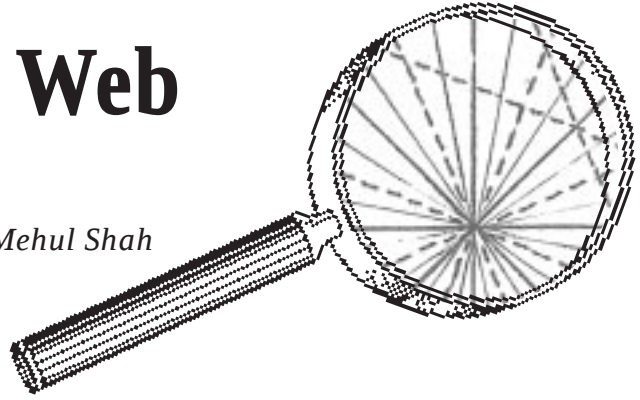


The Hidden Web

Henry Kautz, Bart Selman, and Mehul Shah



■ The difficulty of finding information on the World Wide Web by browsing hypertext documents has led to the development and deployment of various search engines and indexing techniques. However, many information-gathering tasks are better handled by finding a referral to a human expert rather than by simply interacting with online information sources. A personal referral allows a user to judge the quality of the information he or she is receiving as well as to potentially obtain information that is deliberately not made public. The process of finding an expert who is both reliable and likely to respond to the user can be viewed as a search through the network of social relationships between individuals as opposed to a search through the network of hypertext documents. The goal of the REFERRAL WEB Project is to create models of social networks by data mining the web and develop tools that use the models to assist in locating experts and related information search and evaluation tasks.

The vast network of linked documents that make up the World Wide Web (WWW) is only one manifestation of a larger and more profound phenomenon; namely, the social network that links all people. In the 1960s, Stanley Milgram's (1967) pioneering work on the small-world problem showed that any two randomly chosen individuals in the United States are linked by a chain of six or fewer first-name acquaintances. Researchers have replicated this phenomenon in many contexts, ranging from groups of workers within a corporation (Wasserman and Galaskiewicz 1994; Krackhardt and Hanson 1993) to patterns of e-mail communication around the world (Schwartz and Wood 1993). The idea that everyone is connected this way has been popularized by the phrase "six degrees of separation" as well as by games based on tracing the networks around celebrities.

A well-known pastime in the mathematics community is determining an individual's

Erdos number, which is the number of links between the person and the famous mathematician Paul Erdos, where a link corresponds to coauthorship of a published paper (Grossman and Ion 1995). Although this game was inspired by the fact that Erdos was amazingly prolific, similar graphs can be created in many domains that connect any individual, no matter how obscure, with practically all others in only a few links. For example, according to the web site entitled the "Oracle of Bacon at Virginia," over 100,000 actors are within 3 links of Kevin Bacon using the costar relationship.

However, the task of constructing and analyzing social networks is more than just a game. Numerous studies have shown that one of the most effective channels for dissemination of information and expertise within an organization is its informal network of collaborators, colleagues, and friends (Galegher, Krautz, and Edigo 1990; Granovetter 1973). Indeed, the social network is at least as important as the official organizational structure for tasks ranging from immediate, local problem solving (for example, fixing a piece of equipment) to primary work functions (such as creating project teams).

The Hidden Web

Both the WWW and corporate intranets are often cited as tools that can help organizations function more efficiently by making it easier for individuals to directly access information sources. However, almost as soon as the web entered the public consciousness, a major problem was recognized in achieving its promise of democratizing information access, namely, the difficulty a person has in navigating its complex network of hyperlinks. The response has been an arms race in the development of competing web-indexing engines. Although the general problem of

The goal of our REFERRAL WEB Project is to create an interactive tool to help people explore the social networks in which they participate, so that they can quickly find short referral chains between themselves and experts on arbitrary topics.

quickly finding documents on the WWW or an intranet is far from solved, the information-retrieval techniques used by the search engines are good enough to make the web of, at least, some serious use and to maintain the current momentum in its growth.

However, a deeper problem remains that no solution based solely on building a better search engine can address. Much valuable information is deliberately not kept online. Indeed, part of the value of a piece of information resides in the degree to which it is not easily accessible, which is perhaps most obvious with regard to proprietary corporate information. For example, if I am involved in trading General Motors stock, I might be vitally interested in knowing the specifications of the cars to be introduced next year. That such information exists is certain—indeed, factories are already being set up to produce the vehicles—but I am certainly not going to be able to find this information in any database to which I have access. Even when economic imperatives are not apparent, factors of a practical or social nature can play a role in restricting the availability of online information. An expert in a particular field is almost certainly unable to write down all he/she knows about the topic and is likely to be unwilling to make letters of recommendation publicly available that he/she has written for various people. For example, suppose I were trying to decide whether to hire a certain Professor Smith to do some consulting for my company. It is certainly not the case that I would be able to access a database of academics, with entries such as “solid, careful thinker” or “intellectually dishonest.”

