

2 Why Do We Like Networks?

Networks catch hold of you. They are enchanting and contagious. As a first ‘proof’ of these statements let me give you my own example. Just before starting to write this chapter, I sat on a train and watched a charming mother and her little daughter just opposite me. The baby fell asleep playing with her comforter. As I continued to watch, my mind went to work. What could be the periodicity of her suckling motions? Was it perhaps scale-free, showing sudden bursts of activity separated by longer and longer periods of stasis? Did it show self-organized criticality? Was this a punctuated equilibrium? My thoughts continued: What if I looked outside? Would I see fractals instead of trees and clouds? ***“Let me interrupt you here. Why do you assume that we know what ‘scale-free’, ‘self-organized criticality’, ‘punctuated equilibrium’ and ‘fractal’ are supposed to mean? And anyway, what is a network?”*** I am sorry, Spite. Whenever elements are connected with links, we may call them a network. Networks can be formed from atoms, molecules, cells, plants, firms, words, power stations, Internet routers, Web pages, countries, etc. Even your friends, Spite, form a network. The meaning of the other words will be explained later. If you are curious to understand them now, turn to the glossary in Appendix B, at the end of the book.

Returning to the popularity of networks, I am not the only one who has found this field fascinating. Figure 2.1 shows the number of network-related scientific publications in MEDLINE.¹ The arrow points to the publication date of two important network discoveries, the demonstration of the generality of the small-world phenomenon (Watts and Strogatz, 1998) and scale-free behavior (Barabasi and Albert, 1999). Obviously, these data may just reveal a coincidence. They do not directly prove the profound effects of these important discove-

¹Data of in Fig. 2.1 should be treated with caution, since ‘network’ may also refer to a network of authors, for example. An additional non-specific effect arises from the fact that the number of annual publications covered by MEDLINE also increased over the period covered.

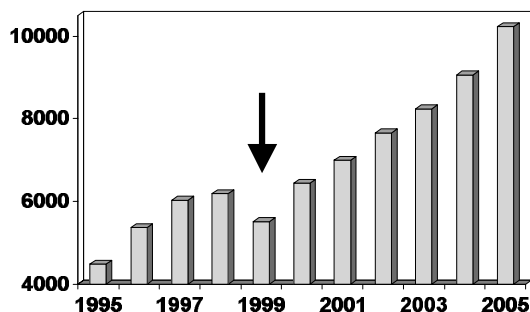


Fig. 2.1. Number of network-related publications in MEDLINE. The number of publications containing the words ‘network’ or ‘networks’ in their title or abstract was collected from MEDLINE (www.pubmed.com). The 2005 data is an extrapolation. The arrow shows the publication date of the two seminal network papers by Watts and Strogatz (1998) and Barabasi and Albert (1999)

ries in the network approach. However, Fig. 2.2 shows the number of annual citations of the above two papers in comparison with the average citations of three randomly selected papers from the same journals having a similar number of total citations. The citations of randomly selected papers peter off after 3 to 4 years. In contrast, citations of the two seminal network papers grow linearly, showing no tendency to decline in this period. No wonder, scientists seem to like networks. But how about the layperson? As a measure of public success, Laszlo Barabasi’s book, *The Linked* was translated into 8 languages in the first two years of its existence.

Having these data to hand, I think we may be quite confident in saying that networks really catch hold of people. People do like networks. However, another question arises: Why exactly do people like networks? This chapter attempts to answer this question and uses the elements of the answer to introduce some important features of networks in general.

Small-worldness, scale-freeness, nestedness, weak-linkedness: these are the titles of the following sections. All these words refer to properties which are general features of most networks around us, and this is why we have acquired a feeling for them. These properties of the networks we either contain or belong to inherently help us to understand the world around us, being basic, underlying elements of our cognition. Therefore small-worldness, scale-freeness, nestedness and weak-linkedness not only mean the actual features of networks (being a small-world network, having scale-free distribution of various proper-

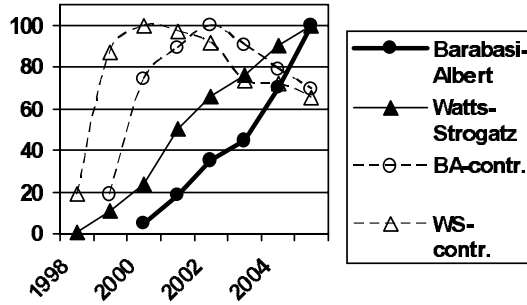


Fig. 2.2. Citations of seminal papers on networks. The numbers of citations for the Watts and Strogatz (1998) and Barabasi and Albert (1999) papers were collected from the Web of Science. Control values show the number of citations of three randomly selected papers from the same journals having a similar total number of citations. Data were normalized to the maximal number of yearly citations. 2005 data is an extrapolation

ties, containing other networks as its elements as well as belonging to higher order networks and having a large number of weak links, respectively), but also refer to the help these network properties give us. What is this help? Please continue, if you would like the answer.

2.1 Small-Worldness

Stanley Milgram did many famous experiments. In his small-world experiment he gave letters to starters, persons, who were asked to pass them to acquaintances known on a first-name basis in order to find an unknown, distant target (Milgram, 1967). Imagine that you have the task of sending a letter to the Reverend Lucas Brown, who lives in the capital of Myanmar, Yangon. It is rather easy. I need the address, ZIP code and a few stamps. But not this time! No address is known, and direct mailing is excluded. You may pass the letter only to one of your friends. The important message of Milgram's work, viz., "we live in a small world, and are only six steps apart from each other", became very popular. There is a good chance that your letter to the Reverend Brown will find its target by passing along a chain of around six friends.

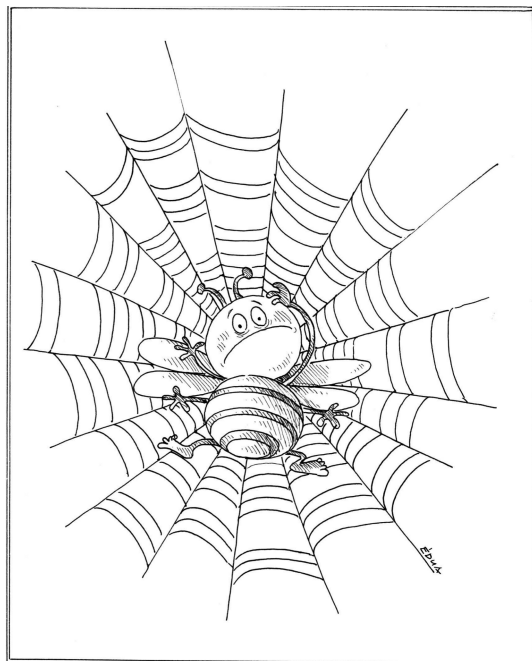


Fig. 2.3. Networks can really catch hold of you



A Hungarian prediction of small worlds from 1929.

Tibor Braun (2004) quotes the following text from a story by the Hungarian writer, Frigyes Karinthy, in 1929: “To prove that nowadays the population of the Earth is in every aspect much more closely interconnected than it ever has been, one member of our gathering proposed a test. ‘Let us pick at will any given existing person from among the one and a half billion inhabitants of the Earth, at any location.’ Then our friend bet that he could establish via direct personal links a connection to that person through at most five other persons, one of them being his personal acquaintance. ‘As people would say, look, you know X.Y. Please tell him to tell Z.V., who is his acquaintance, and so on.’ ‘OK,’ said a listener, ‘then take for example Zelma Lagerlöff’ (Nobel Prize for Literature, 1909). Our friend placing the bet remarked that nothing could be easier. He thought for only two seconds. ‘Right,’ he said, ‘so Zelma Lagerlöff, as a Nobel Laureate, obviously knew the Swedish king Gustav, since the king handed her the prize, as required by the ceremony. Gustav, as a passionate tennis player, who also participated at large international contests, evidently played with Kehrling [Béla Kehrling (1891–1937), Hungarian tennis champion and winner at the Göteborg Olympics 1924], whom he knew well

and respected.’ ‘Myself,’ our friend said (he was also a good tennis player), ‘I know Kehrling directly.’ Here was the chain, and only two links were needed out of the stated maximum of five.” The amazing foresight of Karinty (1929) predicting that we are approximately five steps apart from each other on the global scale was proved decades later by Milgram (1967) and Dodds et al. (2003a).

When I gave a lecture on networks to illustrate the smallness of the small world we live in, I asked my audience how many steps they thought they were from the President of the United States. Some of them guessed around a hundred, others were better informed and said: Six! Then I surprised them with the exact number: Three. ***“How come? Did you know that someone’s parent attended the same school as the President?”*** No, Spite, I knew only my own connections. I happen to know the President of my country, Hungary, who did meet the President of the USA. Since the students knew me, this is exactly three steps for them. Our world is really small. However, there is another message from the Milgram experiment: not only do short paths exist between distant network members, but ordinary people are very good at finding them too (Newman, 2003b). How would you kick off your letter to the Reverend Brown in Yangon? ***“I happen to have a friend who moved to Kuala Lumpur a year ago. If I recall my geography lessons, it is not far from Yangon. I would ask her to look around. She certainly knows many more people in the region than I do.”*** Excellent, Spite. If she happens to know a priest in Kuala Lumpur or Yangon, you might even complete the chain in three steps instead of six.



Why was Milgram lucky? Examining the original numbers, I have to conclude that in spite of the seemingly easy navigation shown above, Milgram was lucky. The final conclusion was based on only 18 letters which actually reached the single target of the Milgram experiment in Boston out of the 96 starters at distant locations in Nebraska. In other studies the success rate was even lower (Kleinfield, 2002). It was often hard to define what was causing the numerous drop-outs. However, a later study (Dodds et al., 2003a), using tens of thousands of emails had the same conclusion: we are about six steps apart even in different parts of the world. Experimenters of robust phenomena are lucky. Their instincts often find the right solution even when the actual proof is shaky. However, if you are a young investigator, let me ask you *not* to rely on this. Unsuccessful examples always outnumber

the few serendipity stories. We do not hear about the failures: most of them never get published. Moreover, our publication habits mean that we mention only the final success stories and not the very important and frustrating path we had to follow to reach them.



The number of dimensions in our brain. How do we select a direction to send our letter towards the unknown target? In fact, we try to get a match between the character of our acquaintances and the known properties of the target. For this matching task, we categorize our friends into social dimensions, as has been shown in the model by Watts et al. (2002). Spite started the letter to the Reverend Brown by finding a friend in Kuala Lumpur. This was a wise tactic, since geographic information is sufficient to perform global routing in a significant fraction of cases (Liben-Nowell et al., 2005). However, I have no friends in the region, and therefore I would probably double check the priests in my circle of friends instead. The number of social dimensions screened for a search is around 5 to 6 (Dodds et al., 2003a; Killworth and Bernard, 1978). This number is actually quite close to our average cognitive dimensionality, which is measured as the number of persons whose intentions towards each other I may still follow (Dunbar, 2005). What should we do if we want to be even more sophisticated? Should we use more social dimensions? I might have bad news here. The dimensionality of our neural network may prevent this. Even if our world grows hopelessly complex, we will still restrict ourselves to half a dozen, or even fewer, social dimensions to describe it, or start to develop more complex neural nets in our brain. Watch out for contact gurus! The evolution of superbrains may actually be happening around you now as you read!

