

Distributed Problem Solving and Multi-Agent Systems: Comparisons and Examples*

Edmund H. Durfee

EECS Department
University of Michigan
Ann Arbor, MI 48109

Jeffrey S. Rosenschein

Computer Science Department
Hebrew University
Jerusalem, ISRAEL

durfee@umich.edu, jeff@cs.huji.ac.il

May 27, 1994

1 Introduction

The term *multi-agent system* is currently in vogue, and has been generally applied to any system that is, or can be considered to be, composed of multiple interacting agents. In the various multi-agent (or, more properly, multiagent) systems that have been proposed or developed, a wide variety of “agents” have been considered, ranging from fully autonomous intelligent agents (such as people) down to relatively simple entities (such as rules or clusters of rules). Thus, as it has become progressively used, “multiagent systems” has come to encompass an increasingly broad variety of issues, approaches, and phenomena, to the point where now there will be a conference on multiagent systems such that one area of interest of the conference is distributed artificial intelligence (DAI).

But this was not always the case. There was a time, when the term was first coined in the AI literature, that multiagent systems referred to a more narrow branch of study

*This work was supported, in part, by the University of Michigan Rackham Graduate School International Partnerships Program, by the National Science Foundation under grant IRI-9015423 and PYI award IRI-9158473, by the Golda Meir Fellowship and the Alfassa Fund administered by the Hebrew University of Jerusalem, by the Israeli Ministry of Science and Technology (Grant 032-8284), and by the S. A. Schonbrunn Research Endowment Fund (Hebrew University) (Grant 034-7272).

within DAI. At the time that it was coined, the term served to distinguish between the prevalent research activities in DAI at the time—distributed problem solving (or cooperative distributed problem solving)—and an emerging body of work. Now, such distinctions have been lost except within a rather small segment of the community.

The overuse and abuse of the term has, unfortunately, made it more difficult to employ it to make useful distinctions of the type that it formerly did. In fact, it is quite possible that many researchers who are relatively new in the field might be only vaguely aware of the distinctions the terms “multiagent” and “distributed problem solving” once meant. Moreover, many who have been aware of these terms might have very different views as to what the distinction really is between these.

In this paper, our principal goal is to revisit these terms, work to clarify what they might mean, and encourage the community to consider useful decompositions of the broader research objectives of DAI. For that reason, the reader is forewarned that, in the bulk of the remaining paper, our use of the term “multiagent system” takes on the more narrow meaning as was first intended, and as derived from the history of the DAI field (Section 2). We then consider several views of how multiagent system (MAS) research differs from distributed problem solving (DPS) research (Section 3). Each of the views provides some insight into important questions in the field, and into different ways of solving problems and designing systems. We conclude by urging the community to not lose track of useful distinctions within the field, and to universally adopt terms to describe distinctive subfields (Section 4).

2 Historical Background

By the middle 1970s, AI research had made significant progress along several fronts, involving both weak methods and strong (knowledge-intensive) methods. Success with approaches such as production systems, where knowledge is encoded into small, manipulable chunks, had ushered in attempts to build modular systems which built off of the metaphor of cooperating specialists. Blackboard systems and ACTORS frameworks captured many of the ideas emerging at that time. Coupled with prior biologically-inspired work on neural networks, the fundamental mindset was ready for the influx of networked technology that made distributed intelligent systems a natural, promising offshoot of AI research.

2.1 The Roots of DAI: Distributed Problem Solving

Building off of the historical roots in AI, early DAI researchers adopted a similar stance to their work. Namely, given a problem to solve, how could they build systems—in this case *distributed* systems—to solve the problem. In many cases, the kinds of problems under consideration were beyond the scope of existing approaches. Early DPS work thus concentrated on harnessing and applying the power of networked systems to a problem, as exemplified by the Contract-Net approach for decomposing and allocating tasks in a network. Early DPS work also addressed harnessing the robustness available from multiple sources of expertise, multiple capabilities, and multiple perspectives. Multiple perspectives generally corresponded to problems that were inherently distributed, as exemplified in air-traffic control and vehicle monitoring domains.

