

Chapter 2

Some Basic Concepts on Complex Networks and Games

Since this thesis is mainly devoted to the study of one particular game, the Prisoner's Dilemma, on complex networks (static ones in the first part of it, and two more sophisticated models that combine the growth with the play in the second), we consider that it is useful to state and explain first some notions on both networks and games. So, in this chapter, we want to provide just a few very basic concepts and definitions on Complex Networks and Game Theory that we will use later on during the full elaboration of this thesis. We hope they will help setting the foundations to understand our work perfectly, so the reader will not need any external help to comprehend, and also it will serve as an introduction to the two fundamental components on which this thesis is based.

2.1 Complex Networks

The study of complex networks is a relatively recent field, and it has been inspired by the observation of many real systems, such as biological, social or technological ones. In the first part of this chapter we want to give a few examples of real networks, just to motivate the study of such structures, by establishing its ubiquity in natural and artificial systems. Then, we will give some of the basic definitions needed in order to properly describe networks [1], such as the degree of a node, the degree distribution of a network, the clustering coefficient or the average path length. Then, we will explain some useful models for building different kinds of graphs, such as the Erdős and Rényi (ER), the Barabási-Albert (BA) or the Small-World by Watts and Strogatz model. Finally, we will mention some of the many possible processes that can take place on top of complex networks.

2.1.1 Examples of Real Networks

As it has been pointed out along the Introduction of this thesis, many real systems [1, 2] can be described as complex networks, and this relatively new approach can provide new insights to better understanding, and tools to deal with unsolved problems. In very different fields, such as biology, immunology, sociology, technology or economics, there are plenty of examples of networks. In every particular field, both the nodes and the links of the networks will represent completely different things, but the fact that this kind of structures are so ubiquitous in Nature, is surprising and very promising.

One can consider technological structures, such as the air transportation networks for a particular region or for the whole planet, where the nodes are airports and the links represent direct flights between them, the road networks connecting cities or the power grids that supply electricity to a country, with its power stations represented by nodes and the links standing for the wires. There is also the WWW, where nodes are web pages connected by hyperlinks, or the Internet (see Fig. 2.1 (Left)), made up of billions of hosts, physically connected among them. Since modern societies depend strongly on these infrastructures, it is obviously very important to have detailed information about them, in order to be able to predict its behavior or act correctly during a crisis.

In biology, there are several examples as well, like food webs on an ecosystem (see Fig. 2.1 (Right)), or on a more basic level, the metabolic networks of different processes. On the other hand, maybe some of the more tangled complex networks one can consider (from the point of view of both number of interconnections and variability over time) are those that describe social relationships, where nodes are people, and links represent some kind of interaction: from groups of mere friends, people with similar interests or collaborators in some particular field [4, 5] (scientific collaborations or citations, or networks of musicians that play together regularly,...), to sexual contact networks or new global phenomena like Facebook, MySpace or Twitter. It is probably because of the complex nature of the human being itself, that such social structures are often so entangled and fascinating.

On the other hand, we want to point out that, when dealing with real networks one has to take into account that the available data can (and probably will) have mistakes: there can be missing or spurious nodes or links. Some effort has been put to try to obtain the ‘real network’ and its topological properties out of the observational data (see for example [6]).

Finally, the kind of processes that will take place on top of them can be very diverse (synchronization, traffic of information or of something else, disease or rumor spreading, games, learning processes...), but it is very useful to be able to characterize them structurally as precisely as possible first, trying to find out what are the main and more relevant features all of them share, if any. Moreover, as we will see later on, the structure will be a key factor in the outcome of any dynamical process that will take place on top of such structured systems. Thus, we will address next the topological characterization of complex networks.

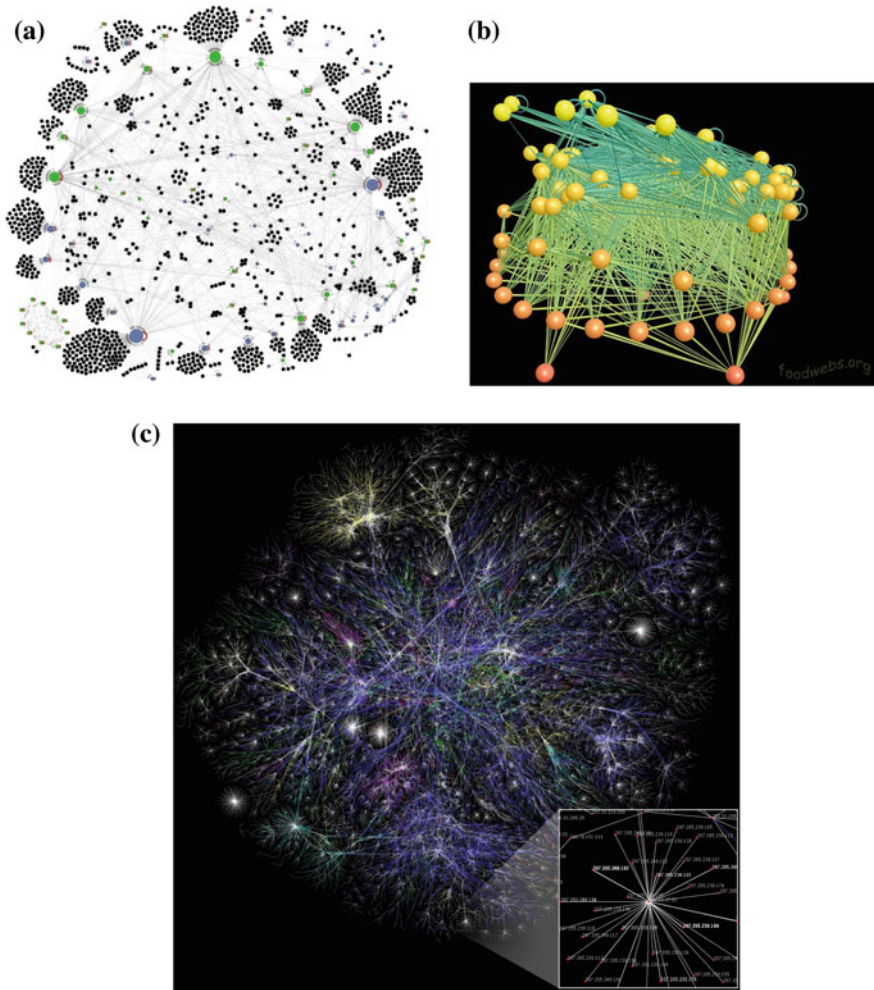


Fig. 2.1 **a** Gene regulation network for the Mycobacterium Tuberculosis. Every node represent a gene, and the links stand for the regulation relationship between a transcription factor and the correspondent regulated gene. Different colors mean different character of the genes, as far as regulation dynamics is concern [2]. **b** Food web of the Caribbean coral reef located in the Puerto Rico Virgin Islands. Node color represents trophic level: *red nodes* represent basal species, such as plants and detritus, *orange nodes* represent intermediate species, and *yellow nodes* represent top species or primary predators. Links characterize the interaction between two nodes, and the width of the link attenuates down the trophic cascade, so a link is thicker at the predator end and thinner at the prey end (Original image from [3], and generated by FoodWeb3D). **c** Visualization of a portion of the Internet, using over $5 \cdot 10^6$ edges. The colors represent different geographical regions. In the inset it is shown a particular node and its neighborhood. (Original image from ‘The Opte Project’: <http://www.opte.org>)

2.1.2 Definitions

A network is a set of items (called *nodes*, points or vertexes), with some connections between them (*links*, lines or edges). A complex network is a network with non-trivial topological features, i.e. its structure is irregular and complex as opposed to lattices, for example, that present total spatial regularity, or they can even evolve, adding and/or losing nodes and/or links over time.

Mathematically, we can represent a network using graph theory. A graph $\mathcal{G} = (\mathcal{N}, \mathcal{L})$, consists of two sets, $\mathcal{N}$ and $\mathcal{L}$, where $\mathcal{N} = \{n_1, n_2, \dots, n_N\}$ are the nodes, and $\mathcal{L} = \{l_1, l_2, \dots, l_K\}$ are the links. Obviously, N is the total number of nodes of the network, and K is the total number of links, which has to be a non-negative number, whose maximum is $N(N-1)/2$ (when the graph is *complete*, i.e. every node is connected to everyone else). A specific node of the network is denoted by a label i in the set $\mathcal{N}$. On the other hand, every link connects a pair of elements of $\mathcal{N}$, i and j , and is denoted by l_{ij} . Thereby, the pair of nodes i and j are called adjacent or neighbors. The usual way of representing a network graphically is by drawing a dot for every node and a line for every link that connects a pair of nodes. In addition to this, we can also define a subgraph $\mathcal{G}' = (\mathcal{N}', \mathcal{L}')$, of the graph $\mathcal{G} = (\mathcal{N}, \mathcal{L})$, if $\mathcal{N}' \subseteq \mathcal{N}$ and $\mathcal{L}' \subseteq \mathcal{L}$. A special case would be the subgraph of all the neighbors of a given node i and its corresponding links, denoted by $\mathcal{G}_i$. On the other hand, a graph is said to be *connected* if, for every pair of nodes i and j , there is a path to go from one to the other. If there is not such a path for at least one pair of nodes, then the graph will be *disconnected* or unconnected, and it will have therefore, two or more disconnected subgraphs.

Besides, another very useful way of representing a network is by using matrix representation. Given a graph $\mathcal{G} = (\mathcal{N}, \mathcal{L})$, the adjacency matrix $\mathcal{A}_{ij}$ is a $N \times N$ square matrix, whose entry a_{ij} ($i, j = 1, 2, \dots, N$) is equal to 1 when the link l_{ij} exists, and zero otherwise. Nonetheless, for implementation or practical purposes, we can use the connectivity matrix $\mathcal{C}_{ij}$ of the graph, that is a $N \times k_{\max}$ matrix, where $k_{\max}$ is the maximum connectivity of the nodes of the graph, and where the row i of it contains all the neighbors of the node i (ordered usually, but not necessarily, from the first to the last to connect with it when constructing the network). And we can also define a matrix of the pairs of neighbors, $\mathcal{D}_{ij}$, which is a $L \times 2$ matrix, whose entries d_{l1} and d_{l2} are the pairs of nodes that are neighbors, with $l = 1, 2, \dots, L$, and being L the total number of links in the network. The definition of these two matrices is not for rigorous mathematical purposes, but nonetheless, they will be very useful in order to implement them on programs and numerical simulations.

Degree of a Node and Degree Distribution of a Network

The *degree* or *connectivity* of a node is the number of neighbors it has. Using the adjacency matrix, we can formally define the degree of a node as:

$$k_i = \sum_{j \in \mathcal{N}} a_{ij} \quad (2.1)$$

If the graph is *directed*, then k_i will have two components: the ingoing links $k_i^{in} = \sum_j a_{ij}$ and the outgoing links $k_i^{out} = \sum_j a_{ji}$, so the total degree will be $k_i = k_i^{in} + k_i^{out}$.

On the other hand, the most basic topological characterization of the network as a whole is the *degree distribution*. We can define the degree distribution of the graph, $P(k)$, as the fraction of nodes in the network that have connectivity k , or equivalently, the probability that a node randomly chosen from the network has k neighbors. For example, random graphs (also known as ‘one-peaked’ or ‘single-scaled’) have a Poissonian degree distribution, while the $P(k)$ for a so-called scale-free network is a power law. For directed graphs, we will have two different distributions, $P(k^{in})$ and $P(k^{out})$.

The mean degree of a graph, $\langle k \rangle$ is the first moment of the degree distribution:

$$\langle k \rangle = \sum_k k P(k) \quad (2.2)$$

Furthermore, the second moment of the distribution, $\langle k^2 \rangle$ is the measure of the fluctuations of the degree distribution. As we will see later on, $\langle k^2 \rangle$ diverges in the limit of infinite graph size for scale-free graphs for certain values of the exponent of the power-law distribution, which is a very interesting property, that affects greatly the outcome of the dynamics that can take place on top of such topologies. For an *uncorrelated graph*, i.e. if the degree of every node is completely independent of its neighbors’, then the degree distribution $P(k)$ is enough to describe the statistical properties of the network. But if the network is *correlated*, as it usually happens in many real systems, then the probability that a node of degree k has a neighbor with connectivity k' , depends on k . In that case, we can define the conditional probability $P(k'|k)$, that a node with connectivity k has a neighbor with connectivity k' . We can also calculate the average degree of the nearest neighbor of nodes with degree k , given by:

$$k_{nn}(k) = \sum_{k'} k' P(k'|k) \quad (2.3)$$

So when the network is uncorrelated, obviously, we have that $k_{nn}(k)$ is independent of k , and equal to $k_{nn}(k) = \langle k^2 \rangle / \langle k \rangle$, but when it is correlated, then we can have *assortative* networks, if $k_{nn}(k)$ is an increasing function of k , or *disassortative* ones, when $k_{nn}(k)$ is a decreasing function of k . The first case implies that nodes tend to be linked with others with similar connectivity, whereas in the second one, the highly connected ones are mostly linked to the poorly connected ones.