Our world is a small world. However, it is not only the expanding circles of friends, the social net, which is a small world. Many other networks, like power nets, the networks of neural cells, etc., are also small-world networks. We live in small-worldness. ***“This is obvious. If I take a hundred people and everyone knows everyone, their world – let me call it Spiteland for short – is really small. Not six steps, but one, separating any one of them from any other.”*** Spite, I appreciate your logic, but real life is not Spiteland. We cannot know everyone. Do you have six billion friends, increasing by dozens every second? I doubt it. However, you have hit upon a good point here. Small worlds are not only small in the sense that their members may reach each other easily. Random graphs, where the connections are made between the elements in a random fashion,

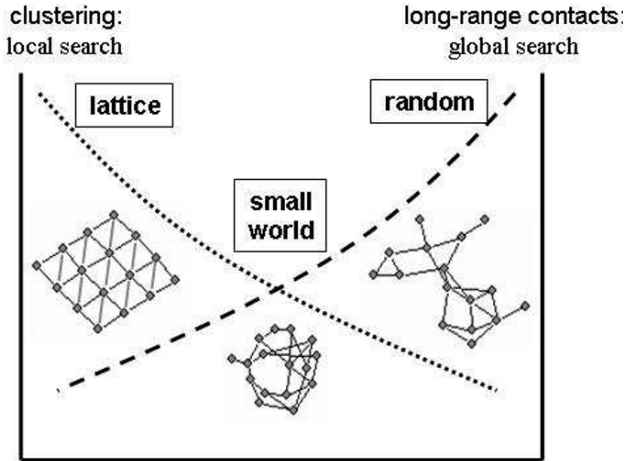


Fig. 2.4. The small-worldness of networks. The figure shows that small-world networks are in-between lattice-type networks and random graphs, having much longer range contacts than the former and much higher clustering than the latter. Note that the measures of both clustering and long-range contacts are purely illustrative

are equally good at this (see Fig. 2.4). In small worlds, your neighbors also know each other. To use a scientific term, this ‘my friend’s friend is my friend’ effect is called clustering. The clustering of small-world networks is high. These networks are lucky mixtures of 100% clustered regular lattices and highly-connected random graphs (see Fig. 2.4) (Watts and Strogatz, 1998). Small-worldness requires both a dense array of local contacts, which is reflected by the high clustering of small worlds, and a good enough number of long-range contacts. The simultaneous presence of both ensures that the small-world network becomes really small, providing easy conditions for finding any of its members. However, this cannot be achieved by extensive cross-linking of the network, since building and maintaining links is costly. Natural small worlds are economical (Latora and Marchiori, 2003). In fact, small worlds are much more economical than either random networks or regular lattices.



Some worlds are not so small. Small-worldness depends on what we consider as a member of the network. As an example, the extent of the small-world status of metabolic networks may vary, if we include relatively simple molecules like water, ATP, use directed links, or restrict the network

to conserved residues of participating molecules (Arita, 2004; Ma and Zeng, 2003).



How many friends do you need to send your message to anyone? In 2000, Jon Kleinberg published an interesting model for message transmission on a two-dimensional lattice, where lattice elements were linked with random shortcuts. The interesting result was that an optimal condition can be defined for the fastest search. If the shortcuts are neither fully random (where lots of short paths exist, but it is extremely time-consuming to find them), nor restricted to short-range contacts (where no short paths exist at all), an optimal condition can be found where the system transmits the messages most efficiently. Under these conditions, you have exactly the same number of friends in your neighborhood, in the rest of your city, in the rest of your country, in the rest of your continent and in the rest of the world. In other words, you only have to worry about how to send your message to someone in the right neighborhood. Once the message has reached the right region, the fine-tuned targeting will rely on the increasingly denser local contacts as the message homes in on the actual target. This makes the search highly efficient and the system behaves like a small world (Kleinberg, 2000). The Kleinberg condition can be reformulated: for an optimal two-dimensional search, if you go to a higher region (neighborhood, city, country, continent, world), the chances of finding a friend after a random selection become an order of magnitude smaller. Kleinberg’s model behaved optimally if the number of connections was the same on all scales, i.e., it was scale-free. Scale-freeness is another important feature of our everyday networks besides small-worldness, and will be discussed in detail in the next section.

Small worlds are easy to navigate. Lattice-type connections with their high clustering ensure the success of the finely-tuned final steps of a target search. Long-range contacts ensure the success of initial steps zooming in on the region of final interest. “The key to generating the small-world phenomenon is the presence of a small fraction of [...] edges, which contact otherwise distant parts of the graph” (Watts, 1999).

Navigation in small-world networks is helped by weak links (Granovetter, 1973; Lin et al., 1978). In many networks such as social networks, most of the long-range links typical of small worlds are weak links (Onnela et al., 2005). As I mentioned above, Dodds et al. (2003a) repeated the Milgram (1967) experiment using more than 60 000 emails. They found that a successful social search was conduc-

ted primarily through intermediate-to-weak strength links and did not require highly connected hubs. Moreover, Skvoretz and Fararo (1989) showed that the more weak links there are in a population, the closer a randomly chosen starter is to all others.

Interestingly, groups with lower or higher socioeconomic status, as well as groups under stress, tend to use strong links instead of weak ones (Granovetter, 1983; Killworth and Bernard, 1978). As a possible consequence of this, people under stress and either on the top or at the bottom of society may belong to a more closed world than those living in relative rather than extreme prosperity.

Why do we like small-worldness? Why do we need it? Humans are cooperative animals (Ridley, 1998). Consequently, our brain has developed to keep an inventory of our contacts (Dunbar, 1998). Connection rules have become essential for our survival. Social surveys show that we get used to circles of 5, 15, 35, 80 and 150 people.² However, we tend to meet more and more people and our cognitive abilities are not prepared for this. In a modern megalopolis, we feel lost. The expanding world has become alienating. Our only escape is to continually redefine and segregate a small world from the contact wilderness outside. Not only does small-worldness provide a key for efficient network search, as well as respecting our cognitive limits, but it is probably also a prerequisite for preserving our safety in an alienating modern world.

2.2 Scale-Freeness

Scale-freeness is popular. We like it, because it resembles the most important aspects of our life. What is scale-freeness? Why is it so general? Scale-freeness refers to a type of distribution. Here, a distribution is a statistic of any property: how many do I find from its various values in my system. Let us take an example from the last chapter. If I have 100 acquaintances in my village (let's call it Budapest, because that is its name),³ I can plot their distribution as a function of the strength of

²These circles correspond to (1) our family and best friends, (2) our close friends, (3) our colleagues and acquaintances, (4) our fellow club members, and (5) our 'village', respectively, (Dunbar, 1998; Hill and Dunbar, 2003). The number of circles is five once again. This is the same number of circles that we saw before in the Kleinberg condition (2000): neighborhood, city, country, continent and world.

³The original saying: "Let's call him Gilberto, because that was his name", is from Leon Lederman's wonderful book *The God Particle* (1993). I am not only extremely privileged to have Leon as a supporter of my scientific research trai-

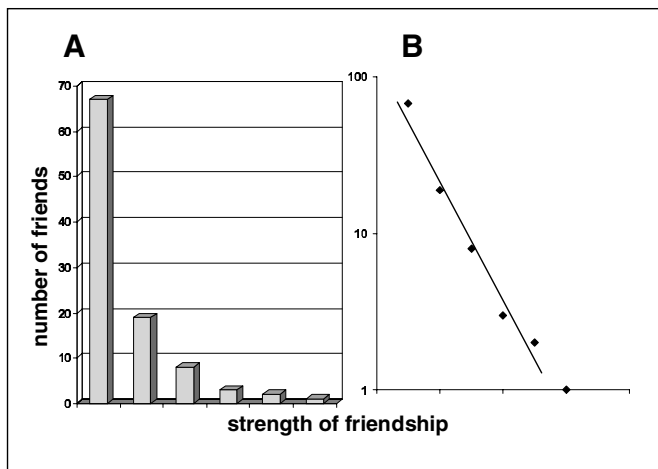


Fig. 2.5. The scale-freeness of networks. The figure shows the approximately scale-free distribution of my hypothetical friends as a function of the strength of our friendship. (A) Histogram. (B) Log-log representation. Note the arbitrary assumption that the strength of friendship increases by an order of magnitude between each of my friendship circles. Moreover, the distribution range is very limited here, which makes the assumption of scale-freeness very inaccurate

our friendship. The histogram of Fig. 2.5A shows that I have one very best friend, two close friends, three friends from the local pub, eight others whom I only see every other week at the bowling club, nineteen from the congregation at my church, and 67 occasional acquaintances from the fitness center, the concert hall and the shopping mall. If the total number of elements in the system is not determined, the actual number may be replaced by probabilities. Taking the above distribution again, from any number of acquaintances, there is a 1% chance that I pick one of my best friends, and there is a 19% probability that I will find a distant acquaintance, whom I see only once a month.

The distribution of scale-free systems follows a power law and can be written as $P = cD^{-\alpha}$, where P is the probability, c a constant, D the distance of our friendship, and α a scaling exponent. Scale-free distributions are best visualized by taking the logarithm of the above equation to get $\log P = \log c - \alpha \log D$, which shows that the logarithm of the probability is a linear function of the distance of our friendship.

ning program for high school students (www.kutdiak.hu/en), but I also treasure a dedicated copy of his book, which became my stylistic template when I started mine.

If we plot the data of the above paragraph in a double-logarithmical way (see Fig. 2.5B), we get a straight line showing that the distribution of my friends as a function of the strength of our friendship follows a power law. It therefore displays the same type of distribution at any scale, i.e., it is scale-free.

2.2.1 Scale-Free Degree Distribution of Networks

The scale-free distribution pattern has been most studied on the degree distribution of networks. What is a degree? The degree of a network element is the number of connections it has. A scale-free degree distribution means that the network has a large number of elements with very few neighbors, but it has a non-zero number of elements with an extraordinarily large number of neighbors. These connection-rich elements are called hubs. If an element has just a few connections, it is often called a node.⁴

Scale-free behavior was long used as an empirical description of experimental data. This model was first developed by Kohlrausch (1854) to explain the discharge in Leiden jars. The first network with scale-free degree distribution was reported a hundred years later by de Solla Price (1965) analyzing the citations between scientific papers. Since then scale-free networks have been reported in all areas of biology, human relations and constructs appearing in every moment of our everyday lives (Barabasi, 2003).

Table 2.1 summarizes the exponents of degree distribution for a few networks. The aim of this list is not to characterize networks. This would be very inappropriate for at least two reasons. On the one hand, networks have a number of important parameters besides their degree distribution characteristics, such as the number of elements, number of links, mean degree, network diameter, clustering coefficient, assortativity, etc. (see Appendix B for definitions). On the other hand, most networks have different exponents for different regions of the degree distribution, or have an exponential cutoff. None of these features are reflected in the data of Table 2.1, which only give a feeling for the multitude and variability of scale-free networks. However, the data of Table 2.1 nicely demonstrate that networks have surprisingly common features in spite of their vastly different constituents and linkage systems.

⁴If the degree distribution of a network is scale free, the network is ‘hub rich’, which means that it has more hubs than a similar random network. Likewise, scale-free networks have no typical degrees.

Table 2.1. Exponents of some scale-free distributions. The exponent refers to the reported range of the exponent α from the equation $P = cD^{-\alpha}$, where P is the probability, c a constant, and D the degree of the network elements. The examples here were assembled from the excellent books and reviews by Albert and Barabasi (2002), Barabasi (2003), Barabasi and Oltvai (2004), Dorogovtsev and Mendes (2002) and Newman (2003). Protein network data are from Bortoluzzi et al. (2003); Chen et al. (2003) and Park et al. (2005a)

Name of distribution	Exponent
Atomic networks	
Occurrence of protein domains	1.6–2.5
Molecular networks	
Prokaryotic protein–protein interaction networks	2.6
Eukaryotic protein–protein interaction networks	2.1
Human protein–protein interaction network	1.7
Gene functional interactions	1.6–2
Yeast gene expression network	1.4–1.7
<i>Escherichia coli</i> metabolic network	1.7–2.2
Biological networks	
Food webs (ecological networks)	1
Social networks	
Scientific collaboration networks	1.2–2.5
Email messages	1.5–2
Zipf’s law for size distribution of cities	2
Phone calls	2.1–2.3
Actors’ appearance in various movies	2.3
Pareto’s law for wealth distribution	2–3
Human sexual contact networks	3.2–3.4
Lotka’s law of scientific productivity	2
Information networks	
WWW (in and out connections)	2.1–2.7
Word co-occurrence	2.7
Citations of scientific papers	3
Technological networks	
Software package parts	1.4–1.6
Internet	2.5
Digital electronic circuits	3
Power grids	4



Not everything is scale-free that seems to be scale-free.

