distinct types of material, which were presented in the following sequence: text, which defines concepts, explains principles and presents ancillary information; examples, which pose a problem and work out its solution; and problems, which the student solves without feedback or help. Chi et al. discovered that students who learned more also explained the examples to themselves more thoroughly. The best learners checked almost every line of the example's solution to see if they could generate it themselves. In particular, Chi et al. found that the number of self-explanation utterances in a protocol was strongly correlated ($r=.81$) with the students' score on the problems. Since all subjects scored roughly the same when tested just prior to studying the examples, this shows that better learning is correlated with more self-explanation of the examples. This phenomenon, called the self-explanation effect, has been replicated twice in other task domains (Pirolli & Bielaczyc, 1989; Ferguson-Hessler & de Jong, 1990).

Cascade 3 is a machine learning program that models the self-explanation effect. This paper discusses its de-

sign and evaluation. (See VanLehn & Jones, in press, for a complete discussion of its design, and VanLehn, Jones & Chi, in press, for a more complete discussion of the evaluation). There were two kinds of constraints on the design. First, the simulation must display the same overt behaviors as the subjects. For instance, subjects often refer to the examples while solving problems, and rarely refer to the text. This should be true of Cascade's performance as well. Second, the simulation should learn what the subjects learn. The first few sections of this paper list the overt behaviors and the types of learning that constrain Cascade's design, and describe how Cascade models each. The last section discusses its evaluation: does it actually reproduce the self-explanation effect?

ORDINARY PROBLEM SOLVING

During the later part of the Chi et. al study, students solve physics problems, such as "A block of mass m is kept at rest on a smooth plane, inclined at an angle of 30 degrees with the horizontal, by means of a string attached to a vertical wall. What is the tension in the string?" A common mode of work is ordinary problem solving, where in the student does not refer to the examples or the text.

Cascade models this overt behavior by using a Prolog meta-interpreter to prove propositions such as

$\text{value}(\text{tension}(\text{string1}), X)$

where $\text{tension}(\text{string1})$ is a sought quantity mentioned in the problem statement. The proof uses equations such as

$\text{tension}(S) = \text{magnitude}(\text{force}(B, S, \text{tension}))$,

which says that the tension in a string S is equal to the magnitude of the tension force exerted on a body B by the string. An equation has conditions that indicate when it applies (e.g., that the string is tied to the body). The combination of an equation and its condition is called a Cascade rule. The meta-interpreter has algebraic knowledge built into it so a single Cascade rule can be used to infer a value for any of the equation's quantities from the other quantities' values. Most of Cascade's knowledge is represented with Cascade rules. Prolog rules are used for visual inference and other kinds of uninteresting, non-learned knowledge.

STUDYING EXAMPLES

A physics example consists of a problem, then a list of lines, where a line can be a force diagram, an equation or a few sentences. For uniformity, Cascade represents all lines as equations, which require inventing some "algebraic" operators for representing diagrams and sentences.

Two fairly distinctive behaviors are common when students study an example line. One is self-explanation. Students try to rederive the line, almost as if they are checking the work of the example's author. Students virtually never reflect on the overall solution and try to recognize a plan that spans all the lines. Their self-explanations are local to a single line. Cascade models self-explanation by proving the line. If the search for a proof succeeds, the derivation is stored as a set of pairs, each consisting of a goal and the Cascade rule used to achieve that goal.

The second behavior observed during example studying consists of simply reading the line, and sometimes paraphrasing it. Students undoubtedly engage in parsing and reference resolution (e.g., determining that "Fa" refers to the force exerted by string A on the block), and these processes also occur as a front end to self-explanation. Instead of modeling them, Cascade presents lines as parsed expressions with all referring terms replaced by their definitions. These data structures model the output of the parsing and disambiguating processes. Thus, to model this second strategy for studying a line, Cascade simply stores the line in memory.

