

Chapter 2

Survey on Community Detection, Validation and Applications

This chapter surveys the related works covering various objectives of the thesis. Generally, two kinds of communities are identified: 1) disjoint communities where nodes can have the membership of only one community and 2) overlapping communities where some nodes can have membership of more than one community. Techniques of identifying both disjoint and overlapping communities are discussed. Various metrics to evaluate both kinds of communities as well as evaluation methodologies are described. Generic datasets that are widely used to validate communities are discussed. Various applications involving post-hoc analysis of communities are also studied.

2.1 Community Detection Techniques

Community detection algorithms are primarily focused on finding disjoint communities (see [54, 189]). However, it is often pointed out that it is common to see overlapping

communities rather than disjoint communities in social networks [98, 162, 186]. Gregory [71] has discussed two kinds of overlap: crisp overlapping and fuzzy overlapping. In crisp overlapping, a node can belong to more than one community, but the node can have membership degree either 0 or 1. In fuzzy overlapping, membership degree of a node is in-between the range $[0, 1]$. The sum of membership degrees of a node to different communities is normalized to 1. Recent surveys have classified both disjoint and overlapping community detection algorithms based on methodological principles, properties incorporated, input and output, or community definition [36, 54, 152, 189, 214]. Fortunato [54] divides disjoint community detection algorithms based on methodological principles, while Coscia et. al [36] categorize overlapping community detection algorithms based on community definition. Xie et al. [214] categorized solely the overlapping community detection algorithms based on the methodological principles covering both crisp and fuzzy overlapping communities. Tang and Liu [189] have divided community detection algorithms based on the properties incorporated in the algorithm covering both kinds of communities. Another recent survey also classified both disjoint and overlapping community detection algorithms from the perspective of input and output [152]. This survey classified community detection algorithms into four categories as discussed below.

2.1.1 Hierarchical Approaches

Hierarchical approaches grow in two directions, divisive and agglomerative approaches. Divisive methods follow top-down approach and agglomerative methods follow bottom-up approach. Agglomerative methods are mostly engaged with node level properties and assimilate nodes into communities starting from ground level entity such as nodes expanding through communities to entire network. Since, agglomerative approaches follow natural process of grouping specially the real world social network, communities detected

are found to be more logical and accurate. Hierarchical approaches have another advantage of not requiring prior knowledge of the number and size of the communities [54].

Agglomerative Techniques

One of the popular agglomerative algorithm is Fast Unfolding (FastU) proposed by Blondel et al. [17], which is widely known as Louvain method. The FastU algorithm aggregates communities at different level to optimize modularity. Initially, individual nodes are considered as communities, aggregate these communities to get next level communities and so on until attains optimal modularity. Another popular agglomerative algorithm is SCAN [216], which also detects hubs and outliers with respect to communities. Unlike FastU algorithm, SCAN does not have clear levels of agglomeration. It merges nodes based on structural reachability and structural similarity. At the initial stage, the core nodes are identified and all those nodes that are structurally reachable from the core nodes are assigned to respective core nodes and form communities. The process continues until all the nodes are either classified or marked as non-member. The nodes those identified as non-members are further classified as hubs and outliers.

Leader based agglomerative approaches such as Top Leader [100], LICOD [97] and Leader-Follower (LeadF) [173] are developed in recent years. Leader based algorithms follow the concept that a community is constituted by a leader node followed by many other follower nodes. Prior to detect communities, these approaches first identifies potential leaders of respective communities based on node and community level properties. Remaining nodes are assigned to one or more leaders depending on requirement of non-overlapping or overlapping communities. Similar to leader based approaches, community expansion from pre-specified seeds is also introduced in recent years. Moradi et al. [135] developed a local seed selection algorithm for overlapping community detection. The algorithm first selects some initial seeds and nodes are assigned based on their locality in

reference to those seeds. Recently, Zhang et al. [222] have proposed another seed based agglomerative approach motivated by label propagation called Membership Degree Propagation (MDP) for fuzzy community detection. Some seed nodes are considered based on local centrality. The algorithm iteratively propagates membership degrees of all nodes and each seed grows community around it. New seeds are considered iteratively and some seeds are deleted to improve the partitioning.

Besides exploring nodes as in leader and seed based approaches, connections are also prioritized in some approaches such as random walk (RandW) [183], where nodes are assigned to the same community if those fall within the randomly performed short walk. A similarity matrix is prepared to keep track of frequency every pair of nodes that appear in these random paths and using that matrix communities are formed. Wang et al. [203] have proposed HC-PIN, where connections are grouped based on node level properties and associated nodes are assigned accordingly to the respective communities. Initially all the nodes are considered as to form different communities. Then Edge Clustering Value (ECV) is calculated for each connection and arranged in a decreasing order in a list. A higher value signifies a greater tendency to be included in the same community.

