

AD-A266 308



2

The Control of Reasoning Under Uncertainty

Paul R. Cohen, *Principal Investigator*
David M. Hart, *Lab Manager*

Experimental Knowledge Systems Laboratory
Department of Computer Science
University of Massachusetts
Amherst, MA 01003

cohen@cs.umass.edu
dhart@cs.umass.edu
413 545 3638

DTIC
ELECTE
JUN 30 1993
S A D

Final Report

Contract N00014-88-K-0009
Period 12/1/87 - 2/14/91

This document has been approved
for public release and sale; its
distribution is unlimited.

93-14773



93 6 29 036

Table of Contents

Numerical Productivity Measures	1
Brief Summary of Technical Results	2
Detailed Summary of Technical Results	3
Publications, Presentations, Reports, Awards/Honors	10
Transitions and DoD Interactions	13
References	16

Approved For	
NHS - CAPI	☒
DDIC - FBI	☐
DDIC - CIA	☐
DDIC - JCS	☐
By <i>per lti</i>	
Distribution /	
Priority Class	
Dist	Area of Interest
	Special
A-1	

COPIES OF THIS REPORT ARE AVAILABLE FROM THE

Numerical Productivity Measures

Refereed papers published: 12

Unrefereed reports and articles: 5

Books or Parts Thereof Published: 2

Patents filed but not yet granted: 0

Patents granted: 0

Invited presentations: 6

Contributed presentations: 12

Honors, including conference committees: 4

Paul Cohen:

Councillor, *American Association for Artificial Intelligence*, 1991-1994.

Chairman, *NSF/DARPA Workshop on AI Methodology*. University of Massachusetts. June, 1991.

Organizing Committee, *AAAI Workshop on Intelligent Real-Time Problem Solving*. Anaheim, CA. July, 1991.

Program Committee, *Sixth International Symposium on Methodologies for Intelligent Systems (ISMIS'91)*. Charlotte, NC. October 1991.

Promotions: 1

Paul Cohen, Associate Professor (1989)

Graduate students supported at least 25% time: 11

Brief Summary of Technical Results

The starting point of our research was to develop declarative representations of problem solving strategies and interpreters to select and execute appropriate strategies. One result was a "strategy frame," a collection of fields or facets of strategies. These included the conditions under which strategies should be evoked, measures of progress, stopping conditions, resource requirements, and so on.

We applied this result in a process control domain with some success, then began to search for a more realistic and demanding task. The Phoenix system, a realtime simulation of forest fire fighting in Yellowstone National Park, was an ideal testbed. In this context, our work on declarative representations of strategies followed two paths:

We developed a declarative representation of the progress of plans, called *envelopes*. Envelopes provide early warning of plan failures, facilitate communication between agents, and help with monitoring and replanning in real time.

We began developing mathematical models of the architectures of agents that facilitate the design of agents for environments with different characteristics.

Envelopes. A valuable component of the original strategy frames was the *measure of progress* slot. It was consulted by the interpreter to decide how and when to reallocate resources. In Phoenix the planner consults more sophisticated but essentially similar structures called *envelopes*. During the term of this contract we accomplished a number of goals in our research with envelopes, including: integrating agent and plan envelopes into Phoenix, collaborating with Dr. Gerald Powell of CECOM to assess the utility of envelopes in operational battlefield planning, publishing numerous papers on the theory and utility of envelopes, and basing collaborations with DARPA and Digital Equipment Corporation on envelopes. More recently we have been formulating a general method for constructing envelopes using principles of Signal Detection Theory.

Modeling, Analysis and Design. During 1989-90 we developed a principled methodology for AI research based on modeling functional relationships between agent architecture and behavior. This emphasis on methodology grew in part out of our work on control using strategy frames. We call this methodology Modeling, Analysis and Design. The methodology divides the design and implementation of agents into seven phases. We have followed these steps in our Phoenix work since developing this approach as a way of testing and refining the methodology. We conducted a survey of AAAI-90 papers that showed how AI could benefit from such an approach, the results of which appeared in *AI Magazine*. We have also engaged in numerous activities to further refine and broaden the methodology, including a workshop on AI methodology held in 1991, the development of a graduate level course in *agentology*, and a AAAI tutorial on experimental methods for AI research.

Detailed Summary of Technical Results

A Declarative Representation of Control Strategy

The starting point of our research was to develop declarative representations of problem solving strategies and interpreters to select and execute appropriate strategies. We began by analyzing the strategies of about a dozen classic AI systems, including HEARSAY-II, PIP, CASNET, MYCIN and MDX. One result was a "strategy frame," a collection of fields or facets of strategies. These included the conditions under which strategies should be evoked, measures of progress, stopping conditions, resource requirements, and so on. In the first year of this contract we implemented a process-control system using strategy frames (Cohen, DeLisio & Hart, 1989). This task was contrived for experimental purposes.