Thus, the search for information often must come down to the search for a person who holds the information privately. Furthermore, it is often not possible to simply identify the individual by searching, for example, personal home pages or internal corporate directories of experts. One needs to have some way to evaluate the expert and motivate the expert to respond. In this situation, searching for a piece of information becomes a matter of searching the social network—the *hidden web*—for an expert on the topic as well as providing a chain of personal referrals from the searcher to the expert. The referral chain serves two key functions: First, it provides a reason for the expert to agree to respond to the requester by explicating their relationship (for example, they have a mutual collaborator). Second, it provides a criteria for the searcher to use in evaluating the trustworthiness of the expert.

Let us consider for a moment how locating an expert on a given topic actually works when it is successful. In a typical case, I contact a small set of colleagues whom I think might be familiar with the topic. Because each person knows me personally, he/she is quite likely to respond. Usually, none of them is exactly the person I want; however, they can refer me to someone they know who might be. After following a chain of referrals a few layers deep, I finally find the person I want.

Manually searching for a referral chain is time consuming and failure prone. When the process fails, it is not because the social network is unconnected but because the chain dies out before reaching an expert. In fact, in Milgram’s original experiments, failed chains often terminated with a person who lived a few blocks from the intended target’s residence. Our own simulation experiments, which we describe here, confirmed that responsiveness is the most important factor in successful referral chaining.

Automating Referral Chaining

How can the efficiency of referral chaining be improved? Milgram’s experiments involved hand delivering a letter without making any copies of it; thus, it involved a search of width one. Electronic mail makes it trivially easy to increase the fan out of the search at each step. Unfortunately, messages that are broadcast to large numbers of people or that appear to be part of a mail pyramid are quickly categorized by most recipients as junk mail and, thus, are ignored. News groups have traditionally functioned as a way to contact a large community in a less obtrusive fashion. However, with the growth of the internet, the social ties between people who are only linked by virtue of being participants in the same news group have become so tenuous that one can no longer count on receiving any responses to posted queries.

The goal of our REFERRAL WEB Project is to create an interactive tool to help people explore the social networks in which they participate, so that they can quickly find short referral chains between themselves and experts on arbitrary topics. We wanted a system that was nonobtrusive and easy to deploy. This desire ruled out approaches based on sending e-mail to large numbers of individuals and those that required all the members of the community to explicitly register information about themselves or install special software on their computer systems.

REFERRAL WEB operates by automatically generating representations of social networks based on evidence gathered from publicly available documents on the WWW. As we describe in the next section, REFERRAL WEB allows users to quickly search for chains in the networks to either named individuals or, more generally, to people who are likely to be experts on a given topic. Typically, a user is only aware of a portion of his or her social network. By instantiating the larger community, the user can discover connections to people and information that would otherwise lay hidden over the horizon.

Before presenting the REFERRAL WEB system in more detail, let us first evaluate the feasibility of any approach to automating referral chaining. The fact that manual referral chaining works even as well as it does indicates that people are generally quite accurate in making referrals, if not directly to the desired expert, at least to someone with a better knowledge of the area. An automated system might be expected to be more responsive, but it is also likely to be less accurate. There will be errors in the model of the social network it creates, both in terms of links between individuals and in terms of the association of individuals with topics of expertise. Therefore, we ran a series of simulation experiments to address the general question of whether responsiveness can effectively be traded for loss of accuracy in an automated referral system.

In our simulation, we consider a communication network consisting of a graph, where each node represents a person, and each link models a personal contact. In the referral-chaining process, messages are sent along the links from node to node. A subset of nodes is designated the *expert nodes*, which means that the sought-after information resides at these nodes. In each simulation cycle, we randomly designate a node to be the requester; the referral-chaining process starts from this node. The goal is to find a series of links that leads from the requester node to an expert node.

First, we model the fact that the further removed a node is from an expert node, the less accurate the referral will be. Why we expect the accuracy of referral to decrease is best illustrated with an example. Assume that one of your colleagues has the requested expertise. In this case, there is a good chance that you know this to be true and can provide the correct referral. However, if you do not know the expert personally but know of someone who knows him or her, then you might still be able to refer to the correct per-

son (in this case, your colleague who knows the expert), but you'll probably be less likely to give the correct referral than if you had known the expert personally. In general, the more steps away from the expert, the less likely it is that you will provide a referral in the right direction. To model this effect, we introduce an accuracy of referral factor A (a number between 0.0 and 1.0). If a node is d steps away from the nearest expert node, it will refer in the direction of the node with a probability of $p(A, d) = A^{\alpha d}$, where α is a fixed scaling factor. With a probability of $1 - p(A, d)$, the request will be referred to a random neighbor (that is, an inaccurate referral).

