

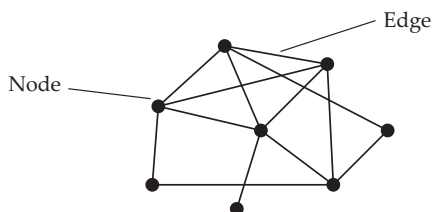
The physics of networks

Mark Newman

Statistical analysis of interconnected groups—of computers, animals, or people—yields important clues about how they function and even offers predictions of their future behavior.

Mark Newman is the Paul Dirac Collegiate Professor of Physics at the University of Michigan in Ann Arbor and a visiting professor at the Santa Fe Institute in Santa Fe, New Mexico.

In its simplest form, a network is a collection of points, or nodes, joined by lines, or edges.



As purely theoretical objects, networks have been the subject of academic scrutiny since at least the 18th century. But they have taken on a new practical role in recent years as a primary tool in the study of complex systems—real-world systems of interacting components, for which networks provide a simple but tremendously useful representation.¹ The internet, for example, can be represented as a network of computers linked by data connections. The World Wide Web is a network of information stored on webpages and connected by hyperlinks. Social networks of friendships between individuals have received a lot of recent attention, and other kinds of social networks, such as those of professional or business contacts, are also attracting their share of interest. Biological networks, such as the interrelated metabolic reactions that run the cell or the food web of predator–prey relations in an ecosystem, are of interest in both experimental and theoretical biology. And networks are increasingly common in the study of epidemiology, computer viruses, computer software, genetics, human transportation and communication, human language, books, movies, music, and many other things. But how do physicists enter the picture?

Physicists' interest in networks is relatively recent. Progress in the first 200 years of the field was mostly the work of mathematicians and social scientists. Leonhard Euler is often credited with the first rigorous result in graph theory (the mathematical study of networks), with his solution of the famous Königsberg bridge problem in 1765. Kenneth Appel and Wolfgang Haken's 1976 proof of the four-color theorem—that four colors are sufficient to color any map in such a way that any two adjacent regions are of different colors—is perhaps the best known recent achievement of graph theory. The empirical study of networks, meanwhile, has its foundations in sociology, in which researchers have been studying social networks since the 1930s.

Interest in networks has, however, seen its most spectacular growth in the past 10 years, with much of the fundamental research in the area being conducted, perhaps sur-

prisingly, by physicists, whose methods turn out to be well suited to the problems of the field. The approach taken by physicists differs from those of mathematicians and sociologists in two important ways. First, unlike most mathematical work, it is founded on and largely inspired by empirical studies of real-world networks such as the internet, friendship networks, and biological networks. One reason for the subject's rise in popularity has certainly been the increased availability of accurate and substantial network data sets.

Second, unlike most sociologists, physicists have been concerned largely with statistical properties of networks—their overall shapes and statistical signatures—rather than with properties of individual nodes or groups of nodes. Whereas a sociologist might have asked, "Which nodes in this network have the most connections?" a physicist might ask, "What is the average number of connections a node has?" or, "What is the distribution of the number of connections?" In asking those questions, physicists have stumbled across a number of intriguing network properties, mostly unremarked in earlier work, that have inspired an impressive array of new theories, techniques, algorithms, models, and measures to describe and illuminate the function of networked systems. Physicists are by no means the only contributors to the outpouring of new work—mathematicians, computer scientists, social scientists, biologists, and many others have made fundamental contributions—but physicists have played a central role, and physics journals have published much of the foundational work in the field.

Figure 1 shows a representation of the internet, the worldwide network of physical data connections between computers. The nodes in the figure represent groups of computers, and the edges represent data connections between those groups. From looking at the network, it is evident that although the pattern of connections is not a regular one, neither is it completely random. The network has clear structure, including the prominent starlike formations at the center and the more filamentary connections around the edges. The central questions that have interested physicists in the study of networks are how to appropriately quantify such patterns and what they mean for the functioning of the system a network represents. In this article I describe some of the approaches that have been developed to tackle those questions.

Degree distributions

The degree distribution is one of the most basic quantitative properties of a network, yet it is a property that until recently received comparatively little attention.