Divisive Techniques

The philosophy of divisive techniques begins with the proposal of Girvan and Newman [61]. They primarily focused on edge between centrality since inter-community connections are supposed to have a large value of the edge betweenness. The idea was to remove all connections from the network starting with highest edge betweenness centrality. Tyler et al. [195] proposed modification of Girvan-Newman algorithm to improve efficiency. The improved algorithm computes the contribution to edge betweenness following a sort of Monte Carlo estimation only from a limited number of centers that are randomly chosen. Rattigan et al. [159] proposed another fast version of Girvan-Newman

algorithm by approximating the edge between centrality. Chen and Yuan [30] pointed out a drawback of Girvan-Newman that counting all shortest paths to compute edge betweenness may lead to unbalanced community sizes and proposed to consider only non-redundant paths. Holme et al. [85] proposed another divisive algorithm, where nodes rather than edges are removed. They introduced a vertex centrality measure similar to edge betweenness and used that in the proposed algorithm. Girvan-Newman approach has been also modified by Gregory for overlapping community detection [69].

2.1.2 Optimization Techniques

The notion of optimization comes up with two pre-occupied questions. First question is about how to design a suitable objective function for obtaining good communities. Follow-up second question is about how to optimize the defined function. One approach for resolving the first question is through rigorous study of network and required communities. Another alternative, easy, and widely utilized approach is to use simply the quality metrics (discussed in section 2.2.2) defined for evaluating communities. Numerous techniques have been developed to optimize the community related objective functions, which are discussed below.

Modularity Maximization

Newman-Girvan defined modularity as a stopping criterion for their community detection algorithm [143]. Afterward modularity rapidly became an essential element of many community detection methods. Modularity is by far the most used and best known quality function for community detection. Newman [141] devised a greedy approach to maximize modularity. Later on, Clauset et al. [34] proposed more efficient data structure like max-heaps to make Newman's algorithm faster. However, Wakita and Tsurumi [202] noticed

that the fast algorithm by Clauset et al. is inefficient due to the bias towards large communities. Schuetz and Caflisch [171] proposed an efficient approach to avoid the formation of large communities. In order to favor small communities, Danon et al. [38] suggested to normalize the modularity variation produced by the merger of two communities by the fraction of edges incident to one of the two communities.

Lancichinetti and Fortunato [108] raised limitation to modularity maximization regarding biases: the tendency to merge small communities and to split large ones. They also pointed out that it is usually very difficult, and often impossible, to tune the resolution for avoiding both biases simultaneously. Arenas et al. [10] adopted multi-level resolution version of modularity to overcome resolution problem of modularity. Cafieri et al. [24] proposed a divisive approach to locally optimize modularity. Recently, Costa et al. [37] proposed another divisive approach to optimize modularity. Besides these, modularity maximization within the framework of mathematical programming is proposed by Agarwal and Kempe [2]. Chen et al. [32] used integer linear programming to transform the initial graph into an optimal target graph consisting of disjoint cliques, which effectively yields a partition. Yazdanparast and Havens [217] proposed positive programming approach to maximize modularity.

Heuristic Approaches

Community detection problem is a NP-Hard problem [54, 156]. Furthermore, Brandes et al. [20] shown modularity maximization is NP-Complete problem. Several heuristic approaches have been developed to deal with the complexity of community detection problem, particularly nature-inspired evolutionary computation techniques are popular (see [26] for the reviews). Tasgin et al. [190] proposed Genetic Algorithm (GA) based approach to optimize network modularity. Pizzuti [150] proposed GA based algorithm GA-net using newly defined function *community score*. Pizzuti [151] proposed another

multiobjective GA based approach that optimizes two objective functions. Gong et al. [67] proposed a multiobjective Evolutionary Algorithm (EA) based on decomposition. Shang et al. [176] proposed hybridized GA with simulated annealing to improve the stability and accuracy of community detection. Wen et al. [209] proposed a maximal clique based multiobjective EA for overlapping community detection. Li and Liu [115] proposed multi-agent based GA for community detection.