Once the experiments were finished, we began to search for a more realistic and demanding task. The Phoenix system, under development since 1987¹, was an ideal testbed. Phoenix is a realtime simulation of forest fires in Yellowstone National Park, and a distributed planning system for controlling fires by the actions of semi-autonomous agents such as bulldozers. In this context, our work on declarative representations of strategies during the second year of the contract followed two paths:

We developed a declarative representation of the progress of plans, called *envelopes*. Envelopes provide early warning of plan failures, facilitate communication between agents, and help with monitoring and replanning in real time.

We began developing mathematical models of the architectures of agents that facilitate the design of agents for environments with different characteristics.

We discuss these developments in turn.

Plans and Envelopes for Phoenix Agents

A valuable component of the original strategy frames was the *measure of progress* slot. It was consulted by the interpreter to decide how and when to reallocate resources. In the Phoenix system, the planner consults more sophisticated but essentially similar structures called *envelopes*.

The idea behind envelopes is to represent explicitly and declaratively the progress of agents. Just as we can explicitly represent the movements of an agent through its physical environment, so can we represent its movement through spaces bounded by failure or other important events. These spaces we call envelopes. The concept is easier to illustrate than it is to define. Imagine you have one hour to reach a point five miles away, and your maximum speed is 5 mph. If you are late, by even a moment, you fail. As long as you maintain your maximum speed, you are *within your envelope*. The instant your speed drops below 5 mph, you *lose or violate* your envelope. This envelope is *narrow*, because it will not accommodate a range of behavior: any

¹ Supported by DARPA and by ONR under the University Research Initiative.

deviation from 5 mph is intolerable. The following problem illustrates a wider envelope. You have one hour to travel five miles, as before, but your maximum speed is 10 mph. You start slowly: your average speed is just 3 mph. After 40 minutes you have traveled just two miles, and you have just 20 minutes to travel the other three. This is possible: If you travel at maximum speed (10 mph), you will achieve your goal with about a minute to spare. On the other hand, if you continue to travel 3 mph for another 171 seconds, you will fail – you will not be able to cover the prescribed five miles in one hour.

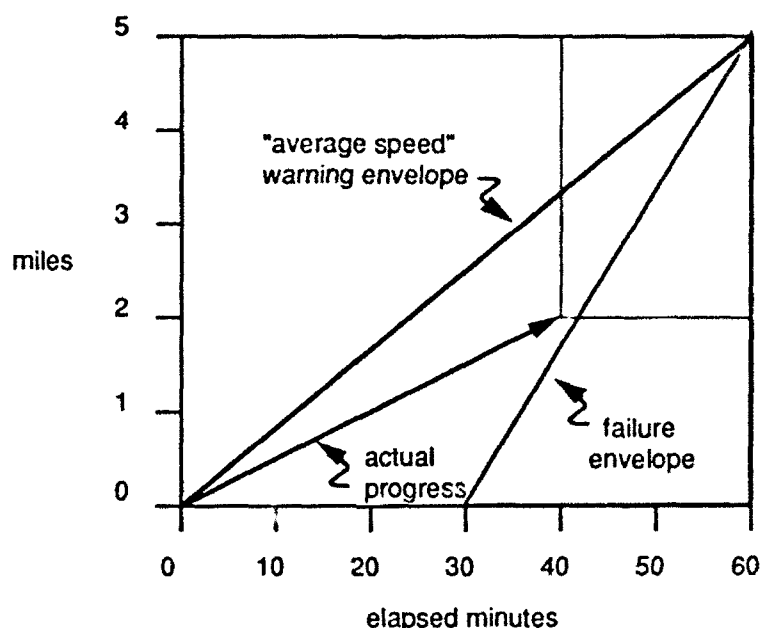


Figure 1. Depicting actual and projected progress with respect to envelopes

Clearly, if the agent waits 40 minutes to assess its progress, it has waited too long, because an heroic effort will be required to achieve its goal. In Phoenix, agents check their progress at regular intervals. They check *failure* envelopes, which tell them whether they will absolutely fail to achieve their goals, and they check *warning* envelopes, which tell them that they are in jeopardy of failure. Typically, there is just one failure envelope but many possible warning envelopes. To continue the previous example, you would violate a warning envelope if your average speed drops below 5 mph, because this is the speed you must maintain to achieve your goal. Violating this envelope says, "You can still achieve your goal, but only by doing better than you have up to this point." These concepts are illustrated in Figure 1. The failure envelope is a line from "30 minutes" to "five miles," since the agent can achieve its goal as long as it has at least 30 minutes to travel five miles. The average speed warning envelope is a line from the origin to the goal, but our agent violated that envelope immediately by at an average speed of 3 mph. In fact, it has moved perilously close to its failure envelope. The box in the upper right illustrates that the agent can construct another envelope from any point in its progress. This new envelope is extremely narrow, as

indicated by the distance from its origin to the point at which the failure envelope intersects the x-axis.