Similarly, we model the fact that the further removed a node is from the requester node, the less likely it is that the node will respond. This aspect is modeled using a responsiveness factor R (again between 1.0 and 0.0). If a node is d links removed from the originator of the request, its probability of responding will be $p(R, d) = R^{\beta d}$, where β is a scaling factor. Finally, we use a branching factor B to represent the average number of neighbors that will be contacted, or referred to, by a node during the referral-chaining process.

Table 1 presents the results of a typical simulation experiment. These results are plotted in figure 1. The general observations here hold across a wide setting of parameter values. The network consists of a randomly generated graph with 100,000 nodes. Each node is linked on average to 20 other nodes. Five randomly chosen nodes are marked as expert nodes. The average branching B , while referral chaining, is fixed at 3. The scaling factors α and β are fixed at 0.5. The table shows the results for a range of values for R and A . The values are based on an average of more than 1000 information requests for each parameter setting.

The final column in table 1 gives the *success rate* of the referral-chaining process, that is, the chance that an expert node is found for a given request. We also included the average number of messages sent when processing a request. Note that because of the decay in responsiveness, at some distance from the requester node, the referral process will die out by itself. Thus, the processing of a request ends when either the chain dies out, or an expert source node is reached. The third column shows how close, on average, a request came to an expert node. The distance is in terms of the number of links, where in a successful referral run, the final distance is zero. On average, in the network under considera-

The fact that manual referral chaining works even as well as it does indicates that people are generally quite accurate in making referrals, if not directly to the desired expert, at least to someone with a better knowledge of the area.

R	A	Average Final Distance	Average Number of Messages	Success Rate
1.00	1.00	0.0	12.1	100.0
1.00	0.70	0.0	42.6	99.7
1.00	0.50	0.0	109.8	99.1
1.00	0.30	0.1	377.2	94.6
1.00	0.10	0.6	1225.2	52.5
1.00	0.00	1.2	1974.6	8.5
0.90	0.90	0.1	21.3	95.2
0.90	0.70	0.2	43.8	92.0
0.90	0.50	0.3	92.1	83.1
0.90	0.30	0.7	204.9	53.2
0.90	0.10	1.5	316.3	13.3
0.70	0.90	0.7	16.0	60.3
0.70	0.70	1.1	22.3	39.0
0.70	0.50	1.7	27.5	19.6
0.70	0.30	2.2	28.3	6.8
0.70	0.10	2.5	31.9	1.2
0.50	0.90	1.4	9.5	26.6
0.50	0.70	1.9	10.5	11.1
0.50	0.50	2.3	10.8	5.2
0.50	0.30	2.6	10.7	3.1
0.50	0.10	2.9	11.1	0.2
0.10	1.00	2.2	3.3	1.3
0.10	0.70	2.7	3.5	0.3
0.10	0.50	2.9	3.5	1.1
0.10	0.10	3.2	3.5	0.1

Table 1. Simulation Results for a 100,000-Node Network.

tion here, the length of the shortest chain between a random node and the nearest expert node is about four links.

Our simulation results reveal several surprising aspects of the referral-chaining process. Let us first consider the extreme cases. With $R = 1.0$ and $A = 1.0$, that is, total responsiveness and perfect referral, we have a 100-percent success rate; on average, it takes about 12 messages to process a request. Now, if we reduce the accuracy of the referral to 0.0 (that is, nodes will simply refer to random neighbors), our success rate drops to 8.5 percent, and almost 2000 messages are sent for each request. (With maximum responsiveness, the referral process would only halt when an expert node is found, giving a success rate of 100 percent by saturating the network. However, in our simulation, we used a maximum referral-chain length of 10. This limit was only reached when R was set to 1.0, and almost 2000 messages are used for each request. There is a reasonably intuitive explanation for the large number of messages and the low success rate. Without any referral information, the search basically proceeds in

a random manner. With 100,000 nodes and 5 expert nodes, we have 1 expert for every 20,000 nodes. In a random search, a request would visit, on average, about 20,000 nodes to locate an expert. Given that the search only involves 2000 messages, the measured value of 8.5 percent is intuitively plausible. Note that our argument is only meant to provide some intuition behind the numbers. A rigorous analysis is much more involved. See, for example, Hedetniemi, Hedetniemi, and Liestman (1988).

