

Letters

■ Editor:

My reaction to Paul Cohen's article in the Spring 1991 issue is quite positive. The research methodology he espouses is one I've tried (often unsuccessfully) to elucidate and follow. My hat is off to him for putting in words things I've tried to say—he does it quite well and backs it up with evidence. However, while I agree strongly with his analysis that we've seen two major (and different) AI methodologies emerge over the past few years, I'm not sure I agree with his main conclusion—that we should all become MAD researchers.

Let me attempt a simple elucidation of an alternative—one which I've expressed on occasion, but never formally in writing. It's basically the same as Cohen's, but instead of focusing on a single methodology (i.e. Cohen's MAD approach) we need some AI practitioners who pick up this third leg—i.e. a few MAD scientists (pardon the pun). That is, I think Cohen has realized that AI has an "excluded middle." That is, when the AAAI went to a theory and systems track, people like Cohen and some of the rest of us formed part of an excluded group that focuses on what I've called the experimental approach. We are a group of AI researchers who are not developing applications, but are not totally concerned with the formalization of their work—in fact, this is the group that Abelson called "the scruffies" in his 1981 Cognitive Science talk. (Nowadays, some folks use neats and scruffies to differentiate applications builders from theorists. However, I believe that the distinction of "neats" and "scruffies" raised at Cog Sci in '81 didn't define scruffies as people who built expert systems [they didn't really exist as a "real" part of AI at the time] but those who implemented to test theories vs. those who theorized. Thus, it was like what Cohen called systems vs. models, but the term systems had less of an applications implication than it does today).

In short, what I believe is an over-

generalization in Cohen's paper is his implication that "the" correct methodology for AI is that we should all be doing MAD. Instead, I believe AI should have three wings: the hard-core theorists (non-mon, uncertainty, etc.) who rarely pay attention to whether the things they hypothesize really exist (theoretical physics being the parallel); an engineering discipline, concerned more with the clean design of artifacts rather than the deeper principles (sort of like the people in nuclear engineering departments—they are very good scientists, but they don't worry about "fundamental principles" of nuclear reactions, but rather about building things to control them); and the middle group whose "job" it is to transition between the two—basically the experimentalists. These are the researchers who read Hawkings and say "gee, if his model of the 10^{-23} second big bang is right, then the distribution of intergalactic gases should be relatively even. Let's go see if that's true. However, to run our experiments we'll need a more sensitive space-based sensing device, so let's work with the engineers to design one." Thus, some of us MUST be practioners of MAD, but all of us had better not be (or where will the innovative theories and challenging application models come from).

Unfortunately, for reasons probably having as much to do with politics and economics as pre-scientific methodology, the current crop of researchers who would like to live in this middle ground have been largely forced into one camp or the other. The few people who live in the middle (and I feel both Cohen and I would be included) have sometimes had a harder time publishing, getting publicity, and otherwise achieving academic kudos than those at either side. There are many reasons for this, not the least is that doing a good job of joining theory and practice is hard! The work must be theoretically strong enough that it cannot be criticized as either "ad hoc" or as "for-

malizing something we all knew," but it must include enough empirical evidence that practitioners cannot shoot it down for not having been tested in a real-world domain. If Cohen is taken (overly) seriously, then every major AI researcher will have to do this—being conversant with theoretical tools (math, logic, etc.), experimental methods (testing theory, statistics, etc.), and implementation strategies (blackboard architectures, real-time systems, etc.)—a pretty tall order.

However, there is a clear solution to this sort of "overkill"—an AI community composed of several different types of practitioners. Consider how Cohen's PHOENIX project could profit if a theorist became interested in the issue of formalizing fireboss communication at the same time that a systems designer decided to provide architectural support for the application. Neither of these approaches would fly without Cohen there to provide guidance to both, but the troika could define a system of use to a user community, but well enough formalized that the usual difficult questions about "how will we know if it works" can be answered.

Finally, of course, it would be nice if the practitioners of MAD got some recognition, instead of being dumped on for being neither fish nor fowl. For that, I hope this article is widely read! I think Cohen has done a good job of defining what the goals of the "experimentalists" should be, and I hope that this will become a research manifesto for MANY projects and labs—just not ALL of them!

Anyway, this response has become longer than it should be and I'm sure that an article that attacks as many sacred cows as this one does will need a lot of space for responses. It's nice to see how clear Cohen's results were, and I congratulate *AI Magazine* for publishing this article. I look forward to seeing what others have to say.