In all of this work, the emphasis was on the *problem*, and how to get multiple agents to work together to solve it in a coherent, robust, and efficient manner. For example, research using the Distributed Vehicle Monitoring Testbed was concerned with how to get distributed problem solvers to work together effectively, where effectiveness was measured based on the external performance of the entire system: How long did it take to generate a map of overall vehicle movements, how much communication was involved, how much could it tolerate loss of or delays in messages, and how resilient was it to lost problem solvers? DVMT research focused on using predesigned organizations and runtime planning, goal exchange, and partial result exchange to increase the coherence of collective activity without increasing the overhead significantly. combine hypotheses channel delay and loss, task complexity, uncertainty, agent failure

2.2 Incorporating New Metaphors: Multiagent Systems

The DPS historical roots were in solving problems with computers, and so it was natural to assume that the individual agents in a DPS system, being programmed computers, could be depended on to take the actions at the right times for which they were built. But this assumption—that individuals would do as they were told—failed to completely model the social systems upon which much of DPS research was built. The literature of economics, game theory, etc. involves individuals that are not so easily programmed, and must just try to determine the conditions and cases where individuals might do as they should, and how to establish those conditions. In other words, while DPS took for granted that agents would be able to agree, share tasks, communicate truthfully, and so on, experiences in the social sciences made it clear that achieving such properties in a collection of individuals is far from simple.

If agents cannot be depended on to share, agree, and be honest, then what are some basic assumptions about agents on which to build? MAS research borrowed from the social science literature the underlying assumption that an agent should be rational: That, whatever it is doing, it should endeavor to maximize its own benefit/payoff. But what can be said about collections of rational agents as a whole? In general, not much, but an ongoing effort in MAS has been to identify under what conditions (what do agents need to know about each other, how do they interact with their environment, etc.) the agents will rationally choose to act such that the society of agents displays certain properties.

Thus, the focus of MAS has been on the *agent*, and getting it to interact meaningfully with other agents. For example, in the work of Rosenschein on deals among rational agents, the research concentrated on how self-interested agents could nonetheless converge on agreement about deals such that each could benefit. protocols different rationality assumptions

3 Relating MAS and DPS

Below are 3 views of the relationship between DPS and MAS. They are not mutually exclusive, and in fact build upon each other to some extent.

3.1 View 1: DPS is a subset of MAS

One view, not inconsistent with the more general definition that MAS has acquired over the years, is that DPS is a subset of MAS. That is, a MAS system is a DPS system when certain assumptions hold. Several such assumptions have been proposed, including the benevolence assumption, the common goals assumption, and the centralized designer assumption.

The Benevolence Assumption. One assumption that has been proposed as a touchstone for whether a system is a DPS system is whether the agents in the system are assumed to be benevolent [8]. Typically, benevolence means that the agents *want* to help each other whenever possible. For example, in the Contract Net protocol [10], agents allocate tasks to do based on suitability and availability, without any sense of agents asking “why should I want to do this task for this other agent.” Upon hearing a task announcement in the Contract Net, an eligible agent will give an honest bid on the task, indicating how well it expects to perform the task, and the agent(s) with the best bid(s) are awarded the task. There is no sense of payment—of transfer of utility—involved. Agents do not need to be bribed, bullied, or otherwise persuaded to take on tasks that others need done; they, in a sense, want to do those tasks.

Even with the benevolence assumption, cooperation and coherent coordination are far from ensured. Even though agents want to do the best they can for each other, difficulties of timing and of local perspectives can lead to uncooperative and uncoordinated activity. In the Contract Net, for example, important tasks could go unclaimed when suitable agents are busy with tasks that others could have performed, or more generally tasks could be improperly assigned so as to lead agents into redundant or incompatible activities. As in the case where people, trying to be helpful, are falling over each other and more generally getting in the way, benevolence is no assurance of cooperation.