Of course, in practice, we have neither perfect referral nor complete responsiveness, and we are particularly interested in the trade-off between the two parameters. In ordinary person-to-person referral chaining, one would expect to have low responsiveness but high accuracy. For example, at $R = 0.5$ and $A = 0.9$, one has a success rate of about 27 percent, with about 10 messages for each request. However, an automated system might have a high responsiveness of $R = 0.9$ but a low accuracy of $A = 0.5$. That is, half the time, the referral (or, equivalently, the link in the model of the network) is inaccurate. In this case,

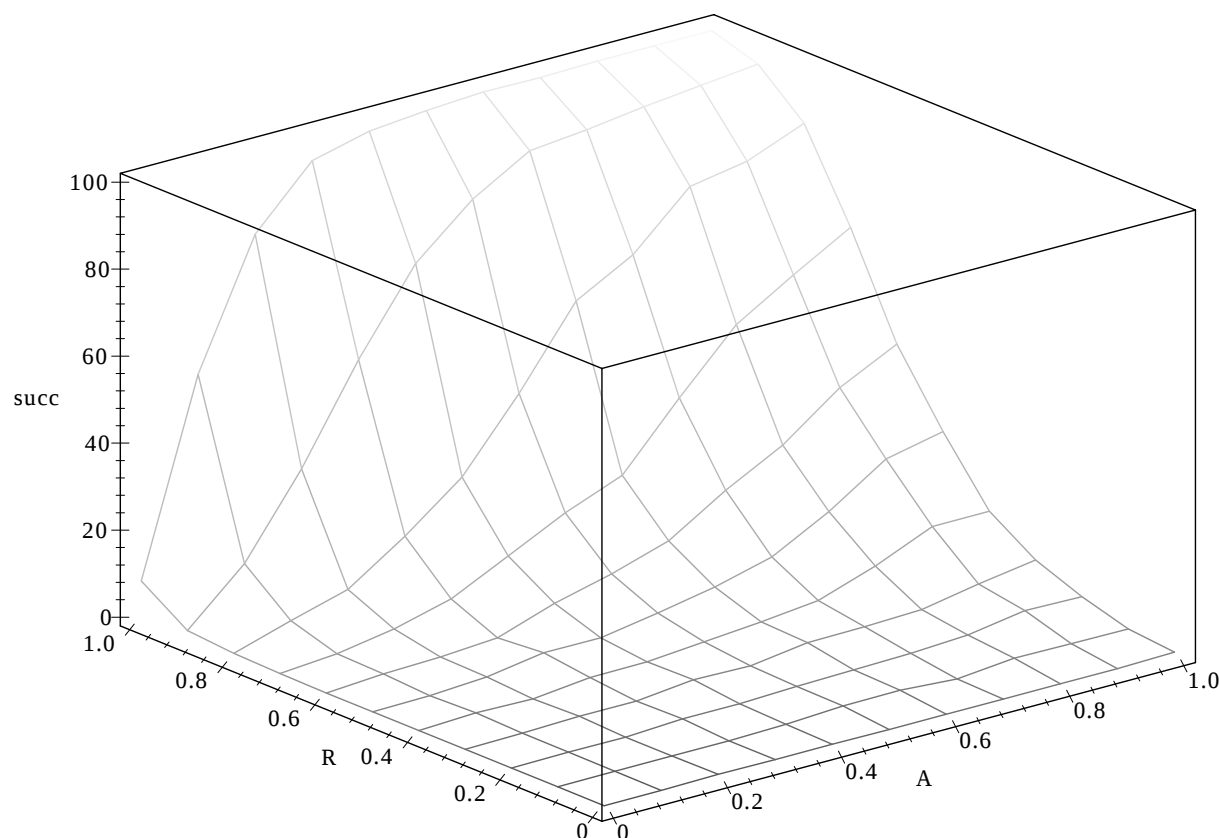


Figure 1. Success Rate as a Function of Responsiveness and Referral Accuracy.

the success rate is 83 percent, and about 92 messages are processed for each request. These numbers suggest that it might be possible to match or exceed the success rate of manual referral chaining by using an automated service, albeit at the cost of a modest increase in the number of messages sent or, equivalently, in the number of nodes searched. However, the 92 messages should be compared to the approximately 2000 messages that would be required in a search without any referral information that has a much lower success rate. If the limitation on the length of referral chains were eliminated, and nodes were simply visited randomly, then a probabilistic argument shows that at least 30,000 nodes would need to be visited to achieve the same 83-percent success rate. Thus, even with limited accuracy, an automated referral-chaining process can be surprisingly successful.

Reconstructing and Querying Social Networks

We now describe the actual REFERRAL WEB system. A social network is modeled by a graph, where the nodes represent individuals, and an edge between nodes indicates that a direct relationship between the individuals has been discovered. An optional weight on the edges indicates the degree of association between the individuals. There are many possible sources for determining direct relationships. At one extreme, which we reject as too burdensome, users could be required to enter lists of close colleagues. Analysis of e-mail logs provides a rich source of relationships, as shown by Schwartz and Wood (1993). In fact, the initial version of our system derived its network by analyzing mail archives (Kautz, Selman, and Milewski 1996). However, the use of such information raises concerns of