Linear regressions are rather tricky. With a little goodwill, a line can be fitted to practically any number of scattered points. Rather careful consideration is needed to judge whether the result has any real meaning. Here are a few hints to double-check your data, if you believe they may show a scale-free distribution:

- **Data should cover many scales.** A large interval – a scale of several orders of magnitude – is needed to demonstrate scale-free behavior convincingly (Avnir et al., 1998; Eke et al., 2002; Malcai et al., 1997). Scale-freeness proves to be especially difficult to judge in ecosystems, where the number of established links is usually very small (Jordan and Scheuring, 2002).
- **Frequency–degree plots can be misleading.** To make a distribution like the one shown in Fig. 2.5A, you need to group individual data values into certain intervals. The selection of these intervals is rather arbitrary and may dramatically change the nature of frequency-based statistics. The use of cumulative, rank–degree plots is therefore suggested (Tanaka et al., 2005).
- **Correct sampling of the original network is needed.** In many cases we do not know, or cannot analyze the entire network. In most assessments, we use a fraction of the Internet or protein–protein interaction networks. Sampling of non-random networks requires great care and good controls. With biased sampling, the archetypal random graph, the Erdős–Rényi random graph (Erdős and Rényi, 1959; 1960) can be shown to have a scale-free degree distribution, which is clearly wrong, since it has a Poissonian, exponential degree distribution (Clauset and Moore, 2005; Egghe, 2005; Lakhina et al., 2003; Stumpf et al., 2005). Incorrect sampling changes many other parameters, such as the average path length, centrality measures, assortativity and clustering coefficient⁵ of networks (Lee et al., 2005b).
- **Distinct parts of networks may display different distributions.** Part of the above sampling bias arises from the inhomogeneities of the network. This is especially pronounced if the network has modules, e.g., parts of the network where the intramodular contacts are denser than the connections to other modules. In such cases the final distribution can be multiscaled (Tanaka, 2005).
- **Various distributions may resemble each other.** If the range of data points is limited, many distributions, such as the log-normal, stretched exponential and gamma distributions may give rather similar fits to the scale-free pattern. A careful analysis is required to discriminate between these options (Stumpf and Ingram, 2005).

⁵For definitions of these network measures, see the glossary in Appendix B.

- **Scale-free distributions may be ‘false positives’.** A recent report by Deeds et al. (2005) proposed a model for unspecific protein–protein interactions. Here the interactions between hydrophobic surface residues developed a scale-free degree distribution and showed other complex features of network structure as well. This study warns that unspecific data may often overshadow the tiny fragment of meaningful data and may give ‘false positive’ scale-free degree distributions.



The scale-free uncertainty principle: A little network quantum mechanics. Degree distributions rely on the fact that networks are constructed on the basis of paired interactions of their elements. Can all interactions be described as a sum of paired interactions in the Universe? Not necessarily. The present-day network description may be analogous to the Newtonian world view in physics, serving well as a first approximation to describe complex systems. But later, when the possibilities of this approximation are exhausted, we may well end up with ‘second generation’ networks, where unpaired, group interactions or interaction halos between many elements will form the basis of the system, parallelling the change in physics when quantum mechanics and wave functions were introduced by Erwin Schrödinger (1935).⁶

2.2.2 Underlying Reason for Scale-Freeness: Self-Organisation

Why is the scale-free degree distribution so general in such a wide variety of networks? The first explanation of scale-free behavior emanates from the work of Herbert Simon (1955). Simon gave an explanation for the empirical law of Pareto (1897) on the scale-free distribution of wealth. The Pareto law, also called the 80–20 rule, says that a small number of people (20%) own a large amount of property (80%) in a given country, originally in Italy, where Vilfredo Pareto lived. Simon (1955) showed that this unequal distribution is a consequence of the ‘the rich get richer’ effect.⁷ The name reflects the fact that the chances of extending existing wealth are bigger than the chances of accumulating it from scratch. This is also quite clear from existence of so many ‘How did I get my first million?’ stories, implying that to get the second, and other millions is easier. As I will describe in detail later,

⁶I am grateful to the LINK group member, István Kovács for this idea.

⁷The ‘rich get richer’ effect is also called the Matthew effect in sociology (after the gospel of Matthew in the Bible: “For to every one that hath shall be given [...]”; Matthew 25:29; Merton, 1968). De Solla Price called it the cumulative advantage (de Solla Price, 1965) and Makse et al. (1995) referred to it as correlated percolation.



Fig. 2.6. The chances of extending existing wealth are bigger than the chances of accumulating it from scratch

networks with scale-free distributions also have the common feature that they are built up from gradual events, which are often elements of self-organization.

In 1999, a mathematical method to generate scale-free networks was devised by Barabasi and Albert (1999). The invention they used was preferential attachment. Instead of starting from an Erdős–Rényi random graph (Erdős and Rényi, 1959; 1960) or from a lattice, and then changing its degree distribution by various rearrangements, they generated a network by attaching each of the new elements preferentially to those that previously had more connections. Here, indeed, the rich node got richer. Barabasi and Albert (1999) also showed the generality of the phenomenon by analyzing three different networks: Hollywood actors, the World Wide Web, and the US power grid. Their work gave another aspect to the ‘rich get richer’ effect: popularity is attractive.



How can one make a scale-free network? Here are a few remarks on the various methods for constructing scale-free networks:

- **Preferential attachment.** After the seminal work of Barabasi and Albert (1999), preferential attachment became a very popular method for constructing scale-free networks. The preference can refer to the degree of existing elements (Barabasi and Albert, 1999) or to their fitness (Bianconi and Barabasi, 2001; Caldarelli et al., 2002). However, preferential attachment does not always result in scale-free distributions. If the fitness difference between the elements is high and there are a few elements which attract the new elements many times more than the others, a ‘winner takes all’ situation may occur, and a star network develops. In this network, some or one element takes most of the connections (Albert and Barabasi, 2002). On the other hand, ‘weak’ preferential attachment, aging effects (where nodes and hubs tend to lose connections as the network gets older) and growth constraints may all cause crossovers to exponential decay (Albert and Barabasi, 2002). Exponential decay means that there are far fewer hubs in the network than in scale-free networks.
- **Duplication and divergence.** Another method to develop a scale-free distribution is the duplication and divergence method. Here the original network is duplicated and then connections are exchanged. Duplication and divergence is an important possibility for the way scale-free distributions may have developed during evolution (Sole et al., 2002; Vazquez et al., 2002).
- **Scale-free degree distribution is not always the optimal solution to the requirement of cost efficiency.** As mentioned before, in small-world networks, building and maintaining links between network elements requires energy. Therefore, network topology is a result of an optimization process. It is optimized with respect to the available resources to ensure optimal communication between different network parts. If the network enjoys unlimited resources, it will have a random distribution with plenty of links. Scale-free degree distribution occurs as a result of optimization in systems with finite resources. If the network experiences even more limited resources, the degree distribution will be steeper than in the case of a scale-free network, and a transition will therefore occur towards a star network (Amaral, 2000; Sole et al., 2003a; Wilhelm and Hanggi, 2003). These phenomena are called topological phase transitions and will be discussed in detail in Sect. 3.4.

What is the advantage of a scale-free degree distribution? As I mentioned in the last section, small-world networks lie somewhere between random graphs and lattice networks. The degree distribution of random graphs (Erdős and Rényi, 1959; 1960) is Poissonian with a maximal characteristic degree and decaying rapidly. In contrast, all nodes in lattice networks have the same degree. In other words, the degree distribution of lattice networks has a single scale only. The scale-free degree distribution lies in-between the two and, similarly to small-

world networks, allows easy navigation and travel (Barabasi, 2003; Bollobas, 2001; Watts and Strogatz, 1998).



Scale-free networks give sensitive responses. Scale-free topologies enable more sensitive responses to various changes than those allowed by random networks (Bar-Yam and Epstein, 2004). This can be a very important property for explaining why scale-free networks have been selected and maintained in many systems.

For a wide occurrence, it is often not enough for a network to be economical and help easy travel with a minimum number of connections, as we have already seen for scale-free networks. It is important for the survival of the network to resist damage. Damage usually occurs in the form of random errors which incapacitate one or other network element. Scale-free networks pass this requirement, too. Albert et al. (2000) have shown that scale-free networks show a much better stability in the face of such errors than random graphs.⁸

We already have a clue as to how scale-free networks may have developed and why they have remained with us. In the next part of the section, I will try to answer the question as to why scale-freeness is so popular.

Is the source of our inherent affection for scale-freeness the fact that we intuitively feel the scale-free degree distribution of networks? We are all aware of the connection capital of our class-mates or colleagues in the social networks around us. “Lucky guy, he has hundreds of friends. Just picks up the phone and all his problems are solved.” Does this description not fit you? Do not worry if you do not have a thousand friends, it is not your fault. Connection heroes are *by nature* rare. They form the thin tail of the scale-free degree distribution. Somehow, deep down, we all know this. However, our cognitive limits (Dunbar, 2005) make it relatively difficult to keep track of the whole connection network. I might know some friends of my friends but I certainly do not know all of them. Thus our intuitive sense for scale-free distribution does not come from the degree distribution, but emanates from different sources. The solution is space and time.

⁸The error tolerance of the scale-free size distribution of forms in space (fractal patterns) has been mentioned by West (1990).



More on fractals in space. Here I will explain a little more about the term fractal, showing its connection to the idea of fractional dimension and explaining the self-similarity of fractals. For a more detailed explanation, the reader should consult Mandelbrot's book (Mandelbrot, 1977).

- **Fractals, fractal dimension and fractional dimension.** The term 'fractal' comes from the Latin 'fractus' meaning broken or fraction (Mandelbrot, 1977). A fractal object has a fractional dimension. What does this mean? Take a two-dimensional object, for instance, a sphere, and multiply its edge length by 3. You can fit 9 of the old spheres into the new, larger sphere. Taking these two numbers, you find that $9 = 3^2$, which means that the sphere has two dimensions. Writing the above equation in general terms, we get the scale-free distribution of $N = (L/l)^d$, where N is the number of smaller objects fitting into the larger object, L/l is the ratio of the characteristic measure of the two objects of different sizes, and d is the exponent, called the fractal dimension. In fractal objects, d is not an integer number (see Fig. 2.7), but is always smaller than the geometric dimensionality of the object (e.g., smaller than 2, if the object is a two-dimensional object).
- **Fractals and self-similarity.** Fractals are self-similar objects. However, not every self-similar object is a fractal, with a scale-free form distribution. If we put identical cubes on top of each other, we get a self-similar object. However, this object will not have scale-free statistics: since it has only one measure of rectangular forms, it is single-scaled. We need a growing number of smaller and smaller self-similar objects to satisfy the scale-free distribution.
- **Limits of self-similarity.** A mathematical fractal is generated by an infinitely recursive process, in which the final level of detail is never reached, and never can be reached by increasing the scale at which observations are made. In reality, fractals are generated by finite processes, and exhibit no visible change in detail after a certain resolution limit. This behavior of natural fractal objects is similar to the exponential cutoff, which can be observed in many degree distributions of real networks.

Recent evidence indicates that many scale-free networks can be simplified, renormalized to a self-similar, fractal hierarchy of network motifs. This is related to the underlying tree of the network, which is composed of edges with high traffic between network components. This tree is also called the skeleton of the network (Alessina and Bodini, 2004; Garlaschelli et al., 2003; Goh et al., 2005; Song et al., 2005a). The fractal transport systems of self-organizing networks may also explain the allometric scaling laws, which are probably the most famous of the empirical scale-free laws (West and Brown, 2004).