Weighted and Directed Networks

Depending on the kind of interaction a link describes within the network, it can be weighted or non-weighted, directed or non-directed, and so will be the network.

If all the interactions in the network are alike, or in other words, when a link only establishes the presence of an interaction between two nodes, then the network is *non-weighted*. Otherwise, if there are different types of interactions, for example, some more important, or more frequent than others, then the links are *weighted*, and so is the graph. In this case, in addition to give the set of nodes and links of the network, we need to specify also the weight of every link in order to properly define a graph. So now we have: $\mathcal{G} = (\mathcal{N}, \mathcal{L}, \mathcal{W})$, where $\mathcal{W} = \{w_1, w_2, \dots, w_K\}$ is the set of weights, that are real numbers attached to the corresponding links. Usually, they will be positive numbers, so the higher the value, the stronger the link between the pair of nodes, but also negative links have been used, describing some kind of repulsive interaction, for example [7]. On the other hand, if a link l_{ij} represents that i interacts with j and vice versa, then it is called *undirected*, but if in a system i can interact with j without j interacting necessarily with i , then in order to describe it correctly, we need *directed links*. In this case, the adjacency matrix will not be symmetric, in general.

Average Path Length, Betweenness and Clustering Coefficient

Given a particular network, it would be interesting to know the minimum distance between every pair of nodes, i.e. the shortest path lengths or geodesics. The knowledge of this information concerning a network can be useful for some processes that could take place on top on it (such as information traffic on the Internet, or rumor spreading on a social club), in order to work the best they can. Thus, we can define a square matrix $\mathcal{D}$, of size $N \times N$, whose entry d_{ij} is the minimum distance between the nodes i and j . On the one hand, the maximum of these d_{ij} is called the *diameter* of the graph, but a more useful magnitude to characterize the network, is the *average path length*, defined as the mean value of the geodesics between every pair of nodes in the network:

$$L = \frac{1}{N(N-1)} \sum_{i,j \in \mathcal{N}, i \neq j} d_{ij} \quad (2.4)$$

One can also ask how important or ‘central’ a particular node is within a graph, meaning how many shortest paths, or geodesics go through it. Thus, we can give a measure of the centrality of a node, by defining its *betweenness*:

$$b_i = \sum_{j,k \in \mathcal{N}, j \neq k} \frac{n_{jk}(i)}{n_{ij}}, \quad (2.5)$$

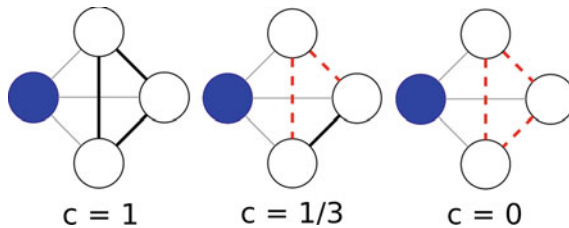


Fig. 2.2 Examples of the local clustering coefficient (for the *dark node*) for different connecting situations. It is computed as the proportion of connections among its neighbors which are actually realized (*thick black lines*) and the number of all possible connections, which in this particular example, is three. For every situation, the missing links are represented with dashed lines

where n_{jk} is the total number of geodesics connecting the nodes j and k , and $n_{jk}(i)$ is the number of geodesics connecting the nodes j and k that go through the node i .

The betweenness is a useful magnitude when constructing community detection algorithms [8, 9].

Clustering, or transitivity of a node, is a measure of how many triangles are there on the graph, or in other words, how likely is that, if a node i has two neighbors, say j and k , then the nodes j and k are also linked to each other. First, given a node i and the subgraph of its k_i neighbors, G_i , we can define the local clustering coefficient of node i as the ratio between the actual number of edges in the subgraph, e_i , and the maximum possible number of them in G_i :

$$c_i = \frac{2e_i}{k_i(k_i - 1)} = \frac{\sum_{j,m} a_{ij}a_{jm}a_{mi}}{k_i(k_i - 1)} \quad (2.6)$$

where a_{ij} are the entries of the adjacency matrix, defined at the beginning of this section. On Fig. 2.2 we show a diagram of how to calculate it for three very simple cases.

And then, we can define the clustering coefficient of the whole network, as the average of c_i over all the nodes in it:

$$C = \frac{1}{N} \sum_{j \in \mathcal{N}} c_i \quad (2.7)$$

Notice that, by definition, both the local and the global clustering coefficient satisfy: $0 \leq c_i \leq 1$ and $0 \leq C \leq 1$. As we will see, SF networks have low values for the average path length, but relatively high values for the clustering coefficient, while random topologies have low values for both magnitudes.

Finally, it is worth mentioning that a power-law dependence of the clustering coefficient with the degree of the node ($C \sim k^{-1}$) is typical of a hierarchical organization on the network, which implies that sparsely connected nodes are part of highly clustered

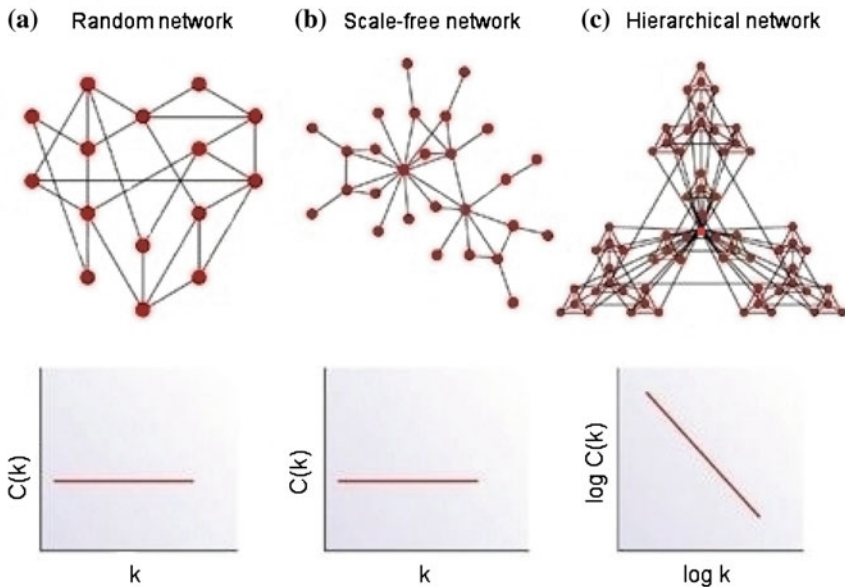


Fig. 2.3 Diagram with some examples of networks, specifically random (a), scale-free (b) and hierarchical ones (c), and its corresponding plots of the clustering coefficient versus the degree of the nodes. This dependence is a power-law for the hierarchical structures, while for the other two types, it is clearly independent. Original figure from [10]

areas, with communication between these different highly clustered neighborhoods being maintained by a few hubs (see Fig. 2.3).

Motifs and Communities on Networks

A *motif* is a n -noded pattern of connections (a subgraph) in a network that appears at a much higher rate than expected in a randomized version of the same network (see Sect. 5.1 for a detailed explanation of the randomizing procedure). Some real networks, such as the metabolic ones, display characteristic motifs, that seem to be specific of each kind of network. On Fig. 2.4 we show as an example, all the possible motifs for a 3-noded directed subgraph. Note that the number of n -noded motifs increases rapidly with n .

On the other hand, we can define a *community* within a network $\mathcal{G} = (\mathcal{N}, \mathcal{L})$, as a subgraph $\mathcal{G}' = (\mathcal{N}', \mathcal{L}')$ or a set of nodes, that are much more connected among themselves than with the rest of the network. Using just the sense that the intra-community connections are denser than the inter-community ones is of course a qualitative way of describing it. Nonetheless, to be able to detect such structures efficiently, a magnitude has been introduced to determine whether or not a partition of a network into communities is accurate enough: the *modularity*.

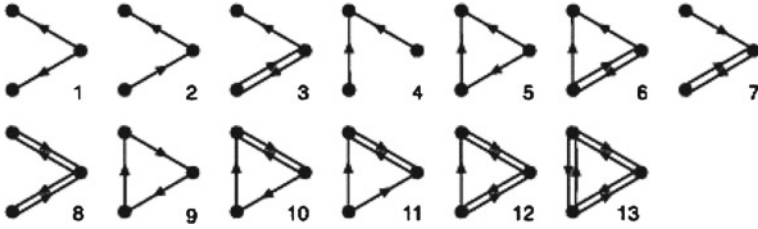


Fig. 2.4 All the possible 3-noded motifs on a directed network

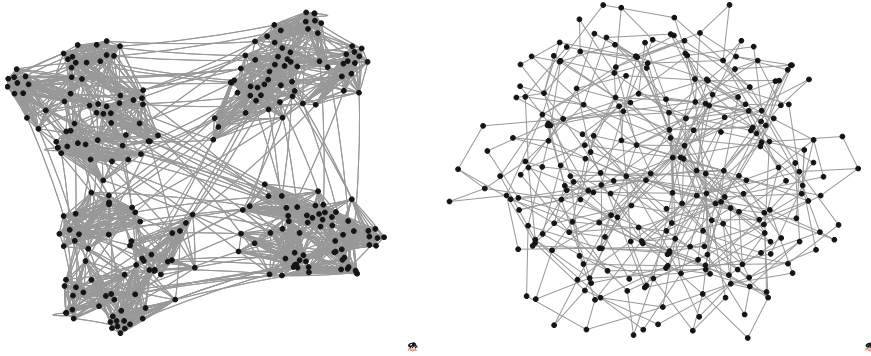


Fig. 2.5 Some examples of a network with (*left*) and without (*right*) community structure, both with $N = 256$ nodes. Original data of the community network created by Dr. L. Izquierdo (<http://luis.izquierdo.org/communities/redes.zip>)

Given an arbitrary network, and an arbitrary partition of it into N_c ‘communities’ (and this time, by this term we mean artificial communities, just a way to part the graph), we can build a $N_c \times N_c$ matrix whose entries e_{ij} are the ratio between the number of links starting at a node in community i and ending at a node in community j , and the total number of links present on the network (so the sum of any row or column, $a_i = \sum_j e_{ij}$, is the fraction of links connected to the community i).

In the case of a random partition of the network i.e., if it does not correspond to the actual community structure, or also if the network itself does not have a community structure (see Fig. 2.5 for some examples of networks with and without community structure), then the fraction of links within communities can be estimated as the probability that a link begins at a node in partition i , a_i , multiplied by the fraction of links that end at a node in partition i , also a_i , so the expected number of intra-community links is just $a_i a_i$. We also know the actual fraction of links exclusively within a partition, e_{ii} , so now we can compare the two values, and thus, we can define the modularity for a specific partition of our network as [8]:

$$Q = \sum_i^{N_c} (e_{ii} - a_i^2) \quad (2.8)$$

Obviously, the closer to 1 the value of the modularity is, the more accurate the partition we have made of the network into communities. It is worth noticing that it is possible to find partitions of random networks that display relatively high values of modularity (up to $Q \sim 0.2$). The reason for this is that random graphs might have some community structure, just due to fluctuations. Moreover, it is important to stress that the presence of communities on a network can not be detected just via its degree distribution, so we can have two graphs with the same $P(k)$, one of them with community structure, and the other one without it.

One can easily realize that the space of possible partitions of a given network into communities is huge, so in order to effectively explore the landscape of values of Q , and find an accurately enough partition, we will need the help of some optimization techniques. For some very nice works on different community detection algorithms, see [8, 9, 11, 12] and references therein.

Finally, we want to mention that it is also possible to consider complex topologies with hierarchical structure, it is to say, networks that have communities within the communities. In this situation, we deal with several levels of description for the structure of the system (multiscale representation) [?]. Also, one can have a system with communities, where there is some degree of overlapping among them, and this fact will make it harder to detect accurately [13].

2.1.3 Some Network Models

In this section we want to present just a few models for growing networks. Specifically, we will address the models to build two of the most used kinds of networks: the ER and the BA model for random and scale-free networks respectively, since we will use them often, later on in this thesis, and also the well-known Small-World model by Watts and Strogatz. On the other hand, we will explain the Gardeñes-Moreno (GM) model, which interpolates between the ER and the BA model, because we will use it also in some chapters to come.