The Utility of Envelopes. In Phoenix and many other distributed planning problems, activities must be coordinated between agents at different levels of a hierarchical command structure, and also between agents at the same level of a hierarchical structure. In Phoenix, several bulldozer agents work under the direction of a fireboss agent. The fireboss tells the bulldozers roughly what to do, and the bulldozers figure out how to do it. *Agent envelopes* monitor the progress of individual agents; for example, the amount of fireline cut by a bulldozer. *Plan envelopes* are maintained by the fireboss, and monitor the progress of several agents. Thus, plan envelopes represent the coordinated activity of several agents (bulldozers) at a single level of the command hierarchy (the subordinate bulldozer level). The coordination among agents at different levels of the hierarchy (i.e., between a bulldozer and the fireboss) is managed by agent envelopes.²

A planner can represent the progress of its plan by transitions within the plan's envelopes. Progress, failures and potential failures are clearly seen from one's position with respect to envelopes, whereas this information is not apparent from one's position in the environment. Just as a planner can project how its actions will propel it through its environment, so it can project how these actions will move it with respect to its envelope. Envelopes function as "early warning" devices: warning envelopes alert the planner to developing problems, and even failure envelopes tell the planner that a plan will fail sometime in the future, so the failure doesn't come as a surprise.

Envelopes Progress. During the term of this contract we accomplished the following goals in our research with envelopes:

- Integrated agent and plan envelopes into Phoenix, as described above.
- Collaborated with Dr. Gerald Powell of CECOM to assess the utility of envelopes in operational battlefield planning.
- Published one paper on the Phoenix planner (Cohen, et al., 1989), one specifically on envelopes for operational planning (Powell & Cohen, 1990), one on envelopes for real time problem solving (Howe, Hart & Cohen, 1989), and one on envelopes in Phoenix (Hart, Anderson & Cohen, 1990).
- Made envelopes the cornerstone of a three-year DARPA real-time planning initiative that started in September, 1989.
- Made envelopes the basis for a collaborative project with Digital Equipment Corporation on planning for competitive computer markets.

Continuing Work with Envelopes and Monitoring. In (Cohen, St. Amant & Hart, 1992) we report on our recent efforts to formulate a general method for constructing envelopes

² In (Hart, Anderson & Cohen, 1990) we describe a Phoenix plan envelope and further discuss the utility of envelopes in Phoenix.

using principles of Signal Detection Theory. In this paper we analyze a tradeoff between early warnings of plan failures and false positives. In general, a decision rule that provides earlier warnings will also produce more false positives. Slack time envelopes are decision rules that warn of plan failures in our Phoenix system. Until now, they have been constructed according to ad hoc criteria. In the paper we show that good performance under different criteria can be achieved by slack time envelopes throughout the course of a plan, even though envelopes are very simple decision rules. We also develop a probabilistic model of plan progress, from which we derive an algorithm for constructing slack time envelopes that achieve desired tradeoffs between early warnings and false positives.

Our work with envelopes has also led to explorations of the general issue of monitoring plan execution in AI planning systems. We are developing a taxonomy of monitoring strategies to serve as a guide to an AI planning system for selecting and parameterizing an appropriate monitoring strategy for each of the plans it executes. A strategy selection mechanism based on this taxonomy will become a component in the *plan steering* architecture we are building for the DARPA Planning Initiative's large transportation planning problem (see Transitions and DoD Interactions).

Modelling AI Systems with Functional Relations

The work on strategy frames was a reaction to a view, prevalent in AI, that control doesn't matter. Most systems have strict forward or backward chaining, or "opportunistic" processing. When we began to work on strategy frames, we thought control was neglected because there were no easy, explicit representations of more sophisticated strategies. Strategy frames were intended to provide such a representation. Eventually, however, we came to believe another explanation: Control is neglected because it really doesn't matter much in the trivial operating environments of AI systems. We need sophisticated control in real-time, dynamic, multi-actor, spatially distributed, unpredictable environments; we probably don't need it for static, predictable, single-agent environments. Today, our research on strategies is a reaction against the methodology of working in trivial operating environments, in which control doesn't matter (Cohen, 1989).

A conceptual breakthrough came about when we characterized the architecture and the environment of an agent as constraints on its behavior. Figure 2 shows the *behavioral ecology triangle* that illustrates these relationships.