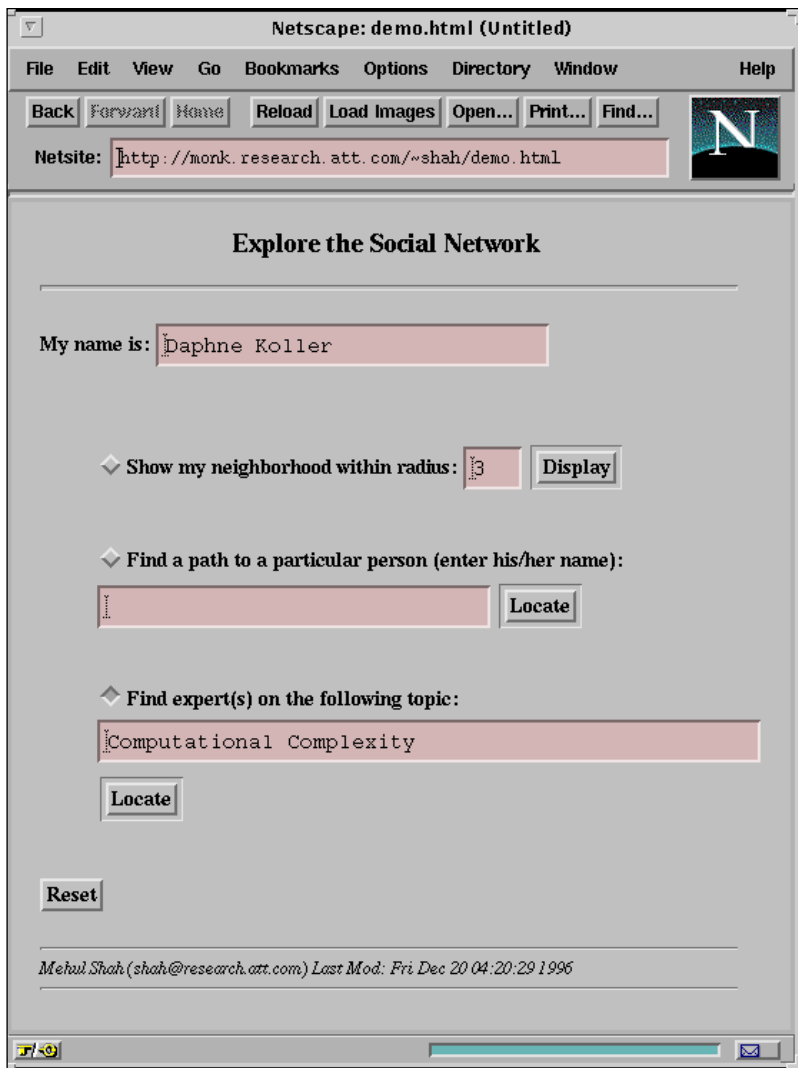


Figure 2. Example of a Query Sent to REFERRAL WEB.

privacy and security that are hard to allay. The current REFERRAL WEB system uses the co-occurrence of names in close proximity in any documents publicly available on the WWW as evidence of a direct relationship. Such sources include links found on home pages, lists of coauthors in technical papers and citations of papers, exchanges between individuals recorded in net-news archives, and organization charts (for example, for university departments).

The network model is constructed incrementally. When a user first registers with the system, it uses a general search engine (currently, Digital Equipment Corporation's ALTAVISTA) to retrieve web documents that mention him or her. The names of other individuals are extracted from the documents,

which can be done with high accuracy (better than 90 percent) using techniques such as those described in Sundheim and Grishman (1995). The result is typically a list of several hundred possible links for an individual. This list is then refined by computing the Jaccard coefficient (Salton 1989) between the user and each member of the list and then applying a threshold to this value. The Jaccard coefficient is simply the number of pages indexed by ALTAVISTA that contain both names divided by the number of pages that contain either name. This process is applied recursively for one or two levels, and the result is merged into the global network model.

Note that the user does not enter any personal information into the system. Furthermore, the network model contains hundreds or thousands of individuals who have never registered with, or even accessed, REFERRAL WEB. Contrast this approach with most other recommender systems on the web, such as FIREFLY from Agents, Inc. These recommender systems can typically only access information regarding the particular community of users who have explicitly interacted with the system.

When started with an empty knowledge base, the REFERRAL WEB spider takes on the order of 24 hours to generate a network of radius three from a given name. The speed of the spider could be increased dramatically if it had direct access to a full-web index; we discuss later other, more limited sources of information that can be processed more rapidly. However, as the network model grows, less and less work is required for each new individual because it becomes more likely that a large part of the neighborhood around the person already has been generated. In fact, if a user appears at all in the current model, he or she can immediately begin to use the system, but the spider runs offline. Because the initial audience for REFERRAL WEB is the AI research community, we have tried to increase the likelihood that the system is immediately useful by "seeding" the system with the names of various AI researchers.