The ER Model

Erdős and Rényi proposed a model (ER) [14] to generate random graphs with N nodes and K links, where the term *random* refers to the disordered nature of the arrangement of links between different nodes. There are two possible ways of constructing such networks: in the first one, we start with N disconnected nodes and choose K pairs randomly, to link them with a probability $0 < p < 1$, avoiding multiple connections between two nodes, and also self-links. The alternative procedure is to start with N disconnected nodes, and link every possible couple with probability $0 < p < 1$. While the first option gets different networks with exactly K links and an average degree of $\langle k \rangle = 2K/N$, the second, gets networks with different number of connections, an average degree $\langle k \rangle = p(N - 1)$, and the probability of having exactly

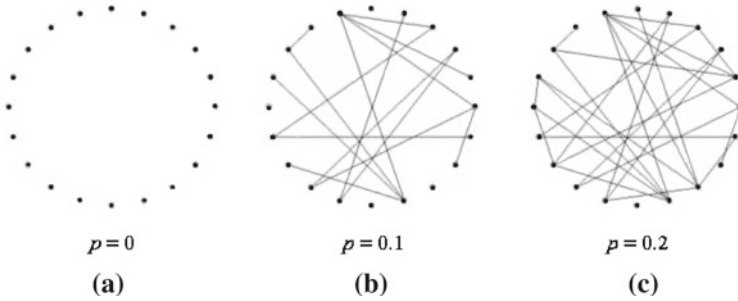


Fig. 2.6 Diagram of the ER model for random networks with $N = 20$ nodes

K links in a particular realization of the network is $p^K (1-p)^{N(N-1)/(2-K)}$. Nonetheless, both models coincide in the limit of large N , or thermodynamic limit. The probability of finding a node with a large connectivity decreases exponentially with K , so vertices with large connectivity, $K \gg 1$, are practically absent (Fig. 2.6).

If one starts increasing the value of the probability of connection, from $p = 0$ (nodes totally disconnected) to $p = 1$ (complete graph), there is an interesting change of behavior at the critical value $p_c = 1/N$. Thus, if $p < p_c$, the graph is not connected (it has no component of size greater than $\mathcal{O}(\ln N)$), if $p > p_c$, then the graph has a component of $\mathcal{O}(N)$, and the transition at p_c displays the typical features of a second phase transition. On the other hand, the probability of having a node with $k = k_i$ connections follows the Binomial distribution:

$$P(k = k_i) = C_{N-1}^k p^k (1-p)^{N-1-k} \quad (2.9)$$

where p^k is the probability of having k edges, $(1-p)^{N-1-k}$ is the probability of the absence of the remaining $(N-k)$ links, and C_{N-1}^k is the number of different ways of selecting the end points of these k nodes. Notice that, since all nodes of the networks are equivalent, this probability $P(k = k_i)$ is also the probability of choosing randomly a node with k_i neighbors. In the limit of large N and fixed $\langle k \rangle$, the degree distribution of the network can be accurately described by the Poisson distribution:

$$P(k) = e^{-\langle k \rangle} \frac{\langle k \rangle^k}{k!} \quad (2.10)$$

Moreover, for this particular topology, the dependence of the clustering coefficient with the size of the system N is given by:

$$\langle C \rangle_{ER} = p = \langle k \rangle / N \quad (2.11)$$

and the average path length, on the other hand shows a dependence given by:

$$\langle L \rangle_{ER} \sim \frac{\ln N}{\ln \langle k \rangle} \quad (2.12)$$

Notice that the value of the clustering coefficient tends to zero in the limit of large N . It is also important to point out that this model produces homogeneous random graphs, which do not share certain topological features with the real networks, for example, they have low values of the clustering coefficient, and do not show any kind of correlations between nodes.

Small-World Networks

A graph in which, although most pairs of nodes are not directly connected to each other, they can nonetheless be in touch by a small number of steps is called Small-world network, since it captures this so-called phenomenon of strangers being linked by a mutual acquaintance (also known as *six degrees of separation* [16–18]). Some properties of real networks can be well modeled using Small-world networks, for example social networks, gene networks or the Internet. Nonetheless, it is important to keep in mind that ‘small-world’ is a concept that includes several kind of topologies: empirical data [19] suggest the existence of three classes of small-world networks, as far as its degree distribution is concern: scale-free networks, broad-scale or truncated scale-free networks, and single-scale or random networks.

The first Small-world network model was proposed by Watts and Strogatz [15], and it interpolates between a regular graph and a random graph, depending on a single parameter $p \in [0, 1]$, without altering neither the number of nodes nor the number of connections per node of the original graph. This is a random graph generation model that produces networks with Small-world properties, possessing short average path length and high clustering coefficient provided the adequate range of the parameter p (see Fig. 2.8).

Departing from a one-dimensional regular lattice or a ring, where each node has exactly the same number of neighbors, z , we rewire every link with a probability p , avoiding multiple connexions between two nodes and self-connections. In another version of the model, we depart from a ring, where each node has exactly z neighbors, and we add a link between every pair of nodes, with probability p , instead of rewiring the existing links. Regarding the degree distribution, for $p = 0$ we have $P(k) = \delta(k - z)$, where z is the coordination number of the lattice ($z = 4$ in the case shown in Fig. 2.7); whereas for finite values of $p \in (0, 1]$, $P(k)$ still has a peak around z , but it obviously gets broader as p increases. For the cases where $p \in (0, 1]$, the probability of finding a node with a large connectivity decreases exponentially with k , as it happen for ER random networks, so vertexes with large connectivity are practically absent as well. For $p = 0$ we keep the initial ring structure, which has high values both for the clustering coefficient ($C \sim 3/4$), but also for the average path length ($L \sim N/(2k) \gg 1$).

On the other hand, for $p = 1$ we have a random network -though, to be rigorous, in the second version, there are not any nodes with connectivity $k < z/2$, as there would

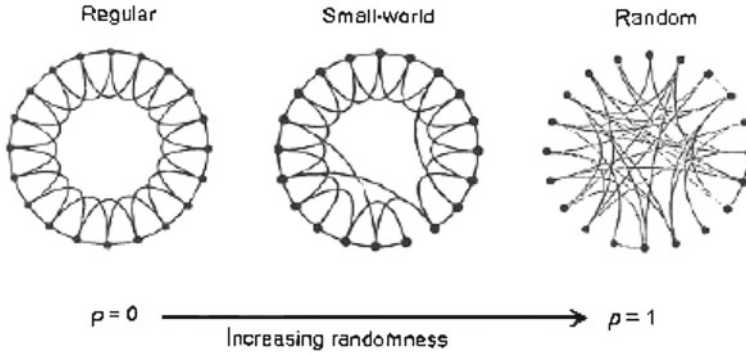


Fig. 2.7 Diagram of the random rewiring procedure for interpolating between a one-dimensional lattice and a random network in the Small-world model. The networks have $N = 20$ nodes and $k = 4$. Original figure from [15]

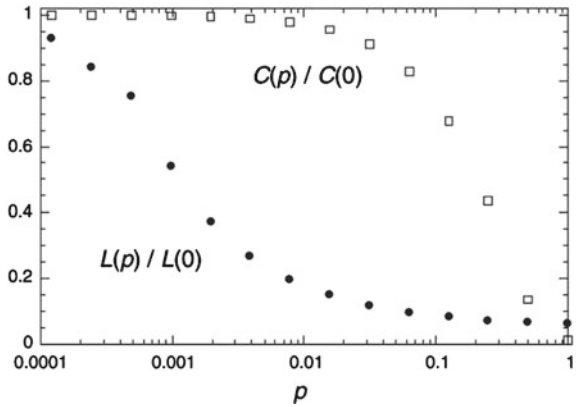
be in a random network built with a mechanism such as ER. Its average path length is short ($L \approx L_{\text{random}} \sim \frac{\ln N}{\ln k}$), but its value for the clustering coefficient is also low ($C \approx C_{\text{random}} \sim k/n \ll 1$). Nonetheless, there is an intermediate region of p where we can get a network with both features: a high value for the clustering coefficient and a short average path length. This is due to the presence of long-range connections or shortcuts introduced by the rewiring procedure. Notice that the introduction of these shortcuts makes the average path length drop, not only for the pair of nodes involved, but for all their neighbors too. Moreover, the removal of some links from a neighborhood due to the rewiring process, does not affect the clustering coefficient too drastically, so it remains unaltered for small values of $p \lesssim 0.01$ (see Fig. 2.8). In other words, during the dropping of $L(p)/L(0)$, the clustering $C(p)/C(0)$ remains almost unaltered, which means that this transition to the Small-world is undetectable on a local level.

Regarding the dependence of the small-world behavior with the size of the system, it has been shown [20] that the emergence of this regime occurs for a value of p that approaches zero as N diverges.

The BA Model

Both the Small-world model and the ER model, explained previously, although are most undoubtedly useful and insightful, display two important features that make them very different from the real networks. The first one is the assumption that the whole system is present from the very beginning, it is to say, that the network has a fixed size N and it does not grow, no new nodes are added. In contrast, it has been observed that most real networks are open systems, and they get new vertexes that connect with the ones already present, so the number N keeps increasing throughout the lifetime of the graph. The second one is the assumption that the probability

Fig. 2.8 Average path length and Clustering coefficient for the Small-world model, as a function of the probability of rewiring p , normalized by their respective values for the ring, i.e. when $p = 0$. Notice that the x -axis is shown in logarithmic scale. The graphs have $N = 10^3$ nodes and $\langle k \rangle$. The data shown is the average over 20 different rewiring procedures. Original figure from [15]



that two vertexes are connected is uniform. Again, in contrast, most real networks show clearly a preferential attachment: usually, the more connected a node is, the more easily it will get even more neighbors due to connections from new nodes.

The Barabási-Albert (BA) [20] is a model for building scale-free networks that is based on two fundamental ingredients: *preferential attachment*, i.e. the assumption that the likelihood of receiving new edges increases with the node's degree, and *growth*. Actually, variants of the model, with just one of the two ingredients have been tried, but neither of them get networks with power-law distributions. This was a model originally inspired on the growth of the World Wide Web, and as we have already mentioned, the idea behind it is that the highly connected nodes get new links at a higher rate than the lower connected ones or, in other words, the catchphrase 'rich get richer' [21]. This is a phenomenon easily found on real systems (and it is known in sociology as the Matthew effect [22]).

Thus, we start with a little core of m_0 disconnected nodes, and at each time step $t = 1, 2, 3, \dots, N - m_0$, a new node i is added to the system with $m \leq m_0$ links to existing nodes. The probability that an existing node j gets one of the links from the newcomer is proportional to its own connectivity, k_j , in a linear way:

$$\Pi_j = \frac{k_j}{\sum_l k_l} \quad (2.13)$$

Since every new node links to m other nodes, at any given moment t , the network has $N(t) = m_0 + t$ nodes and $K(t) = mt$ links. Besides, for long times, the average degree of the network is $\langle k \rangle = 2m$. The degree distribution of these networks is a power law, $P(k) \sim k^{-\gamma}$, with $\gamma = 3$. A scale-free degree distribution implies that there are a lot of nodes with just a few connections, and a small number of nodes with a very high connectivity. These highly connected nodes are called *hubs* and they usually play an important role in most dynamical processes that can take place on the system, as we will see with some detail during this thesis. Besides, the degree

distribution $P(k)$ of the BA networks is independent of time, and thus independent of the size of the system, indicating that despite its continuous growth, the system organizes itself into a scale-free stationary state.

The dependence of the clustering coefficient with the size of the system N is approximately a power law, given by:

$$\langle C \rangle_{BA} \sim N^{-0.75} \quad (2.14)$$

The average path length, on the other hand shows a dependence given by:

$$\langle L \rangle_{BA} \sim \frac{\ln N}{\ln(\ln N)}. \quad (2.15)$$

The value of the average path length in BA networks is smaller than in ER networks for any value of N , so obviously, the heterogeneous topologies help bringing the nodes together more than the homogeneous ones. On the other, hand, comparing the values for the clustering coefficient, the corresponding values for the BA networks are about five times higher than for ER networks, and this factor even increases slightly with the size of the system. Moreover, it is worth pointing out the existence of the so-called *age correlations* [23–25] among nodes for the scale-free topologies, which means that the older nodes, i.e. the ones that appear first on the system, are more likely to end up being hubs, just by construction, while the later a node appears, the lower connectivity it will get.