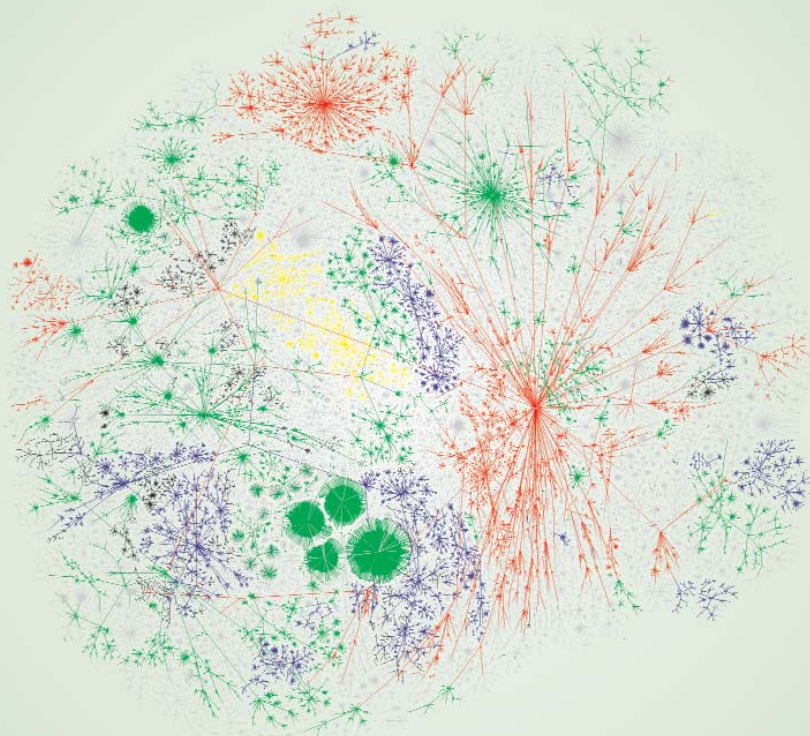


Figure 1. A network representation of the internet. Each node represents an autonomous system—a group of computers under single administrative control, such as the computers at a university or a corporation. Edges (connecting lines) represent direct peering relations between autonomous systems, which are a rough indicator of where optical fibers and other data connections run. (Patents pending; © Lumeta Corp 2007. All rights reserved.)

The degree of a node in a network is the number of edges connected to that node. In a social network of friendships, for example, your degree would be the number of friends you have. As figure 1 shows, many nodes in the internet have low degree—just one or two connections—but a few (often called hubs) have very high degree, as much as a thousand or more in some cases.

Figure 2 shows a plot of the distribution of degrees on the internet and reveals an interesting pattern: The distribution roughly follows a straight line on the logarithmic scales used in the figure, meaning that the number $P(k)$ of nodes with degree k follows a power law $P(k) \propto k^{-\alpha}$, an observation first made in 1999 by Christos Faloutsos, Michalis Faloutsos, and Petros Faloutsos,² three brothers who are all professors of computer science. The value of the exponent α for the internet is about 2.1.

When first discovered, the power-law distribution was a surprise to many researchers. If edges in the network were placed between nodes simply at random, the resulting degrees would follow a Poisson distribution, which is very different in shape from a power law, with most nodes having degrees close to the mean value and no high-degree hubs at all. The observation of a power-law distribution thus indicates that the placement of edges in the network is, in a sense, far from being random.

A number of other networks have also been shown to have degree distributions that follow power laws; such networks are now often referred to as scale free. The earliest report of a scale-free network that I'm aware of was given in 1965 by Derek de Solla Price, who studied citation networks, in which the nodes represent learned papers and the edges represent citation of one paper by another. Price showed that the distribution of the number of citations a paper receives—the degree distribution of the citation network—follows a power law.³

Recent interest in degree distributions has been particularly sparked by the work of Réka Albert, Hawoong Jeong,

