

Classifying and Detecting Plan-Based Misconceptions for Robust Plan Recognition

Randall J. Calistri-Yeh

A Robust Plan-Recognition Algorithm

The traditional plan-recognition problem can be posed as a search problem, where the object is to find a path in the plan hierarchy that explains the user's query and is consistent with what we know about the world. This problem can be handled by a best-first-search algorithm such as A* (Hart, Nilsson, and Raphael 1968; Slagle and Bursky 1968; Pearl 1984). However, some modifications are needed to handle the various classes of plan-based misconceptions. The 10 simplified types of misconceptions can be divided into 2 large groups, based on how they affect the A* algorithm. *Constraint misconceptions* (involving bindings, preconditions, and orderings) are quite easy to handle because they do not affect the structure of the user's plan. The user's plan can still be found in the plan hierarchy, but there might be certain violations that need to be explained. *Structural misconceptions*, on the other hand, are much harder because the user has incorrect beliefs about how the steps of a plan are put together. In these cases, the user's actual plan will not be found in the correct plan hierarchy at all, and we need to modify the structure of the hierarchy itself to identify the plan.

Two extensions to the A* algorithm allow us to handle situations where the user has flawed plans (Calistri-Yeh 1991a). The extension for constraint misconceptions does not add to the asymptotic complexity of the algorithm; so, plans with these misconceptions can be recognized just as efficiently as correct plans. However, the extension to handle structural misconceptions results in a very severe combinatorial explosion. These misconceptions are fundamentally more difficult to handle, and in order to deal with them in a reasonable way, we need powerful heuristics to guide the search.

A Probabilistic Interpretation

Probabilistic methods can be used to handle the three main difficulties of robust plan recognition: dealing with the hard classes of misconceptions, controlling the combinatorial explosion, and selecting the most likely plan from among multiple competing explanations.

My Ph.D. dissertation (Calistri 1990) extends traditional methods of plan recognition to handle situations in which agents have flawed plans.¹ This extension involves solving two problems: determining what sorts of mistakes people make when they reason about plans and figuring out how to recognize these mistakes when they occur. I have developed a complete classification of plan-based misconceptions, which categorizes all ways that a plan can fail, and I have developed a probabilistic interpretation of these misconceptions that can be used in principle to guide a best-first-search algorithm. I have also developed a program called Pathfinder that embodies a practical implementation of this theory. Pathfinder is a probability-based plan-recognition system based on the A* algorithm that uses information available from a user model to guide a best-first search through a plan hierarchy.

Most plan-recognition systems assume that the agent will have perfect plans, and sometimes this approach is acceptable. However, intelligent interfaces are specifically designed to interact with people, who can make mistakes. In fact, it is the least experienced people, the people who make the most mistakes, who need intelligent interfaces the most. The more mistakes a person is likely to make, the more important it is to be able to recognize and correct these mistakes. Ambiguity already shows up quite often in the traditional plan-recognition problem, but flawed plans are inherently ambiguous, and the number of alternative explanations is much larger than with correct plans. Although probabilistic methods have already been applied to plan recognition to handle ambiguity (Goldman and Charniak 1989), they have not addressed the

difficulties encountered with faulty plans. In addition, the work that has been done with plan-based misconceptions has generally ignored the problem of ambiguity (Pollack 1986; Quilici, Dyer, and Flowers 1988; and van Beek 1987).

Classifying Misconceptions

By analyzing the structure of a plan, I have isolated 16 distinct ways that a plan can improperly be constructed. Unlike some of the earlier, unprincipled classifications, this approach is based on the structure of plans and provides a complete classification of

Pathfinder is a probability-based plan-recognition system based on the A algorithm...*

all types of plan-based misconceptions for plans that can be represented within a particular framework. Based on further analysis of the information available to an outside observer (the plan recognizer), I have simplified the classification to 10 detectable and distinguishable types of plan-based misconceptions.

There are two important advantages to this approach. First of all, it is completely independent of the domain. Second, by associating misconceptions directly with plan components, I impose a structure on the mistakes that allows me to handle misconceptions and plan recognition simultaneously.

With certain reasonable assumptions, the probability of a particular plan can be calculated incrementally from the probabilities of the steps and misconceptions that make up the plan. The step probabilities are fairly easy to estimate, but the probability of the user having a particular misconception seems much more difficult. My approach is to break the probability of a misconception into a number of independent features that can be calculated at run time when the misconception is identified (Calistri-Yeh 1991b). One example of a feature is *similarity*: Generally, the more similar two things are, the more likely it is that the user might have confused them. Each class of misconception has associated with it a set of functions for calculating the relevant features based on information provided by a user model.

These probabilities can form a heuristic function suitable for best-first search. While constructing a potential explanation for the user's query, the probability that this is the correct explanation can be calculated incrementally by factoring in the steps and misconceptions as they are proposed by the algorithm. We can use the probability of the partial explanation constructed so far as a heuristic estimate of the likelihood of the entire explanation.

Evaluation and Results

In order to test both the soundness of the theory and the practicality of the Pathfinder program, I compiled a large corpus of naturally occurring dialogues, comprising over 100 exchanges that contain examples of plan-based misconceptions (Calistri 1989). As far as is known, this is an order of magnitude larger than any other collection to date. The examples come from a wide range of sources (a Usenet newsgroup, a radio talk show, online computer dialogues, and so on) and cover a large number of topics, such as income taxes, financial investments, the UNIX operating system, and the Emacs text editor.