We consider that it is important to stress again that SF networks built via this BA procedure have very low values for the clustering coefficient, when comparing with real networks, so we must admit that this kind of topologies might reproduce the degree distribution of those systems, but can not do the same for the clustering coefficient. Along these lines, there have been some other models that, based on BA, tried to put a remedy to this fact. For example, the work by Holme and Kim [26], presents a model for constructing SF networks with tunable clustering coefficient. In few words, this model starts with a set of m_0 unconnected nodes and adds a new one to it every time step, up to N . Each one of the new nodes launches $m \leq m_0$ links. The probability of an existing node i to receive the first link of a newcomer j is proportional to its connectivity k_i , but for the remaining $m - 1$ links that the new node j has to establish, there is a probability p to launch them to a (randomly selected) neighbor of i , and a probability $(1 - p)$ to launch them following the original preferential attachment rule. In this way, the family of networks we obtain have all exactly the same power-law degree distribution $P(k) \sim k^{-3}$, but the higher the value of the probability p , the higher the value of the clustering coefficient (it can easily achieve values of 0.5, when we recall that for BA, it tends to zero as N increases, so the order of magnitude of a typical value can be around 10^{-2} for $N = 10^3$). For the particular case $p = 0$, we recover the original BA model, obviously. Moreover, with this Holme-Kim model, the clustering coefficient is independent of the size of the system, as opposed to what happens with BA, where it decreases with N , as we have seen. On the other hand, it is also worth mentioning that, one may think

that by increasing p the average path length of the final structure will decrease, since some links that would help shortening it by linking to nodes far apart, are now linking nodes in the same neighborhood. As it turns out, the value of the average path increases slightly with the probability p , but the dependence with the size of the system remains logarithmic, so we do not lose the 'small-world' property with this model.

Finally, we also want to remark two points regarding preferential attachment. First, other mechanisms for building SF networks have been proposed [27], that are not based on growth and preferential attachment like the BA model is. Instead, an intrinsic fitness (from a given probability distribution) is assigned to each node in the system, and then pairs of them are linked together, according to a function of their fitness. And second, if one combines growth, preferential attachment and some aging mechanism or introduces a cost per link, then one will obtain SF topologies with a cutoff on the degree distribution, or even make the scale-free regime disappears altogether [19].

The GM Model

The Gardeñes-Moreno is a model [28] that interpolates between Erdős-Rényi random networks and Barabási-Albert scale-free networks as far as the degree distribution is concerned, through a tunable parameter α , so it generates a one-parameter family of networks. This parameter $\alpha \in [0, 1]$ determines the degree of heterogeneity of the network, whose final size will be Ω . Thus, $\alpha = 0$ gives rise to scale-free networks and $\alpha = 1$ to random graphs, and for in-between values, the topology will have an intermediate degree of heterogeneity.

The procedure to generate these networks is as follows: we start with a small fully connected core of m_0 nodes, and a set $\mathcal{U}(0)$ of $(\Omega - m_0)$ disconnected nodes. At each time step, a new node j from the set $\mathcal{U}(0)$ is chosen, and it makes a link in two possible ways: with a probability α , it attaches to any other node i from the whole set of $\Omega - 1$ nodes with uniform probability:

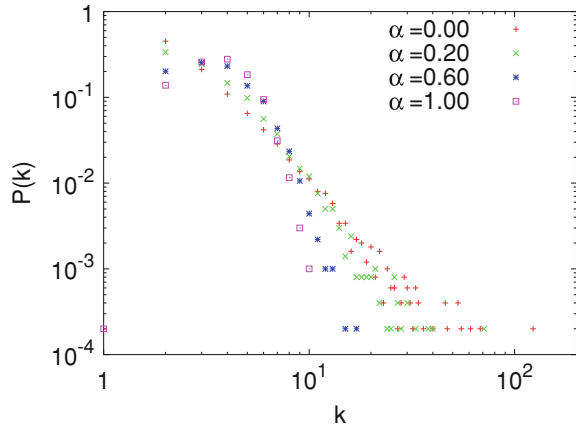
$$\Pi_i^{uniform} = \frac{1}{\Omega - 1} \quad (2.16)$$

and with probability $1 - \alpha$, it establishes a link following a preferential attachment (PA) strategy. This means that the probability for any other node i to get attached to node j is a function of its connectivity, in a way given by:

$$\Pi_i^{PA} = \frac{\hat{k}_i^{pa} + A_i}{\sum_{l \in \Omega} (\hat{k}_l^{pa} + A_l)} \quad (2.17)$$

where $\hat{k}_i^{pa}$ is the incoming PA degree of the node i , that is, those links received by i when other node launches (in average) $(1 - \alpha)m$ links following the PA rule. On

Fig. 2.9 Degree distributions for several networks obtained for the shown values of the parameter α with the GM model, interpolating between random ($\alpha = 1.0$) and scale-free ($\alpha = 0.0$) graphs. The size of the system is $\Omega = 5 \cdot 10^3$ and $\langle k \rangle = 2m = 4$. Every point is the average over 10^3 different realizations



the other hand, A_i is an initial attractiveness (or fitness) the new node has when it is introduced in the connected component (either because it is chosen at random by any node or because it is launching its m outgoing links over the rest of nodes). This associated parameter is zero if the node i is not in the connected set and is $A_i = A$ if it is linked to other nodes, i.e., if it belongs to $N(t)$. Thus, the preferential attachment is strongly correlated with the simultaneous uniform random linking, and, on the other hand, it is linear with the incoming PA degree of the node $\hat{k}_i^{pa}$. Next, we repeat the linking procedure for another $m - 1$ times for the same node j , and then we repeat the whole process altogether for the rest of the nodes, i.e., for another $\mathcal{U} = \Omega - m_0$ more time steps.

On Fig. 2.9 we show the degree distribution for some networks obtained with the GM model, for several values of the parameter α but the same size Ω and average connectivity k . Notice that the transition between heterogeneous and homogeneous topologies is smooth, as α increases.

2.1.4 Processes on Networks

So far in this chapter, we have studied some general topological properties of networks, as well as some well known and widely used models to generate them, and some real examples too. Nonetheless, we have to keep in mind that the ultimate goal of studying these structures, is to finally be able to model, describe and predict the different dynamics that can take place on top of them. Those include a wide and varied collection, such as disease [1, 2, 29–35] or rumor spreading, synchronization [1, 36–40], diffusion, traffic information and congestion, network search and navigation, percolation, robustness against random failures or targeted attacks [41, 42], cultural dissemination, opinion formation or language dynamics [43], and games [44]. In this section, it is not our intention to go exhaustively through all of them at

Fig. 2.10 Schematic representation of the SIR model



all (for some very nice reviews on the subject, see [1, 2, 40, 44]), but just to briefly examine a few of them, as an example, describing some the most popular models or approaches that have been proposed, and also pointing out the differences introduced by the underlying topology on the outcome of the dynamics, in comparison to well-mixed situations or lattices.

Disease Spreading

Epidemic spreading is a very interesting and obviously very important object of study [1, 2, 29–35]. The aim in this field is not only to understand the mechanisms through which diseases spread on a population, but also to design strategies to control them, and to be able to protect the population from pandemics.

Specifically, *Compartmental Models* in epidemiology stand for some models that, in order to describe the progress of an epidemic in a large population comprising many different individuals, reduce such population diversity to a few key characteristics which are relevant to the infection under consideration. For example, for most common childhood diseases that confer long-lasting immunity, such as the chickenpox, it makes sense to divide the population into those who are susceptible to the disease, those who are infected and those who have recovered and are therefore immune. Thus, one can ignore the rest of the information about the population, such as age distribution or race, because it is irrelevant for the model. These subdivisions of the population are called *compartments*.

In particular, one of the more used (and at the same time simple) models to study disease spreading is the SIR model. It considers that the population is compartmentalized into three possible states: Susceptible, Infected (and infectious), and Recovered (or removed). Thus, a susceptible individual can get infected with a certain probability if it is in direct contact with an infected one, and in turn, an infected individual recovers (or dies) with a different probability, not being able to get infected again in any case. This simple model describes many infectious diseases, such as measles, mumps and rubella. On Fig. 2.10 we show a simple scheme for the dynamics of this model. Of course, there are other models much more sophisticated, that take into account other intermediate states in the infectious process, such as latency, infected asymptomatic individuals or vaccination (see for example [44, 45]).

As a first approximation, one can consider the homogeneous mixing hypothesis, which assumes that people with whom a susceptible individual has contact are chosen at random from the whole population. This is a strong and somehow questionable assumption, since it does not take into account local details, such as individual diversity on the number of acquaintances, community structure or geographic constrictions. And, on the other hand, one should take into account that some illness like

the common cold, can be modeled accurately enough as a random-contact process, ignoring the social structure underneath, while it has been proved than for some others, such as the venereal diseases, one can not even describe them using a random degree distribution for the population, but a scale-free, so in these cases, the structure is essential.

Nonetheless, this approximation made by the SIR model allows us to describe analytically the behavior of the models simply by using ordinary differential equations for the densities of individuals in each compartment:

$$\begin{aligned}\frac{ds(t)}{dt} &= -\lambda \bar{k} \rho(t) s(t) \\ \frac{d\rho(t)}{dt} &= -\mu \rho(t) + \lambda \bar{k} \rho(t) s(t) \\ \frac{dr(t)}{dt} &= \mu \rho(t),\end{aligned}\tag{2.18}$$

where $s(t)$, $\rho(t)$ and $r(t)$ are respectively, the fraction of susceptible, infected and recovered individuals on the population at time t , so $s(t) + \rho(t) + r(t) = 1$. On the other hand, one susceptible individual becomes infected (if in contact with another infected one) with a probability λ , an infected individual recovers (or dies) with a probability μ , and $\bar{k}$ stands for the connectivity of the population, assumed exactly the same for everyone.

The most relevant prediction of this model is the existence of a non-zero *epidemic threshold*,

$$\lambda_c = 1/\bar{k}\tag{2.19}$$

so if $\lambda > \lambda_c$, the disease spreads and infects a finite fraction of the population, and if $\lambda < \lambda_c$, the total number of infected individuals (the so-called *epidemic incidence*, defined as $r_\infty = \lim_{t \rightarrow \infty} r(t)$) is infinitesimally small in the limit of a large population.

On the left panel of Fig. 2.11 we show an example of time evolution of the dynamics for a meaningful set of the parameters, namely, for $\lambda = 0.94$, $\mu = 1.0$, $\bar{k} = 6$ and using as initial conditions: $s(0) \simeq 1$, $\rho(0) \simeq 0$ and $r(0) \simeq 0$. On the right panel, it is shown the dependence of the epidemic incidence with the infection probability λ .

To deal with situations where the population is not well-mixed, or as we have mentioned before, the nature of the disease itself does not allow us to treat the pattern of interactions as homogeneous, we will need to represent the system as a graph, where nodes are the individuals (belonging to one of the three possible states: Susceptible, Infected or Recovered), and links are the interactions through which a susceptible node can become infected, if it has another infected node as a neighbor. So now, we want study the SIR process on an uncorrelated heterogeneous network (with generic degree distribution $P(k)$ and a finite average connectivity $\langle k \rangle$). We will study $s_k(t)$, $\rho_k(t)$ and $r_k(t)$, meaning the time evolution of the fractions of susceptible, infected and recovered individuals, respectively, within a connectivity class k , and with the normalization condition: $s_k(t) + \rho_k(t) + r_k(t) = 1$ for any given connectivity

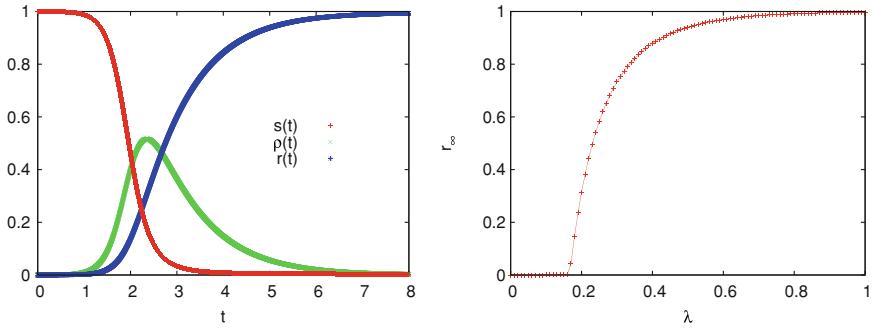


Fig. 2.11 Time evolution of the SIR dynamics (*left*) for $\lambda = 0.94$, $\mu = 1.0$, $\bar{k} = 6$ and taking $s(0) \simeq 1$, $p(0) \simeq 0$ and $r(0) \simeq 0$ as initial conditions, and the dependence of the epidemic incidence (*right*) with the probability of infection, λ for $\mu = 1.0$ and $\bar{k} = 6$

class and time instant. The global magnitudes are now given by the average over all the classes of connectivity present on the graph, so for example, the total fraction of infected individuals on the population at a given time t is: $\rho_k(t) = \sum_k P(k) \rho_k(t)$. Here it is important to notice that the network is considered static, so $P(k)$ does not change over time.