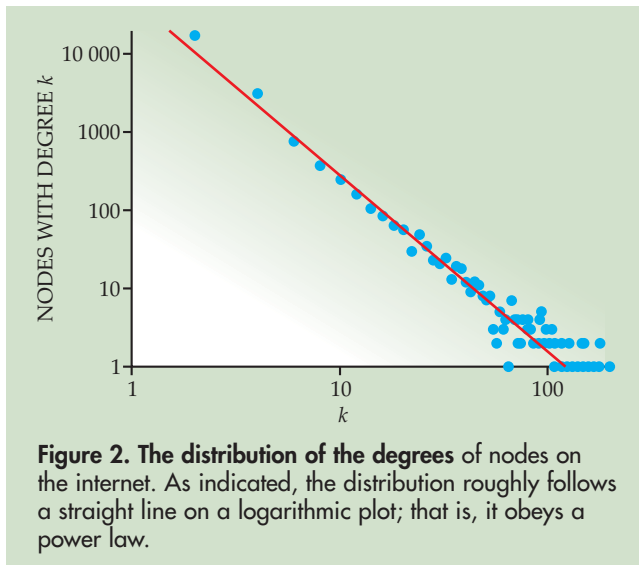
and Albert-László Barabási (all at the University of Notre Dame), who discovered power laws in a number of networks, including the World Wide Web.⁴ Subsequent studies by various researchers have shown that many other networks, though not always following the power-law pattern precisely, do tend to have skewed degree distributions with a lot of low-degree nodes and a small number of high-degree hubs.⁵

Those findings turn out to be enormously important because many of the behaviors of networked systems are dominated by their hubs. Though the hubs may be few in number—sometimes just a few percent of the total number of nodes, depending on how you define them—they are nonetheless the principal factor determining many aspects of the behavior of the overall system. The next two sections explore examples of this phenomenon.

Resilience to the removal of nodes

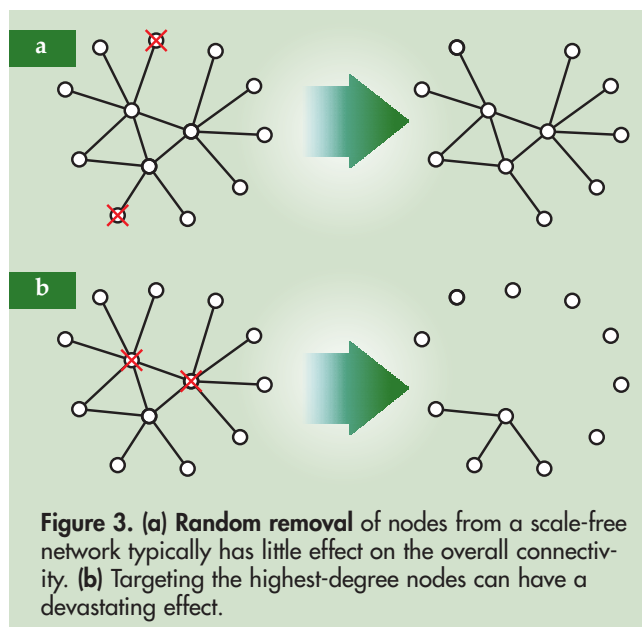
Suppose that some of the nodes in a network disappear for some reason. Routers on the internet, for instance, fail all the time: By some estimates, as many as 3% of routers worldwide may be nonfunctional at any given time. Nonetheless, most people remain able to reach the websites and e-mail servers they want because the network offers—by design—more than one route from most points to most others. Building on that idea, Albert and colleagues asked what fraction of nodes in a network would have to be removed before performance *was* affected.⁶ They found that the answer depends crucially on the network's degree distribution.

The researchers considered two different schemes for removing nodes. In the first, they removed nodes uniformly at random, whereas in the second, they deliberately targeted the highest-degree nodes for removal, on the assumption that the second scheme would cause the network to fail faster. It turns out, however, that on a network in which the nodes are connected at random, there is little difference between the two schemes. As discussed above, a random network has a



Poisson degree distribution, with most nodes having degrees close to the mean. In such a network, a process that removes the highest-degree nodes is not much different from a process that removes nodes at random: The degrees of the nodes removed are about the same either way, and the network is roughly equally resilient to both attacks.

But real-world networks are not random. Many have highly skewed degree distributions, and for them the difference between random and targeted removal of nodes is striking. If you remove nodes purely at random, then most of them have low degree, since most of the nodes in the network have low degree. Thus the removed nodes are connected to few others and have little effect when they are removed, as shown in figure 3a. One of the more remarkable theoretical results to emerge in recent years is that if nodes are removed uniformly at random from a network with a power-law degree distribution, then the network typically remains connected and functional no matter how many nodes are removed.⁷ Of course, the nodes that are removed are themselves no longer connected to the network, but of the nodes left behind, an extensive fraction remain connected.