Ninety-seven transcript examples were used to test the classification theory, and all but 3 could successfully be classified. During this exercise, I discovered that there are certain inherent ambiguities in classifying specific misconceptions. Pathfinder can exploit this ambiguity in order to increase its efficiency (by assigning misconceptions to "easier" classifications). The Pathfinder program was tested on 58

examples (39 directly from the transcripts and 19 variations of transcript examples). In 53 cases (91 percent of the test cases), Pathfinder produced completely reasonable and believable interpretations. In 4 other cases, it still preferred the "correct" explanation but assigned unreasonable probabilities to the remaining interpretations. There was only 1 case where it produced the wrong interpretation.

Summary

Virtually all the plan-recognition techniques of the past (and the present) have ignored the problem of human fallibility. This approach creates a serious problem when systems that use these techniques are faced with a flawed plan; most of these systems are completely unable to understand what the agent is doing in these cases. Intelligent interfaces, intelligent tutoring systems, and any other systems that endeavor to interact intelligently with people need the capacity to recognize, diagnose, and correct the mistakes that are bound to occur.

System designers have typically dealt with the problem of recognizing faulty plans by including bug lists of common mistakes, but there are many studies that have shown that bug lists are not robust enough to handle the full range of mistakes that people make (for example, Carroll and Aaronson [1988]). Instead of trying to predict these mistakes one at a time, I have developed a new classification of plan-based misconceptions that covers all ways that a plan can be constructed or executed improperly.

In order to handle the ambiguity problem, I employ new probabilistic methods. It is possible to isolate the primary "features" that determine whether a plan-based misconception is reasonable in a given situation. These features can be calculated from a model of the user, and once the class of misconception is known, the features can be combined to form the probability that the user has this misconception. Both the theoretical classification and the Pathfinder program have been tested on a large corpus of naturally occurring plan-based misconceptions, which were extracted from real conversations in a variety of domains.

Acknowledgments

The research described in this thesis was performed at Brown University under the guidance of Eugene Charniak. Some work

was done while I was at the GE Corporate Research and Development Center. Support at Brown was provided in part by National Science Foundation grants IST 9416034 and IST 8515005 and Office of Naval Research grant N00014-88-K-0589.

Note

1. For information on obtaining this dissertation, either write to the author at ORA Corporation, 301A Harris B. Dates Drive, Ithaca, NY 14850-1313, <calistri@oracorp.com>, or contact Librarian, Computer Science Department, Brown University, P.O. Box 1910, Providence, RI 02912, (401) 863-7600.

References

- Calistri, R. 1990. Classifying and Detecting Plan-Based Misconceptions for Robust Plan Recognition. Ph.D. diss., Dept. of Computer Science, Brown Univ.
- Calistri, R. 1989. An Annotated Compendium of Naturally Occurring Plan-Based Misconceptions, Technical Report, CS-89-37, Dept. of Computer Science, Brown Univ.
- Calistri-Yeh, R. 1991a. An A* Approach to Robust Plan Recognition for Intelligent Interfaces. In *Applications of Learning and Planning Methods*, ed. N. G. Bourbakis, 227–251. World Scientific Series in Computer Science, volume 26. Teaneck, N.J.: World Scientific.
- Calistri-Yeh, R. 1991b. Utilizing User Models to Handle Ambiguity and Misconceptions in Robust Plan Recognition. *User Modeling and User-Adaptive Interaction* 1(4). Forthcoming.
- Carroll, J., and Aaronson, A. 1988. Learning by Doing with Simulated Help. *Communications of the ACM* 31(9): 1064–1079.
- Goldman, R., and Charniak, E. 1989. A Semantics for Probabilistic Quantifier-Free First-Order Languages, with Particular Application to Story Understanding. In *Proceedings of the Eleventh International Joint Conference on Artificial Intelligence*, 1074–1079. Menlo Park, Calif.: International Joint Conferences on Artificial Intelligence.
- Hart, P.; Nilsson, N.; and Raphael, B. 1968. A Formal Basis for the Heuristic Determination of Minimum Cost Paths. *IEEE Transactions on Systems Science and Cybernetics* SSC-4(2): 100–107.
- Pearl, J. 1984. *Heuristics: Intelligent Search Strategies for Computer Problem Solving*. Reading, Mass.: Addison-Wesley.
- Pollack, M. 1986. Inferring Domain Plans in Question-Answering. Ph.D. diss., Dept. of Computer and Information Science, Univ. of Pennsylvania.
- Quilici, A.; Dyer, M.; and Flowers, M. 1988. Recognizing and Responding to Plan-Oriented Misconceptions. *Computational Linguistics* 14(3): 38–52.
- Slagle, J., and Bursky, P. 1968. Experiments with a Multipurpose, Theorem-Proving Heuristic Program. *Journal of the ACM* 15(1): 85–99.
- van Beek, P. 1987. A Model for Generating Better Explanations. In *Proceedings of the Twenty-Fifth Annual Meeting of the Association for Computational Linguistics*, 215–220. Stanford, Calif.: Association for Computational Linguistics.

Randall Calistri-Yeh is a computer scientist at ORA Corporation in Ithaca, New York. His research interests include plan recognition, misconceptions, probabilistic reasoning, intelligent tutoring systems, user modeling, and natural language.