The equations for the evolution of the three compartments are similar to Eq. 2.18, but now we differentiate among connectivity classes:

$$\begin{aligned} \frac{ds_k(t)}{dt} &= -\lambda k s_k(t) \Theta(t) \\ \frac{d\rho_k(t)}{dt} &= -\mu \rho_k(t) + \lambda k s_k(t) \Theta(t) \\ \frac{dr_k(t)}{dt} &= \mu \rho_k(t), \end{aligned} \quad (2.20)$$

where $\Theta(t)$ is the probability of a given link to point towards an infected node, and is given by:

$$\Theta(t) = \frac{\sum_k k P(k) \rho_k(t)}{\langle k \rangle}. \quad (2.21)$$

Notice that this probability is the same for any node we consider, so it does not take into account any possible correlations between the connectivity of the nodes.

Again, one can get that there is an epidemic threshold, given by:

$$\lambda_c = \frac{\langle k \rangle}{\langle k^2 \rangle} \quad (2.22)$$

below which the epidemic incidence is zero, and above which it has a finite value. As we can see, this threshold depends inversely on the connectivity fluctuations of

the network the disease is spreading on, so for a system whose topology has a finite value, $\langle k^2 \rangle$, such as a random graph, then we get a threshold with a finite value as well (and therefore, a standard phase transition scenario). However, for scale-free networks we know that their connectivity fluctuations $\langle k^2 \rangle$ diverge when $N \rightarrow \infty$, which implies a vanishing epidemic threshold for increasingly larger systems.

The absence of a threshold in scale-free topologies is an important result that differs drastically from the one obtained for random networks or well-mixed scenarios, and it should be taken into account, for instance for prevention or vaccination strategies to be used by the health authorities, in order to efficiently fight off epidemics.

On the other hand, it is also worth noticing that real networks, even when they present some degree of heterogeneity on the connections, do have a finite size, and thus an effective threshold, depending on its $\langle k \rangle$ and $\langle k^2 \rangle$. Nonetheless, this value is usually very small for a large enough population, and is considerably smaller than the one for a random graph of the same size.

With regard to immunization strategies on scale-free topologies, we can point out that random vaccination is not effective, since there is always a non-zero epidemic incidence, even for very high vaccination ratios among the population. Nonetheless, targeted immunization, i.e., vaccinating the most connected individuals in a population, can give better results. On the other hand, is not always realistic to assume that the number of connections of a node on a real network can be known. A possible solution to this problem is the vaccination of random acquaintances of random chosen individuals, since the probability of reaching a particular node by following a randomly chosen edge is proportional to its degree.

Finally, we will say that for correlated networks it has been found that the qualitative behavior is the same as for uncorrelated networks, although there are some quantitative differences: on the one hand, while the likelihood of an epidemic outbreak is not modified when taking into account positive correlations, the epidemic incidence is smaller than in networks without correlations, and on the other hand, the diseases can live longer in assortative topologies.

Synchronization

Synchronization [1, 36–40] is a self-organized phenomenon where a set of individuals, initially acting on their own, gradually become more similar in their deeds, without any appointed leader or environmental external signal to guide them. In this way, after some time, they start behaving under the same pattern, showing, if not total, at least some identifiable level of clocking: they became somewhat ‘in sync’. There are many examples of synchronization in natural and human systems: crickets chirping in a summer night, neurons firing at the same pace, kids playing or singing along on spur of the moment, or groups of women living together, whose periods synchronize,...

A simple model has been frequently used in order to address synchronization: the Kuramoto model. It approaches the problem considering a mean field approximation, where every individual is an oscillator, and they are all supposed to interact to

everyone else through a purely sinusoidal coupling, so the governing equations for each one of them is given by:

$$\dot{\theta}_i = \omega_i + \frac{K}{N} \sum_{j=1}^N \sin(\theta_j - \theta_i) \quad (2.23)$$

where K is the coupling constant, ω_i is the natural frequency of the oscillator i , and the factor $1/N$ is incorporated to make sure that the system behaves correctly in the thermodynamic limit. The natural frequencies are assumed to be distributed according to some unimodal and symmetric function, whose mean frequency is Ω .

The collective behavior of the whole system is described by the macroscopic complex order parameter:

$$r(t)e^{i\phi(t)} = \frac{1}{N} \sum_{j=1}^N e^{i\theta_j(t)} \quad (2.24)$$

so the modulus $0 \leq r \leq 1$ measures the *phase coherence* of the population, whereas $\phi(t)$ is the average phase. The value $r \simeq 0$ corresponds to the lack of synchronization (the oscillators move incoherently) and $r \simeq 1$ to the case where almost the whole system is in sync (their phases are locked). The existence of a critical value, K_c , can be derived for the coupling, which separates a ‘disordered’ from an ‘ordered’ regime. In this second regime (when $K \geq K_c$), there are two types of long term behavior: a group of oscillators for which $|\omega_i| \leq Kr$, that are phase-locked at frequency Ω , and the rest of them, with $|\omega_i| > Kr$, that are drifting around the circle, sometimes accelerating and sometimes rotating at lower frequencies.

If one should include some kind of structure in the population in order to give an account of the complex interaction patterns among individuals, then, instead of Eq. 2.23, one needs to consider an extension of it:

$$\dot{\theta}_i = \omega_i + \sum_{j=1}^N \sigma_{ij} a_{ij} \sin(\theta_j - \theta_i) \quad (2.25)$$

where σ_{ij} accounts for the specific coupling strength between individuals i and j , and a_{ij} is the adjacency matrix of the network.

The mean field approach for complex networks considers that every oscillator is influenced by the local field created in its neighborhood, so the local order parameter is proportional to the connectivity of the node, k_i . It can be obtained the critical coupling for this situation:

$$\sigma_c = K_c \frac{\langle k \rangle}{\langle k^2 \rangle}. \quad (2.26)$$

It is to say, we get a rescaled critical value for the all-to-all topology, K_c , by the ratio between the mean connectivity of the particular network and its fluctuations.

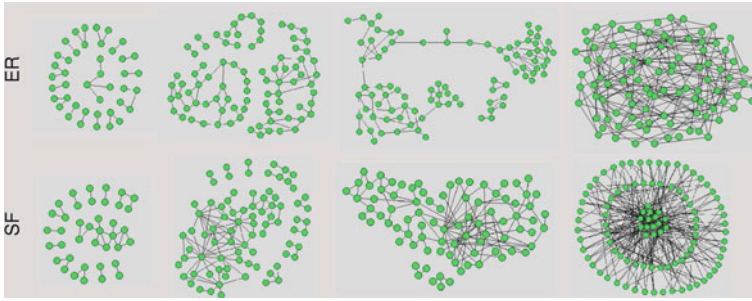


Fig. 2.12 Schematic representation of the different paths to synchronization displayed for SF (bottom) and ER (top) networks (higher values of the coupling strength are shown from left to right). Original figure from [36]

So once again, it is clear that for random networks there will be a threshold, but for (infinite) SF networks, this critical value will tend to zero.

Besides, it is important to point out that no exact analytical results for the Kuramoto model on general complex networks are available up to date, but one can always numerically simulate its dynamics. These simulations [36, 39] confirm the theoretical predictions, since they have shown that the onset of synchronization first occurs for SF, and as the topology becomes more homogeneous, the critical point moves to larger values, and the system seems to be less synchronizable. On the other hand, the particular paths to synchronization [36, 40] are also very different depending on the underlying structure (see Fig. 2.12): in SF networks, links and nodes are incorporated together to the largest of the synchronized clusters, while for homogeneous topologies, what are added are links between nodes already belonging to such cluster, making the route to complete synchronization a ‘sharper’ process, somehow. In other words, in the presence of hubs, a giant component of synchronized pair of oscillators forms and grows by recruiting nodes linked to them, while on the contrary, in homogeneous structures, many small clusters first appear and then grow together.

Cultural Dissemination

A very interesting aspect of human interactions is how people from different cultures can relate to each other, changing some of their own cultural traits when they meet. Nonetheless, if two individuals do not share any cultural features to begin with, it will be probably very hard for them to communicate and interact, but if they do have initially something in common (like some interests, hobbies, goals or even an aversion against something), they may start some kind of relationship. Moreover, it makes sense to assume that the more similar they are before meeting each other, the more likely it is for them to interact and become even more similar after that (this phenomenon is known as *homophilia*). As a result, not only individuals, but also societies change over time due to this mechanism of cultural influence. Thus,

one would expect these societies to become homogeneous (global) over time, as far as culture is concern, but as it turns out, sometimes they do not. Instead, such interactions can give rise to different groups with practically nothing in common, surprisingly enough.

Since Axelrod proposed his agent-based model [46] to address the issue of cultural dissemination in 1997, much effort has been put on studying these kind of processes [43, 47–52]. Under this paradigm, we generally consider that an individual's culture can be represented in terms of a set of attributes, such as language, religion, technology, dressing style, literary and musical preferences, sport preferences, and so on. Thus, an individual can be represented using a vector $\vec{V}_i = (v_i^1, v_i^2, \dots, v_i^F)$, with $i = 1, 2, \dots, N$, and where F is the total number of features that define a culture. Each one of these components can take only Q integer values, or cultural traits, and we assume that Q is the same for the F features. It is worth noticing that within this model, we do not consider as 'cultural' those features an individual can not change, for example skin color or physical constitution. Besides, we will consider our society to be placed in a lattice of size $L \times L = N$, where individuals will interact only with their neighbors.

Once we have randomly distributed the initial values for all the features of every individual in the system, the cultural interaction dynamics is defined as follows: every time step, an individual i is randomly chosen and one of its neighbors j , is also randomly selected. One measures the overlap between their cultural vectors, given by:

$$S_{ij} = \frac{1}{F} \sum_{l=1}^F \delta(v_i^l - v_j^l) \quad (2.27)$$

where $\delta(x) = 1$ if $x = 0$ and $\delta(x) = 0$ otherwise. If these two individuals are totally different ($S_{ij} = 0$) or exactly the same ($S_{ij} = 1$), then nothing happens, since the link between them is *blocked*. But if that is not the case, and $S_{ij} \in (0, 1)$, then the link is '*active*', and we then consider the value of the overlap S_{ij} as the probability that one of them imitates the other in one of the other features they have different. Obviously, the more similar they are, the higher the probability of becoming even closer through social interaction.

Letting the system evolve, it will eventually reach a *frozen state*, meaning that all the links between individuals are blocked. A useful order parameter is the relative size of the largest cultural cluster, $S_{\max}$, it is to say, the largest group of individuals that share the values for all their cultural features. According to some studies on lattices [47, 49, 51, 53], when $F > 2$, a non equilibrium first-order phase transition from order to disorder is observed as a function of the number of traits Q (the control parameter). There is a critical value, so if $Q < Q_c$, the final state of the system corresponds to $S_{\max} \sim 1$, a *global*, homogeneous state, while if $Q > Q_c$, then $S_{\max} \ll 1$, a *polarized* state with different *cultural domains* arises (see Fig. 2.13 (left)). This transition gets sharper as the size of the system increases.

If we analyze the time evolution of the relative number of blocked links (see Fig. 2.13 (right)), it can be seen that there is a non-zero initial value, due to just

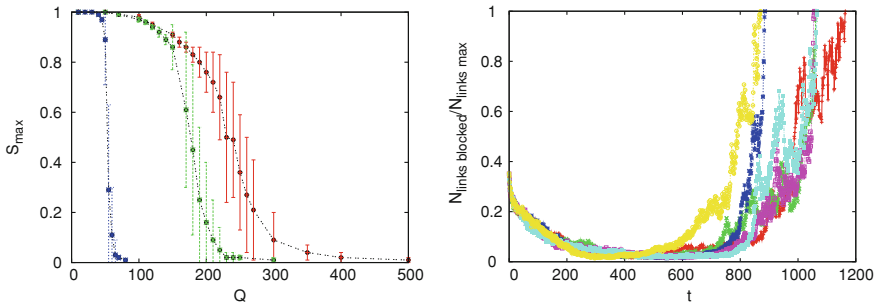


Fig. 2.13 *Left* Dependence of the largest cluster of global cultural consensus with the number of traits per feature for a 50×50 node square lattice with a 4-node neighborhood (*left*), ER random network (*center*) and BA scale-free network (*right*), and always for $F = 10$. The last two topologies have $\langle k \rangle = 6$ and $N = 10^3$ nodes. Every point is the average of 100 independent realizations. *Right* Several examples of time evolution of the relative number of blocked links. The underlying topology is a SF network made up of $N = 10^3$ nodes and $\langle k \rangle = 6$ and for a fixed value of $F = 10$

random assignment of the traits, that drops quickly as the dynamics starts, and individuals begin to interact. Then, this magnitude remains very low for a considerable amount of time, to finally rise up to the final value, corresponding with the rapid rise of $S_{\max}(t)$. This reflects the fact that, while the individuals have almost nothing in common, the system seems to spend a lot of time in that state, unable to get to an agreement, but once the individuals share some values for the features, then the final state is rapidly achieved. Notice that every realization shown in Fig. 2.13 (*right*) reaches its final state at its particular ‘consensus time’, since it is a stochastic process.