Besides GA, another nature inspired optimization technique called Particle Swarm Optimization (PSO) is widely used for community detection [6, 41, 181]. Xiaodong et al. [213] proposed a PSO based web community detection approach. Gong et al. [64] identifies communities by multiobjective discrete PSO. Cai et al. [25] used discrete PSO to identify communities in signed networks. Rees and Gallagher [160] explored overlapping communities using PSO with decentralized multi-threading processing and label propagation. Another swarm intelligent technique called Ant Colony Optimization (ACO) is also gaining attention in community detection [81, 86, 93, 94]. Extremal optimization [18] is also used to maximize modularity for community detection by Duch and Arenas [46]. Furthermore, memetic algorithm is also used for community detection. Gong et al. [66] used memetic algorithm for community detection. Later on, they have proposed an improved version by incorporating level propagation and elitism strategy [65]. In the same line of work, Gach and Hao [56] proposed another memetic algorithm based community detection approach. Recently, Ma et al. [123] developed a multi-level learning based memetic approach for community detection.

2.1.3 Spectral Clustering

Spectral clustering is one of the traditional clustering technique. The algorithms that are designed to cluster set nodes by using spectrum of matrix (better known as eigenvalues)

or matrices derived from eigenvalue are referred as spectral clustering techniques. Often, graph Laplacian matrices are considered as key tools for spectral clustering [200]. Normalized graph Laplacian matrices are used extensively in spectral clustering. Shi and Malik [179] developed two-way normalized spectral clustering algorithm. Minimization of normalized cut is the basis of their algorithm. Following same principle, Ding et al. [42] formulated min-max cut algorithm. Ng et al. [145] proposed k-way normalized spectral clustering algorithm. However, Nadler and Galun [137] discussed the limitations of these approaches that they cannot successfully cluster datasets that contain structures at different scales of size and density. Newman [142] proposed recursive two-way spectral clustering algorithm considering modularity maximization instead of various graph cuts. Verma and Meila [198] argued multiway spectral clustering algorithms are expected to perform better if easily traceable structures are present in the data. Meila et al. [130] explored spectral clustering in the view of random walks. Filippone et al. [52] studied the unification of kernel and spectral clustering. Recently, Joseph et al. [95] carried out wide range of study addressing the impact of regularization in spectral clustering.

2.1.4 Statistical Inference

Statistical inference in graph data aims fitting statistical models to the actual graph based on hypotheses on how nodes are connected to each other. Community detection problem is expressed as an inference or maximum likelihood problem then preferably Bayesian inference [212] or other models such as Potts model [161], blockmodeling [45], information theory [124] etc. are adopted.

Generative models

Bayesian inference has been used extensively in community detection, where the best fit is obtained through the maximization of a likelihood (known as generative models). Hastings [78] chose a planted partition model of network with communities. Reichardt and Bornholdt [161] proposed a fuzzy community detection algorithm based on Potts model. Hofman and Wiggins [84] proposed a generic Bayesian approach to identify communities by inferring modules of the network. Newman and Leicht [144] proposed a method based on mixture model and expectation-maximization technique. Ren et al. [163] developed another similar method based on the group fractions. Copic et al. [35] defined an axiomatization approach for community detection problem using maximum likelihood estimation. Psorakis et al. [153] proposed a method based on Bayesian non-negative matrix factorization for detecting overlapping communities.

Information theoretic approach

The modular structure of a graph can be treated as a compressed description of the graph. Rosvall and Bergstrom [165] utilized this concept to approximate the whole information contained in the adjacency matrix that represents graph. They envisioned a communication process in which community the graph represents a synthesis of the full structure that tries to infer the original graph topology. Sun et al. [188] also followed the same idea for bipartite graphs evolving in time. Rosvall and Bergstrom [166] utilized further the notion of describing a graph by using less information to optimally compress the information needed to describe the process of information diffusion across the graph. Chakrabarti [27] has applied the minimum description length principle to put the adjacency matrix of a graph into the (approximately) block diagonal form for a good compression of the graph

topology, and having very homogeneous blocks, for a compact description of their structure. Ziv et al. [228] proposed a principled information theoretic community detection approach on basis of module discovery in the network. Buhmann [23] proposed an information theoretic model validation approach for clustering.

2.1.5 Other Approaches

Some algorithms that do not fit into the above mentioned categories are discussed below. Zhang et al. [224] proposed an iterative process that reinforces the network topology and propinquity that is interpreted as the probability of a pair of nodes belonging to the same community. The propinquity between two vertices is defined as the sum of the number of direct links, number of common neighbors and the number of links within the common neighborhood. Gregory [70] proposed an overlapping community detection algorithm, where each node updates its belonging coefficient by averaging the coefficients from all its neighbors at each time step in a synchronous fashion. Xie et al. [215] developed another overlapping community detection algorithm based on general speaker-listener information propagation process. A game-theoretic framework was proposed by Chen et al. [31], in which a community is associated with a Nash local equilibrium. Long et al. [138] devised an interesting technique that is able to detect various kinds of groups, not necessarily communities. Traag and Bruggeman [193] studied communities in the networks, where connection weights are both negative and positive.