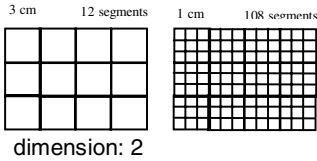
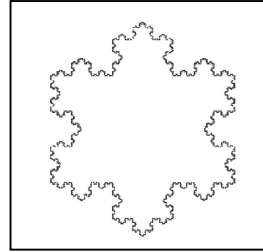
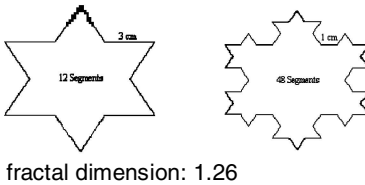
Squares**Koch's curve**

Fig. 2.7. Comparison between non-fractal objects (squares) and a fractal object (Koch's curve). If we decrease the characteristic measure of a square to one third of the original, we can put 108 small squares into the 12 old ones. Taking the equation $N = (L/l)^d$, where N is the number of smaller objects fitting into one of the larger objects, L/l is the ratio of the characteristic measure of the two objects of different sizes, and d is the dimension, we get a value of 2 for d ($108/12 = 3^2$) showing that the square was a two-dimensional object. In contrast, if we do the same with the Koch curve, we obtain $48/12 = 4$ for N , which gives approximately 1.26 for d , showing that Koch's curve has non-integer dimensionality, i.e., it is a fractal. In the *inset*, Koch's curve is shown after a few more recursive steps. With the increase in the number of repetitive steps, the number of scales increases, and the image begins to look like a fractal

**Allometric scaling laws: the mouse-to-elephant curve.**

Allometric scaling laws are probably the most famous of the empirical scale-free laws. These laws cover a wide range of empirical scaling relationships, which show the self-similar behavior of various complex systems such as cells, organs, organisms, etc., over a wide range of masses. In the scale-free relationship here, $P = cM^\alpha$, P refers to the property, c is a constant, M is the mass of the organism or organelle, and α is a scaling exponent, which varies depending on the nature of P . The expression 'allometric' was first used by Julian Huxley (1932) to show the generality of the concept, which describes the mass dependence of the metabolic rate, lifespan, growth rate, heart rate, DNA nucleotide substitution rate, lengths of aortas and genomes, tree height, mass of cerebral grey matter, density of mitochondria and concentration of

RNA, to name but a few (West and Brown, 2004). The most studied of these is the basal metabolic rate, which obeys Kleiber's law (1932). Here P is the basal metabolic rate, i.e., the amount of energy per unit time required by a living organism to remain alive, and the scaling exponent is $3/4$. A limited fraction of this empirical scaling law is often called the mouse-to-elephant curve, describing the universality it represents, since it is equally valid for all mammals from mice to elephants. However, recently, the applicability of the formula has been extended even further, from the largest animals down to cells, and individual enzymes (West et al., 2002). The law has become remarkable, since from geometrical considerations (area-dependent metabolism per volume-dependent mass), one would arrive at a value for the α exponent equal to $2/3$, rather than $3/4$. The deviation from $2/3$ to $3/4$ has been explained by fractal-type transport systems (West et al., 1997), as will be described in Sect. 7.2. In the other examples, the value of the exponent is different: the mass dependence of the heart rate ($\alpha = -1/4$), lifespan ($\alpha = 1/4$), the radii of aortas and tree trunks ($\alpha = 3/8$), unicellular genome length ($\alpha = 1/4$), and RNA concentration ($\alpha = -1/4$), all have different exponents in their scaling relationship $P = cM^\alpha$. However, one has to bear in mind that models always have limitations. Hence the debate concerning a possible over-interpretation of the allometric scaling laws (Dodds et al., 2001).

2.2.3 Scale-Free Distribution in Time: Probabilities

Systems show scale-free behavior in time, too. The probability of the occurrence of a highly connected hub in a network follows similar statistics to the probability of an unusual event. The archetypal example is an earthquake. The Gutenberg–Richter law states that both the occurrence and the magnitude of earthquakes follow a power law (Gutenberg and Richter, 1956).⁹ We know this, although we do not feel it. Fortunately, earthquakes do not happen often enough for us to have any inherent feeling for their statistics, without the meticulous records of decades and centuries. There is a similar scale-free event on a shorter time scale: rain. The dry spells between two rainfalls and the magnitude of the rainfall both follow scale-free statistics (Peters

⁹The Gutenberg–Richter law states that $N = a10^{-bM}$, where N is the number of earthquakes of magnitude M , a is a constant, and b is the exponent. The value of b seems to vary from area to area, but worldwide it seems to be around unity. Here the magnitude M of an earthquake is gauged on the Richter scale, and it is proportional to the logarithm of the maximum amplitude of the Earth's motion. What this means is that, if the Earth moves one millimeter in a magnitude 2 earthquake, it will move 10 millimeters in a magnitude 3 earthquake, and 10 meters in a magnitude 6 earthquake.

and Christensen, 2002). Our nervous ancestors, when watching out for rain over their drying and dying crops, did acquire a sense for the way this heavenly force behaves. Understanding scale-freeness was crucial for life.



The Noah effect. Not surprisingly, Mandelbrot (1977) called the low, but definite probability of extreme events the Noah effect, referring to the great flood in the Bible. In fact, Noah's case may not be the best example of an extreme rainfall, since the hypothesis of Ryan and Pitman (1998) proposes that the Biblical flood was more probably caused by the sudden burst of the Mediterranean Sea into the empty basin of the Black Sea through the Bosphorus Strait. ***"I have an objection. Recent data suggest that the gigantic spill was not completed in the Biblical 40 days but lasted for 33 years, if not longer (Schiermeier, 2004)."*** Well done, Spite! But just imagine that your favorite lake grows 150 meters above your head before your children really grow up. A rather frightening prospect, is it not?



More fractals in time: mono- and multifractals. The functions $f(t)$ typically studied in mathematical analysis are continuous, and have continuous derivatives. Hence, they can be approximated in the vicinity of time t_i by a so-called Taylor series or power series:

$$f(t) = a_0 + a_1(t - t_i) + a_2(t - t_i)^2 + \cdots + a_h(t - t_i)^h + \cdots ,$$

where h is an integer (a whole number). In contrast, most time series found in 'real' experiments cannot be approximated by the above formula. If a non-integer number h is enough to quantify a local singularity in the noisy time series, we call it a fractal series. If we find a single value $h = H$ for all singularities t_i in the signal, then the signal is a monofractal. If we need several distinct values to describe the time series, then the signal is multifractal. (If none of them works, then the signal is not fractal, but has some different behavior.)

Most people in the developed countries do not have to worry about dying crops. World trade provides a safety net against droughts. There should be rain somewhere around the globe. We have a better example to show that scale-freeness is really close to us: the Bernoulli law. No, not those faint high-school memories, if they exist at all, on the fundamental laws of hydrodynamics, which help us to understand and

design airplanes that can actually fly, or sails and boats, or evacuation systems that get the sewage safely out of our flat. I am referring to the St. Petersburg paradox, which was also conceived by Daniel Bernoulli, the most talented of the three Bernoullis from 18th century Basel in Switzerland.

While staying as the guest of the czar with his friend, the great mathematician, Leonhard Euler in Saint Petersburg, Bernoulli was thinking about the chances of winning when tossing a coin. *“Now I know why I hated the Bernoulli law in physics. This Bernoulli fellow was a rather dumb guy. Why even think about this? A coin has two sides. The chance of winning is clearly 50%. You do not need to be famous to figure this out.”* Spite, you are right, but I am afraid you were not patient enough. Bernoulli was contemplating the cumulative wins and losses. The rule in this game is to play until you have a winning series of consecutive heads (or tails): you flip a new coin each time, and you win all the coins you can flip in a row. And what is the probability that you win in a row? In fact, it is scale-free again. If a billion people play the game, then on the average, half will win one coin, a quarter will win two coins, an eighth will win four coins, etc. In this game we always have a chance of winning an order of magnitude higher, but this chance is an order of magnitude smaller (Bernoulli, 1738; Shlesinger, 1987).



The real St. Petersburg paradox and bet-hedging funds. The well-known form of the St. Petersburg paradox is a bit more complicated than the scale-free distribution of successive wins and losses. It is related to the price one is willing to pay for the game which offers an algebraic payment (2, 4, 8, etc.) for a series of successive heads when tossing a coin. The expected value of the game is $2 \times 1/2 + 4 \times 1/4 + 8 \times 1/8$, etc., which is clearly infinite. However, no one is willing to pay even a modest amount for this game, and this is why it is called the St. Petersburg paradox. In order to solve the paradox, Bernoulli introduced the *relative* value of any amount of money, in relation to one's total wealth. This is now called the principle of marginal utility and has become a central element of economics. Analyzing coin tosses, Bernoulli also introduced bet-hedging as a tactic for spreading risk (Bernoulli, 1738). Bet-hedging is achieved by splitting resources, a tactic which makes better (smoother) use of the winning series than putting all one's money up at once. Diversity of action pays well. Bet-hedging has already had an unbelievably successful career. It has been demonstrated that it has been widely used throughout evolution, and it has proved its worth again since the idea was introduced into modern economics. I will return to

bet-hedging as a reason for diversity in Sect. 6.3.



The Joseph effect. Mandelbrot (1977) called the clustering of probabilities the Joseph effect, referring to the seven years of great plenty and the seven years of great famine predicted by Joseph in Biblical Egypt. In fact, the successive high and low flood levels of the Nile and many other rivers are also persistent (Mandelbrot, 1977). The Joseph effect has been observed recently in email communications (Barabasi, 2005), and has been explained by a model based on a decision-based queuing process. However, the underlying cause may be more general, as shown below.

Probabilities give a novel interpretation of preferential attachment (Barabasi and Albert, 1999). To illustrate the link, let me quote the famous saying by Benjamin Franklin:

A little neglect may breed mischief: for want of a nail, the shoe was lost, for want of the shoe, the horse was lost, for want of the horse, the rider was lost, [which was later continued by] for want of the rider, the battle was lost, for want of the battle, the kingdom was lost, and all for the want of a horseshoe nail!

This shows the extremely strong natural sense for chains of unlikely events. Indeed, the successful (or unsuccessful) completion of many sub-tasks will result in a scale-free probability and clustering for the overall success of the complete task (Montroll and Shlesinger, 1982; Shockley, 1957). This is one of the main reasons behind the emergence of Pareto's law, which shows a scale-free distribution of wealth among citizens (Pareto, 1897).



Scale-freeness is related to the self-organisation of the Universe. The surprisingly general occurrence of scale-free properties in space and time and the scale-free distribution resulting from the completion of successive steps raise the idea that this property is tightly linked to the self-organization of matter in the Universe. The scale-free distribution is related to the emergence and maintenance of life (Kauffman, 2000).

2.2.4 Scale-Free Survival Strategies: Levy Flights

The scale-free probability of the several unlikely events mentioned above gives us the impression that we may be able to predict the

unpredictable. This would help us to survive in an increasingly unpredictable modern world. Do we have other scale-free survival strategies? Yes, we do. These are the Levy flights (Levy, 1937). When an albatross, an ant, a bumble bee, a deer, a jackal or a monkey makes a search,¹⁰ the length of individual trips follows a scale-free distribution (Atkinson et al., 2002; Cole, 1995; Ramos-Fernandez et al., 2004; Viswanathan et al., 1998; 1999). In most cases we explore the immediate neighborhood by making a number of small trips, because it is cost-efficient. However, from time to time we make a bigger jump, and on very rare occasions, we go really far to find our target, whether it be fish, pollen, grass, fruit, or the citation beginning the next paragraph.