The Common Goals Assumption. A motivation for benevolence among agents is having a common goal. That is, if the agents all value the same outcome of group activity, they will each attempt to contribute in whatever way they can to the global outcome. This assumption could be considered to be at the heart of Contract Net, and is also arguably at the core of cooperation in inherently distributed tasks, such as distributed interpretation tasks, where agents each value the development of a global result. In the DVMT [5], for example, the system-wide goal was for the distributed sensing nodes to integrate their local maps of vehicle movements into a global map of vehicle movements. Since each node is trying to help the system converge on the global solution, each is trying to help the others form good local interpretations as quickly as possible.

However, once again, local views of the problem to be solved can lead to local decisions that are incoherent globally. Without strict guidelines about responsibilities or interests, agents can inundate each other with superfluous information. Worse, agents can work at cross purposes and send information that can distract others into pursuing unimportant tasks [2].

Also unclear is the level at which goals should be common to make a system a DPS system. If agents are meeting to hold a competition, then they might share a high-level goal

of holding the competition while having opposing goals as to who is supposed to win the competition. Similarly, in a situation like that studied by Sycara in her PERSUADER system [11], where the agents are representing opposing sides in a labor contract, the agents share a goal of reaching an agreement (forming a contract) while having very diverse preferences in rating candidate contracts. Is this a DPS system?

The Centralized Designer Assumption Most recently, the argument has been put forward that a DPS system is a system with a centralized designer. This perspective subsumes the previous assumptions, since the central designer’s goals would be embodied in the agents (giving them common goals) and the designer, being concerned about getting the parts to work as a whole, would likely make each agent benevolent. Moreover, since the designer has the “big picture”, the preferences/foci of the agents could be calibrated and the mechanisms for expressing and acting on those preferences can be standardized.

The open question here, as in the case of the common goals assumption, is to what detail must the common designer specify the agent design to make them into a DPS system? Is any commonality in design sufficient? Is identical design down to the smallest detail necessary? For example, in the case of social laws [9], if we assume that all agents are constrained in their design to follow the laws laid down by a common designer (a legislative body, perhaps), then does this make them a DPS? Or is the fact that the society is “open” in the sense that very different agents could come and go, so long as each is a law-abider, and thus agents could have very different (law-abiding) plans and goals, grounds for defining the system as MAS? Until we define exactly what aspects of agents need be dictated by a common designer to make the system a DPS, this approach to categorizing systems will likely remain arbitrary.

Evaluation of View 1. In summary, this first view, that DPS is a subset of MAS, has arguments in its favor but suffers from some disadvantages. One of these, described above, is that there appears to be a slippery slope between the two types of systems: That, as the commonality of goals and/or designers is developed to increasingly detailed levels, the system is more of a DPS system, but there is no clear dividing line for exactly when the transition to DPS occurs. A second, related disadvantage is that, because agents in DPS systems can behave a cross purposes, distract each other, and commit other such non-cooperative acts, an observer of a system could not be able to tell whether the system is DPS or MAS just by watching the agents’ behavior. Thus, without being able to either look inside of the agents to identify common goals, or being able to see the design process leading up to the system, classifying DPS and MAS using this first view is not possible.

Of course, it could be that the DPS and MAS character really is not an inherent property of the system, but rather of the context in which the system has been developed. We return to this perspective in View 3.

3.2 View 2: MAS provides a substrate for DPS

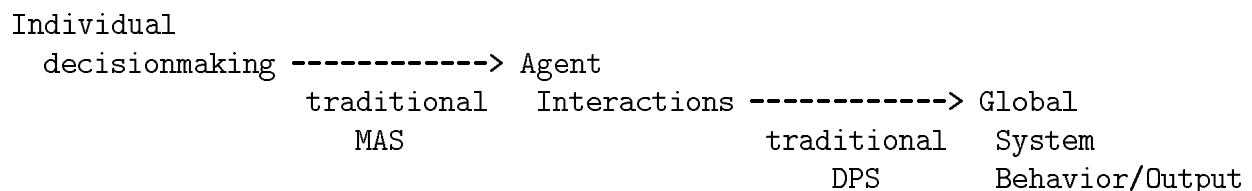
Traditional DPS research has taken, as its starting point, that internal properties of the system can be assumed (generally, designed in). These properties can include that agents will be truthful in their communications, that they will follow defined communication protocols,

that they will perform tasks as promised, that they will promise to accomplish tasks when asked to and when they are able to, and so on.