2.2 Validation Metrics

Evaluation of predicted community structure is another aspect of community detection. If ground truth community structure of a network is available, then simply predicted communities are compared with that of ground truth to measure accuracy. However, it is difficult to collect ground truth communities for most of the real-world networks. Therefore, community evaluation has to rely on the quality related measures that are designed based on connectivity pattern of communities. In this section, first various accuracy metrics are detailed, and then briefed some popular quality metrics.

2.2.1 Accuracy Metrics

Communities predicted by the algorithm are evaluated in terms of various accuracy metrics as follows. Let $C = \{C_1, C_2, \dots, C_k\}$ be the set of communities obtained with a community detection algorithm on a network of n nodes. Let $R = \{R_1, R_2, \dots, R_m\}$ be the real community structure. Here, C_k and R_m are interpreted as set of nodes in the respective communities. Overlapping of nodes in C and R can be summarized with a contingency table as presented in 2.1. Contingency table can also be presented with 2×2 matrix as shown in 2.2. Entries in the tables are decision pairs of two different nodes. If any pair of nodes are predicted to belong in the same community C_i and these two nodes are actually lie in same community R_j then entry for the pair of nodes will be True Positive (TP). If nodes pair is predicted to lie in same community, but actually the pair belong to different community, then the entry will be False Positive (FP). If nodes pair is predicted to lie in different community and the pair actually belong to different community, then the entry will be True Negative (TN). If nodes pair is predicted to lie in different community, but the pair actually belongs to the same community, then entry will be False Negative (FN).

Table 2.1: Contingency table I. Here, each n_{ij} denotes the number of nodes in common between communities C_i and R_j : $n_{ij} = |C_i \cap R_j|$.

$R \setminus C$	C_1	C_2	...	C_k	Sums
R_1	n_{11}	n_{12}	$\cdots$	n_{1k}	a_1
R_2	n_{21}	n_{22}	$\cdots$	n_{2k}	a_2
$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\vdots$
R_m	n_{m1}	n_{m2}	$\cdots$	n_{mk}	a_m
Sums	b_1	b_2	$\cdots$	b_k	$\sum_{ij} n_{ij} = n$

Adjusted Rand Index: Adjusted Rand Index (ARI) is corrected version of Rand Index [157] that measures the degree of overlapping between two partitions. As Rand Index suffers scaling problem [183] and its expected value between two random partitions does not take a constant value [199]. Hubert and Arabie [89] proposed a corrected version as follows:

$$AdjustedIndex = \frac{Index - ExpectedIndex}{MaxIndex - ExpectedIndex}. \quad (2.1)$$

More specifically, with overlapping entries of different communities of C and R in contingency table as shown in Table 2.1, ARI is computed as follows:

$$ARI(R, C) = \frac{\sum_{ij} \binom{n_{ij}}{2} - \frac{[\sum_i \binom{a_i}{2}] \sum_j \binom{b_j}{2}}{\binom{n}{2}}}{\frac{1}{2} \left[\sum_i \binom{a_i}{2} + \sum_j \binom{b_j}{2} \right] - \frac{[\sum_i \binom{a_i}{2}] \sum_j \binom{b_j}{2}}{\binom{n}{2}}} \quad (2.2)$$

where, n_{ij} is the number of nodes that are present both in communities C_i and R_j , a_i is the summation of all n_{ij} corresponding to any R_j of R and all C_i of C , and b_j is the summation of all n_{ij} corresponding to any C_i of C and all R_j of R . Rand Index can take a value between the range $[0, 1]$, while the ARI may take negative values if the index is less than the expected index [201]. Even though ARI takes negative values, generally considered range for ARI is $[0, 1]$. ARI value 0 indicates real and predicted communities does not agree on pairing, ARI value 1 indicates real and predicted communities both represents the same communities.

Table 2.2: Contingency table II.