Scale-free networks are thus extremely resilient against random removal or failure of their nodes. The same is not true of purely random networks.

For the case of targeted removal of the highest-degree nodes, on the other hand, the reverse is true. The high-degree nodes in a scale-free network are hubs with connections to many other nodes, so the removal of just a few of them can have a substantial effect, as shown in figure 3b. Analytic calculations, for instance, indicate that no matter what the exponent of the power law, no more than 3% of nodes need to be removed before the entire network becomes disconnected, meaning that the average probability that there is a path connecting any two nodes vanishes.⁸

In the context of a communication network such as the internet, that kind of fragility to a targeted attack could be a bad thing: It's certainly not desirable for critical infrastructure to be susceptible to failure of, or attacks on, just a few central hubs. In other domains, however, fragility can be good. One reason for the current high level of interest in social networks is their importance in the spread of disease. Diseases travel over networks of contact between individuals just as information travels over the internet, and nodes can be "removed" from those networks by vaccination, assuming a vaccine is available for the disease in question. If vaccination procedures could be targeted toward the highest-degree nodes in a social network, it might in theory be possible to disconnect the network and thus prevent the spread of disease while vaccinating only the tiniest fraction of the population. This effect, in which a vaccination campaign ends up protecting not only those vaccinated but also many others as well, is known to epidemiologists as "herd immunity." In practice, unfortunately, it can be hard to exploit because of the difficulty of finding the high-degree individuals.

Spreading processes on networks

The spread of disease also provides the second example of the importance of the degree distribution in the functioning of networked systems. It is clear that the degrees of network nodes must play some role in the spread of disease—or other disease-like elements, such as rumors or fads, that pass over networks of contacts between individuals. If no one in a network has any connections, for instance, then diseases of course cannot spread. If everyone has many connections, diseases can spread quickly. One might guess that the speed at which a disease spreads over a network would be determined by the mean degree of a node. Although that is the right basic idea, it turns out to be wrong in detail: The crucial parameter is not the mean degree but the mean squared degree, as is shown by the following qualitative argument.

Consider a hub in a social network: a person having, say, a hundred contacts. If that person gets sick with a certain disease, then he has a hundred times as many opportunities to pass the disease on to others as a person with only one contact. However, the person with a hundred contacts is also more likely to get sick in the first place because he has a hundred people to catch the disease from. Thus such a person is both a hundred times more likely to get the disease and a hundred times more likely to pass it on, and hence 10 000 times more effective at spreading the disease than is the person with only one contact. That is not a rigorous argument, but it can be made rigorous, and the basic result is correct: It is not a node's degree but its degree squared that is the crucial parameter.

This result is again particularly important when applied to scale-free networks. For a power-law degree distribution with an exponent less than 3, the mean squared degree formally diverges, and with it the rate of growth of a disease epidemic.