If we consider now that the pattern of interactions is given by a finite complex network [47], instead of by a lattice, the general picture of the phase transition remains unaltered (see Fig. 2.13 (*left*)), but with a higher value for Q_c (even higher for SF than for random networks, but qualitatively similar).

On the other hand, recent studies [52] have shown that, one can analyze the cultural evolution process towards the final state, from a global point of view (it is to say, considering the macroscopic level of consensus in the system though $S_{\max}$), but also from a feature level. It means that at any given time, we consider F layers or subgraphs of the original graph G . In the subgraph $G_f(t)$, two individuals are connected if they are physically connected in G , and if they share the value of the feature f at that precise instant of time. In this way, we can observe how cultural consensus evolve in every layer, $S_{\max}^f$, and we get to discover that there are some relevant differences between the two approaches: while for the global consensus point of view, the system remains apparently unordered for a large fraction of the simulation time, to finally get organized very quickly (Fig. 2.14 (*left*)), the organization at a feature level starts much earlier. Actually, $S_{\max}^f$ increases monotonously over time from the very beginning (Fig. 2.14 (*right*)).

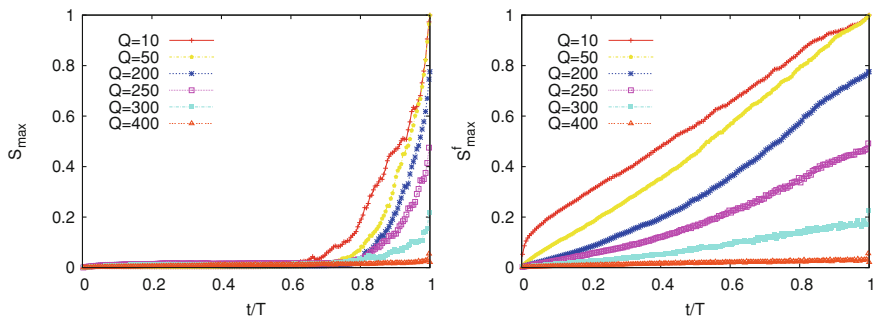


Fig. 2.14 Time evolution (relative to the final consensus time T) of the largest cluster of cultural consensus at global (*left*) and at feature (*right*) level for a value of $F = 10$ in SF networks made up of $N = 4 \cdot 10^3$, with average connectivity $\langle k \rangle = 6$

Finally, it is also worth mentioning that there are many other works with different variations of the Axelrod model [43], including for example noise [54], an external field [55], rewiring of the connections between nodes [56], even movility of the individuals [57], or even a combination of the original Axelrod model for cultural dissemination with the original Schelling model of social segregation [58].

2.2 Games

A *game* can be considered as a formal abstraction of social interactions between individuals. There must be at least two decision makers (or *players*), who can choose between at least two different actions (also called *strategies*). It is worth stressing that a player does not need a brain in order to adopt a strategy, on the contrary, they can be very simple agents: bacteria, for example, have the basic capacities to play games, since they are highly responsive to certain aspects of their chemical environment, and they can respond differently depending on the actions of their neighbors, the behavior can affect the fitness of others and vice versa, and finally, the conditional strategies can be inherited by the offspring [59]. The outcome of the interaction depends on the strategy every player adopts. Thus, *Game Theory* is a branch of applied Mathematics that tries to capture these situations and it is usually considered to have its origin with the work of von Neumann and Morgenstern [60] in 1944. Historically, Game Theory has been used in very different fields, such as economics, biology, political science or sociology, and there are two main different approaches: Classic Game Theory and Evolutionary Game Theory, which made different assumptions about the systems they model.

Classic Game Theory formally studies how rational players should behave in order to obtain the maximum possible benefit or payoff. Nonetheless, one could easily object to the concept of ‘rational player’ as an accurate representation of real

individuals in a social or biological context. ‘Rational player’ means that its only goal and motivation is to maximize its benefits, given its belief about its opponent’s strategy, but there are plenty of real situations where the actions of the players do not seem to aim a maximum payoff.

Evolutionary Game Theory [61–63] was originated in 1973 by the work of Maynard Smith and Price works. It studies the time evolution of large populations of individuals who repeatedly play a game and are exposed to selection and replication (with or without mutation). Their strategies are fixed, and usually, the encounters between the individuals are supposed to happen at random, in a ‘well-mixed’ situation, so there is no social structure behind it (everyone interacts with everyone else), and it allows for the analytical treatment of the problem. Thus, the probability of interacting with an individual that uses strategy i is proportional to the fraction of individuals that are using that particular strategy in the system at the moment, x_i . The payoffs from all these interactions are added up, and success in the game is interpreted as reproductive success. Thus, payoff means fitness in the Darwinian way: the strategies that perform better, reproduce faster, which can be straightforwardly interpreted as natural selection.

In this section we intend to establish just a few useful concepts and results in Classical Game Theory, always keeping in mind that our goal is to understand the problem of cooperation. Then we will move on to the approach given by Evolutionary Game Theory, and finally, we will point out some mechanisms that have been introduced to explain the survival of cooperation observed in several natural and social systems, specially, the differences in the outcome of a game when dealing with a structured population, it is to say, when we have an underlying topology.

2.2.1 Classical Game Theory

In Classical Game Theory (CGT), we consider that interacting individuals can choose a strategy -or a way to act- among a well-defined set of them. A game is called *normal-form* if it is determined by a payoff matrix. Thus, for instance in a 2×2 game, we have two players and two different strategies A and B , and then depending on their particular choices, the benefits the players will obtain are given by the payoff matrix:

$$\begin{array}{cc} & \begin{array}{cc} A & B \end{array} \\ \begin{array}{c} A \\ B \end{array} & \begin{pmatrix} a & b \\ c & d \end{pmatrix} \end{array} \quad (2.28)$$

This means that, for instance, when a player uses strategy A against a player using also A , it get a payoff equal to a , when a player uses strategy A against a player using a strategy B , it get a payoff equal to b , and so on. We say that strategy A *dominates* strategy B , if $a > c$ and $b > d$. In that situation, no matter what strategy your opponent uses, it is better always to use A . Conversely, B *dominates* A , if $a < c$ and $b < d$.

Now, in a general case of a $N \times N$ payoff matrix U , if we denote the N *pure strategies* by $R_1, R_2, \dots, R_N$, then the simplex S_N of the linear combinations of pure strategies:

$$S_N = \left\{ p = (p_1, p_2, \dots, p_N) : p_i \geq 0 \text{ and } \sum_i p_i = 1 \right\} \quad (2.29)$$

is the set of *mixed strategies*. A mixed strategy can be seen as the one used by a player that chooses strategy R_i with a probability p_i , where $i = 1, 2, \dots, N$. The N vertexes of the simplex S_N are the N pure strategies, while the interior of the simplex is the set of completely mixed strategies, it is to say, those for which $p_i > 0 \forall i$. The boundaries of the simplex, on the other hand, correspond to mixed strategies that must have necessarily one of the probabilities set to zero. We can calculate the benefit of a p -strategist against a q -strategist as:

$$pUq = \sum_{i,j} p_i u_{ij} q_j \quad (2.30)$$

and the set of strategies for which the application $p \rightarrow pUq$ achieves its maximum value is called *best responses* to q .

A strategy q is called a *Nash Equilibrium* (originally called ‘equilibrium for n -person games’ by Nash in 1950 in [64]) if it is the best response to itself. This means that if two individuals are both using a strategy that is a Nash Equilibrium, then neither of them can unilaterally deviate from that strategy and increase its payoff. Moreover, a Nash Equilibrium is called *Strict* if it is the only best response to itself, therefore $\forall p \neq q$ it is fulfilled that $pUq < qUq$. If q is a Nash Equilibrium, then there is a constant c that satisfies that $(Uq)_i \leq c$, and from this result can be derived that a Nash Equilibrium is always a pure strategy.

A strategy $\hat{p}$ is *Evolutionary Stable* if $\forall p \in S_N$ with $p \neq \hat{p}$ the inequity:

$$pU(\epsilon p + (1 - \epsilon)\hat{p}) < \hat{p}U(\epsilon p + (1 - \epsilon)\hat{p}) \quad (2.31)$$

is fulfilled $\forall \epsilon > 0$, as long as it is smaller than a certain appropriate invasion threshold $\bar{\epsilon}(p)$. It can be proven the following logic chain:

Strict Nash Equilibrium $\rightarrow$ Evolutionary Stable Strategy $\rightarrow$ Nash Equilibrium.

Let’s now consider again a particular set of 2×2 games. We can analyze the possible outcomes within the CGT framework. We consider two different strategies: cooperate (C) and defect (D), and the correspondent payoff matrix:

$$\begin{array}{cc} & \begin{array}{cc} C & D \end{array} \\ \begin{array}{c} C \\ D \end{array} & \begin{pmatrix} R & S \\ T & P \end{pmatrix} \end{array} \quad (2.32)$$

Depending on the relative ordering of the parameters, we can define three games:

- The Hawks and Doves (or Snow Drift or Chicken) game [65–68] fulfills: $T > R > S > P$. Players are referred to as greedy, since they prefer unilateral defection to mutual cooperation ($T > R$). In this situation, C is the best response for D , and vice versa, so one should always try to choose the opposite of what the opponent does, in order to maximize the benefits.
- The Stag Hunt game [69, 70] satisfies $R > T > P > S$. Players prefer mutual defection to unilateral cooperation ($S < P$), resulting in an intrinsic fear of individuals to cooperate. In this situation, C is the best response for C , and D is the best response for D , or in other words, both are Nash equilibria, so it is better always to try to play the same strategy as your opponent.
- The Prisoner's Dilemma game [59, 63, 71–73], for which $T > R > P > S$, both tensions described above are incorporated at once, so is the most difficult situation for cooperation to arise. In this scenario, D dominates C . No matter what strategy your opponent uses, it is better always to defect.

2.2.2 Evolutionary Game Theory

Within the Theory of Evolution, the central actor of an evolutionary system is the *replicator*. A replicator is an entity that possesses the ability of making copies of itself. It can be a gene, an organism, a strategy in a game, a particular belief or opinion, a technique or any other cultural trait in general. A *replicator system* is a set of replicators in a particular environment, with some kind of interaction among the individuals. An evolutionary dynamics of a replicator system is a process of change over time on the replicator frequency distribution, in such a way that the strategies with higher benefits reproduce at a faster pace.

Let us consider that the population is divided into n types of individuals $E_1, E_2, \dots, E_n$ with frequencies (or relative abundances) $x_1, x_2, \dots, x_n$ respectively. The fitness (or expected number of descendants) f_i of the type E_i will be assumed to be a function of the composition of the whole population. If the population is big enough, and the individuals of a generation are supposed to meet and interact continuously and at random (*well-mixed* scenario), then we can consider that the state of the system $x(t)$ evolves in the simplex S_n as a derivable function of time. The increase of the rate $\dot{x}_i/x_i$ of the type E_n is a measure of its success, in the Darwinian evolutionary sense of the term. Then, we can express this success as the difference between the fitness f_i of this type and the average fitness of the population, $\bar{f}(x) = \sum_i x_i f_i(x)$, and thus describe the evolution of every type in the population using the *Replicator Equation* [62, 74–76]:

$$\dot{x}_i = x_i[f_i(x) - \bar{f}(x)] \quad (2.33)$$

with $i = 1, 2, \dots, n$. It is easy to see that the simplex S_n is invariant under these equations, so if $x(0) \in S_n$, then $x(t) \in S_n \forall t > 0$. Moreover, the faces of the simplex are also invariant: if one or several strategies are not present at a given moment t_0 of the evolution of the system, then they will never be for any t_1 , such as

$t_1 > t_0$. In the case of having mixed strategies, we can also obtain the correspondent Replicator Equation. If there is a game with N pure strategies $R_1, R_2, \dots, R_N$ and a $N \times N$ payoff matrix U , then a strategy is a point in the simplex S_N , and the $E_1, E_2, \dots, E_n$ types of individuals present in the system correspond to n points $p^1, p^2, \dots, p^n \in S_N$.