Real \ Predicted	Same	Different	Total
Same	TP	FN	P
Different	FP	TN	N
Total	$\hat{P}$	$\hat{N}$	P+N

Normalized Mutual Information: Normalized Mutual Information (NMI) [8, 184] is an information theoretic approach to measure shared information between two data distribution. In information theory, the information contained in a distribution is called entropy [189]. In perspective of communities of a network with n nodes, the entropy $H(R)$ of real partitioning R and the entropy $H(C)$ of predicted partitioning C can be expressed as follows:

$$H(R) = - \sum_{i=1}^m \frac{n_i^R}{n} \log \left(\frac{n_i^R}{n} \right) \quad (2.3)$$

$$H(C) = - \sum_{j=1}^k \frac{n_j^C}{n} \log \left(\frac{n_j^C}{n} \right) \quad (2.4)$$

where, m and k are the number of communities present in R and C respectively, n_i^R represents the number of nodes in community $R_i \in R$ and n_j^C represents the number of nodes in community $C_j \in C$. Again mutual information share between R and C is expressed as:

$$I(R, C) = \sum_{i=1}^m \sum_{j=1}^k \frac{n_{ij}^{RC}}{n} \log \left(\frac{\frac{n_{ij}^{RC}}{n}}{\frac{n_i^R}{n} \times \frac{n_j^C}{n}} \right) \quad (2.5)$$

The definition of mutual information demonstrates that $I(R, C) \leq \frac{(H(R)+H(C))}{2}$ [8]. Thus, NMI between R and C is defined as:

$$NMI(R, C) = \frac{2 \times I(R, C)}{H(R) + H(C)}. \quad (2.6)$$

After simplifying Equation 2.6 by placing values of $H(R)$, $H(C)$ and $I(R, C)$ NMI is derived as follows:

$$NMI(R, C) = \frac{-2 \times \sum_{i=1}^m \sum_{j=1}^k n_{ij}^{RC} \log \left(\frac{n_{ij}^{RC} \times n}{n_i^R \times n_j^C} \right)}{\sum_{i=1}^m n_i^R \log \left(\frac{n_i^R}{n} \right) + \sum_{j=1}^k n_j^C \log \left(\frac{n_j^C}{n} \right)}. \quad (2.7)$$

NMI takes values in the range $[0, 1]$. Value 1 indicates maximum mutual information share between R and C or in other words both represents same set of communities. Value 0 indicates no information sharing between R and C or in other words both represents two different set of communities.

Purity: Purity is a very simple and clear measure, considers only correctly assigned nodes to any community. Each community is assigned a community label which is most frequent, and then count correctly assigned nodes with respect to ground truth communities. This number is then divided by the total number of nodes in the network. Purity of R and C is calculated as follows [126]:

$$Purity(R, C) = \frac{1}{n} \sum_m \max_k (R_m \cap C_k) \quad (2.8)$$

Actually purity is the average of total correctly assigned nodes to different communities. Purity takes values in the range $[0, 1]$. Higher purity value indicates high accuracy and lower value indicates inaccurate communities.

Precision: Precision is a measure of exactness of the predicted communities with respect to real communities. It evaluates the proportion of true positives against all the positive results. Precision can be expressed in terms of elements of contingency table

(Table 2.2) as follows:

$$Precision = \frac{TP}{TP + FP} = \frac{TP}{\hat{P}} \quad (2.9)$$

Recall: Recall is a measure of completeness of the predicted communities with respect to real communities. Recall also known as the true positive rate, or the recall rate, or sensitivity. It evaluates the proportion of true positives against actual true positives results. Recall can be expressed in terms of elements of contingency table (2.2) as follows:

$$Recall = \frac{TP}{TP + FN} = \frac{TP}{P} \quad (2.10)$$

F-measure: Harmonic mean of Precision and Recall is referred as F-measure, also known as F-score or F_1 -score.

$$F - measure = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (2.11)$$

This is the balanced version F-measure, where Precision and Recall given equal weight. Actually this is balanced version of general F_β -measure, where Precision given specific non negative weight as given below:

$$F_\beta - measure = (1 + \beta^2) \times \frac{Precision \times Recall}{\beta^2 \times Precision + Recall} \quad (2.12)$$

As F_1 -score takes its best value at 1 and worst score at 0, also more popular in evaluating communities than F_β -score, so we consider F_1 -score. When we use F-measure it will refer to F_1 -score only.

Entropy: Similar to NMI, entropy is another information theoretic concept which measures the information content of messages [103]. In community evaluation, entropy measures how different predicted communities share the distribution of nodes of real communities. Entropy of predicted communities C with respect to real communities R is defined as [225]:

$$H(R, C) = \sum_{i=1}^k \frac{n_i^C}{n} \left(-\frac{1}{\log m} \sum_{j=1}^m \frac{n_i^j}{n_i^C} \log \frac{n_i^j}{n_i^C} \right) \quad (2.13)$$

where, n is total number of nodes, k and m are number of communities present in predicted and real partitioning, n_i^C is the number of nodes in the community i of predicted communities, n_i^j is the number of nodes of real community j is assigned to the predicted community i . Lower value of entropy indicates better partition of predicted communities.