Viswanathan et al. (1999) showed that the scale-free pattern of Levy flights has a reason. And the reason is simple (see Fig. 2.8). Scale-free Levy flights are the best strategy to minimize the probability of returning to the same site again (a disadvantage of random search) and also to maximize the number of newly visited sites (a disadvantage of the lattice, Brownian-type search). Once again neither random, nor regular brings us the optimum, only something in-between, which is scale-freeness. Note that this is neither rain, nor gambling, i.e., it is neither something like rain, which you have to observe but cannot influence (with the known exceptions), nor something like gambling, which you may avoid (with the known exceptions). If our ancestors, probably right back at the unicellular stage, had not learnt how to plan a search obeying the scale-free rule of Levy flight, they would not have survived. Non-Levy life died out from Earth early on. Scale-freeness may be embedded in our genes.



Definition and comparison of Brownian and Levy-flight search strategies. Both the Brownian and Levy-flight search strategies have the same probability distribution for the step lengths, obeying the equation $P = cL^{-\alpha}$. Note that this is the same equation as we had in Sect. 2.2, where P is now the probability of the given step, c is a constant, L is the length of the step, and α is a scaling exponent. If the search is in two dimensions and the value of α is three, we speak about a Brownian search, whereas if it is two, a Levy flight is performed. For simple models, Levy flights confer a significant advantage over Brownian search in realistic situations, when the

¹⁰Levy flight may not require conscious action. As the writer sits here and types this book, and as the reader reads this line, millions and millions of cells may be making a Levy flight in our body, hunting for food, infectious intruders, or wounds to be healed.

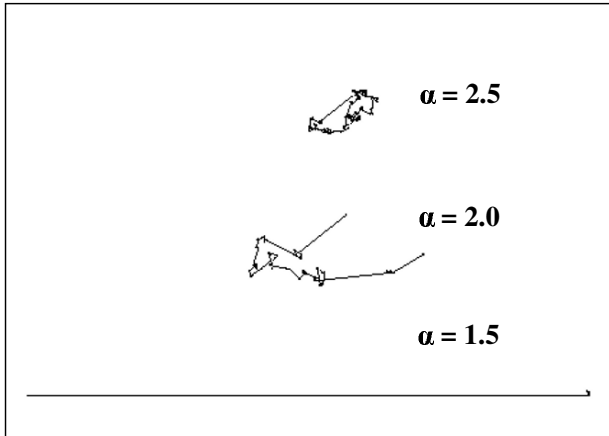


Fig. 2.8. Levy flights are optimal search strategies. In Levy flights, the probability P of the length of a given step obeys the equation $P = cL^{-\alpha}$, where c is a constant, L is the length of the step, and α is a scaling exponent. The search patterns of 1 000 steps are shown for various exponents α . In the model system of Vishnavathan et al. (1999), the case $\alpha = 2.0$, which corresponds to Levy flight, was shown to be optimal. Adapted from Visnawathan et al. (1999) with kind permission

searcher is larger or moves rapidly relative to the target, and when the target density is low (Bartumeus et al., 2002; Viswanathan et al., 1999).

“I do not understand something here. How does a cell make a search strategy? Does she learn it from older cells?” A cell is not complex enough to be conscious. Moreover, Levy flights are characteristic of turbulent vortices (Solomon et al., 1993), anomalous diffusion in polydisperse systems (like Knudsen diffusion; Gheorghiu and Coppens, 2004; Stapf et al., 1995) and electron trajectories (Geisel et al., 1985; Micolich et al., 2001). The common mechanism most probably requires some form of arrest in relaxation processes. This is useful in networks in a low-resource environment since it does not allow dissipation of hard-earned energy. Relaxation-arrest may keep the available resources and also lead to the development of scale-free self-organized criticality (Bak et al., 1987; Bak and Paczuski, 1995; Bak, 1996; Bonn and Kegel, 2003), which may ‘automatically’ induce scale-free Levy flights. Levy flights may characterize nested networks at many levels.



Levy flights of macro-networks. *“If all these nested networks have Levy flight, then should social groups and ecosystems also make Levy flights, when they try to survive or evolve?”* Spite, your imagination is working well today. This is a very good question for future research and will be addressed in detail in Chap. 12.

2.2.5 Scale-Free Pleasures

If a sense for scale-freeness is so crucial for life, have we figured out any way of practising it?¹¹ Yes, one form of scale-free training is called music. Loudness and pitch fluctuations are both scale-free in all classical and folk music from Bach to Mozart, from the pygmies to native Americans. Modern music also follows this trend from Sergeant Pepper by the Beatles to blues and jazz. Here again, random fluctuations produce white noise, the hiss from an unused loudspeaker, which is rather boring. In contrast, synthetic music, obeying regular patterns, produces ‘Brownian’ music, which is too correlated and becomes boring again. Excitement and beauty come with scale-freeness (Gardner, 1978; Voss and Clarke, 1975).

There are a few exceptions, however. The atonal music of Schoenberg and Stockhausen do not fit this rule (Voss and Clarke, 1975). We may have discovered by the 20th century that no more practice of scale-freeness is needed. Are we right? If anyone reads this book in a few hundred years (if anyone reads anything in a few hundred years), they may be able to give the answer.



Bach abridged. The scale-free structure of music gives us an interesting chance to make a musical sample which is similar to the original. Scale-free systems in space are fractals and are self-similar as mentioned above. Scale-free systems in time, like music, behave likewise. If we reduce a Bach composition to half of its original notes by doubling the scale, the ‘half-Bach’ and even the ‘quarter-Bach’ will still sound like original Bach music (composed perhaps when the master was exhausted by the horde of little Bach kids running around, and increasingly sparing with the ornamentation). Interestingly, the final reduction of, e.g., Johann Sebastian Bach’s Invention No. 1 in C Major gives the basic three notes on which the whole piece was

¹¹Levy flights may be part of the ‘inherent response set’ brought with us when we are born. However, any helpful reflex can be overridden by a misdirected mind. We probably need practice to keep our Levy flights alive.

composed (Hsu and Hsu, 1991).



Is art scale free? After discovering the scale-freeness in music, one is tempted to ask: How general is this? Can we find scale-free elements in the composition of paintings or sculptures? How about the composition of Shakespeare dramas, novels or Hollywood films? Should they not remind us of the Levy flights and other evolutionary lessons? I will return to all of these in Sect. 9.2.



Is play scale free? Going one step further, if exciting music and gambling are both scale-free, how about play? Is it only those parlor games and card games where the probability of winning is scale-free that are popular? Are the chances of placing cards scale-free in the game solitaire? Can a careful series of changes in game rules be discovered in football, basketball, tennis or other games, adjusting the excitement of the game to scale-free behavior? All these are exciting questions for further research.¹²



Good schools are scale-free schools. A good school should have a scale-free probability distribution. Without a few strict rules, the identity of the school cannot be defined and the school lacks the discipline helping the socialization of the students (over-democratic, anarchic schools). Random schools are not optimal. However, the other extreme is probably even worse. An over-disciplined school having almost exclusively strong links efficiently kills all creativity and playfulness (like the Prussian-type schools). Lattice schools are not optimal either. In a good school, not only link strength but also unusual events should follow a scale-free distribution. What does this mean? Most of the time, nothing extraordinary happens. However, from time to time a good school needs an unusual hour or day. The more unusual it is, the more infrequent it can be. But unusual situations should go quite far from the normal, even to extremes at times. And they must occur in a rather random fashion, as in all of our play, as in art, and as in all human activities and institutions, which prepare us for the scale-free probabilities of life itself. A school which either gives too much freedom or is not free enough to create this scale-free behavior serves this key purpose poorly. It is no wonder that students dislike both. Prison schools are disliked instantly, while schools with extreme freedom leave no memories. There is one more important aspect here. Schools socialize us with respect to our society. We learn about the norms of science and social life in schools. Schools tell us what has become acceptable and what has not during the development of

¹²I am grateful to the LINK group member, István Kovács for this idea.

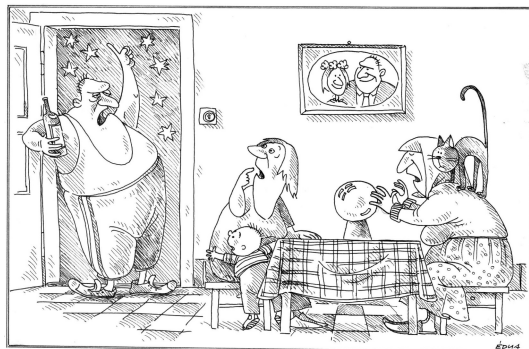


Fig. 2.9. The unpredictable becomes domesticated: the scale-free distribution of unpredictable events gives the impression that they are understandable

human knowledge over the past hundred thousand years. Schools pinpoint certain elements of our knowledge network as reliable and important, and suppress others as superficial, outdated, ridiculous or forbidden. With this labeling, schools strengthen certain links in our knowledge network. If schools do not add many more weak links to the knowledge network by building up creative, playful associations, and providing a rich environment of social interactions, the whole knowledge network will become unbalanced and rather ineffective.¹³ Besides teaching the disciplines, a good school also gives ample time for the discussion of the ‘non-scientific’, ‘soft’ part of the world. *“I will print this out and pin it on the wall in my class. My geography teacher will experience a moment of truth, if he understands it.”*

In conclusion, scale-freeness not only helped our ancestors to survive though the eons of evolution by Levy-flight searches for food and any other goodies we needed, but helps us today to cope with our growing uncertainty about the globalized world. The unpredictable becomes domesticated: the scale-free distribution of otherwise unpredictable events gives the impression that they are understandable. Scale-freeness surrounds us in our environment and appears in all kinds of art to remind us of this task: to be prepared! So there is nothing left but to ask you to be prepared for the next reason why we like networks: nestedness.

¹³I am grateful to Gergely Hojdák for this idea.

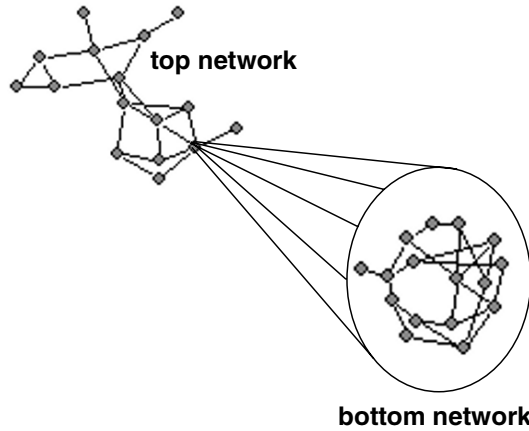


Fig. 2.10. Nested networks. In real, self-organizing networks, each element of the (top) network usually consists of an entire network of elements at a lower (bottom) level. The elements of the lower level form bottom networks. As an example, a top network can be a network of neurons, where the elements are neural cells. Here each top network element, each neural cell, is a network of proteins. Proteins of a single neural cell form a bottom network in this example. Other examples of top–bottom network pairs can be found in Table 2.2. In real, self-organizing networks, both the top and the bottom networks have more elements than the simplified examples in the figure

2.3 Nestedness

Networks are like Russian dolls. They lie inside each other. My memorable physics teacher, László Holics, used to tell us in class: “Look at this dot. If you examine it closely enough, it is infinite. If you look at it from far enough way, it is a dimensionless point.” Networks behave just like this. Most of them can be regarded as structureless elements of a higher order network which I will call the top network in the rest of the book. Similarly, most elements of the top network are themselves whole networks, with a complex structure and stability. These networks, which are elements of a top network, will be called bottom networks in the rest of the book (see Fig. 2.10 and Table 2.2.). In other words, bottom networks are nested in their top network. Our world is an onion, which can be peeled and peeled again. Reduction of complex networks to simple elements and the rediscovery of complex networks in these elements are recurrent features of our thinking. From the moment a child disassembles her first watch and encounters the complexity within it, nestedness has been born.