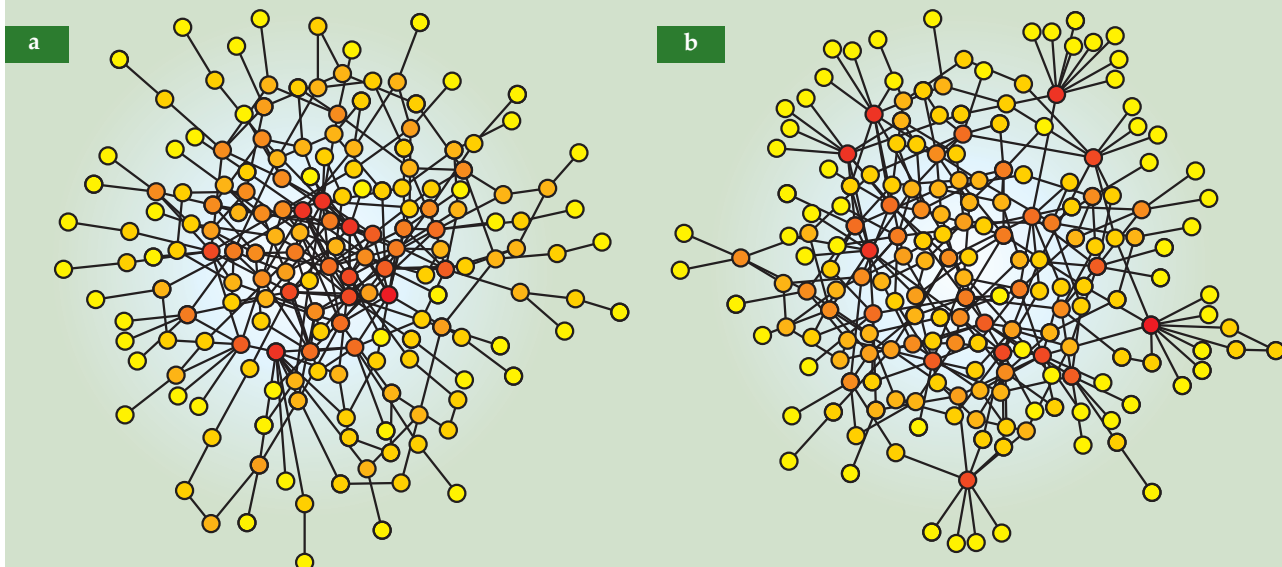


Figure 4. Two networks with the same distribution of node degrees but different degree correlations. Network (a) is positively correlated, whereas (b) is negatively correlated. The nodes are color coded to emphasize their degrees (yellow for low degree, red for high degree). The high-degree nodes clump together in network a but are more spread out in network b. A number of the hub-and-spoke formations characteristic of negatively correlated degrees—created when a high-degree node connects to a lot of low-degree ones—are visible around the edges of network b.

(In a network of finite size, it does not actually diverge, but it does become very large.) Although some experts have debated whether real contact networks have power-law degree distributions, there is little doubt that the existence of network hubs—or “superspreaders,” as they are sometimes called in epidemiology—plays a big role in determining whether and how fast diseases spread.⁹

The small-world effect

Perhaps the best known discovery in the study of networks is the so-called small-world effect, the finding that most pairs of people, no matter how distant they may be, are connected by at least one and probably many short chains of acquaintances. The idea was explored mathematically by Ithiel Pool and Manfred Kochen¹⁰ starting in the 1950s and more recently by Duncan Watts and Steven Strogatz in a widely cited 1998 paper.¹¹ But it is most strongly associated with the work of Stanley Milgram, a social psychologist who at Harvard University in the 1960s conducted a now-famous experiment in which he asked participants to pass letters from acquaintance to acquaintance in an effort to get them to a chosen target person. Milgram found that the letters that reached the target took an average of just six steps to get there from a starting point chosen roughly at random,¹² a result that has passed into folklore under the name “the six degrees of separation” and has become the starting point for many after-dinner conversations and parlor games.

The small-world effect is not confined to social networks and seems to apply to almost all kinds of networks. There are exceptions—networks with special regularities such as low-dimensional networks or lattices, for example—but all the networks mentioned in this article seem to be small worlds in the sense discussed by Milgram. The average number of degrees of separation varies from network to network—and indeed the number six found by Milgram is only a rough estimate, given the limitations of his experimental method—but the basic principle, that you can get from almost any node to any other in just a small number of steps, is well documented in a wide array of systems.

Although many people—Milgram included, by his own report—find the small-world effect surprising, it is not mathematically unexpected. The fundamental explanation is that the number of people you can reach by taking a given number of steps in your social network increases exponentially with the number of steps you take—a result known as the expander property—so the number of steps needed to reach anyone in the world increases only logarithmically with world population. Since the logarithm is a slowly increasing function, the typical number of steps between any two people in the world is relatively small, even though the population numbers in the billions. The exponential behavior has been confirmed empirically for a wide variety of networks and appears well established.