The state of the whole population is given by the frequencies x_i of the types E_i . The benefits of a p^i -strategist playing against a q^j -strategist is given by $a_{ij} = p^i U p^j$, and thus, the fitness f_i of the type E_i is $f_i(x) = \sum_j a_{ij} x_j = (Ax)_i$. A state $\hat{x} \in S_n$ is a Nash Equilibrium if $x A \hat{x} \leq \hat{x} A \hat{x}$, $\forall x \in S_n$, and it can be proven that if $\hat{x}$ is a Nash Equilibrium, then it is an equilibrium point of the Replicator Equation. A state $\hat{x} \in S_n$ is said *evolutionary stable* if $\forall x \neq \hat{x}$ in an environment of $\hat{x}$ it is fulfilled that $\hat{x} A x > x A x$. The same way, it can be proven that if $\hat{s}$ is an evolutionary stable state, then it is a point of asymptotically stable equilibrium of the Replicator Equation (but the reciprocal result is not necessarily true).

Replicator Equation for 2×2 Games

For the particular case of a 2×2 symmetric game, we will have again that the generic payoff matrix is given by:

$$\begin{matrix} & \begin{matrix} A & B \end{matrix} \\ \begin{matrix} A \\ B \end{matrix} & \begin{pmatrix} a & b \\ c & d \end{pmatrix} \end{matrix} \quad (2.34)$$

And according to the Evolutionary Game Theory, we should consider that the fitness of an individual playing a certain strategy depends on the fraction of individuals that play every strategy (it is to say, the so-called *frequency-dependent selection*), so if the vector $\vec{x} = (x_A, x_B)$ represents the composition of the population, in terms of the two possible strategies, and we denote respectively, $f_A(\vec{x})$ and $f_B(\vec{x})$ the fitness of both of them. The selection dynamics can be written as

$$\begin{aligned} \dot{x}_A &= x_A [f_A(\vec{x}) - \phi] \\ \dot{x}_B &= x_B [f_B(\vec{x}) - \phi] \end{aligned} \quad (2.35)$$

where $\phi = x_A f_A(\vec{x}) + x_B f_B(\vec{x})$ is the average fitness of the entire population. Obviously, since $x_A + x_B = 1$, we can consider $x \equiv x_A$ and $1 - x \equiv x_B$, and then we can rewrite the previous differential Eq. 2.35 in a simpler way as:

$$\dot{x} = x(1 - x)[f_A(x) - f_B(x)] \quad (2.36)$$

It can be easily shown that $x = 0$ is a stable equilibrium if $f_A(0) < f_B(0)$, and conversely, $x = 1$ is a stable equilibrium if $f_A(1) > f_B(1)$. On the other hand, any interior value of $x \in (0, 1)$ is a stable equilibrium x^* if the first derivative of the fitness functions satisfies $f'_A(x^*) < f'_B(x^*)$.

In particular we can calculate the expected fitness of an individual playing A or B respectively, in the well-mixed scenario explained before as:

$$\begin{aligned} f_A &= ax_a + bx_b \\ f_B &= cx_a + dx_b \end{aligned} \quad (2.37)$$

so if we again introduce this expression for the fitness in 2.35 we obtain:

$$\dot{x} = x(1-x)[(a-b-c+d)x + b-d] \quad (2.38)$$

Depending on the relative ordering of the coefficients of the payoff matrix, we can have different situations for the selection dynamics [73, 77, 78]:

- (a) A *dominates* B , if $a > c$ and $b > d$. No matter what strategy your opponent uses, it is better always to use A , and selection will lead to a final state where all players are A .
- (b) B *dominates* A , if $a < c$ and $b < d$. No matter what strategy your opponent uses, it is better always to use B , and selection will lead to a final state where all players are B .
- (c) A and B are *bistable*, if $a > c$ and $b < d$. In this situation, A is the best response for A , and B is the best response for B , so it is better always to try to play the same strategy as your opponent. There is an unstable equilibrium at $x^* = \frac{d-b}{a-b-c+d}$, and depending on the initial fraction of every strategy, the system will converge to all- A (if $x(0) > x^*$) or all- B (if $x(0) < x^*$).
- (d) A and B *coexist*, if $a < c$ and $b > d$. In this situation, A is the best response for B , and vice versa, so one should always try to choose the opposite of what the opponent does. Selection will make the system converge to the interior equilibrium $x^* = \frac{d-b}{a-b-c+d}$.
- (e) A and B are *neutral*, if $a = c$ and $b = d$. No matter what action you choose, you will always win exactly the same as your opponent, so selection will not modify the initial fraction of every strategy, but this scenario is obviously not very interesting for us.

And some other useful concepts are:

- (a) Strategy A is called *risk-dominant* if $a + b > c + d$, and then strategy B has a basin of attraction smaller than $1/2$.
- (b) Strategy A is called *pareto-efficient* if $a > d$.
- (c) Strategy A is *advantageous* if $a + 2b > c + 2d$, and then strategy B has a basin of attraction smaller than $1/3$.

As a particular example of 2×2 game, we have the Prisoner's Dilemma (see 2.32), that has been widely used to study the phenomenon of cooperation in very different fields, from biology to sociology or economics. It is obvious that defection is the best response, regardless the opponent's (it is in fact, the only Nash equilibrium), despite the fact that, if both cooperate, then they will win more than if both defect.

Thus, both in a Classic Game Theory approach, and in an Evolutionary context using the Replicator Equation we obtain straightforwardly an all-D state, since defectors have higher payoff than cooperators. Cooperation can not survive in a well-mixed situation, it is inevitable. In fact, there are a great deal of examples of this well-mixed or transitory-pairing environments in Nature, which lead to non-cooperative or exploiting situations for the individuals, on the contrary to what usually happens with stable pairing, or even mutualism between different species [59].

Finite Populations

Additionally, one can wonder what happens to the dynamics in the very realistic case of finite populations (notice that we still do not take into account an internal structure). In this case, in order to describe the evolution of a N -sized population, a stochastic theory is needed, and we calculate fixation probabilities for the different possible strategies [73, 79], instead of equilibrium states of the system. The *probability of fixation* of strategy B is the probability of a single mutant B to invade an entire population of A -players.

In order to approach this situation, we can use, among other stochastic processes, the *Moran process* [80], which could be a finite- N analogue to the Replicator Equation. It is a birth-death process that describes the probabilistic dynamics in a finite population of constant size N in which two strategies A and B are competing for dominance. In each time step, a random individual is chosen for reproduction and a random individual is chosen for death; thus ensuring that the population size remains constant. To model selection, one type has to have a higher fitness (considered constant) and is thus more likely to be chosen for reproduction. The same individual can be chosen for death and for reproduction in the same step. It is worth mentioning that in finite populations, even if all different strategies had the same fitness, all but one type will eventually go extinct. This principle is called *neutral drift*. Thus, since coexistence is not possible, there are as many absorbing states as different strategies at the beginning. In a population on size N made up of A individuals, we can calculate [73] the probability of fixation of another strategy B (it is to say, the probability for a single neutral mutant to take over the entire population), and it is given by $1/N$. It means that when dealing with finite populations, just due to random drift, a mutant (with the same fitness as the majority strategy) can invade the system, which is a very different outcome from the infinite-population scenario, where having the same fitness meant coexistence of different strategies. In the same way, the probability of ending up in an all- B state, just due to random drift, when starting with $i \leq N$ individual playing B in a population of A is i/N . On the other hand, if a mutant B has a relative fitness r , with respect to the A players, it can be proven [73] that its probability of fixation is then $\rho = \frac{1-1/r}{1-1/r^N}$. Notice that in this scenario, there is always a non-zero probability that a mutant strategy will invade and take over the whole population, even though it is opposed by selection [81].

2.2.3 Evolution of Cooperation

As we have seen previously, neither within the Classic or the Evolutionary Game approach, can cooperation survive. Nonetheless, there are plenty of examples of real situations where cooperators arise and thrive, so there must be some mechanisms behind it. Over the years, five main ideas [77] have been proposed to help understand this phenomenon: kin selection, direct reciprocity, indirect reciprocity, group selection and network reciprocity.

According to Hamilton [72], natural selection can favor cooperation if the donor and the recipient of an altruistic act are genetic relatives. More precisely, Hamilton's rule establishes that the coefficient of relatedness, r , must exceed the cost-to-benefit ratio of the altruistic act, it is to say: $r > c/b$. This coefficient r is defined as the probability of sharing a gene (it is equal to $1/2$ for siblings, equal to $1/8$ for cousins,...). This theory is called *Kin Selection*, but obviously it can not help understand cooperation among unrelated individuals, or even members of different species.

Trivers proposed the *Direct Reciprocity* mechanism. Let us assume that there are repeated encounters [71] of a the Prisoner's Dilemma Game between the same two individuals, and every time they can choose to be cooperators or defectors. The idea is that if I cooperate in this round of the game, maybe you will cooperate in the next one. When considering the repeated game on a whole population, it can be proven that direct reciprocity leads to the evolution of cooperation only if the probability of another encounter between the same two individuals, w , exceeds the cost-to-benefit ratio of the altruistic act: $w > b/c$.

Let us now consider the following scenario: among a population, two individuals meet once, one of them is in the position of helping the other one (this help is supposed to be less costly for the donor than beneficial for the receiver), and although there is no possibility for direct reciprocation, helping others will establish a good *reputation* which will be rewarded by others. In this way, when deciding how to act, one will take into consideration the consequences for their reputation. Moreover, the next step can be to take into consideration the opponents' reputation, in order to decide whether or not he/she deserves our help, and how it will affect our own. This theory constitutes *Indirect Reciprocity* [82, 83], and when applied to human behavior, it can help understand the origin of moral and social norms.

We can take into account that selection not only acts on the individual level, but also on the group level. A simple model for *Group Selection* is as follows [84]: the population is divided into different groups, and individuals cooperate inside its own group, while defectors do not help anyone. Individuals reproduce proportional to its fitness and the offspring belongs to the same group as the ancestors. When a group reaches certain size, it can split in two, making another group disappear, in order to preserve the total size of the population constant. In a mixed group, a defector reproduces faster than a cooperator, but groups of pure cooperators split faster than those of pure defectors. For the limit of weak selection and considering the case of rare group splitting, it can be obtained that, if n is the maximum group size and

m is the number of groups, then Group Selection allows evolution of cooperation, provided that: $b/c > 1 + (n/m)$, where b/c is the cost-to-benefit ratio.

Finally, one can realize that the Evolutionary approach for the PD game always leads to all-D situations, but it considers a well-mixed scenario, it is to say, at any given time, every individual has equal probabilities to interact with everyone else. Nonetheless we know that this is a very unrealistic assumption, since groups and societies have usually some kind of internal structure. In other words, there is a well defined pattern of interactions among individuals, so every one of them has a fixed number of neighbors. It has been shown that spatial structure affects greatly the outcome of an evolutionary dynamics, allowing cooperators to survive in many situations. Specifically, cooperators form network clusters, where they help each other. The analytical treatment of this problem is hard, and many times, even impossible, but it has been found that this *Network Reciprocity* can favor cooperation if $b/c > k$, where k stands for the average number of connections of the individuals in the population.

Prisoner's Dilemma Game on Structured Populations

According to what we have seen previously, one of the mechanisms that helps promote cooperation is Network Reciprocity, and it happens to be also the one we will be interested during this thesis. Thus, the natural next step for us in order to build more realistic models of social or biological interactions, is to consider some sort of underlying structure, in account for the particular pattern of relationships between individuals. The first attempts to model such social structure for the Prisoner's Dilemma game considered the individuals placed in a regular lattice [85–91]. Those studies found that spatial structure affects greatly the outcome of such dynamics. Specifically, by making the agents play just with a small number of fixed neighbors, we can make cooperation and defection coexist, or even enhance cooperation. In fact, when dealing with games in spatial structure populations, the equilibria among strategies are no longer necessarily characterized by their having equal average payoff. Instead, the asymptotic equilibrium properties are now determined by 'local relative payoffs', and not by global averages [88]. It was also found for the PD in lattices, that under certain symmetrical initial conditions for the distribution of strategies, certain values of the temptation to defect b , and as long as we use deterministic updating rules, kaleidoscopic carpet-like chaotically-changing spatial patterns arise [86, 87]. Moreover, it has been found that there is a critical phase transition in the Prisoner's Dilemma game in lattices that falls into the same universality class than directed percolation [91].