The network model is then used to guide the search for people or documents in response to user queries. Most simply, a user interactively explores a graphic representation of the portion of the network centered on himself or herself. Alternatively, a person can ask to find the chain between himself or herself and a named individual. For example, a query might be, "What is my relationship to Marvin Minsky?" To search for an expert, the user can specify both a topic and a social radius; for example, a query might be, "What

colleagues of mine, or colleagues of colleagues of mine, know about simulated annealing?”

Figure 2 shows the interface for entering a query on our prototype system. (The uniform resource locator [URL] shown in the browser and the details of the interface might differ from the version that will be deployed on the web at the time this article is published.) We entered as the user name *Daphne Koller*, a professor at Stanford University who was not involved in this project. Her name already appears in the network model because we started the spider with the names of the authors of this article. We assume that Koller chooses to explore the network model immediately, without waiting for the spider to explicitly explore her local neighborhood.

The query asks to find an expert on computational complexity. Figure 3 shows REFERRAL WEB's response. Two experts were found: (1) Hector Levesque and (2) Robert Constable. The chains from Koller to the experts, Koller-Halpern-Selman-Levesque and Koller-Halpern-Clarke-Constable, are highlighted.

Most of the links that appear in this response are accurate, although the graph is incomplete. For example, most of the linked individuals have coauthored papers, and one has served as the Ph.D. adviser of the other. Halpern and Selman have never coauthored a paper, but their names have often appeared together on program committees and near each other in bibliography files. In fact, the names Koller and Selman also co-occur on the web, but the strength of the link between them (as measured by the Jaccard coefficient) is below the threshold; so, the link does not appear in the graph.

This information could be used in several ways: For example, the user could directly contact one of the experts and mention how he or she is linked to encourage a response. Alternatively, the chain could be used to evaluate which expert should be contacted. For example, if Koller knows that Selman is in AI, but Clarke is in theory (a fact she could also learn by contacting Halpern), then she could guess that Levesque would be a better expert to contact for applications of computational complexity to problems in AI.

We are currently exploring several mechanisms for associating individuals with topics of expertise. In the example shown in figure 3, the topics were taken to be capitalized phrases that appeared in documents that were retrieved by the spider for the individual but were not proper names. Alternatives we are experimenting with include full-text



indexing of all the retrieved pages and using ALTAVISTA to dynamically search for experts that the network models indicate are close to the user.

One can imagine other ways to use the social network model. Another kind of query that we plan to implement takes advantage of a designated, known expert to control the search. For example, one might ask to “list documents on the topic *computational complexity* by people close to Bart Selman.” It is important to emphasize that REFERRAL WEB does not replace generic search engines such as ALTAVISTA but instead uses the social network to make the search more focused and effective. For example, in the Selman example, the requirement of proximity to Selman helps make the query more specific to AI.

Exploring Specialized Communities

The version of REFERRAL WEB we described here creates a generic social network model based on whatever information can be found on the web. More accurate and detailed models of various communities can be extracted from more specialized information sources. Future releases of REFERRAL WEB will let the user explore some of these specialized domains.

For the academic research community, online bibliographies provide a rich source of relationships using the coauthor relationship.