However, another aspect of the small-world effect really is surprising and was noted only recently by computer scientist Jon Kleinberg. He pointed out that Milgram’s experiment reveals not only the existence of short paths between pairs of individuals in social networks but also that people are good at finding those paths.¹³ He demonstrated how non-trivial that statement is by using a simple network model to show that in most cases short paths are difficult to find unless you know the entire network—which, of course, the people in Milgram’s experiment didn’t. Kleinberg proposed a possible mechanism to explain how people navigate around social networks in practice: Although nobody knows the entire network, he suggested that everyone has friends at various distances and has an intuitive feeling for roughly how close each of those friends is to any given target person. In an experiment like Milgram’s, people would then pass the letter in a series of steps that close in on the target by stages: A random starting person in, say, New York aiming for a target in Los Angeles gives the letter a long first step to a friend in California. From there that friend takes a shorter step, passing it to an acquaintance in Los Angeles, and the process is repeated with shorter and shorter steps until eventually the letter reaches someone who knows the target personally.

Kleinberg’s crucial discovery was that a process of this type, simple though it sounds, only works if people have the

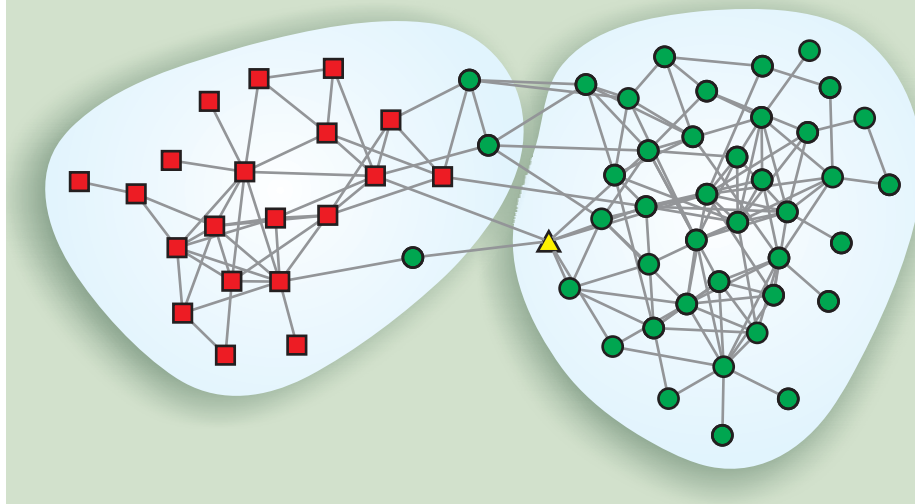


Figure 5. The social network of a community of bottlenose dolphins. The node shapes and colors indicate the two groups into which the network split upon the departure of the dolphin denoted by the triangle. The two circled regions indicate the division of the network found by an automated network analysis technique that considered only the pattern of connections.

right distribution of friends at different distances. If they do not, then either the letter stalls in the hands of an individual with no friends closer to the target than himself, or the steps it takes toward the target are all small, and a huge number of them are needed to reach the destination. Thus the success of Milgram's experiment implies that the social network has a special form. Although Kleinberg phrased his results in terms of geographic distance, variations have been proposed that employ various forms of social distance or that work in non-Euclidean spaces.

In addition to their scientific interest, network navigation processes have practical applications too. One is in the use of so-called spiders to perform customized searches of the Web or other distributed databases. Traditional Web searches are performed by crawling the Web on a huge scale for information on every conceivable topic, indexing it all, and then searching that index to locate webpages of interest. Spiders take a more customized approach, performing a limited perusal of the network for information on just one topic. For example, the simplest spider is a blind "hill climber" that starts at a random point and follows links across the network, jumping to pages of successively greater relevance to the topic of interest (as measured, for instance, by the occurrence of key words) until no further such steps are possible; then the spider reports the last page found. That strategy works because the spider, like the participants in the small-world experiment, has a measure of how close it is to the target. It only works, however, if pages on the Web are connected to an appropriate distribution of relevant and less relevant neighbors, so that the spider can take both long and short steps when needed, just as Milgram's letters did. Thus the network must again have a special structure if rapid navigation is to be possible. Luckily, the Web does appear to have that kind of structure, and efficient spiders have been created that find uses in various kinds of specialized Web search.