Table 2.2. Examples of top–bottom network pairs

Top network	Elements of top network	Bottom network	Elements of bottom network
World economy	Countries	Social network	People
Social network	People	Cellular network	Cells
Cellular network	Cells	Protein network	Proteins
Protein network	Proteins	Atomic network	Atoms

In fact, nestedness is a rather old idea even in the network field. System hierarchy was emphasized by von Bertalanffy as early as 1950. Feibleman (1954) pointed out that the mechanism of any level is found at lower levels, while the purpose of any level is found at levels above. Koestler and Smythies (1969) called nested networks a nested hierarchy of holons. Holons have a Janus behavior. Each higher level is a holon to the lower level. However, the same holon breaks into parts and behaves as an assemblage of them if seen from one level lower. Eldredge (1985) arrives at nestedness (what he calls genealogical, historical and ecological hierarchy) from the genealogy of nested taxa. Nestedness was also described by Oltvai and Barabasi in 2002.



My first encounter with nestedness. For me nestedness did not come in the form of the incapacitation of my father’s favorite watch. I had cheese. I still remember the endless minutes I spent at the breakfast table as a three-year old naturalist. I was perplexed. In my hands I had a piece of cheese. However, it was not the cheese but the box which showed me one of the great wonders of the world. On the box there was a bear holding up the same box of cheese. Even on the second box you could see the silhouette of the next bear holding up the next box. Then the resolution gave up and no more bears, no more boxes were seen. My imagination soared. I suddenly saw thousands, millions of bears and boxes inside each other. At the time (after some parental explanations), I thought I had discovered the infinite. Today I know that this was my first encounter with nestedness. Nestedness is an enchanting principle. Once you have felt it, it stays forever.

How was nestedness born? How do networks start to behave as elements and assemble into top networks? Symbiosis-driven nestedness, also called integration by Sole et al. (2003a), occurs if a network enjoys a longer period of relative stability and then, by associating with other

networks, all become elements of a top network (Sole et al, 2003a). This self-development is not a purely theoretical assumption. The major transitions in evolution, e.g., formation of the first cells and the first eukaryotes, the first sex, the first multicellular organisms, the first social groups (Margulis, 1998; Maynard-Smith and Szathmary, 1995) may all be regarded as instances of symbiosis-driven nestedness. Life itself requires a higher level of organization, since it is possible only in a macroscopic system. On small, molecular scales, the order required for life would be destroyed by microscopic fluctuations (West and Deering, 1994).

Symbiosis-driven nestedness is the bottom-up, self-organizing, assembling approach to forming a top network. However, there is another way. The top-down approach, the self-structuring, discriminating segregation of elements in the top network. This may be called modularization-driven nestedness. This segregative process is referred to as parcellation by Sole et al. (2003a). Here the already-assembled top network first forms modules. Later on, modules can develop into elements of a top network. Module formation and the discrimination between modules and elements of a top network can teach us many lessons. So let us have a detailed look at these processes.

2.3.1 Network Modules

Module formation is a general feature of most networks. Modules are also often called communities (see Fig. 2.11). According to their most widespread definition, modules contain elements which have more links inside the module than outside (Hartwell et al., 1999; Radicchi et al., 2004; Wasserman and Faust, 1994).



More definitions for modules. When we look at a graphical representation of a network, it is often easy to recognize its modules. However, the exact definition of modules is rather difficult. Here I will list some of the available definitions (Radicchi et al., 2004; Wasserman and Faust, 1994). *“I like your approach, Peter, trying to give us as much information as you can in an extremely condensed format, but do you really think we need your survey on the available module determination methods? How about just drawing the network and simply seeing these modules? Would this not be easier?”* If you check the network visualization tools, Spite (a short list of available Web sites is given in Appendix A), you will see that visualizations of networks often inherently contain the modularization protocols. If you want to draw a network ‘nicely’, you

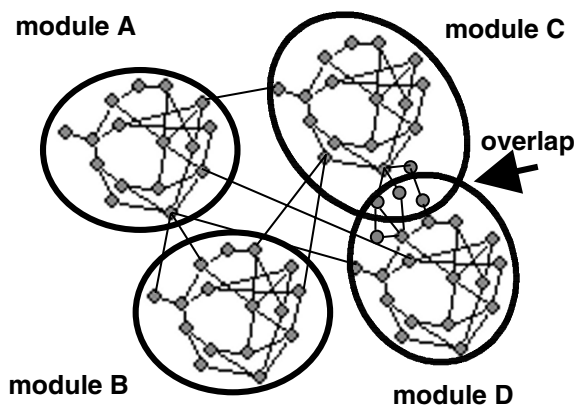


Fig. 2.11. Example of a modular network. The hypothetical network shown here can be divided into 4 modules. Modules A and B, B and D, and A and C do not overlap. However, modules C and D have an overlap region, also called the fringe area. Overlaps play an important role in the defense and regulation of networks, as will be discussed later. In real, self-organizing networks, the modules have more elements than the simplified examples shown here

should draw the highly-connected parts close to each other and the outlying elements further from the rest. What is this, if not module formation? We have no other choice, we need to know the module definitions in detail. So here they are:

- **k -cliques.** In the strongest sense, the elements of the module form a clique, meaning that each of them is connected to all the other elements belonging to the same module. As a modification of this definition, the elements of the module form interconnected k -cliques, where k is the number of elements in each clique.
- **LS-sets.** In the strong sense, each element of the module has more connections with other elements of the module than with the outside. In an even stronger sense, this criterion is extended to each subset of the module. Such a module is called an LS-set.
- **Modified LS-sets.** The above criterion can be applied to the sum of the links, meaning that the sum of the outward links of the whole module should be smaller than the sum of the intramodular connections.
- **k -cores.** In a weak sense, the module is a k -core, meaning that the elements of the module are connected to at least k other elements of the same module (Radicchi et al., 2004; Wasserman and Faust, 1994).
- **Co-sets.** A different type of definition arises from the dynamical properties of networks. In this context, co-sets are reaction sets that always occur simultaneously, and these have been defined as network modules (Papin et al., 2004).



How can modules be discriminated from one another?

Here is a brief survey of the methods available for module determination:

- **Clustering methods.** Clustering methods are often called agglomerative since, in one of their representations, all the initial links are deleted and then rebuilt again, starting from the most closely linked communities. In one of these methods, the original network structure is ‘renormalized’ in such a way that groups of tightly connected elements are considered as one cluster, and this step is then repeated. In this method the network is finally transformed into a dendrogram (Wasserman and Faust, 1994). A very similar procedure has already been mentioned when we talked about the fractal properties of scale-free networks (Alessina and Bodini, 2004; Garlaschelli et al., 2003; Goh et al., 2005; Song et al., 2005a).
- Clustering can be performed by random association of network elements, where the resulting cluster is probed by a modularity parameter derived from a module definition above. Modules emerge as the optimization of this modularity parameter proceeds by successive selection of elements (Newman, 2004).
- A special version of clustering methods uses the Potts model of superparamagnetic clustering, where clusters are identified by assigning a ‘spin’ to each element. In this method neighboring spins interact with each other, and tend to form clusters, where elements have the same spins (Spirin and Mirny, 2003).
- **Divisive methods.** The gold-standard of these methods utilizes the special feature of intermodular, bridging links, which comes from the definition of modules. If the number of intermodular links is much smaller than the number of intramodular links, then intermodular links should be parts of a much larger number of shortest routes linking two network elements. In other words, the betweenness centrality of intermodular links is larger than that of the intramodular links. The method finds the link with the largest betweenness centrality and removes it. This procedure is repeated until the original network is split into two modules. The method is costly in computation time, since the betweenness centrality changes by the removal of any links, so that it has to be recalculated after each removal (Girvan and Newman, 2002; Newman and Girvan, 2003).
- The edge-clustering method uses the feature that intermodular links are only very seldom parts of link triangles. Deleting links with the lowest amount of edge clustering rather effectively dissects the network at modular borders (Radicchi et al., 2004).
- Another divisive approach starts the procedure with a random split of the original network. Using the module definitions above, the method calculates the ‘fitness’ of each element after the initial split, and moves the element with the lower fitness to the other module. When an optimal overall fitness is reached, the two modules are ultimately separated and

the process starts again with the two modules (Duch and Arenas, 2005). A similar method has been worked out recently by Guimera and Amaral (2005) using simulated annealing to find the maximal number of modules (reasonable splits) in the network.

- Clustering similarities (shared common neighbors) and shortest path similarities have been combined in a recent method to provide a more complex grouping of network elements into appropriate modules (Poyatos and Hurst, 2004).
- **Fuzzy clustering methods.** Neither clustering, nor divisive methods give module overlaps, if executed in their original sense. Module overlaps are also called fringe areas and are important elements of the regulation and evolution of complex networks (Agnati et al., 2004). Reiherdt and Bornholdt (2004) worked out a fuzzy version of the Potts model described above. Here, not only is the agglomeration of identical spins rewarded, but the final fitness parameter also has a second term which favors identical spins over the entire network. A proper ratio of these counterbalancing terms identifies overlapping network modules. Fuzziness may be introduced into the network description by adding a random noise to the weights of the network (Gfeller et al., 2005).
- **Fuzzy divisive methods.** Fuzziness can be introduced into the ‘traditional’ Girvan–Newman method (Girvan and Newman, 2002) by a random selection of the deleted link from those having the highest betweenness centrality. This method is based on the observation that overlapping elements end up in different modules, if we vary the order of deletion of links with a large betweenness (Wilkinson and Huberman, 2004).
- **Topological overlap methods.** These methods utilize the feature that elements of the same module tend to have a much larger overlap of higher order neighbors than elements from different modules (Ravasz et al., 2002; Yip and Horvath, 2005).
- **Network walk methods.** Overlapping modules can be obtained in the strong sense by a k -clique method, where overlapping subnetworks of completely linked k -element subgraphs are identified by a continuous exploration (a network walk) as modules. The method starts with the deletion of a fraction of weak links which usually connect modules, and hence gives overlapping modules in their strongest sense (Palla et al., 2005). Methods using network walks in their weaker sense are currently being developed (Kovacs and Csermely, unpublished work).
- **Dynamical methods.** A new class of approach will take into account the dynamical properties of networks. As an initial attempt, co-sets, i.e., reaction sets that always occur simultaneously, have been determined as modules (Papin et al., 2004).

Modularization is a spontaneously occurring property of networks, where the links are gradually reorganized. Module formation is related to the fractal growth of networks. Modules may also result from

the duplication of a network segment and the subsequent divergence of the two arising modules in further evolution. These parallel processes may give rise to modular structures ‘automatically’, and make them extremely widespread in the organization of life on Earth (Guimera et al., 2004; Sole and Fernandez, 2005; Song et al., 2005b). As an example of this, 74% of known metabolic enzymes of the bacterium *Escherichia coli* are clustered together in modules (von Mering et al., 2003).



Modules as fossils of the past. Symbiotic network modules may be a good source for uncovering the history of network formation. Modules of symbiotic, top networks often conserve an isotemporal cluster of elements which can be identified by detecting the amount of random errors (mutations) in them or knowing their evolutionary history (Kunin et al., 2004; Qin et al., 2003).

Formation of modules seems to be a natural consequence of network development. What is the driving force? What are the benefits of the modular structure?

- Modules localize the effects of deleterious perturbations (Maslov and Sneppen, 2002).
- Modules can evolve rather independently (Hermisson et al., 2003; Kirschner and Gerhart, 1998) and give an advantage, if design specifications (environmental conditions) change from time to time (Alon, 2003).
- Modules allow a separation of various network functions and decrease the cross-talk between them (Maslow and Sneppen, 2002). Modularization usually induces divergence and may be a good source of diversity. All of these things stabilize networks and increase their chances of faster evolution (Kirschner and Gerhart, 1998).
- Modules are not always separated and the module connection is flexible. The intermodular boundary is often not a line but an area, called the fringe area. Fringes can either facilitate or prevent contact between neighboring network elements. The facilitative or occlusive behavior of fringe areas may change from time to time (Agnati et al., 2004).