Assuming these internal properties, DPS is generally concerned with how the system can demonstrate certain desirable external properties. Typically, these external properties are generating appropriate solutions to instances of "the problem" (that motivated the construction of the system in the first place). These instances could involve different tasks/environment combinations, including initial distributions of tasks, task arrival times, agent/communication failures, communication delays, etc. For example, in the DVMT, the property most often measured for the system was the response time: how long the system of agents took to generate a correct hypothesis of vehicle movements through the sensed area. A coordination strategy was generally considered successful if the network of agents could successfully track vehicles with good response time despite losses of messages, varying input data (including noisy data), and even failures of some agents.

While DPS thus (generally) assumes that whatever internal properties desired of the system can be instilled, MAS (generally) is concerned with how to instill these properties in the first place. That is, MAS generally only makes assumptions about the properties of individuals (most typically, that they are rational utility-maximizers), and considers what properties will emerge internally among agents given the incentives (payoffs) and features of their environment. MAS research can thus define incentive structures (as in Clarke Tax mechanisms [4]) or environmental features (as in the ability to discover or conceal lies [12]) that either exist naturally or can be imposed such that desired internal properties (such as truth telling or fair access to resources) are achieved.

Thus, this view takes a divide and conquer approach to DAI research. Rather than trying to jump all the way from how individual, self-interested decisionmaking on the part of each agent could lead to an overall system that accomplishes some desirable task, we can divide the problem up. MAS studies how individual, self-interested decisionmakers might discover (or be coerced into) stable, predictable, and desirable ways of interacting among themselves. DPS then considers how these dependable, desirable interactions can be initiated, controlled, and otherwise exploited to yield a system that accomplishes some externally-defined goal. This can be graphically depicted as:



Note, finally, that (once again) specific research projects might blur this decomposition. For example, a MAS approach might, while concentrating on internal properties of the collection of agents, could also have some overall external behavior that can be measured and evaluated (such as maximizing global efficiency in the Postal problem [12]). Typically, though, the external problem is an abstracted/simplified version of a real problem because it is not the main emphasis of the goals of the mechanisms. Similarly, a DPS system, while concentrating on robustly accomplishing an externally-imposed task, might also allow variations on internal properties, such as Corkill's work on externally-directed versus internally-directed nodes in the DVMT [1]. The variations, however, are generally quite limited.

3.3 View 3: MAS and DPS are complementary research agendas

Implicit in View 2 is that the kinds of questions/problems asked by MAS researchers are somewhat different from those asked by DPS researchers. This leads to the view that MAS and DPS are really labels not for particular kinds of systems, but rather for research agendas. As mentioned before, chances are that an observer of a system would not be able to classify it as MAS or DPS based on its observable behavior. But a system could be part of either stream of research depending on how it is experimented with.

As one example among many, consider the problem of generating and executing plans for multiple agents in a decentralized manner. Several systems have been developed for this problem, and there are many similarities among them, but what makes them part of different research agendas is the kinds of questions the researchers ask when developing, analyzing, and experimenting with these systems.

As one example, in the partial global planning approach, agents dynamically and reactively construct local plans in response to changes in their task environments and to changes in what they know about other agents' plans. Thus, at any time, an agent will have a model of the collective plans of subsets of agents (called partial global plans) and will modify its own activities appropriately, assuming that others will modify their activities in compatible ways. The research questions focused on in this work concerned issues such as how efficiently the agents could coordinate their plans, how robust the approach was to agent or network failures, how performance is impacted by message delays or losses, and so on.