James Hendler
University of Maryland

■ Editor:

Paul Cohen has addressed a deep, well-recognized dichotomy in AI, i.e., the battle of the scruffies versus the neats. However, he has given it a bit of a new twist, claiming that the problem is essentially methodological, and that the two camps can find a common ground in his MAD methodology. I find it quite odd to cast the differences as primarily methodological, as if each camp started with some firm methodological beliefs, and then defined progress in AI to be just whatever happens to be generated by this method. It seems much more plausible to me that the differences stem from substantive beliefs about the nature of intelligence, and about what sort of knowledge we are striving for.

In any event, I would like to suggest an alternative to Cohen's MAD methodology, and also address the question of whether the goal of a single method is a worthy one. The approach I want to suggest is the one that has been taken by the SOAR group and others working within the bounds of a cognitive architecture. The computational model instantiated by the SOAR architecture, for example, ensures that the various programs which control mobile robots, design algorithms, model visual attention, etc., have some sort of meaningful "computational common ground." The research within this community is clearly cumulative, can be described in an appropriately abstract and precise way, and addresses significant tasks. In addition to helping to bridge the neat/scruffy gap, there is the additional advantage of being able to address the thorny issues of integrating components, e.g., learning while problem-solving while interacting with the external world.

This is not intended as an advertisement for SOAR, but an argument for a research strategy centered around cognitive architectures (or, to use Allen Newell's phrase with a more psychological slant, unified theories of cognition). A given architecture could perhaps be usefully thought of as a point in a space of computational models that support intelligent behavior. This space may have many local optima, and many architectures must evolve if we are to explore any significant part of it. Research on many AI issues might benefit from this sort of strategy

(although I don't imagine that this, or any other, strategy is adequate for the whole field).

Finally, I want to express my reservations about any attempt to legislate a single method for a field of investigation. If model-builders and system-builders do not interact now, it is unclear that they can be made to do so by fiat. I tend to believe that whatever shortcomings exist in published AI research have more to do with real difficulties inherent in the subject itself rather than with the recalcitrance of researchers. I know of no other field of inquiry where it has seriously been maintained that there is a single sufficient methodology, appropriate for all problems and purposes. This is hubris.

In any event, I applaud this effort to draw attention to these important issues. As a springboard for discussion, I have little doubt of this paper's success.

James Herbsleb
Cognitive Science and Machine
Intelligence Lab
University of Michigan

■ Editor:

Paul Cohen's analysis of the Proceedings of AAAI-90 is interesting and full of insights. His feat of stamina in conducting the survey is also quite impressive. However, I do not believe that his overall impressions about the state of the field nor his concluding prescriptions necessarily follow from the survey results. But even if it isn't broke exactly as purported, many of his suggested fixes would be beneficial, and it certainly would not hurt to pay more attention to many of the methodological issues raised.

As Cohen acknowledges, a serious problem with the survey is that the AAAI conference proceedings do not accurately represent the field. However, it is not that researchers are not submitting their best work or that the reviewers are rejecting it, but the general constraints of the conference forum and length of proceedings papers. (Perhaps the field is too conference-oriented, as Paul Rosenbloom suggested at a recent conference.) One can conduct perfectly balanced model- and system-oriented research and yet exclusively write conference papers emphasizing one or the other, relying on cross-reference to link the mutual lessons drawn from the work. It is a fact of communication that in

a limited time and space it is best to package and deliver a small polished nugget if one wishes to convey any scientific lesson at all to a conference audience and proceedings readership. To get the big picture, one has to read the book or dissertation.

The benefits of combining models and systems work accrue from the results of the two enterprises informing and influencing the direction of the other. This does not require that they be conducted perfectly in tandem or even by the same researchers. Indeed, it is not at all clear to me that it would be a desirable end for all or even many AI researchers to be performing both types of research. To take Knowledge Representation as an example, I am genuinely appreciative of the work of both Lenat and Levesque, and think it would be absolute folly for either to adopt much of the approach of the other. Of course, they should (and I expect, do) read each other's papers and be influenced by them in their subsequent efforts. Division of labor is a remarkable enhancer of productivity when individuals have varying types of skills, and absolutely essential for progress in specialized scientific fields like AI.

I think one could make the case (although not from the data collected in Cohen's survey) that the two methodologies are not informed and influenced by each other to the extent they should or could be. Perhaps when Cohen reads the AAAI-91 proceedings, he can also read all the referenced papers, in order to measure the cross-influences of the varying methodologies. But while it is reasonable and laudable to call for less empty theory and more analyzable systems work, insisting on and expecting a combined methodology at the level of individual research efforts is neither realistic nor desirable.