Given this view, we characterized AI as a kind of design. We don't design graphics, or VLSI circuits, or mechanical devices: we design intelligent agents. The agents are evaluated by how they behave, determined by their environments and their architectures. Once we adopt this view, we see immediately that we do not know enough about the relationships between agent architectures and behaviors to design intelligent agents in a principled way. We cannot answer the question, "How would the behavior of this AI program change if you change its architecture this way: ... ?" But until we can answer this question, AI system design will remain ad hoc. Since this

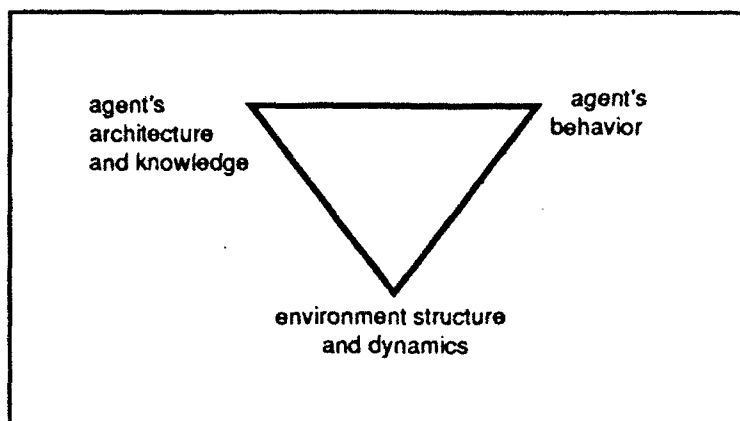


Figure 2. The three components of an agent's behavioral ecology.

breakthrough we have been developing mathematical models of the functional relationships between an agent's architecture and its behavior.

Modeling, Analysis and Design. During 1989-90 we developed a principled methodology for AI research based on modeling functional relationships between agent architecture and behavior. We call this methodology Modeling, Analysis and Design (MAD). We also conducted a survey of AAAI-90 papers that showed how AI could benefit from such an approach. The survey appeared in *AI Magazine* (Cohen 1991).

The model-based design and analysis methodology is based on these premises:

- The goal of AI is the design and implementation of autonomous agents; that is, programs whose behavior is not completely determined by their relatively complex and dynamic environments.
- The behavior of an agent is determined by the interactions between its architecture and its environment.
- It is possible to build formal models of these interactions that are sufficiently predictive to support design and analysis, despite the inherent complexity of AI architectures and environments.

The methodology itself divides the design and implementation of agents into seven phases. In practice, we cannot push a project through each phase in one pass, but must iterate over the experiment/explain/redesign phases:

- 1. Environment assessment:** Determining which aspects of the environment must be represented in a model for design and analysis
- 2. Modelling:** Formally specifying the functional relationships from which to predict behavior, given the architecture and environment of an agent.
- 3. Design:** Inventing or adapting architectures that are predicted to behave as desired in particular environments. In addition, redesign involves modifying a design when it is shown, by way of a model, to perform less well than it might.

4. Prediction: Inferring from the functional relationships in a model how behavior will be affected by changing the architecture of the agent or its environment.

5. Experiments: Testing the veracity of predictions by running the agent in its environment.

6. Explanation: Finding the source of incorrect predictions in a model, and revising the model, when unexpected behaviors emerge from the interactions between an agent and its environment.

7. Generalization: Whenever we predict the behavior of one agent in one environment, we should ideally be predicting similar behaviors for agents with related architectures in related environments. In other words, our models should generalize over architectures, environmental conditions and behaviors.

Applying Modeling, Analysis and Design to Phoenix. Since developing the MAD methodology we have applied it to all of our work in Phoenix. The first step in MAD is environmental assessment (Cohen, 1991) -- determining which aspects of the environment must be represented in a model for design and analysis. During this contract we developed several models to assess the Phoenix environment:

- A model of the way fires spread in the Phoenix simulator has been derived from statistical analysis of hundreds of randomly set fires in our simulated Yellowstone (Silvey, 1990).
- The assumption of constant weather conditions over the life of each fire made by the model above is unrealistic, since changing weather is one of the primary impediments to long-range planning in this domain. To introduce this environmental variability we have developed an analytical model that gives realistic changes in global wind speed and direction (Hansen, 1990).

The next steps in MAD call for the construction of predictive models relating the agent's architecture to its environment, along with empirical verification of the models. We have followed these steps as we studied several research issues in Phoenix:

- Modeling optimal fire-fighting strategies. Models show that when fighting multiple fires sequentially in Phoenix, the best strategy is to fight the youngest fire first (Cohen, 1990a). In (Cohen, Hart & deVadoss, 1991) we report on experiments that verify these models.
- Failure recovery analysis. In (Howe & Cohen, 1991) we report on an extended model of error-recovery in Phoenix and on experiments we conducted to test the predictions this model makes about the cost of error-recovery.
- The Phoenix agent architecture provides two components that read data from sensors and program effectors, thereby providing two sense-act loops for a single agent. We argue in (Anderson, Hart & Cohen, 1991) that having two components is justified because of differences between the kinds of tasks assigned to each component and the resulting interruption of one task by another if they were to be assigned to a single component.