ANALOGICAL PROBLEM SOLVING

There are two rather distinct types of analogical problem solving evident in the protocols. In the first type, students read through an example, starting at the beginning and stopping when they find a line they can use for the problem they are working on. Cascade models this with a kind of transformational analogy (Carbonell, 1986). The givens of a problem are represented as a set of ground literals, such as

tied-to(string1,b1).

To perform an analogy, Cascade retrieves an example, partially matches its givens to the givens of the problem, and saves the resulting mapping. [Some details: Retrieval is not currently modelled in Cascade. It is simply told which examples go with which problems. Partial matching enumerates all constant-to-constant mappings, subject to some type constraints (e.g., constants denoting situations may not be mapped to constants denoting physical objects). Partial matching then picks the mapping that maximizes the number of literals in the givens that are put into 1-1 correspondence.] Example lines are represented as equations. To see if an example line is useful, transformational analogy applies the saved mapping to it, which yields an equation expressed in terms of the problem instead of the example. If this equation matches the current goal, it is used. If not, Cascade moves on to the next

example line.

The second type of analogy is much more focussed than the one just described. At the beginning of a problem solving attempt, students say something like "This is just like example 6" and open the book to the appropriate page. Although they probably look at the diagrams, they generally do not read much of the example. As they work through their problem, they sometimes refer directly to a line in the example without reading the preceding lines.

Cascade models this kind of analogy with a mechanism called analogical search control. At the beginning of a problem, it retrieves an example and forms a mapping just as in transformational analogy. Whenever Cascade needs to choose among several rules for achieving a goal, it uses the map to find an equivalent goal in the example's derivation, and if it succeeds, it applies the rule paired with that goal during self-explanation. If it fails to find such a goal, perhaps because the example's lines were not self-explained, then Cascade chooses the first rule that matches the goal. If this choice fails, will back up and try the next rule that matches. Analogical search control is thus a heuristic for choosing rules. It is also a form of learning, because every time an example is explained or a problem is solved, its derivation is saved and thus adds more heuristic power for choosing rules. Analogical search control is similar to the search control learning of Eureka (Jones, 1989).

LEARNING NEW DOMAIN RULES

The textbook does not mention all the physics knowledge needed for explaining the examples and solving the problems. Many previous simulations of textbook learning (e.g., Cohen, 1990; VanLehn, 1987) have found the same incompleteness. In order to model this incompleteness, we first constructed a knowledge base sufficient to solve all the problems and explain all the examples. We asked two people not involved in the construction of the rules to judge whether each rule was mentioned in the textbook. The judges agreed 95% of the time, and disagreements were settled by a third judge. Of the 62 physics rules, only 29 (47%) were judged to be present in the textbook. Because some students get most of the problems correct, either they knew the missing knowledge already or they acquired new rules as the studied examples and solved problems. There is some empirical evidence against the prior knowledge account (VanLehn, Jones & Chi, in press), so Cascade is based on the hypothesis that students learn the missing rules.

Cascade uses explanation-based learning of correctness (VanLehn, Ball & Kowalski, 1990) to model domain knowledge acquisition. When Cascade cannot achieve its present goal with any known domain rule, it uses a "non-domain" rule. Such rules are distinctively marked in order to represent the belief that they are part of some other task domain (e.g., common sense physics) or are incorrect. For instance, one non-domain rule is, "The property value of a quan-

tity is a function of the property value of one of its defining parts." This rule is overly general and hence incorrect. However, Cascade uses it to prove the example line, "The incline of the normal force on block-1 by surface-1 is perpendicular to the incline of surface-1." A new domain rule is created from this proof by specializing the non-domain rule. In this case, the new domain rule is, "The incline of a normal force on B by S is perpendicular to the incline of S." This rule is correct.