Correlations and communities

Even after you allow for skewed degree distributions, the connections in typical networks are, unsurprisingly, far from random. One indication of deeper patterns in network structure can be found in the correlations between the degrees of different nodes. You can look at adjacent nodes in a network and ask whether, statistically speaking, their degrees are correlated. Are the high-degree nodes connected to other high-degree nodes, for example, and low to low? Are the party people hanging out with other party people, or are they hanging out with hermits? Correlations can be quantified

using a joint probability distribution $P(k_1, k_2)$ for the degrees k_1 and k_2 of adjacent nodes, but such joint distributions are hard to measure. A simpler approach, and one natural for physicists, is to define a correlation coefficient r whose value, positive or negative, quantifies the level of correlation:

$$r = \frac{\langle k_1 k_2 \rangle - \langle k_1 \rangle \langle k_2 \rangle}{\sigma_k^2}$$

where the averages are taken over all edges and σ_k^2 is the variance of the node degree k . Values of r have been calculated for various networks and reveal an intriguing pattern: Most social networks, it turns out, have positive correlations between the degrees of adjacent nodes, whereas most nonsocial ones, including technological and biological networks, have negative correlations. Although there are exceptions to the pattern, it seems to be quite a reliable rule of thumb. However, its origins are, for the moment at least, not entirely clear.

What is clear is that degree correlations have a strong effect on the structure of networks. Networks with positive correlations tend to have a core-periphery structure: The nodes with high degree are attracted to one another and so coagulate into a highly interconnected core surrounded by a periphery of lower-degree nodes, as shown in figure 4a. In negatively correlated networks, by contrast, the high-degree nodes tend to be scattered more broadly over the network, as shown in figure 4b.

Those structural differences can have substantial effects on the way a networked system behaves. A disease, for instance, can persist more easily in a positively correlated network by circulating in the dense core where there are many opportunities for it to spread. On the other hand, the below-average density of the periphery makes it harder for the disease to leave the core. In a negatively correlated network, the same disease finds it harder to persist, but if it does persist, then it typically spreads to the whole network.

Degree correlations are an example of a more general phenomenon known as assortative mixing, in which the probability of two nodes being connected by an edge depends on some property of those nodes. The property in that example was degree, which has a special importance because it is itself a network measure, but network connections can depend on all sorts of other properties too. In social networks, for instance, the probability of a connection between two individuals has been observed to depend on their ages,

incomes, races, and the languages they speak, among other characteristics. On the internet, the probability of a connection depends on the geographical locations of nodes and other node properties such as bandwidth.

These observations lead in turn to another recent point of interest in the study of networks, the appearance of community structure. In networks in which like nodes associate preferentially with one another, connected clusters of like nodes tend to form—think of ethnic communities in cities, for example. Many networks are found to have such groups or communities within them, although it's not always clear what, if anything, is the underlying commonality among a group's members. Nonetheless, the discovery of groups within a network can in many cases lead to a better understanding of the behavior of a system.

Figure 5 shows an example from a collaboration I recently undertook with marine biologist David Lusseau on a study of bottlenose dolphins.¹⁴ Dolphins are highly social animals that demonstrate attachments by performing “dances,” whose observation can serve as raw data for the construction of social networks like the one in the figure. By an appropriate analysis of the network, we discovered two clear subcommunities within it, indicated by the circled regions. Interestingly, during the course of Lusseau's multiyear observation of the community, one dolphin, denoted by the triangle in the figure, disappeared. (It wasn't dead, it turns out; it just disappeared unexpectedly for a couple of years, to who knows where, before returning equally unexpectedly.) Upon its disappearance, the remaining dolphin population split into two subpopulations, denoted by the circles and squares, that then went their separate ways. Apparently, the absent dolphin was the crucial glue holding together the larger community, and with its departure the community split. But notice that the communities discovered by the analysis of the network before the split provide a good, though not perfect, prediction of how the split would occur. In a sense, the analysis quantitatively predicted the future behavior of the social system. (Some studies of human social networks, and of nonsocial networks as well, have given similar results.)

The development of methods for finding communities within networks is a thriving sub-area of the field, with an enormous number of different techniques under development. Methods for understanding what the communities mean after you find them are, by contrast, still quite primitive, and much needs to be done if we are to gain real knowledge from the output of our computer programs.

Further directions

There are many other interesting areas in the study of networks. A huge volume of theoretical work has been done on network models; a growing field of experimental analysis and data mining has produced both new techniques and a host of interesting new results in the last few years; there are studies of dynamical networks, which change over time, and theories that describe them, and of dynamical systems on networks, such as synchronization phenomena or the dynamics of metabolic reactions; and many interdisciplinary projects are bringing ideas from statistics, machine learning, algorithms, and other fields to bear on networked systems.