Continuing Methodological Development. Paul Cohen was invited to deliver keynote addresses on methodological issues at a conference and a AAAI Spring Symposium (see Invited Presentations). He also participated in the recent Workshop on Research

in Experimental Computer Science, the goal of which was to identify issues and problems arising in experimental work in the entire field of Computer Science. Sponsored by ONR, DARPA, and NSF, this workshop was held in Palo Alto, CA, October 16-18, 1991.

We can also report a number of encouraging developments growing out of this work since the contract period:

- Workshop on AI Methodology. Held in June of 1991, this workshop brought together a group of leading AI researchers to discuss growing methodological concerns and develop a consensual strategy for addressing them (see Transitions and DoD Interactions).
- Agentology Curriculum. During the summer of 1991 we conducted a summer school designed develop the skills in our graduate students needed to conduct MAD research, and believe that this effort has laid the groundwork for a curriculum in *agentology* -- the principled design of autonomous agents for complex environments. From that summer school we have developed a research methods course for AI graduate students and are working on an accompanying textbook on Experimental Methods for AI Research (tentatively scheduled for release at AAAI-93).
- AAAI-92 Tutorial on Experimental Methods for AI Research. This tutorial was offered jointly with Prof. Bruce Porter from the Univ. of Texas, Austin.

Publications, Presentations, Reports, Awards/Honors

Refereed Papers Published

- Cohen, P.R., 1991. A survey of the *Eighth National Conference on Artificial Intelligence*: Pulling Together or Pulling Apart? *AI Magazine*, 12(1): 16-41.
- Cohen, P.R., 1990. Methodological problems, a model-based design and analysis methodology, and an example. Keynote address and paper at the *International Symposium on Methodologies for Intelligent Systems*. Knoxville, Tennessee.
- Cohen, P.R., 1990. Designing and analyzing strategies for Phoenix from models. *Proceedings of the Workshop on Innovative Approaches to Planning, Scheduling and Control*. Katia Sycara (Ed.). Morgan Kaufman. Pp. 9-21.
- Cohen, P.R., Greenberg, M.L., Hart, D.M. & Howe, A.E., 1989. Trial by fire: Understanding the design requirements for agents in complex environments. *AI Magazine* 10(3): 34-48.
- Cohen, P.R. & Howe, A.E., 1989. Toward AI research methodology: Three case studies in evaluation. *IEEE Transactions on Systems, Man and Cybernetics*. 19(3): 634-646.
- Cohen, P.R., DeLisio, J.L. & Hart, D.M., 1989. A declarative representation of control knowledge. *IEEE Transactions on Systems, Man and Cybernetics*. 19(3): 546-557.
- Hart, D.M., Anderson, S.D. & Cohen, P.R., 1990. Envelopes as a vehicle for improving the efficiency of plan execution. *Proceedings of the Workshop on Innovative Approaches to Planning, Scheduling and Control*. Katia Sycara (Ed.). Morgan Kaufman. Pp. 71-76.
- Howe, A.E. & Cohen, P.R., 1991. Failure recovery: A model and experiments. *Proceedings of the Ninth National Conference on Artificial Intelligence*. AAAI Press. Pp. 801-808.
- Howe, A.E., Hart, D.M. & Cohen, P.R., 1990. Addressing real-time constraints in the design of autonomous agents. *The Journal of Real-Time Systems*, 1(1/2): 81-97.
- Howe, A.E., Hart, D.M. & Cohen, P.R., 1990. Designing agents to plan and act in their environments. Abstract for *The Workshop on Automated Planning for Complex Domains* at the *Eighth National Conference on Artificial Intelligence*. Boston, MA.
- Howe, A.E. & Cohen, P.R., 1988. How evaluation guides AI research. *AI Magazine*, 9(4): 35-43.
- Powell, G.M. & Cohen, P.R., 1990. Operational planning and monitoring with envelopes. *Proceedings of the IEEE Fifth Annual AI Systems in Government Conference*, Washington, DC.

Unrefereed Reports and Articles

Cohen, P.R. (Ed), 1992. *Working Papers of the Workshop on Artificial Intelligence Methodology*. Forthcoming.

Cohen, P.R., 1991. AI methodology: A position paper. *Working Papers of the Workshop on Artificial Intelligence Methodology*. Northampton, MA. June 2-4, 1991. Pp. 19-25.