The work of Ephrati and Rosenschein [4] addresses a similar task, namely, how can agents converge on a joint plan to achieve goals when each is constructing pieces of that plan locally. Like partial global planning, their approach involves an exchange of information about aspects of local plans and the integration of that information to identify relationships which in turn lead to changes in local plans. But at this point, the research agenda diverges from that of partial global planning. The main questions that Ephrati and Rosenschein ask are questions about the extent to which their approach will perform appropriately in the face of manipulative, insincere agents, ignoring environmental concerns such as communications failures and delays.

Thus, in contrast to the second view above which sees the MAS and DPS fields as fitting together into a chain that leads from individuals to an overall system, this view sees them as starting from a common beginning point but varying different parameters in the exploration of the field. At the risk of oversimplification, let's try to nail this down more precisely. Let's say that distributed AI involves *agents* who act in an *environment* to comprise a *system*. So for each of these three, we can talk about their properties:

- Agent properties: What can we say about an individual agent? Is it rational? Are its preferences common knowledge? Are they even shared? What are its capabilities? Are these known to others? And so on.
- Environment properties: What can we say about the environment? Is it static? Closed? Benign? Are outcomes of actions taken by agents predictable? Temporally bounded?
- System properties: What can we say about the overall agent/environment system? Does the system assure certain internal properties, such as fair access to resources

	Agent Properties	Environ Properties	System Properties
MAS	variable	fixed	fixed (internal)
DPS	fixed	variable	fixed (external)
?	fixed	fixed	variable

Table 1: Matrix of Research Agendas and Properties to Vary.

or honesty among agents? Does it assure certain external properties, such as timely output responses to system inputs?

With these three kinds of properties, we can summarize MAS and DPS as in Table 1.

In words:

MAS corresponds to a research agenda that has focused on getting certain internal properties in a system of agents whose individual properties can vary. Thus, MAS has been concerned with how agents with individual preferences will interact in particular environments such that each will consent to act in a way that leads to desired global properties. MAS asks how, for a particular environment, can certain collective system properties be realized if the properties of agents can vary uncontrollably.

DPS has focused on getting external properties such as robust and efficient performance, under varying environmental conditions, from agents with established properties. DPS asks how can a particular collection of agents attain some level of collective performance if the properties of their environment are dynamic and uncontrollable.

Again, this categorization should be taken with a grain of salt. Most systems will not fall neatly into one pure camp or the other. Nevertheless, it appears likely that any DAI research project will fall predominantly into one side or the other, since varying too many parameters at once prohibits a systematic scientific investigation.

Finally, it is an open question as to how to label the remaining category of research, where agent and environment properties are fixed but system properties vary. We might speculate that some work in artificial life, or even neural networks, can fall into this category, but this deserves more careful thinking.

4 Discussion

Now we have identified three possible views of the relationship between MAS and DPS. To summarize, what these views have in common is that none of them talk about observable properties of the systems so much as the systems in a larger context of endeavor:

view 1 focuses on *who* made the system (by a single designer or designers with shared goals?);

view 2 focuses on *how* the system was made (were individuals thrown together or were team interactions relied upon?); and

view 3 focuses on *why* the system was made (was it made to ask questions about the impact of changing environment or of changing agent population?).

Which of these views is correct? That is a question for the community as a whole to debate and answer. Is the distinction important? Again, that is a matter of opinion, but the advantages of being able to identify relevant technology based on the desired properties of a target system seem strong. For example, when faced with the task of designing a distributed AI system for monitoring and managing a computer network, where the prime measure of performance is to minimize user complaints, and where the implementation is under the control of the designer(s), techniques borrowed from the DPS side of the line can be fruitfully employed [6, 7]. On the other hand, when faced with an open system where standard task-level protocols among agents are brittle or undefined, allowing interaction patterns and protocols to emerge from first principles (agent preferences, abilities, and rationality) in an MAS manner is a promising approach [3].

If we conclude that the distinction is important, how can we build on this distinction to help map out the larger field of study, to encourage the systematic exploration of issues in the DAI field?