Some effort was put also on the analytical study of how different kind of structures can favor fixation of the strategies or, on the contrary, favor neutral drift, explicitly calculating to that end the corresponding probabilities of fixation of the strategies on some networks with very particular topologies, such as stars, paths, downstreams, upstreams or funnels [73, 76, 92]. Moreover, striking results in terms of survival of cooperation were found for random and SF networks, but for such general structures,

no explicit calculations can be performed, so one needs to rely totally on simulations. In this area, a great deal of effort has been put too, and as a very general remark, it can be said that the complex topologies behind the interactions among a given population affect the outcome of any process [29, 30, 36, 40–42, 93] not only games [44, 76, 86, 92] to a large extent. Specifically, as we will see with some detail in Chap. 3, when it comes to the Prisoner’s Dilemma game on complex networks, a large number of studies [66, 94–98] have pointed out that cooperation benefits from heterogeneity. It is to say, it has much better chances to survive in scale-free than in random topologies, for the same given values of the parameters of the game.

References

1. S. Boccaletti, V. Latora, Y. Moreno, M. Chavez, and D. U. Hwang, *Phys. Rep.* **424**, 175 (2006).
2. M. E. J. Newman, *SIAM Review* **45**, 167 (2003).
3. J. Dunne, R. Williams, and N. Martinez, *Marine Ecological Press Series* **273**, 291 (2004).
4. M. E. J. Newman, *Proc. Natl. Acad. Sci. USA* **98**, 404 (2001).
5. A. Arenas, L. Danon, A. Díaz-Guilera, P. Gleiser, and R. Guimerà *European Physical Journal B* **38**(2), 373 (2004).
6. R. Guimerà and M. Sales-Pardo, *Proc. Natl. Acad. Sci. USA* **106**, 22073 (2009).
7. S. Gómez, P. Jensen, and A. Arenas, *Phys. Rev. E* **80**, 016114 (2009).
8. M. Newman and M. Girvan, *Phys. Rev. E* **69**, 026113 (2004).
9. L. Danon, A. Díaz-Guilera, J. Duch, and A. Arenas, *J. Stat. Mech.* p. P09008 (2005).
10. A. Barabási and Z. Oltvai *Nature Reviews, Genetics* **5**, 101 (2004).
11. J. Duch, and A. Arenas *Phys. Rev. E* **72**, 027104 (2005).
12. L. Danon, J. Duch, A. Arenas, and A. Díaz-Guilera, *World Scientific* p. 93 (2007).
13. M. Sales-Pardo, R. Guimerà, A. Moreira, and L. Amaral *Proc. Natl. Acad. Sci. USA* **104**, 15224 (2007).
14. P. Erdős, and A. Renyi, *Publicationes Mathematicae Debrecen* **6**, 290 (1959).
15. D. J. Watts and S. H. Strogatz, *Nature* **393**, 440 (1998).
16. S. Milgram, *Psychol. Today* **2**, 60 (1967).
17. J. Guare, *Six degrees of separation: a play*. (Vintage Books, New York, 1990).
18. J. Travers and S. Milgram, *Sociometry* **32**, 425 (1969).
19. L. Amaral, A. Scala, M. Barthélémy, and H. Stanley, *Proc. Natl. Acad. Sci. USA* **97**, 11149 (2000).
20. M. Barthélémy and L. Amaral, *Phys. Rev. Lett.* **82**, 3180 (1999).
21. H. Simon, *Biometrika* **42**, 425 (1955).
22. R. Merton, *Science* **159**, 56 (1968).
23. R. Albert, and A. L. Barabási, *Rev. Mod. Phys.* **74**, 47 (2002).
24. S. N. Dorogovtsev, J. F. Mendes, and A. N. Samukhin, *Phys. Rev. Lett.* **85**, 4633 (2000).
25. S. N. Dorogovtsev, and J. F. F. Mendes, *Evolution of networks. From biological nets to the Internet and the WWW*. (Oxford University Press, Oxford, UK, 2003).
26. P. Holme, and B. Kim, *Phys. Rev. E* **65**, 026107 (2002).
27. G. Caldarelli, A. Capocci, P. D. L. Rios, and M. A. M. noz, *Phys. Rev. Lett.* **89**, 258702 (2002).
28. J. Gómez-Gardeñes and Y. Moreno, *Phys. Rev. E* **73**, 056124 (2006).
29. R. Pastor-Satorras, and A. Vespignani, *Phys. Rev. Lett.* **86**, 3200 (2001a).
30. R. Pastor-Satorras and A. Vespignani, *Phys. Rev. E* **63**, 066117 (2001b).
31. Y. Moreno, R. Pastor-Satorras, and A. Vespignani, *European Physical Journal B* **26**, 521 (2002).
32. R. Pastor-Satorras, and A. Vespignani, *Phys. Rev. E* **65**, 036104 (2002).

33. M. Boguñá, R. Pastor-Satorras, and A. Vespignani, *Phys. Rev. Lett.* **90**, 028701 (2003).
34. R. May, and A. Lloyd, *Phys. Rev. E* **64**, 066112 (2001).
35. M. Newman, *Phys. Rev. E* **66**, 016128 (2002).
36. J. Gómez-Gardeñes, Y. Moreno, and A. Arenas, *Phys. Rev. Lett.* **98**, 034101 (2007a).
37. L. Donetti, P. I. Hurtado, and M. A. Muñoz, *Phys. Rev. Lett.* **95**, 188701 (2005).
38. J. Gómez-Gardeñes, and Y. Moreno, *Int. J. Bifurcation Chaos* **17**, 2501 (2007).
39. J. Gómez-Gardeñes, Y. Moreno, and A. Arenas, *Phys. Rev. E* **75**, 066106 (2007b).
40. A. Arenas, A. Díaz-Guilera, J. Kurths, Y. Moreno, and C. Zhou, *Phys. Rep.* **469**, 93 (2008).
41. R. Cohen, K. Erez, D. ben-Avraham, and S. Havlin, *Phys. Rev. Lett.* **86**, 3682 (2001).
42. R. Cohen, K. Erez, D. ben-Avraham, and S. Havlin, *Phys. Rev. Lett.* **85**, 4626 (2000).
43. C. Castellano, S. Fortunato, and V. Loreto, *Rev. Mod. Phys.* **81**, 591 (2009).
44. G. Szabó and G. Fáth, *Phys. Rep.* **446**, 97 (2007).
45. V. Colizza, A. Barrat, M. Barthélemy, A. Valleron, and A. Vespignani, *PLoS Medicine* **4**, e13 (2007).
46. R. Axelrod, *J. Conflict Res.* **41**, 203 (1997b).
47. K. Klemm, V. Eguíluz, R. Toral, and M. San Miguel, *Phys. Rev. E* **67**, 026120 (2003a).
48. D. Vilone, A. Vespignani, and C. Castellano, *European Physical Journal B* **30**, 399 (2002).
49. K. Klemm, V. Eguíluz, R. Toral, and M. San Miguel, *Physica A* **327**, 1 (2003b).
50. J. González-Avella, V. Eguíluz, M. Cosenza, K. Klemm, J. Herrera, and M. San Miguel, *Phys. Rev. E* **73**, 046119 (2006).
51. C. Castellano, M. Marsili, and A. Vespignani, *Phys. Rev. Lett.* **85**, 3536 (2000).
52. B. Guerra, J. Poncela, J. Gómez-Gardeñes, V. Latora, and Y. Moreno, *Phys. Rev. E* **81**, 056105 (2010).
53. K. Klemm, V. Eguíluz, R. Toral, and M. San Miguel, *J. Economic Dynamics and Control* **29**, 321 (2005).
54. K. Klemm, V. Eguíluz, R. Toral, and M. San Miguel, *Phys. Rev. E* **67**, 045101(R) (2003c).
55. J. Gonzalez-Avella, M. Cosenza, V. Eguíluz, and M. San Miguel, *New J. Phys.* **12**, 013010 (2010).
56. A. Grabowski, and R. A. Kosinski, *Phys. Rev. E* **73**, 016135 (2006).
57. C. Gracia-Lázaro, L. Lafuerza, L. Floria, and Y. Moreno, *Phys. Rev. E* **80**, 046123 (2009).
58. T. Schelling, *The American Economic Review* **59**, 488 (1969).
59. R. Axelrod, and W. Hamilton, *Science* **211**, 1390 (1981).
60. J. von Neumann, and O. Morgenstern, *Theory of Games and Economic Behavior*. (Princeton University Press, Princeton, NJ, 1944).
61. J. Maynard Smith, and G. Price, *Nature* **246**, 15 (1973).
62. H. Gintis, *Game theory evolving*. (Princeton University Press, Princeton, NJ, 2000).
63. J. Hofbauer, and K. Sigmund, *Evolutionary games and population dynamics*. (Cambridge University Press, Cambridge, UK, 1998).
64. J. Nash, *Proc. Natl. Acad. Sci. USA* **36**, 48 (1950).
65. T. Killingback, and M. Doebeli, *Proc. R. Soc. Lond.* **263**, 1135 (1996).
66. C. P. Roca, J. A. Cuesta, and A. Sánchez, *Phys. Rev. E* **80**, 046106 (2009).
67. M. Tomassini, L. Luthi, and M. Giacobini, *Phys. Rev. E* **73**, 106 (2006).
68. C. Hauert, and M. Doebeli, *Nature* **428**, 643 (2004).
69. L. E. Blume, *Games and Economic Behavior* **5**, 387 (1993).
70. G. Ellison, *Econometrica* **61**, 1047 (1993).
71. R. Axelrod, *The Evolution of Cooperation*. (Basic Books, New York, 1984).
72. W. Hamilton, *J. Theor. Biol.* **7**, 1 (1964).
73. M. Nowak, *Evolutionary dynamics: exploring the equations of life*. (Harvard University Press, Cambridge, MA, 2006b).
74. J. Hofbauer, P. Schuster, and K. Sigmund, *J. Theor. Biol.* **81**, 609 (1979).
75. P. Taylor, and L. Jonker, *Math. Biosci.* **40**, 145 (1978).
76. H. Ohtsuki, and M. A. Nowak, *J. Theor. Biol.* **243**, 86 (2006).
77. M. Nowak, *Science* **314**, 1560 (2006a).
78. M. A. Nowak, and K. Sigmund, *Science* **303**, 793 (2004).

79. A. Traulsen, M. Nowak, and J. Pacheco, *Phys. Rev. E* **74**, 011909 (2006).
80. P. Moran, *Mathematical Proceedings of the Cambridge Philosophical Society* **54**, 60 (1958).
81. M. Nowak, A. Sasaki, C. Taylor, and D. Fudenberg, *Nature* **428**, 646 (2004).
82. M. Nowak, and K. Sigmund, *Nature* **393**, 573 (1998).
83. M. Nowak, and K. Sigmund, *Nature* **437**, 1291 (2005).
84. A. Traulsen, and M. A. Nowak, *Proc. Natl. Acad. Sci. USA* **103**, 10952 (2006).
85. M. Nowak and K. Sigmund, *Games on Grids*, in: *The Geometry of Ecological Interactions*. (Cambridge University Press, Cambridge, UK, 2000).
86. M. A. Nowak and R. M. May, *Nature* **359**, 826 (1992).
87. M. Nowak, S. Bonhoeffer, and R. May, *Int. J. Bifurcation Chaos* **4**, 33 (1994).
88. M. Nowak, S. Bonhoeffer, and R. May, *Proc. Natl. Acad. Sci. USA* **91**, 4877 (1994).
89. G. Szabó and C. Tóke, *Phys. Rev. E* **58**, 69 (1998).
90. M. Nakamaru, H. Matsuda, and Y. Iwasa, *J. Theor. Biol.* **184**, 65 (1997).
91. C. Hauert and G. Szabó, *Am. J. Phys.* **73**, 405 (2005).
92. E. Lieberman, C. Hauert, and M. A. Nowak, *Nature* **433**, 312 (2005).
93. D. Callaway, M. Newman, S. Strogatz, and D. Watts, *Phys. Rev. Lett.* **85**, 5468 (2000).
94. G. Abramson, and M. Kuperman, *Phys. Rev. E* **63**, 030901(R) (2001).
95. F. C. Santos, and J. M. Pacheco, *Phys. Rev. Lett.* **95**, 098104 (2005).
96. F. C. Santos, F. J. Rodrigues, and J. M. Pacheco, *Proc. Biol. Sci.* **273**, 51 (2006a).
97. H. Ohtsuki, E. L. C. Hauert, and M. A. Nowak, *Nature* **441**, 502 (2006).
98. F. C. Santos, and J. M. Pacheco, *J. Evol. Biol.* **19**, 726 (2006).

Evolutionary Games in Complex Topologies
Interplay Between Structure and Dynamics

Poncela Casasnovas, J.

2012, XIV, 158 p., Hardcover

ISBN: 978-3-642-30116-2