2.2.2 Quality Metrics

Modularity: Modularity (Q) [143] is the most widely used metric designed specially for the purpose of measuring quality of predicted communities. This metric has got wide acceptance over the years for evaluating communities when ground truth is not known. Modularity is computed as follows. Let us consider an algorithm predicted community structure C with k communities on a network. Define a $k \times k$ symmetric matrix e whose element e_{ij} is the fraction of all edges in the network that connect nodes in the community i to nodes in community j . The trace of this matrix $tr(e) = \sum_i e_{ii}$ gives the fraction of edges in the network that connect nodes in the same community. Whereas, the row (or column) sums $a_i = \sum_j e_{ij}$ represents the fraction of edges that connect to nodes in community i . With these $tr(e)$ and a_i modularity is defined as:

$$Q(C) = \sum_i (e_{ii} - a_i^2) = tr(e) - \|e^2\| \quad (2.14)$$

where, $\|x\|$ represents the sum of the elements of matrix x . Thus, Q effectively measures the fraction of edges in the network that connect nodes in the same community by subtracting the expected value of this quantity from it if the edges were placed at random. If the number of connections between nodes of same community is less than random, we will get $Q = 0$ and value approaches $Q = 1$ if connections between nodes of same community is higher.

Coverage: Coverage [21] of any community structure C is the fraction of edges that connect two nodes of same community within the total no of edges present in the network. Thus, Coverage of C is computed as follows:

$$Coverage(C) = \frac{\sum_{i=1}^k \sum_{j=1, i=j}^k e_{ij}}{m} \quad (2.15)$$

where, k is the number of communities present in C and m is total number of edges present in the network. Intuitively, higher value of coverage indicates better quality communities. If all communities predicted has only one node (i.e. $n = k$, n is number of nodes in the network) coverage will become 0. If number of communities is one, coverage takes value 1, since all edges of the network will lie within the community.

External Density: External density [169] of a network partitioning is defined as the ratio of edges that connect two different communities to the maximum number of edges possible that connects two different communities. It is the ratio of inter community edges to the maximum number of inter community edges possible. External density is computed as follows:

$$ExtD(C) = \frac{\{(u, v) | u \in C_i, v \in C_j, i \neq j\}}{n(n-1) - \sum_{i=1}^k (|C_i|(|C_i| - 1))} \quad (2.16)$$

where, u and v are any pair of nodes, C_i and C_j are communities, n is the number of nodes in the network, C is the predicted community structure that comprises k communities. Lesser value of external density indicates better quality communities.

Fuzzy Modularity: Membership of nodes with different belongingness to multiple communities are handled using fuzzy membership of nodes. To measure the quality of communities having fuzzy membership of nodes, fuzzy version of modularity has been developed [119, 223]. Havens et al. [79] proposed a generalized fuzzy modularity as follows:

$$Q_g = \frac{\text{tr}(UBU^T)}{\|W\|} \quad (2.17)$$

where, W is weighted adjacency matrix of the graph, U is partition matrix for membership of nodes to communities, $B = \left[W - \frac{(m^T m)}{\|W\|} \right]$, and $m = (m_1, m_2, \dots, m_n)^T$.

2.3 Datasets

Evaluation of communities requires benchmark networks. There are widely used real-world networks compiled by network scientists specifically for community detection. Besides these networks, synthetic networks are also used by imitating the properties of real-world networks. Some popular real-world networks and synthetic networks used for community evaluation are discussed below.

2.3.1 Real-world Networks

Karate: This data set is about study of a karate club network by Zachary [218]. The network consists of 34 members of a karate club as nodes and 78 connection among members representing friendships in the club which was observed over a period of two years. Due to a disagreement between members of the club's administrator and the club's instructor, the club was split up into two groups. Originally, Zachary constructed friendship network among members of the club with various measure of ties. Unweighted version of the network was used by Girvan and Newman [61] for evaluating communities.

Strike: Michael [132] studied labor strike patterns in a wood-processing facility after new management proposed compensation packages. The study was based on age and ethnic group. Set of 24 labors were grouped into three such groups depending on 34 association among labors during different schedules of strikes.

Football: United States college football network is a popular dataset used for community evaluation, which was used firstly by Girvan and Newman [61]. Network was represented with the schedule of Division I games for the 2000 season. Nodes in the network represent teams and connections represents regular-season games between two connected teams. The network contains 115 teams and 616 games were played among different teams. The teams are divided into conferences containing around 8-12 teams. However, generally considered 12 conferences for community evaluations.