Modules often have a hierarchical organization, which has been shown to occur in metabolic networks (Ravasz et al., 2002) and in protein folding (Compiani et al., 1998). The organization of the modular hie-

rarchy may follow the scale-free pattern. However, a central module may collect most links to other modules, behaving like a dictator or a black hole. This corresponds to the star structure of networks and reveals an increased level of environmental stress. We shall return to this when describing the topological phase transitions of networks in Sect. 3.4.



Rich clubs and VIP clubs. The hierarchical structure of networks often leads to the development of an inner core. The yeast protein interaction network has a core network of essential proteins, which has an exponential degree distribution (Pereira-Leal et al., 2004). The inner core becomes a rich club if it is formed by the hubs of the network. Rich clubs are formed by Internet routers (Zhou and Mondragon, 2003), large cities in transportation networks (de Montis et al., 2005), etc. Rich clubs are typical of assortative networks, such as social networks, where hubs tend to associate with hubs. However, the rich-club phenomenon extends to disassortative networks, where the non-typical links between hubs become those with the highest traffic. VIP clubs are different from rich clubs. In VIP clubs the most influential members have a low number of connections. However, many of these connections lead to hubs, which can mobilize the rest of the network. Thus VIPs might serve as an interconnected, closed group of the elite, which may even become masterminds of the whole network, as has been demonstrated by Masuda and Konno (2005) for the network of the world's best tennis players.

A rather interesting and hitherto poorly-studied question concerns the nature of the discriminative step between module formation and the behavior of these modules as elements of a top network. *“I guess I have a reason why this interesting question was ‘hitherto poorly-studied’: it is not appropriate. You said just a few lines above that it is only a question of perception how we define an element of a larger network. If we look at it from a distance, it is a point, but if we go closer, it becomes a whole bottom network. I presume we will have the same inherent difficulties in the discrimination between modules of the bottom network and elements of the top network.”* Spite, in a way you are right, discrimination is indeed difficult. However, we have to achieve this discrimination if we want to decide how to describe the network with our one-dimensional information flow, with our words. Before starting to talk about it, we need to know whether it is

a module or an element. Let me list a few discriminating features. The higher the values, the more likely it is that the modules of the bottom network will become elements of the top network:

- number of modules,
- network size per average module size,
- structure of modules,
- average number of intramodular per intermodular links,
- weight of intramodular per intermodular links,
- difference of modules from each other,
- availability of module as an independent network,
- transient and unnecessary intermodular links.

Do you know, Spite, if you think over all these steps, what do you imagine? A fantastic step in the self-organization of matter in the Universe: the development of a new layer of nested networks. The gradual enrichment of nestedness is one of the most beautiful events in our past. Without this, neither you Spite, nor me, nor the reader would be sitting here enjoying this book.



Are modularization and network nestedness evolutionary tools to restore the stability of overgrown networks? Are the above examples of spontaneous module development (Guimera et al., 2004; Sole and Fernandez, 2005; Song et al., 2005b) general? In other words, do self-organization and growing network complexity automatically destabilize networks? Is modularization and the subsequent development of a top network an evolutionary tool to restore the stability of overgrown networks? ***“This is a pseudo-question. You nicely showed that modularization and the subsequent development of top networks are unavoidable steps of self-organization. You even assured us that we could not enjoy your book without this. I beg your pardon, but why are you asking this?”*** In spite of the seemingly obvious answer, it is worth asking this question, Spite, since the automatic affirmative has a number of very interesting consequences. If all these processes are spontaneous, the driving force behind cell division might not be just the well-known relationship between cell surface area and cell volume,¹⁴ but cell division may also be triggered by

¹⁴The growth of the cell surface area is proportional to the second power of the number of constituent molecules, while the growth of the cell volume is proportional to the third power. Thus the cell volume occasionally outgrows the maximal possible transport provided by the cell surface.

an increase in protein network size and complexity. Similarly, the Parkinsonian modularization of social groups¹⁵ may reflect a stabilization of growing social networks.



Is modularization related to the extent of network resources? In Sect. 3.4, I discuss the topological phase transitions of networks described by Tamas Vicsek and coworkers (Derenyi et al., 2004). The preferential random configuration of networks enjoying a resource-rich environment may inherently mean that modularization is triggered by relatively lower resources of network development. However, a formal proof for this is missing.



Should we expect the modularization of the Internet, world economy and Gaia?¹⁶ To take this a little further, could we envision the parcellation of other growing networks, like the Internet, world economy, Gaia, and so on? Is there any rule for this? The Internet and the world economy certainly grow. Is it possible that their growth will lead to their modularization, if the resources cannot support the current growth rate? If we think about a sudden modularization of the world economy, it may easily lead to rather violent events, such as an economic crash or war. An even more interesting question concerns the ecosystem of the whole Earth, Gaia. Does this grow? Currently, we cannot even guess the answer. However, a sudden modularization of Gaia may lead to cascading extinctions of several species like ourselves.



Is modularization a threat? *“Peter, first you described modularization as a joy, which led to our appearance on Earth. Then, in the last few paragraphs, you described the same modularization as a life-threatening event, which we should avoid. I am confused. Is modularization good or bad for us?”* If we are the top network, modularization of our bottom networks is generally a joy for us,

¹⁵The Parkinsonian modularization of social groups was first demonstrated by the cyclic growth in the number of members of the British Crown Council, which was accompanied by the development of inner circles of real power (Parkinson, 1957). A possible reason behind this particular example might be that the number of persons able to conduct a meaningful conversation leading to an efficient decision-making process is rather limited (Dunbar, 2004). However, there may be a general explanation behind the same phenomenon.

¹⁶I am grateful to Bálint Pató for these questions.

since it shows their growing complexity.¹⁷ However, if we are the bottom network, the modularization of our own network or the top network above us may be a threat, especially if it happens as a fast transition. We need time to adapt to such a change. We need a better understanding of these events to predict and control them. *“Wow, this sounds like something really urgent. What if our understanding lags behind? Do you have any suggestions for NOW?”* My instinctive answer to your question, Spite, is that we should be very cautious when we grow. Until we know more about the sudden phase transitions of network modularization, it is better to avoid them.

2.3.2 Intermodular Weak Links: The Role of Creative Elements

In the last section, network modules emerged as rather enigmatic groups of the network organization. On the one hand, they do have a certain level of identity and isolation from the rest of the network, while on the other, the network should integrate its constituent modules. Such seemingly contradictory tasks can be accomplished by intermodular weak links. These long-range connections are THE weak links in the sense of the original Granovetter concept (see Chap. 1; Granovetter, 1973; Kossinets et al., 2008). These bridging contacts have also been called bottlenecks, referring to their key position with regard to information spreading within the network. The importance of bottlenecks was demonstrated by their enriched content of essential proteins in yeast cells. Bottlenecks often form transient weak links, since these proteins are significantly less well co-expressed with their neighbors than other members of the protein–protein interaction network (Yu et al., 2006).

Weak links are essential to provide stability in a dynamic network. Contrary to expectations, Palla et al. (2007) found that, in a social network of telephone users, group stability was increased by the dynamic change of several group members over a longer time. In contrast, those groups that were absolutely stable in the sense of having a constant, unchanging membership, were quite often dissolved, and vanished in the long run. This is once again a dynamic example of the stabilizing role of weak links within and across network modules. A further example comes from the analysis of human warfare. Zhao et al. (2009)

¹⁷It is an exciting question in itself, if the numerical measures of complexity which I describe in Sect. 4.4 do indeed grow through growing modularization.

showed that internal network dynamics, i.e., the emergence of inter-modular weak links, greatly prolongs the minority group's survival time. This behavior is an underlying reason for the extraordinary dynamism of guerilla or terrorist networks. The general model of Zhao et al. (2009) may be relevant in the explanation of the long-term survival of latent diseases, or metastatic cancer cells.



Intermodular weak links as killers: The cases of epilepsy and schizophrenia. Intermodular contacts are commonplace in our brains. As the reader is going through these lines, thousands of modules are formed and dissolved again in the reader's brain. A certain nerve cell may find itself in the middle of a module at one moment, and as a bottleneck between several modules at the next. However, a hyperconnectivity of the so-called resting network in the brain has been observed in schizophrenia by Whitfield-Gabrieli et al. (2009), and has been shown to contribute to the misdirection of attentional resources, increased mind-wandering, hallucination, etc., features typical to this disease. If intermodular connections became extremely dense, the even more severe symptoms of epilepsy may also develop (Chavez et al., 2008). The destabilizing role of an extreme excess of intermodular weak links is similar to the destabilizing role of abundant links in general, which will be shown in Sect. 3.3 (Watts, 2002).



Intermodular weak links: Key players in stress management. We have seen in the previous box that a great surplus of intermodular weak links makes the network overconnected and unstable. However, just the appropriate amount of intermodular weak links is life-saving in stress. Stressed organisms decouple their modules from each other. As an example, the modules of the yeast protein-protein interaction network exhibit a much smaller overlap when the hosting yeast cell experiences any of a wide range of stressful events (Szalay et al., 2007; Palotai et al., 2008; Mihalik et al., 2008). Module decoupling is beneficial, since it helps a more specialized, economical life during the energy shortage in stress. Moreover, various kinds of damage cannot propagate so easily through decoupled modules, providing an additional safety measure.



Intermodular weak links: A possible mechanism for modular evolution. Decreased intermodular contacts help organisms to survive during stress (Szalay et al., 2007; Palotai et al., 2008; Mihalik et al., 2008). What happens once the survival is completed? What does a post-

crisis world look like? Regretfully, we have only rather sparse experimental evidence regarding the recovery of cellular or other networks after stress. One of the reasons is human: crisis is abrupt, recovery is prolonged. Many scientists believe that, if they want to make a spectacular career, they need to do it fast. This usually prevents specifically long-term endeavors, such as observation of lengthy stress recovery processes. Therefore we may only make an educated guess at the moment. If the remaining links between the modules are few and weak during stress, than the probability of getting back exactly the same modular coupling after stress is rather low. This is not actually a disaster, but a great piece of luck. Formulating this in a different way, after each stressful event, the organism has a chance to rebuild itself slightly differently than it was previously. How great is the difference? Most of the time it is infinitesimally small, but sometimes it is larger, and in very special cases, one could not even guess that the two entities, the pre-stress and post-stress organisms, were actually the same. These latter scenarios provide a network-based mechanism for modular evolution (Korcsmaros et al., 2007).

A special type of intermodular position at the molecular level is provided by the active centers of proteins and enzymes. The amino acids of these active centers:

- occupy a central position in amino acid networks;
- often have many neighboring amino acids;
- have non-redundant connection sets with their neighborhood;
- integrate the communication of the whole amino acid network;
- are individual, and do not participate in the dissipative motions of ‘ordinary’ residues; and
- collect most of the energy of the whole amino acid network (Chennubhotla and Bahar, 2007; del Sol et al, 2006; Liu et al., 2007; Piazza and Sanejouand, 2008).

In summary, active centers influence the communication of all other network elements while maintaining their individuality.

The above summary of protein active center properties sounds like the characterization of a mastermind, broker, innovator or network entrepreneur of the social networks of human societies. Indeed, we may find similarly central and individual elements in other complex systems. In protein–protein interaction networks a good analogue of an active centre is a date hub. Date hubs are proteins forming different complexes with different subsets of their partners at different times (Han et al., 2004). Date hubs are enriched for structural flexibility, and this makes them different from other proteins, which are structu-

rally more constrained (Kovács et al., 2006). Date hubs are preferentially located in the overlaps of multiple modules, so their connections are non-redundant and unique (Wilkins and Kummerfeld, 2008). As an example of date hubs, molecular chaperones have an intermodular localization and are amongst those date hubs that constitute the true central coordinators of the cellular network (Komurov and White, 2007; Palotai et al., 2008). Several chaperones are also called stress proteins or heat shock proteins, and their importance increases in protein–protein interaction networks after stress (Palotai et al., 2008). In other words, chaperones become key integrators when the cell experiences an unexpected situation, i.e., a stress. Central elements of signal transduction networks have been termed ‘critical nodes’ by Ronald C. Kahn and coworkers (Taniguchi et al., 2006). Critical nodes often have many isoforms, which shows the importance of the need for ‘back-ups’ for active centers as well as the need to extend the variability of these key elements of signal transduction further.