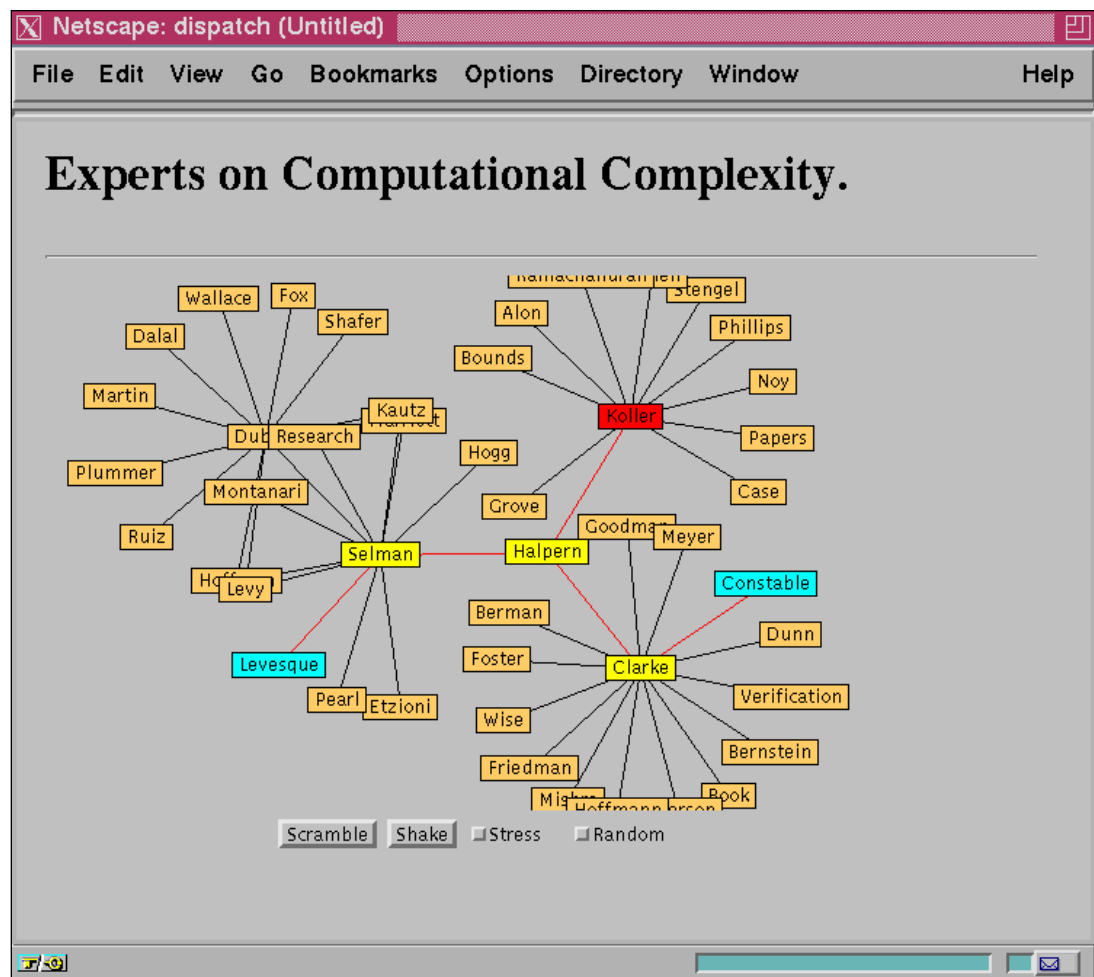


Figure 3. Graphic View of Response to Query for Experts on Computational Complexity.
Two chains were found: Koller-Halpern-Selman-Levesque and Koller-Halpern-Clarke-Constable.

Common academic affiliation also provides a source of evidence for a relationship between individuals.

Frequent contributors to news groups form another, looser kind of community. Archives of many news groups are easy to obtain and can be mined to obtain both the names of contributors and some evidence that particular contributors have a social tie. For example, if two individuals frequently quote from each other's postings, then this public correspondence reveals some tie between the two. (They might, of course, be antagonists, rather than colleagues!) News-group postings are also a good source of evidence of an individual's interests and areas of expertise.

Limiting the network to a specialized community partially overcomes one of the major technical problems in automating the con-

struction of social network models—distinguishing between different individuals with the same name. REFERRAL WEB does not attempt to deal with this problem: Ambiguous names are simply a source of noise that can lead to errors in the network model. As we argued earlier, there is reason to believe that referral chaining can be robust even against a significant degree of inaccuracy in the model. However, we will continue to explore this empirical question as use of the system grows.

As mentioned previously, a number of social-filtering (Malone et al. 1987) and recommendation systems (Resnick 1996) already exist, such as FIREFLY. REFERRAL WEB has a number of features that sets it apart from such systems:

First, REFERRAL WEB attempts to uncover

existing social networks rather than provide a tool for creating new communities. Although new communities might be appropriate for recreational uses of the web, we emphasize helping individuals make more effective use of their large, existing networks of professional colleagues.

Second, although recommender systems are often designed to provide anonymous recommendations, REFERRAL WEB is based on providing referrals by way of chains of named individuals. It is critical to provide the names because not all sources of information are equally desirable.

Third, some recommender systems require the user to manually enter a personal profile of interests, preferences, or expertise. Recommendations are generated by matching profiles that exist within the system. By contrast, because REFERRAL WEB builds its network model from public documents, the model includes many more individuals than those models based on names of individuals who explicitly register with the service.

Fourth, users of REFERRAL WEB are not limited to any set of topic areas determined in advance. The PHOAKS system (www.phoaks.com; Hill and Terveen 1996) has some features in common with REFERRAL WEB: It mines information (in particular, recommended URLs) from general net-news postings and to improve access to the resources of an existing community, and one can view the named individuals who made the recommendations. However, PHOAKS does not attempt to make connections between individuals.

A Test Bed for Referral Chaining

The REFERRAL WEB Project can be reached at www.research.att.com/~kautz/referralweb/.

REFERRAL WEB is an evolving system both because its network model grows with use and because we are experimenting with different algorithms and sources of information for hypothesizing social links. As mentioned previously, the target audience for our initial prototype is the community of AI researchers; therefore, we are concentrating our efforts on tuning the system for such users. For example, we are adding heuristics to give higher weight to co-occurrences of names that appear to indicate coauthorship of scholarly papers. For other groups, other sources of information, such as home pages or news groups, might be more important. In a corporate intranet application, the company-staffing database would be an important (although not the only) source of informa-

tion about relationships. We welcome feedback from users on both their experiences with the system and new ways they would like to be able to use the social network model to aid in the search for information and expertise.