Most of the overly general rules required for learning correct domain rules have the same characteristics as the one just mentioned. They all move property values from objects to related objects. However, there are 3 occasions where new types of forces are used without any introduction or explanation, and these could not be handled with such property-value manipulations. For instance, one problem asserts that an object is prevented from sliding down an inclined plane by a spring that is at the bottom of the plane. This requires inventing a new object (a force), so manipulating property values will not handle this case. In order to handle such learning events, we gave Cascade non-domain rules for common sense physics, including a rule that says that compressed springs push back. We also gave it overly general rules that convert common sense physics to proper physics. One such rule is, "If object X pushes on object Y, then there is a force on Y due to X."

These force-invention cases demonstrate that the success of explanation-based learning of correctness is not a foregone conclusion. It may seem like a learning method that cannot fail, as the only constraint on the content of the overly general rules is the inventiveness of the rules' writer. However, because we could not handle the force-invention cases by just inventing overly general rules, we had to adopt a different approach, which is based on the intuitively plausible hypothesis (espoused by di Sessa, 1983, among others) that formal physics is acquired by refining and adjusting one's naive physics.

Because Cascade applies explanation-based learning of correctness (EBLC) as soon as it fails to achieve a goal, it is deadly for Cascade to stray from a solution path, for it has no way to tell whether the failure at the end of a dead end is caused by missing domain knowledge or by having made a wrong search control decision earlier. In the latter case, EBLC will still be tried and may learn an incorrect rule. Avoiding failure paths is easier when explaining an example, for the goals mention values as well as sought quantities, and the form of the value often eliminates some of the false choices. We found that during problem solving, when goals usually mention only sought quantities and not their values, some other form of search control was necessary in order to prevent EBLC from learning incorrect rules. We were surprised (and delighted) to find that analogical search control provided enough constraint to allow EBLC to learn only correct rules during problem solving.

EBLC is similar to other knowledge-level learning methods whose non-domain rules take the form of explanation patterns (Schank, 1986), causal attribution rules (Pazzani, Dyer & Flowers, 1986; Lewis, 1988) and terminations (Widmar, 1989). EBLC uses the same format for both domain and non-domain rules, which should simplify acquisition of non-domain rules via syntactic generalization. Also, EBLC operates during both problem solving and example explaining, whereas most of its predecessors operate only while explaining examples, and thus did not need the strong search control that analogical search control provides. The discovery of the interaction between search control learning and this class of knowledge level learning may be the most important technical contribution of the Cascade research.

TAKING THE EXAMPLE'S WORD FOR IT

One of the original 62 domain rules could not be learned via EBLC from overly general rules nor from common sense physics. The rule applies first in an example and then in a series of problems, all of which have three strings converging on a knot, with various forces pulling on each of the strings. The rule simply says that the knot is the body (which means that forces will be resolved on it). The text leads students to believe that bodies are physical objects with mass, but the strings are massless and so presumably are the knots. All the subjects that self-explained the example line stating that the knot was the body found the line confusing. After searching in vain for an explanation, they all wound up "taking the example's word for it." That is, they accepted the line as correct even though they could not derive it themselves. On subsequent three-string problems, they would again make the assumption that the knot is the body, often referring to the example by name.

To model this behavior, we tried syntactic techniques for learning by completing explanations (e.g., VanLehn, 1987; Hall, 1988), but could not get the rule learned during the example to apply to all the problems. We relaxed transformational analogy enough to allow it to transfer the knowledge from examples to problems, but this made it too liberal, and it often transferred incorrect knowledge. We solved the problem with an ad hoc but interesting method. When the example explainer has to "take the example's word for it," it builds a special rule that says, "If the goal is to find a body, try a transformational analogy to the knot of example sx," where example sx is the one with the three strings in it. Thus, Cascade learns which kind of situations are good ones for applying transformational analogy.

EVALUATION

The mechanisms of Cascade were constructed to simulate behavior directly visible in the protocols and to model learning inferred from the protocols. It is not

clear whether these mechanisms will actually reproduce the self-explanation effect found by Chi et al. Although the effect consists of correlations between five measures (VanLehn, Jones & Chi, in press), we will here discuss only the most important correlation: Students who utter more self-explanatory statements while studying examples also answer more problems correctly.