Going several levels of integration higher to the level of mammalian networks, dolphins occupying intermodular positions in dolphin communities were shown to act as brokers of social cohesion for the whole group. If a key, connecting dolphin leaves the group, the original group breaks into two subgroups. If the dolphin returns, the two subgroups become unified again (Lusseau and Newman, 2004). In social networks the archetype of the above unique, intermodular element is the ‘stranger’ described by one of the forefathers of sociology, George Simmel, a hundred years ago (1908). The stranger is different from anyone else. The stranger belongs to all groups, but at the same time does not belong to any of them. A later, well-known example came from Ronald S. Burt (1995), who proved that innovators and successful managers occupy ‘structural holes’, which are exactly the non-redundant, centrally connecting positions of the active centers in protein structure networks. People bridging structural holes have ‘weak links’, e.g., they often change their contacts (Burt, 1995).

Going one level higher in the hierarchy again, having a central position also offers a great advantage to groups. As an example of this, biotech companies with diverse portfolios of well-connected collaborators were found to have the fastest access to novel information, and governed the evolution of the field. This was only possible in the long run if most of these connections were transient (Powell et al., 2005). The transient, far-reaching, exploratory contact structure helps performance only in those cases where the tasks are novel (e.g., those emerging in uncertain environments or in crisis) and require creative

thinking to solve. Conversely, if the task is typical, and the expertise that is already present within the group is enough to solve it, the maintenance of exploratory contacts is costly and hinders performance (Hansen et al., 2001; Krackhardt and Stern, 1988).

The above analogies enrich the characteristics of intermodular, creative elements. These creative elements are not only central, having a unique set of properties and integrating the communication of the entire network, but they also perform a partially random sampling of the whole network, and connect distant modules. Creative elements have transient, weak links leading to hubs in the network, and become especially important when the whole system experiences an atypical situation requiring a novel, creative solution. The properties of creative elements both require and predict each other, and therefore make an integrated set of assumptions as shown in the following:

- **Autonomy and transient links.** Creative elements are the least specialized, and are the best among all network elements to conduct an individual, autonomous life, independent from the rest of the network – this independence explains why they might, and should, continuously rearrange their contacts.
- **Transient links and structural holes.** Creative elements must connect elements that are not directly connected to each other. If creative elements introduced their new and unexpected content to multiple sites of a densely connected region, they would create extremely large cumulative disorder, which would be either intolerable or would lead to a permanent change instead of a transient change. For the same reason, creative elements must connect to hubs to allow either the dismissal, or the fast dissipation of their novel content.
- **Structural holes and network integration.** If an element connects distant modules (with transient weak links leading to the generation of important positions of the modules involved), this element performs a continuous sampling of key information of the entire network, and therefore has a central and integrating role in network function.
- **Network integration and creativity.** If an element accommodates key and representative information of a whole network, it (a) may easily invent novel means to dissipate an unexpected, novel perturbation, or (b) may connect distant elements of the network with ease and elegance, helping them to combine their existing knowledge to cope with the novel situation. The reformulation of the original problem (by translating it from one distant element to

another), the generation of novel associations and novel solutions, flexibility, divergence, and originality are all well-known hallmarks of human creativity (Csikszentmihalyi, 2008).



The transient life of creative elements. In most networks the status of the creative element is, by itself, transient. Creative elements may well be transformed into task-distributing party hubs (Han et al., 2004; Komurov and White, 2007; Wilkins and Kummerfeld, 2008) or bridges which preferentially connect two modules with strong links, or into problem-solving, specialized elements. These transformations of creative elements usually happen after repeated stress, showing that the network ‘learned’ the novel response by reorganizing its topology, and provides the first unusual, creative solution in a regular, reliable and highly efficient manner (Szalay et al., 2007). This ‘commercialization of creativity’ may explain why signaling networks have isoforms of their critical nodes (Taniguchi et al., 2006) which may replace each other in a redundant fashion when one has become engaged continuously with a specific task.



Our brain, where creativity is probably present everywhere. Our brain might be an especially unique place of creativity. The continuous remodeling of neuronal modules most probably gives an opportunity to all neurons to occupy a very special, central, intermodular position. Thus, the brain very likely offers ‘five minutes of fame’ to all of its members.¹⁸ *“Peter, this assumption is really nice – I do feel the competition of neurons in my brain to be the creative element of the next moment, and to enjoy the five minutes fame. By the way, your ‘five minutes of fame’ seems to me too much. My brain is certainly faster than that. But what happens with creativity in our other tissues? Are all our muscles, heart, eyes, ears, and face boring and dumb?”* Not at all, Spite, these tissues also have their creative cells. These cells are called stem cells. Stem cells maintain tissue renewal. Regretfully, as we grow old, the creativity of our stem cells becomes diminished. But an extremely large, uncontrolled level of creativity is also not helpful, as you will see in the next box.



When creativity becomes dangerous. Creative elements add random elements to network behavior, inducing an increase in noise. This is highly beneficial to a certain extent as we saw in the previous box, but be-

¹⁸I am thankful for Zoltán Nusser for this idea.

comes intolerable if it exceeds a certain threshold. This threshold is high if the hosting network lives an individual life and often meets unexpected situations. However, the same threshold becomes low if the hosting network is part of a higher level organization which provides a stable environment. An excess of creative cells, e.g., violently proliferating tumor cells, significantly disturbs the regular functions of the hosting network and finally causes its disintegration, then death. As another example, symbiotic organisms which have become engulfed by another shed off a large section of their network variability (Pál et al., 2006), probably including most of their creative elements.

Creative elements are the luxury of a network operating in ‘business as usual’ situations. Therefore, the number of creative elements is usually very small. This situation may be characteristic of most man-made networks, such as the internet, traffic networks, or power grids. However, creative elements are the ‘life insurance’ of complex systems, helping them to survive during any unexpected damage. Therefore, the number and importance of creative elements should increase if the complex organism experiences a fluctuating environment (Kovács et al, 2006; Palotai et al, 2008).

The adaptation of a large group of competing organisms to fluctuating environments is described by the process of evolution. The capacity of an organism to generate heritable phenotypic variation is called evolvability (Krischner and Gerhart, 1998). Evolvability is a selectable trait, which assumes that it is modulated by specific mechanisms (Earl and Deem, 2004).

As a summary of the above ideas, creative elements play a crucial role in the development, inheritance, and regulation of evolvability (Csermely, 2008). Creative elements are key players in the evolution of complex systems. They help to disassemble and reassemble modules, and most probably they had an important function in building the hierarchical, nested layers of the multiple networks on Earth.

In summary, nestedness is a central element in the way complexity evolved on Earth, and probably in the whole Universe. At the same time, nestedness helps us to explain the complexity of the world around us. If I open it, if I go close to it, all of a sudden a miracle happens: it is not a point, but a whole network! Nestedness serves our cognition requirements as well. Nestedness allows the simplification of a whole network to a point, when the complexity of the whole network would disturb our generalization. Human beings have to enjoy a certain level

of isolation to develop their consciousness, and at the same time, they have to be connected to preserve their safety and to enjoy the benefits of labor division. Nestedness is also important to formulate and preserve our own image, our self.

2.4 Weak-Linkedness

The task is already completed. We have come to the end of the list of reasons why people like networks. Small-worldness helps to preserve our safety in an alienated world. Scale-freeness helped our ancestors to survive, domesticates the unpredictable, and brings excitement and beauty into our everyday life. Nestedness helps us to explain the complexity of the world around us, and makes us feel at home on Earth. Is there anything left? Weak links, i.e., contacts between network elements which have a low probability, intensity or affinity, do not add very much (yet) to our excitement about networks. However, I hope to show in this book that weak-linkedness must also be added to the all-time network favorites, including small-worldness, scale-freeness and nestedness. This section presents a few introductory links between these four properties.

Weak Links and Small-Worldness

Long-range contacts, which are what make small worlds small, are usually formed by weak links. Weak links are necessary for the establishment of small worlds (Dodds et al., 2003a; Granovetter, 1973; 1983; Onnela et al., 2005; Skvoretz and Fararo, 1989).

Weak Links and Scale-Freeness

In real networks, elements are not identical. This leads to the emergence of strong and weak links. Indeed, modeling of Ethernet traffic (Leland et al., 1994), data transport (Goh et al., 2001; Ghim et al., 2004), the metabolism of the bacterium *Escherichia coli* (Almaas et al., 2004), air traffic, scientific collaboration (Barrat et al., 2004a), emails (Caldarelli et al., 2004), and market investments (Garlaschelli et al., 2005a) shows that widely different natural networks develop a scale-free distribution with regard not only to the degrees of various elements in space (self-similarity, fractals) and in time (event probabilities), but also in the distribution of the weights of the links between network elements. Recent modeling of network development has shown that

preferential attachment may explain the parallel emergence of scale-free properties of both the degree and link strength distribution of networks (Barrat et al., 2004b; Li and Chen, 2004; Yook et al., 2001). Scale-freeness thus involves the scale-freeness of link strength distribution. What does this mean? The message is rather simple: weak links always accompany strong ones. And in fact, in most networks, we have far more weak links than strong. Somehow networks cannot exist without weak links. What can be the reason of this? This book tries to find the answer.

Weak Links and Nestedness

Network stability may be a key element in the development of multilevel, nested networks. The formation of nested networks obviously requires at least a few contacts between the bottom networks. However, evolutionary selection requires the independence and at least temporary isolation of the bottom networks themselves. Weak links are probably the only tools for solving this apparent paradox. Now you see them, now you don't. Weak links can be broken and reformed again. This feature makes them what they are – weak. Indeed, modules are connected by weak links:

- in proteins (here weak links are provided by water molecules; Csermely, 2001a; Kovacs et al, 2005),
- in cells (here weak links are provided by low affinity protein bridges; Maslow and Sneppen, 2002; Rives and Galitski, 2003; Spirin and Mirny, 2003),
- in societies (here weak links are provided by superficial acquaintances; Degenne and Forse, 1999; Granovetter, 1973).



Date hubs are probably weak. A recent analysis by Han et al. (2004) showed the existence of two types of hubs in the yeast protein network. One of them had its partner proteins around all the time. This was called the party hub. The other hub, which was called the date hub, kept changing its partners. Date hubs were connecting functional and protein modules together and thus most probably contained weak links. Date hubs have also been identified in transcriptional networks (Luscombe et al., 2004). It will be an exciting task in the future to analyze the stabilizing role of date and weak hubs.



Weak links in cultural evolution. The role of weak-link-induced temporary isolation is not only observed in biological evolution. It

also plays an important role in cultural evolution. Cultural innovations often come from a module which is temporarily segregated from the giant component of the society.¹⁹ This relative isolation allows time for the novel idea to develop undisturbed and makes the competition and selection of these ideas by the majority of society possible. If such an idea breaks isolation by weak links, it has a good chance of conquering the society. Blues, pop, the Founding Fathers of the USA were all more or less isolated subcultures originally. ‘Incubators’ of larger firms also provide an excellent example of isolation-accelerated innovation (Sabel, 2002).

In summary, weak links seem to be necessary for the development of small-worldness, are a consequence of scale-freeness, and make a key contribution to the formation of nestedness. Weak links are also general and important elements of networks, forming most of their contacts. When we talk about the reason why we like networks, we have to talk about weak-linkedness.

¹⁹I am grateful to Viktor Gaál for these ideas.

Weak Links

The Universal Key to the Stability of Networks and
Complex Systems

Csermely, P.

2009, XIX, 392 p., Softcover

ISBN: 978-3-642-01192-4