Dolphin: Bottlenose dolphins network study [122] of dolphins living in Doubtful Sound, New Zealand is also used widely for evaluating communities. They had studied behavior of about 62 bottlenose dolphins for over seven years. The network was divided into two groups depending on the association patterns of dolphins.

Sawmill: Data were collected from a mid-sized softwood sawmill to examine the network of communications among employees [133]. This data was collected in order to analyze the communication structure among the employees after a strike. An edge in the network means that the two connected employees have discussed the strike with each other very often. The network of communications among employees consists 36 employees and 62 communicating links.

PollBooks: It is network of books about US politics sold by the online bookseller Amazon.com. Edges between books represent frequent co-purchasing of books by the same buyers. The network was compiled by V. Krebs [104]. The network contains 105 nodes and 441 connections.

Jazz: List of edges of the network of Jazz musicians. Gleiser and Danon [62] studied the collaboration network of jazz musicians. Two different levels of networks were prepared. First the collaboration network between individuals, where two musicians are connected if they have played in the same band. Second the collaboration between bands, where two bands are connected if they have a musician in common. The later network contains 198 bands and 2742 connections, and is used widely to evaluate community detection algorithms.

LesMis: Les Miserables is a French historical novel, written by Victor Hugo and published in 1862 [90]. The co-appearance of weighted network of characters in the novel has been extracted by Knuth [102]. The nodes are the characters of the novel and an edge indicates that the two characters appear together in the same chapter of the novel, at least once. There are 77 nodes representing different characters and 254 connections in the network.

Email: Gleiser et al. [73] studied a social network constructed from email communications within a medium sized university with employees. This is the email communication network at the University Rovira i Virgili in Tarragona in the south of Catalonia in Spain. Nodes are users and each connection represents that at least one email was sent. The direction of emails or the number of emails are not stored. The network was constituted with 1133 users and 5451 interchanges of emails as connections.

Words: Newman [142] compiled a network representing juxtapositions of words in a corpus of English text, in this case the novel David Copperfield by Charles Dickens. To construct this network, the 60 most commonly occurring nouns in the novel and the 60 most commonly occurring adjectives were considered. The nodes in the network represent words and an edge connects any two words that appear adjacent to one another at any point in the book. Eight of the words never appear adjacent to any of the others and are excluded from the network, leaving a total of 112 vertices and 425 connections.

GR-QC: Arxiv GR-QC (General Relativity and Quantum Cosmology) collaboration network [112] was prepared from the e-print arXiv and covers scientific collaborations between authors and papers submitted to General Relativity and Quantum Cosmology category. The data covers papers published within the period from January 1993 to April 2003 (124 months). Network was constructed with authors as nodes and co-authorship among authors were represented with connections.

HEP-TH: Arxiv HEP-TH (High Energy Physics - Theory) collaboration network [112] was prepared from the e-print arXiv and covers scientific collaborations between authors and papers submitted to High Energy Physics-Theory category. Network was constructed with authors as nodes and co-authorship among authors were represented with connections.

Wiki-Voter: Administrator selection in Wikipedia was done through voting among users and volunteers around the world. Using the Wikipedia page edit history in [111] extracted all administrator elections and vote history data. Nodes in the network represent Wikipedia users and voting of any user for other users are represented with a directed connection.

2.3.2 Synthetic Networks

LFR Graphs: Lancichinetti, Fortunato and Radicchi defined benchmark graphs for evaluating community detection algorithms, which are referred as LFR graphs [109]. These benchmark graphs are generated based on power law distribution of both node degree and community size. The graphs are generated according to various parameters such as number of nodes (n), average degree, maximum degree, exponent degree distribution and community size distribution, range of community sizes and mixing parameter μ for controlling the neighbors in other communities. Mixing parameter μ determines intra-community and inter-community connections. To evaluate performance of community detection algorithms variety of LFR graphs are generated. LFR graph generated with variation of mixing parameter, number of nodes and average node degree are popularly utilized in community evaluation. Source code of LFR graphs available online ¹.