Our hypothesis is that the correlation is caused by a strategy difference rather than a prior knowledge difference. Thus, we ran Cascade twice and gave it the same prior knowledge both times (the 29 rules judged to be mentioned in the text). On the so-called Good Student run, we had Cascade self-explain every line of all 3 examples from the Chi et al. study. On the Poor Student run, we had it explain no lines. As expected, this strategy difference during example studying caused a large difference during problem solving. The Good Student solved all 23 of the Chi et al. problems correctly, while the Poor Student solved only 9 problems correctly. Because Cascade is deterministic, explaining an intermediate number of example lines would have caused an intermediate score on the problem solving phase. Thus, Cascade does reproduce the observed correlation.

According to Cascade, there are several sources for this correlation. First, when more lines are explained, there is more opportunity for Cascade to stumble across gaps in its domain knowledge. These gaps cause EBLC, which fills in the gaps with new domain rules. The Good Student learned 7 rules during example studying. Of course, the Poor Student learned none, as it didn't even process the examples.

The Good Student learned 15 rules during problem solving. Of these, 3 were also learned by the Poor Student. In order to determine why the Good student learned 12 more rules than the Poor student, we ran Cascade on the problems with analogical search control turned off. It was given the 36 Good Student rules (29 rules from initial knowledge and the 7 rules learned during example studying). It learned 3 of the 12 rules that were learned by the Good student and not the Poor student. This shows that the knowledge acquired during example studying set up contexts during problem solving that enabled further rule acquisition. In order to control search enough to learn the remaining 9 rules, analogical search control was necessary. To put it another way, of the 22 rules learned by the Good student, 3 can be learned during problem solving even if the examples are not self-explained, 10 can be learned by self-explaining the examples, and 9 require both self-explanation and analogical search control.

DISCUSSION

The research so far has shown that Cascade is qualitatively similar to human behavior in that it can solve problems, both analogically and non-analogically, and it can explain examples. Moreover, both major strate-

gies for studying examples are modelled, as well as several types of analogical reference. The research has also shown that Cascade reproduces the general correlations observed in aggregate data, especially the most important correlation, which indicates that amount of learning is related to amount of self-explanation.

The next step in the research is to try to fit the protocols of individual subjects. Given a student's protocol, Cascade will be forced to explain exactly those lines that the student explained. Will this cause Cascade to make the same errors as the subject? If the subject makes errors that Cascade does not and Cascade uses a rule that it learned via EBLC, then EBLC may be too powerful. If the student makes an error that Cascade does not and Cascade uses only the 29 prior knowledge rules, then our assumption that students all have the same prior knowledge is inaccurate. If Cascade makes errors that the subject does not, then the subject may have used some kind of learning that Cascade does not model. Thus, this further work should test our hypothesis that the self-explanation effect is caused by a strategy difference and not by a prior knowledge difference, and it should help determine whether Cascade's learning models are appropriate and complete.

Although it would be premature to draw strong conclusions about the psychological adequacy of Cascade from the present research, the work taught us much about the technical difficulties of integrating several types of learning and problem solving. EBLC looked like it could not fail to learn, as we could write any overly general rule we liked. Moreover, it seemed certain that it would reproduce the self-explanation effect, because only the self-explaining students processed the examples and we thought all learning occurred there. However, we found it impossible to get EBLC to learn with just overly general rules; we had to embrace diSessa's hypothesis that formal physics evolves from naive physics. We also found that most rules are not learned during example explaining, but during problem solving. So why does explaining the examples help? It turns out the examples help because they control the search (via analogical search control) during problem solving. Without such search control, Cascade wanders down false paths and EBLC learns incorrect rules. In short, we learned that the self-explanation effect is not due solely to knowledge-level learning, but that search control learning is necessary in order to properly bias knowledge-level learning. This may be a general principle that applies to all EBLC-like learning methods (e.g., Lewis, 1988; Pazani, Dyer & Flowers, 1986; Schank, 1986; Widmar, 1989) and perhaps even to all impasse or failure driven knowledge-level learning methods.