And yet we have large holes in our understanding, with many fundamental questions still unanswered. The current state of the field is perhaps analogous to quantum mechanics before the discovery of the Schrödinger equation. Or maybe it's not even that good. For instance, we currently have little idea of how to reliably estimate the properties of net-

works for which we have incomplete structural data. Unfortunately, that includes many, perhaps most, of the networks we're interested in, which means that we have reliable measurements of quantities of interest for few of the systems mentioned in this article. Moreover, it's hard to know whether we are even measuring the right things in many cases. I've described a variety of network measures—degree distributions, correlation coefficients, and so forth—and we have discovered much by our study of those quantities. But there may well be other equally important quantities we should be looking at but haven't thought of yet. In some fields of physics, there are fundamental theories that state that if you measure certain things—correlation functions, for example—then in principle you know everything about the measured system, at least up to some given level of approximation. But no such theory exists yet for networks, nor even an idea of how to develop one.

And even if we can solve those problems, we will have only just begun to tackle the real questions at the heart of the field. In the end the goal is to understand how networked systems behave. Characterizing their structure is a good first step, and studies of things like robustness and spreading processes certainly address elements of the behavior question. But, ultimately, the reason we study the internet is to understand how it works, responds, and evolves. We study social networks because we want to understand and perhaps predict the behavior of societies. We study the World Wide Web to understand how it stores information and how people use it, with the hope perhaps of building better search engines or better Web browsers. And although we have made enormous progress toward those goals, big questions of fundamental importance are still begging for answers. Of course, it has taken a century to arrive at our still-incomplete understanding of quantum mechanics, next to which the mere 10 years or so that physicists have devoted to the study of networks is the blink of an eye. What has been achieved in that short time is extensive and invaluable. Still, it seems certain that crucial results are waiting to be discovered by experimenters and theorists alike, and there is plenty of room for those with ideas to contribute.

References

1. For further reading, see S. H. Strogatz, *Nature* **410**, 268 (2001); A.-L. Barabási, *Linked: The New Science of Networks*, Perseus, Cambridge, MA (2002); D. J. Watts, *Six Degrees: The Science of a Connected Age*, Norton, New York (2003); M. Newman, A.-L. Barabási, D. J. Watts, *The Structure and Dynamics of Networks*, Princeton U. Press, Princeton, NJ (2006).
2. M. Faloutsos, P. Faloutsos, C. Faloutsos, *Comput. Commun. Rev.* **29**, 251 (1999).
3. D. J. de S. Price, *Science* **149**, 510 (1965).
4. R. Albert, H. Jeong, A.-L. Barabási, *Nature* **401**, 130 (1999).
5. L. A. N. Amaral, A. Scala, M. Barthélémy, H. E. Stanley, *Proc. Natl. Acad. Sci. USA* **97**, 11149 (2000); A.-L. Barabási, R. Albert, *Science* **286**, 509 (1999); S. N. Dorogovtsev, J. F. F. Mendes, *Adv. Phys.* **51**, 1079 (2002).
6. R. Albert, H. Jeong, A.-L. Barabási, *Nature* **406**, 378 (2000).
7. R. Cohen, K. Erez, D. ben-Avraham, S. Havlin, *Phys. Rev. Lett.* **85**, 4626 (2000).
8. D. S. Callaway, M. E. J. Newman, S. H. Strogatz, D. J. Watts, *Phys. Rev. Lett.* **85**, 5468 (2000).
9. R. Pastor-Satorras, A. Vespignani, *Phys. Rev. Lett.* **86**, 3200 (2001); A. L. Lloyd, R. M. May, *Science* **292**, 1316 (2001).
10. I. de S. Pool, M. Kochen, *Soc. Netw.* **1**, 1 (1978).
11. D. J. Watts, S. H. Strogatz, *Nature* **393**, 440 (1998).
12. S. Milgram, *Psychol. Today* **1**(1), 60 (1967).
13. J. M. Kleinberg, *Nature* **406**, 845 (2000).
14. D. Lusseau, M. E. J. Newman, *Proc. R. Soc. B* **271**, S477 (2004). ■