Finally, returning to the point that began this paper, given that the term multi-agent systems has now been generalized to be a superset of DAI systems (and, seemingly, most systems), we have lost a useful distinguishing term in the field. A goal of this paper, therefore, is to stimulate the field to identify and adopt terms that have concrete meaning in the field, beginning with the DPS and MAS distinctions, as well as to prompt more general discussion about the utility of characterizing the field in this way and to incite debate over which viewpoint on DPS and MAS, if any, is correct.

References

- [1] Daniel D. Corkill and Victor R. Lesser. The use of meta-level control for coordination in a distributed problem solving network. In *Proceedings of the Eighth International Joint Conference on Artificial Intelligence*, pages 748–756, Karlsruhe, Federal Republic of Germany, August 1983. (Also appeared in *Computer Architectures for Artificial Intelligence Applications*, Benjamin W. Wah and G.-J. Li, editors, IEEE Computer Society Press, pages 507–515, 1986).
- [2] Daniel David Corkill. *A Framework for Organizational Self-Design in Distributed Problem Solving Networks*. PhD thesis, University of Massachusetts, February 1983. (Also published as Technical Report 82-33, Department of Computer and Information Science, University of Massachusetts, Amherst, Massachusetts 01003, December 1982.).
- [3] Edmund H. Durfee, Piotr J. Gmytrasiewicz, and Jeffrey S. Rosenschein. The utility of embedded communications and the emergence of protocols. In *Submitted to the 1994 Distributed AI Workshop*, 1984.
- [4] Eithan Ephrati and Jeffrey S. Rosenschein. Multi-agent planning as a dynamic search for social consensus. In *Proceedings of the Thirteenth International Joint Conference on Artificial Intelligence*, August 1993.
- [5] Victor R. Lesser and Daniel D. Corkill. The Distributed Vehicle Monitoring Testbed: A tool for investigating distributed problem solving networks. *AI Magazine*, 4(3):15–

- 33, Fall 1983. (Also published in *Blackboard Systems*, Robert S. Englemore and Anthony Morgan, editors, pages 353–386, Addison-Wesley, 1988 and in *Readings from AI Magazine: Volumes 1–5*, Robert Englemore, editor, pages 69–85, AAAI, Menlo Park, California, 1988).
- [6] Young pa So and Edmund H. Durfee. Distributed big brother. In *Proceedings of the Eighth IEEE Conference on AI Applications*, March 1992.
 - [7] Young pa So and Edmund H. Durfee. A distributed problem-solving infrastructure for computer network management. *International Journal of Intelligent and Cooperative Information Systems*, 1(2):363–392, 1992.
 - [8] Jeffrey S. Rosenschein and Michael R. Genesereth. Deals among rational agents. In *Proceedings of the Ninth International Joint Conference on Artificial Intelligence*, pages 91–99, Los Angeles, California, August 1985. (Also published in *Readings in Distributed Artificial Intelligence*, Alan H. Bond and Les Gasser, editors, pages 227–234, Morgan Kaufmann, 1988.).
 - [9] Yoav Shoham and Moshe Tennenholtz. On the synthesis of useful social laws for artificial agents societies (preliminary report). In *Proceedings of the Tenth National Conference on Artificial Intelligence*, July 1992.
 - [10] Reid G. Smith. The contract net protocol: High-level communication and control in a distributed problem solver. *IEEE Transactions on Computers*, C-29(12):1104–1113, December 1980.
 - [11] Katia Sycara-Cyranski. Arguments of persuasion in labour mediation. In *Proceedings of the Ninth International Joint Conference on Artificial Intelligence*, pages 294–296, Los Angeles, California, August 1985.
 - [12] Gilad Zlotkin and Jeffrey S. Rosenschein. Cooperation and conflict resolution via negotiation among autonomous agents in non-cooperative domains. *IEEE Transactions on Systems, Man, and Cybernetics*, 21(6), December 1991. (Special Issue on Distributed AI).