Acknowledgements

This research was supported by the Cognitive Science Division of ONR (N00014-88-K-0086) and the Information Sciences Division of ONR (N00014-86-K-0678). We thank Micki Chi for the use of her group's data and for many insightful conversations.

References

- Carbonell, J. (1986). Derivational analogy: A theory of reconstructive problem solving and expertise acquisition. In R.S. Michalski, J.G. Carbonell & T.M. Mitchell (eds) *Machine Learning, An AI Approach: Vol. 2*. Los Altos, CA: Morgan-Kaufman.
- Chi, M.T.H., Bassok, M., Lewis, M.W., Reimann, P. & Glaser, R. (1989). Self-explanations: How students study and use examples in learning to solve problems. *Cognitive Science*, 13, 145-182.
- Cohen, W. W. (1990) Learning from textbook knowledge: A case study. *Proceedings of AAAI-90*. Los Altos, CA: Morgan Kaufman.
- diSessa, A. (1983) Phenomenology and the evolution of intuition. In D. Bentner & A. Stevens (Eds.), *Mental Models*, Hillsdale, NJ: Erlbaum.
- Ferguson-Hessler, M.G.M. & de Jong, T. (1990). Studying physics texts: Differences in study processes between good and poor solvers. *Cognition and Instruction*, 7, 41-54.
- Jones, R. (1989). A model of retrieval in problem solving. PhD. Thesis, Information and Computer Science, University of California at Irvine.
- Hall, R.J. (1988). Learning by failing to explain: Using partial explanations to learn in incomplete or intractable domains. *Machine Learning*, 3, 45-78.
- Lewis, C. (1988). Why and how to learn why: Analysis-based generalization of procedures. *Cognitive Science*, 12, 211-256.
- Pazzani, M., Dyer, M. & Flowers, M. (1986) The role of prior causal theories in generalization. *Proceedings of AAAI-86*. Los Altos, CA: Morgan-Kaufman.
- Pirolli, P. & Bielaczyc, K. (1989). Empirical analyses of self-explanation and transfer in learning to program. *Proceedings of the 11th Annual Conference of the Cognitive Science Society*, Hillsdale, NJ: Erlbaum.
- Schank, R. (1986). *Explanation Patterns: Understanding mechanically and creatively*. Hillsdale, NJ: Erlbaum.
- VanLehn, K. (1987) Learning one subprocedure per lesson. *Artificial Intelligence*, 31, 1-40.
- VanLehn, K. (1989) Problem solving and cognitive skill acquisition. In M.I. Posner (Ed.) *Foundations of Cognitive Science*, Cambridge, MA: MIT Press.
- VanLehn, K., Ball, W. & Kowalski, B. (1990). Explanation-based learning of correctness: Towards a model of the self-explanation effect. *Proceedings of the 12th Annual Conference of the Cognitive Science Society*. Hillsdale, NJ: Erlbaum.
- VanLehn, K. & Jones, R.M. (in press) Integration of explanation-based learning of correctness and analogical search control. In S. Minton and P. Langley (Eds.) *Proceedings of the Symposium on Learning, Planning and Scheduling*. Los Altos, CA: Morgan Kaufman.
- VanLehn, K., Jones, R.M. & Chi, M.T.H. (in press). A model of the self-explanation effect. *Journal of the Learning Sciences*.
- Widmar, G. (1989) A tight integration of deductive and inductive learning. *Proceedings of the Sixth International Workshop on Machine Learning*. Los Altos, CA: Morgan-Kaufman.
- Wilkins, D.C. (1988). Knowledge base refinement using apprenticeship learning techniques. *AAAI-88*, Los Altos, CA: Morgan-Kaufman.