Cohen, P.R., Hart, D.M. & deVadoss, J.K., 1991. Models and experiments to probe the factors that affect plan completion times for multiple fires in Phoenix. EKSL Memo #17, Department of Computer Science, University of Massachusetts, Amherst.

Cohen, P.R. & Howe, A.E., 1991. Benchmarks are not enough; Evaluation metrics depend on the hypothesis. *Collected Notes from the Benchmarks and Metrics Workshop*. Technical Report FIA-91-06, NASA Ames Research Center. Pp. 18-19.

Cohen, P. R., 1989. Why knowledge systems research is in trouble and what we can do about it. Technical Report #89-81, Department of Computer Science, University of Massachusetts, Amherst.

Books or Parts Thereof Published

Cohen, P.R., 1990. Architectures for reasoning under uncertainty. *Readings in Uncertain Reasoning*. Glenn Shafer and Judea Pearl (Eds). Morgan Kaufmann.

Howe, A.E. & Cohen, P.R., 1990. How evaluation guides AI research. Reprinted in *A Sourcebook of Applied Artificial Intelligence*. Gerald Hoppole and Stephen Andriole, Eds. TAB Books, Inc. (Originally published in *AI Magazine*, Winter, 1988.)

Invited Presentations

Cohen, P.R.

- A brief report on a survey of AAAI-90, some methodological conclusions, and an example of the MAD methodology in Phoenix. Keynote address, *AAAI Spring Symposium on Implemented AI Systems*. Palo Alto, CA. March, 1991.
- Methodological problems, a model-based design and analysis methodology, and an example. Keynote address at the *International Symposium on Methodologies for Intelligent Systems*. Knoxville, TN. October 25-27, 1990.
- Modelling for AI system design. Imperial Cancer Research Fund, London, England. June 25, 1990.
- Modelling for AI system design. Digital Equipment Corporation, Galway, Ireland. June 25, 1990.
- Fire will destroy the pestilence, or, How natural environments will drive out bad methodology. Texas Instruments, Dallas. May 24, 1990.
- Designing autonomous agents. Thayer School of Engineering, Dartmouth College. November 21, 1989.

Contributed Presentations

Cohen, P.R.

- Welcoming address (untitled) at the NSF/DARPA *Workshop on Artificial Intelligence Methodology*. Northampton, MA. June 2, 1991.
- The Phoenix project: Responding to environmental change. *Workshop on Innovative Approaches to Planning, Scheduling, and Control*. San Diego, CA. November 1990.
- What is an interesting environment for AI planning research? Panel moderator, *Workshop on Automated Planning for Complex Domains* at the *Eighth National Conference on Artificial Intelligence*. Boston, MA. July 30, 1990.
- Intelligent real-time problem solving: Issues and examples. *Intelligent Real-Time Problem Solving Workshop*. Santa Cruz, CA. November 8-9, 1989.
- Evaluation and case-based reasoning. Member of panel on Evaluation Issues in CBR at the *DARPA Case-Based Reasoning Workshop*. Pensacola Beach, FL. May 1989.
- Why knowledge systems research is in trouble and what we can do about it. *Workshop on the Foundations of Expert Systems*. Ponte de Lima, Portugal. March 23, 1989.

Greenberg, M.L.

- Real-time problem solving in the Phoenix environment. *Workshop on Real-time AI Problems* at *IJCAI-89*. Detroit, MI. August 20, 1989.

Hart, D.M.

- The Phoenix project: Simulating a complex, real-time environment for autonomous agents. *Fourth AI & Simulation Workshop* at *IJCAI-89*. Detroit, MI. August 21, 1989.

Howe, A.E.

- Failure recovery: A model and experiments. *Ninth National Conference on Artificial Intelligence*. Pasadena, CA. July, 1991.
- Designing agents to plan and act in their environments. *Workshop on Automated Planning for Complex Domains* at the *Eighth National Conference on Artificial Intelligence*. Boston, MA. July 30, 1990.
- The Phoenix project: Problem solving in a crisis management environment. *First Workshop on Symbolic Problem Solving in Noisy, Novel and Uncertain Task Environments* at *IJCAI-89*. Detroit, MI. August 20, 1989.

Powell, G.M.

- Operational planning and monitoring with envelopes. *IEEE Fifth Annual AI Systems in Government Conference*. Washington, DC. May 9, 1990.

Transitions and DoD Interactions

1989. Work on envelopes was transferred in several ways during 1989. Two organizations had a direct interest in envelopes for their own applications. The Army, via Dr. Gerald Powell of CECOM in Fort Monmouth, was exploring the applications of envelopes to battlefield planning [Powell & Cohen]. Digital Equipment Corporation also sponsors some of our research on Phoenix and envelopes. DEC was interested in having Phoenix agents play the roles of competitors in simulations of market dynamics. Phoenix was also chosen as a testbed for AFOSR's Intelligent Real Time Problem Solving Initiative.