Michael P. Wellman
Wright Laboratory AI Office
Wright-Patterson Air Force Base

■ Editor:

I'm always amused to find myself described as an "advocate" of artificial intelligence, though by now I shouldn't be, especially if by advocate is meant someone who's glad to see the field achieve partial successes, and who's still waiting hopefully for the field's scientific maturity, even

Applied AI News

Hitachi Data Systems (Santa Clara, CA) has added a download microcode enhancement to its Hi-Track expert system. The enhancement will allow Hi-Track to remotely identify and solve potential problems in a customer's storage subsystem, over the telephone.

AT&T Network Systems (Oklahoma, OK) has developed System Test History Analysis, an expert system to lower circuit pack repair expenditures and to isolate and resolve intermittent problems prior to shipment to customers. The system reviews the test history on multiple switching module configurations of digital telecommunications systems equipment.

A number of expert systems were used in support of **Operation Desert Storm**, including PRIDE (Pulse Radar Intelligent Diagnostic Environment), SABRE (Single Army Battlefield Requirements Evaluator), TOPSS (Tactical Operation Planning Support System), TACOS (The Automated Container Offering System), and AALPS (Automated Airlift Planning System).

The Knowledge Worker System, contracted by the **US Army's Construction Engineering Research Lab**, can provide assistance both to military personnel and civilians who are frequently relocated and face the challenges of searching through mountains of paperwork and manuals. The system was originally designed by Georgia Tech developers to reduce the learning curve for office workers managing the Army's \$1 billion annual construction program.

Echelon (Palo Alto, CA), in partnership with Motorola and Toshiba, has introduced the Neuron Chip, which will make intelligent control networks possible for a number of "smart" applications in homes and office applications, such as thermostats, lights and sprinkler systems. Motorola and Toshiba will produce the multiprocessor chip in support of Echelon's Local Operating Network, or LON.

Alamo Rent A Car (Fort Lauderdale, FL) has developed an expert system to set its prices nationwide for Alamo's rental cars. The embedded system analyzes the competition's prices, compares them to Alamo's, and then suggests a suitable pricing alternative.

InterVoice (Dallas, TX) has developed a help desk system that disperse the knowledge of a few experts to a group of end-user support staff through a case-based reasoning expert system. The system was developed using CBR Express, a new case-based reasoning software product from Inference (El Segundo, CA).

Cogensys (San Diego, CA), a developer of expert system software for the financial services industry, has been acquired by Cybertek (Dallas, TX), a publicly-traded life insurance software company. Cybertek is expected to incorporate Cogensys' expert system technology into a new line of expert workstations currently under development.

SRI International (Menlo Park, CA), under contract from the **Gas Research Institute** (Chicago, IL), is developing petrochemical neural network applications to analyze gas well data and predict such matters as porosity, water saturation and permeability. Preliminary results have been good: The neural networks can predict parameters with a degree of accuracy that matches the best experts.

Computer Recognition Systems (Ayer, MA) has developed a machine vision system that recognizes automobile license plates at speeds of up to 100 mph, allowing the plates to be read and specific plates to be detected. The License Plate Reading System consists of a series of cameras, mounted to cover the area to be monitored, and the software to parse the data into characters.

Letters (continued)

(someday... perhaps...) the unified theory.

However, in a 1985 book, *The Universal Machine*, I raised (much more playfully) one of the questions David M. West and Larry E. Travis raise in their important article, "The Computational Metaphor and Artificial Intelligence". There I suggested that with the notion of Thinking, AI might have gone off on the wrong track, rather like Columbus believing he'd discovered the Indies. Columbus hadn't discovered the Indies; in fact he'd stumbled on something as least as interesting; but thanks to Columbus's monomania we have here to this day a Federal Bureau of Indian Affairs. I was afraid then—I am sometimes afraid now—that artificial intelligence research is prone to limit itself the same way.

But the appearance of West and Travis's article, along with Paul Cohen's thoughtful survey of the 8th National Conference, and Stephen Sloane's case study of a failure, suggest a field that is taking a good look at itself—and incidentally make the Spring 1991 issue of *AI Magazine* one of the most interesting ever.

Pamela McCorduck
Princeton, New Jersey

■ Editor:

Stephen Slade's recent article "Case-based Reasoning: A Research Paradigm," (Spring, 1991), correctly attributes the flowchart in Figure 1 to an unpublished 1987 report by William Bain and myself.

To transfer credit where credit is due, however, I would like to note that our flowchart derives in turn from Kris Hammond's flowchart for the CHEF system, which can be found in his *Case-based Planning: Viewing Planning as a Memory Task*, from Academic Press, and Riesbeck and Shank's *Inside Case-based Reasoning*, from Lawrence Erlbaum Associates. The comment on page 50 that "CHEF's case-based planning process closely follows the flowchart in Figure 1" would be more correct with the opposite direction of causality.

Christopher K. Riesbeck
The Institute for the Learning
Sciences
Northwestern University

(Continued on page 32)