2.4 Evaluation Methodologies

Popularly used community evaluation method is value based analysis, where predicted communities are evaluated in terms of various metrics discussed above (see section 2.2). Value based comparison of community detection algorithms is rather common strategy in

¹<https://sites.google.com/site/andrealancichinetti/>

order to determine performance of algorithms in the ground of quality and accuracy. Accumulation of indications of different metrics is a major difficulty in value based analysis. Another disadvantage is value based analysis cannot determine whether communities are good or bad when there is a trade-off between quality metrics and accuracy metrics. Kou et al. [103] developed MCDM technique to evaluate community detection algorithms, where a single comparative score is generated accumulating the values obtained for different metrics. Although with the methodology proposed by Kou et al. resolves the trade-off between quality and accuracy metrics, but it cannot express explicitly how likely the algorithm will identify accurate communities or good quality communities.

There are some other aspects of community evaluation, which are discussed in different literature. Visual inspection of predicted community structure in reference to ground truth is often considered [13, 117, 183, 216]. However, the visual inspection of community structure is suitable only for small networks. Sensitivity of community detection algorithm with variation of certain network parameters such as size of network (i.e. number of nodes), inter-community connections, intra-community connections are analyzed [113, 141]. Use of mixing parameter μ of LFR graph for variation of inter-community and intra-community connections is popular in community evaluation [107]. Statistical measures such as mean and standard deviation of different metrics are used when community detection algorithms are randomized i.e. the algorithm detects different community structures in different execution of algorithm [185, 186]. With this methodology, information about the distribution of metric values computed for various community structures can be obtained. However, it cannot express good (or bad) communities are how much good (or bad) in comparison to the communities identified by other algorithm.

2.5 Applications of Community Structure

In this section, various applications that incorporate predicted communities are studied and discussed how the community information are utilized in those systems.

2.5.1 Community Structure in Link Prediction

The information of community of nodes is leveraged in the task of link prediction. Ding et. al [43] have proposed a multi-resolution community division based link prediction approach. They considered the property of hierarchical organization of communities to extract different levels of communities. Then, a simple frequency statistical model is used to compute frequency of node pairs in different resolution of community, and likelihood of missing links are generated. Recently, Ding et. al [44] proposed another link prediction algorithm by using the latent information between different communities that introduces a community similarity feature called community relevance. Kuang et. al [105] also proposed a similarity based algorithm exploring herd phenomenon in different communities to predict missing links. Mimno et. al [134] proposed a generative model for text documents using the notion of community to identify missing links. Sachan and Ichise [167] proposed to build a link predictor in a co-authorship network, and showed that the knowledge of a pair of researchers lying in the same dense community can be used to improve the accuracy of predictor further. Shahriary et. al [174] have explored communities to predict links and sign of the links in signed networks.

2.5.2 Community Structure in Information Diffusion

Communities are central for efficiently disseminating news, rumors, and opinions in human social networks. Dissemination of information in the network in the context of community structure has been studied extensively. Belak et al. [16] studied the role of communities in target oriented information diffusion, and they showed exploring community information highly efficient strategy can be achieved. Epidemic spreading in weighted scale-free networks with community structure is investigated by Chu et al. [33]. The impact of overlapping community structure on susceptible-infected-susceptible (SIS) epidemic information diffusion model is investigated by Chen et al. [29]. Kimura et al. [101] used community analysis within the independent cascade model to find influential nodes for information diffusion on a social network. Nematzadeh et al. [138] proposed a linear threshold model for systematic analysis of community structure influence on global information diffusion. Weng et al. [210] developed predictive models of information diffusion based on interactions among communities. Shang et al. [175] considered both overlapping and non-overlapping communities to examine the affect in epidemic spreading.

2.5.3 Community Structure in Recommendation Systems

Community detection algorithms are also useful in the development of network based recommendation system. Kamahara et al. [96] explored user preference projected by the community in which the user belongs and introduced the notion of partially similar interest of users for recommendation system. Musto et al. [136] proposed a recommender system based on user community behavior to recommend a set of relevant keywords for the resources to be annotated. Zhuhadar et al. [227] proposed visual recommender system using community information for recommending resources to cyberlearners that belong to same community. Lisboa et al. [118] proposed another approach for community specific

recommendation system using Bayes modeling consideration changes in the behavior of users over time. A community based social recommender system is proposed by Fatemi and Tokarchuk [51], where social data is utilized for person specific recommendation based on the communities constructed from user interactions over the time.

2.6 Summary

Community detection problem is studied in three perspectives: identification communities, evaluation of detected communities, and post-hoc analysis leading to applications. First the study covers various community detection algorithms categorizing those based on their methodological principles. Second discussed aspects of community evaluation, which include validation metrics, datasets and evaluation methods. Lastly, various applications of community structure are discussed. Community detection problem has been studied several inter-disciplinary domains, yet the problem is not solved satisfactorily and leaves us with number of challenging issues.