The work on modelling and functional relationships can potentially provide the basis for design and analysis of AI systems. The Department of Defense spends enormous resources on AI systems and AI research, but must rely on the intuitions of systems builders. Moreover, AI systems are notoriously difficult to test, much less analyze. We envision a time when AI systems are designed, analyzed, and modified not by intuition, but with the guidance of models. Our preliminary results are modest, but they are significant: We can answer design questions analytically when previously they were answered, "Try it and see what happens."

1990. Gerald M. Powell was a visiting faculty member in 1990 under the Secretary of the Army Research and Study Fellowship Program. Dr. Powell, who worked then for the Center for Command, Control, and Communications Systems, CECOM, Ft. Monmouth, New Jersey, was investigating computational approaches to various problems in battlefield planning for the previous five years, and is very interested in the present capabilities and further design and development of Phoenix. He previously worked with Paul Cohen applying envelopes to an operations planning problem in battlefield management (see above), and during his visit studied the application of approximate processing techniques for real-time control in Phoenix.

Paul Cohen and David Hart visited the Decision Systems Laboratory at Texas Instruments in Dallas, May 24-25. Cohen presented a talk entitled "Fire will Destroy the Pestilence, or, How Natural Environments will Drive Out Bad Methodology." Phoenix was demonstrated for the DSL, and we looked at a number of their projects, including CACTUS, a battlefield planning system that is conceptually similar to Phoenix, although implemented differently. We discussed doing a comparative analysis of these two systems to show they fall within an equivalence class with respect to the task environments and design of agents for those environments. Such an analysis would attempt to show that both systems can be represented using the same underlying model for the task environment and agent design, thus substantiating the methodological approach we advocate.

We also discussed at length with TI the use of visualization techniques to aid in the interpretation and analysis of a system that simulates of shop floor activities in a semi-conductor fabrication plant. The simulation allows experimentation with vari-

ous scheduling strategies to improve plant throughput. However, the volume of data it produces overwhelms the capabilities of traditional data analysis techniques. Our discussions focused on ways of visualizing pathologies that arise during (the simulation of) shop floor processing that cause the operant scheduling strategy to perform poorly or fail. These visualizations would allow the user to intervene as problems develop and explore the causes by pausing and interacting with the system graphically. These ideas are based on our work in Phoenix with simulation, graphical interfaces, and envelopes, and led ultimately to a proposal (currently a contract) to apply these techniques to the development of a plan steering architecture in the DARPA Planning Initiative.

We also joined TI in a proposal to use models of autonomous agents developed at TI and in Phoenix to improve semi-autonomous forces technology for battlefield training simulations such as SIMNET. This is an effort to extend the Joint Training Simulation Concept to support Joint Service and International (NATO) training operations.

Paul Cohen presented a talk entitled "Modelling for AI System Design" at Digital Equipment Corporation in Galway, Ireland, and at the Imperial Cancer Research Fund in London on June 25. These talks led to plans to hold a workshop sponsored by NSF and DARPA in early 1991 on methodology in AI research (see below).

1991.

Workshop on AI Methodology. We held a Workshop on Artificial Intelligence Methodology in Northampton, Massachusetts on June 2-4, 1991. This workshop, while funded by the National Science Foundation and the Defense Advanced Research Projects Agency, grew directly out of research started under this contract in 1989-90. Thirty AI scientists from laboratories, universities, and funding agencies throughout the U.S. met to assess the current state of AI methodology and discuss specific methodological tactics to improve the current state. The motivation for the workshop came from our recent survey of 150 papers in the *Proceedings of the Eighth National Conference on Artificial Intelligence* (Cohen, 1991), in which we found relatively few papers that offered hypotheses, predictions, experiments, and replicable, additive, general results. Instead, we found methodological barriers between theorists and empiricists, rendering theories largely inapplicable and systems largely ad hoc. AI methodology raises structural, endogenous impediments to scientific research (Cohen, 1990b). The goal of the workshop was to identify these impediments and take the first steps to remove them.

Although the workshop focused on AI's status as a science, its weak methodology also affects our ability to produce technology. We literally do not know whether a technique, demonstrated once on a handful of examples, will work in a different task environment; we do not even know why it worked in the first place. We cannot turn techniques into technologies unless we understand why the techniques work and when they are expected to fail. In other words, we require a scientific understanding – not merely a limited empirical one – based on hypotheses and predictions about per-

formance, and replications of performance in different environments, in order to turn AI techniques into technologies. Because AI methodology does not foster this kind of understanding, AI technology suffers.