Acknowledgments

Portions of this article appeared in Kautz, Selman, and Milewski (1996) and Kautz, Selman, and Shah (1997). We thank Will Hill for useful discussions and Ellen Spertus for the reference to Stanley Milgram's original paper.

References

- Galegher, J.; Kraut, R.; and Edigo C. 1990. *Intellectual Teamwork: Social and Technological Bases for Cooperative Work*. Hillsdale, N.J.: Lawrence Erlbaum.
- Granovetter, M. 1973. Strength of Weak Ties. *American Journal of Sociology* 8:1360–1380.
- Grossman, J., and Ion, P. 1995. On a Portion of the Well-Known Collaboration Graph. *Congressus Numerantium* 108:129–131.
- Hedetniemi, S. T.; Hedetniemi, S. M.; and Liestman, A. L. 1988. A Survey of Broadcasting and Gossiping in Communication Networks. *Networks* 19:319–349.
- Hill, W., and Terveen, L. 1996. Using Frequency-of-Mention in Public Conversation for Social Filtering. In *Proceedings of the ACM Conference on Computer-Supported Cooperative Work (CSCW-96)*, 106–112. New York: Association of Computing Machinery.
- Kautz, H.; Selman, B.; and Milewski, A. 1996. Agent-Amplified Communication. In *Proceedings of the Thirteenth National Conference on Artificial Intelligence*, 3–9. Menlo Park, Calif.: American Association for Artificial Intelligence.
- Krackhardt, D., and Hanson, J. 1993. Informal Networks: The Company behind the Chart. *Harvard Business Review* 71:4.
- Malone, T. W.; Grant, K. R.; Turbak, F. A.; Brobst, S. A.; and Cohen, M. D. 1987. Intelligent Information Sharing Systems. *Communications of the ACM* 30(5): 390–402.
- Milgram, S. 1967. The Small-World Problem. *Psychology Today* 1(1): 60–76.
- Resnick, Paul, ed. 1996. *Communications of the ACM* (Special Section on Recommender Systems) 30(3).
- Salton, R. 1989. *Automatic Text Processing*. Reading, Mass.: Addison-Wesley.
- Schwartz, M. F., and Wood, D. C. M. 1993. Discovering Shared Interests Using Graph Analysis. *Communications of the ACM* 36(8): 78–89.
- Sundheim, B., and Grishman, R., eds. 1995. *Proceedings of the Sixth Message-Understanding Conference (MUC-6)*. San Francisco, Calif.: Morgan Kaufmann.
- Wasserman, S., and Galaskiewicz, J., eds. 1994. *Advances in Social Network Analysis*. Thousand Oaks, Calif.: Sage.



Thirteenth Conference on Uncertainty in Artificial Intelligence

*August 1-3, 1997
Brown University
Providence, Rhode Island, USA*

Visit the UAI '97 home page at
<http://cuai97.microsoft.com>

Eric Horvitz, *Conference Chair*
Dan Geiger and Prakash Shenoy, *Program Cochairs*

UAI '97 will be held right after the National Conference on Artificial Intelligence (AAAI '97) at nearby Brown University.

For additional information contact the
UAI '97 conference chair at
horvitz@microsoft.com



Henry Kautz is head of the Artificial Intelligence Principles Research Department at AT&T Labs. He holds a Ph.D. and an M.Sc. in computer science from the University of Rochester and the University of Toronto, respectively, and an M.A. in writing from the Johns Hopkins University. He joined AT&T Bell Labs in 1987. In 1989, he received the Computers and Thought Award from the International Joint Conference on Artificial Intelligence in recognition of his work on plan recognition, temporal constraint satisfaction, and the complexity of inference. His work with Bart Selman on stochastic search for planning won the best paper award at the Thirteenth National Conference on Artificial Intelligence. His current research interests include software agents, planning, theory approximation, and stochastic algorithms for propositional reasoning. His e-mail address is kautz@research.att.com.



Bart Selman is a principal scientist in the Artificial Intelligence Research Department at AT&T Labs. He holds a Ph.D. and an M.Sc. in computer science from the University of Toronto and an M.Sc. in physics from Delft University of Technology. His research has covered many areas in AI, including tractable inference, knowledge representation, search, planning, default reasoning, constraint satisfaction, and natural language understanding. He has received best paper awards at both the American and Canadian national AI conferences and at the International Conference on Knowledge Representation. His current research projects are on efficient reasoning, stochastic search methods, planning, knowledge compilation, and software agents. His e-mail address is selman@research.att.com.



Mehul Shah received a B.S. in computer science and engineering and a B.S. in physics in June 1996 from the Massachusetts Institute of Technology. During this time, he was a co-op student at AT&T Labs and the former AT&T Bell Labs. He received his masters of engineering in June 1997 from the Electrical Engineering and Computer Science Department at the Massachusetts Institute of Technology. His current research interests include information retrieval, databases, and parallel algorithms. His e-mail address is shah@mit.edu.