In the course of the workshop, working groups developed techniques for specific types of AI research, speakers described experimental programs in subareas of AI, and the group debated and achieved consensus on half a dozen general recommendations to the field. Considering the number of approaches and areas represented by the participants, this consensus was remarkable. The recommendations, position papers [4], talks, and other material from the workshop will be published in the near future, probably as a book.

DARPA Planning Initiative. We have just come under contract to participate in the new DARPA Planning Initiative built around a large-scale transportation planning problem. This is an ambitious program focusing the work of numerous AI researchers on a common problem. For technical guidance and organization the initiative relies on a set of Issues Working Groups, each of which is chaired by two researchers with broad experience with the issues addressed by that group. One of the groups is tasked with methodological concerns and the specification of a Common Prototyping Environment. This group will define the critical experiments that will be used to assess the initiative's progress. As a result of his longstanding concern for methodological issues (as well as considerable experience with prototyping environments), Paul Cohen has been asked to serve as co-chair of this Working Group.

References

- Anderson, S.D., Hart, D.M. & Cohen, P.R., 1991. Two ways to act. Working notes of the 1991 AAAI Spring Symposium on Integrated Intelligent Architectures. Also in *SIGART Bulletin*, 2(4): 20-24.
- Cohen, P.R., Hart, D.M. & St. Amant, R., 1992. Early warnings of plan failure, false positives and envelopes: Experiments and a model. *Proceedings of the Fourteenth Annual Conference of the Cognitive Science Society*. Lawrence Erlbaum Associates. Pp. 773-778.
- Cohen, P.R., 1991. A survey of the *Eighth National Conference on Artificial Intelligence*: Pulling together or pulling apart? *AI Magazine*, 12(1), 16-41.
- Cohen, P.R., Hart, D.M. & deVadoss, J.K., 1991. Models and experiments to probe the factors that affect plan completion times for multiple fires in Phoenix. EKSL Memo #17, Dept. of Computer Science, Univ. of Massachusetts, Amherst.
- Cohen, P.R., 1990a. Designing and analyzing strategies for Phoenix from models. *Proceedings of the Workshop on Innovative Approaches to Planning, Scheduling and Control*. Katia Sycara (Ed.). Morgan Kaufman. Pp. 9-21.
- Cohen, P.R. (Ed.), 1990b. *Working Papers of the Workshop on Artificial Intelligence Methodology*. Northampton, MA. June 2-4, 1990. Forthcoming.
- Cohen, P.R. & Howe, A.E., 1989. Toward AI research methodology: Three case studies in evaluation. *IEEE Transactions on Systems, Man and Cybernetics*, 19(3): 634-646.
- Cohen, P.R., J.L. DeLisio & Hart, D.M., 1989. A declarative representation of control knowledge. *IEEE Transactions on Systems, Man and Cybernetics*, 19(3): 546-557.
- Cohen, P.R., 1989. Why knowledge systems research is in trouble and what we can do about it. Technical Report #89-81, Dept. of Computer Science, Univ. of Massachusetts, Amherst.
- Cohen, P.R., Greenberg, M.L., Hart, D.M., Howe, A.E., 1989. Trial by fire: Understanding the design requirements for agents in complex environments. *AI Magazine*, 10(3): 34-48.
- Hansen, E.A., 1990. A Model for Wind in Phoenix. EKSL Memo #12. Dept. of Computer Science, Univ. of Massachusetts, Amherst.
- Hart, D.M., Anderson, S.D. & Cohen, P.R., 1990. Envelopes as a vehicle for improving the efficiency of plan execution. *Proceedings of the Workshop on Innovative Approaches to Planning, Scheduling and Control*. Katia Sycara (Ed.). Morgan Kaufman. Pp. 71-76.
- Howe, A.E. & Cohen, P.R., 1991. Failure recovery: A model and experiments. *Proceedings of the Ninth National Conference on Artificial Intelligence*. Morgan Kaufman. Pp. 801-808.

Howe, A.E., Hart, D.M. & Cohen, P.R., 1990. Addressing real-time constraints in the design of autonomous agents. *The Journal of Real-Time Systems*, 1(1/2): 81-97.

Howe, A.E. & Cohen, P.R., 1988. How evaluation guides AI research. *AI Magazine*, 9(4): 35-43.

Powell, G.M. & Cohen, P.R., 1990. Operational planning and monitoring with envelopes. *Proceedings of the IEEE Fifth Annual AI Systems in Government Conference*, Washington, DC.

Silvey, P.E., 1990. Phoenix baseline fire spread models. EKSL Memo #13. Dept. of Computer Science, Univ. of Massachusetts, Amherst.