

RESEARCH ARTICLE

Clustering Scientific Publications Based on Citation Relations: A Systematic Comparison of Different Methods

Lovro Šubelj^{1*}, Nees Jan van Eck², Ludo Waltman²

1 University of Ljubljana, Faculty of Computer and Information Science, Ljubljana, Slovenia, **2** Leiden University, Centre for Science and Technology Studies, Leiden, Netherlands

* lovro.subelj@fri.uni-lj.si



OPEN ACCESS

Citation: Šubelj L, van Eck NJ, Waltman L (2016) Clustering Scientific Publications Based on Citation Relations: A Systematic Comparison of Different Methods. PLoS ONE 11(4): e0154404. doi:10.1371/journal.pone.0154404

Editor: Lutz Bornmann, Max Planck Society, GERMANY

Received: December 30, 2015

Accepted: April 13, 2016

Published: April 28, 2016

Copyright: © 2016 Šubelj et al. This is an open access article distributed under the terms of the [Creative Commons Attribution License](https://creativecommons.org/licenses/by/4.0/), which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Data Availability Statement: The data have been obtained from Thomson Reuters' Web of Science database. Our license agreement with Thomson Reuters does not allow us to make the data freely available. Readers can contact Thomson Reuters to obtain the data (<http://thomsonreuters.com/thomson-reuters-web-of-science/>). All interested parties will be able to obtain the data in the same manner as we did. Most research organizations have a Web of Science license and therefore have access to the Web of Science database. Readers that do not have access to the Web of Science database can contact Thomson Reuters to obtain a license. Relevant

Abstract

Clustering methods are applied regularly in the bibliometric literature to identify research areas or scientific fields. These methods are for instance used to group publications into clusters based on their relations in a citation network. In the network science literature, many clustering methods, often referred to as graph partitioning or community detection techniques, have been developed. Focusing on the problem of clustering the publications in a citation network, we present a systematic comparison of the performance of a large number of these clustering methods. Using a number of different citation networks, some of them relatively small and others very large, we extensively study the statistical properties of the results provided by different methods. In addition, we also carry out an expert-based assessment of the results produced by different methods. The expert-based assessment focuses on publications in the field of scientometrics. Our findings seem to indicate that there is a trade-off between different properties that may be considered desirable for a good clustering of publications. Overall, map equation methods appear to perform best in our analysis, suggesting that these methods deserve more attention from the bibliometric community.

Introduction

There is an extensive literature on the topic of graph partitioning and community detection in networks [1]. This literature studies methods for partitioning the nodes in a network into a number of groups, often referred to as communities or clusters. The general idea is that nodes belonging to the same cluster should be relatively strongly connected to each other, while nodes belonging to different clusters should be only weakly connected.

Which methods for graph partitioning and community detection perform best in practice? The literature does not provide a clear answer to this question, and if the question can be answered at all, then most likely the answer will be dependent on the type of network that is being studied and on the type of partitioning that one is interested in.

information can be found at <http://thomsonreuters.com/en/products-services/scholarlyscientific-research/scholarly-search-and-discovery/web-of-science.html>.

Funding: This work has been supported in part by the Slovenian Research Agency Program No. P2-0359. The funders had no role in study design, data collection and analysis, decision to publish, or preparation of the manuscript.

Competing Interests: The authors have declared that no competing interests exist.

In this paper, we therefore address the above question in one specific context. We are interested in grouping scientific publications into clusters and we expect each cluster to represent a set of publications that are topically related to each other. Clustering scientific publications is a problem that has received a lot of attention in the bibliometric literature. In this literature, publications have for instance been clustered based on co-occurring words in titles, abstracts, or full text [2, 3], based on co-citation or bibliographic coupling relations [4–6], and sometimes even based on a combination of different types of relations [4, 7–9]. Following Waltman and Van Eck [10] and Boyack and Klavans [11, 12], our interest in this paper is in clustering publications based on direct citation relations. Direct citation relations are of special interest because they allow large sets of publications to be clustered in an efficient way. Waltman and Van Eck for instance cluster ten million publications from the period 2001–2010 based on about hundred million citation relations between these publications. In this way, they obtain a highly detailed classification system of scientific literature covering all fields of science.

The analysis presented in this paper focuses on systematically comparing the performance of a large number of clustering methods when applied to the problem of clustering scientific publications based on citation relations. The following clustering methods are included in the analysis: spectral methods [13, 14], modularity optimization [15–18], map equation methods [19, 20], matrix factorization [21], statistical methods [22], link clustering [23], label propagation [24–28], random walks [29], clique percolation [30] and expansion [31], and selected other methods [32, 33]. These are all methods that have been proposed during the past years in the literature on graph partitioning and community detection.

To evaluate the performance of the different clustering methods, we perform an in-depth analysis of the statistical properties of the clusterings obtained by each method. On the one hand we focus on general properties of the clusterings, but on the other hand we also consider a number of properties that are of special relevance in the context of citation networks of publications. However, to obtain a deep understanding of the differences between clustering methods, we believe that analyzing the statistical properties of clusterings is not sufficient. Understanding the differences between clustering methods also requires an expert-based assessment of different clusterings. This is a challenging task that involves a number of practical difficulties, but in this paper we nevertheless make an attempt to perform such an expert-based assessment. The expert-based assessment is performed for publications in the field of library and information science, focusing on the subfield of scientometrics.

This paper is organized as follows. We first discuss the data and methods included in our analysis. We then present the results of the analysis. We conclude the paper by providing a detailed discussion of our findings.

Methods

Below we first discuss the citation networks of publications that we consider in our analysis. We then discuss the clustering methods included in the analysis. Finally, we discuss the criteria that we use for comparing the clustering methods. These criteria relate to the following four properties of a clustering method:

Cluster sizes. Ideally the differences in the size of clusters should not be too large. For instance, the largest cluster preferably should be no more than an order of magnitude larger than the smallest cluster.

Small clusters. For practical purposes, it is usually inconvenient to have a large number of very small clusters. Therefore the number of very small clusters should be minimized as much as possible.

Clustering stability. Running the same clustering method multiple times may yield different results (due to random elements in many clustering methods), but the results should be reasonably similar. Likewise, when small changes are made to a citation network, this should not have too much effect on the results of a clustering method.

Computing time. Preferably, a clustering method should be fast. Especially in applications to large citation networks the issue of computing time is of significant importance.

In addition to the above four properties, a fifth property for comparing clustering methods is the intuitive sensibility of the results provided by a method. Experts should be able to interpret the clusters obtained from a clustering method in terms of meaningful research topics. We do not evaluate this fifth property using quantitative criteria. Instead, our expert-based assessment of the results of different clustering methods is focused on this criterion.

Citation networks of scientific publications. Citation relations between scientific publications are represented as a simple undirected and unweighted graph by first discarding the directions of citations, any multiple citations and citations from a publication to itself. Publications neither citing nor cited by any other are also discarded. Let n be the number of nodes N , $n = |N|$, and m the number of links in such citation network. Denote k to be the average node degree, i.e. the number of links incident to a node, $k = 2m/n$, and LCC the largest connected component, i.e. the largest subset of mutually reachable nodes.

We analyze four citation networks representing publications in the fields of Scientometrics, Library & Information Science and Physics, and also the entire science (see Table 1). Publications and their citations were collected from the Web of Science bibliographic database produced by Thomson Reuters. More specifically, we used the in-house version of the Web of Science database of the Centre for Science and Technology Studies of Leiden University. This version of the Web of Science database is very similar to the one available online at www.webofscience.com. However, there are some differences, notably in the identification of citations between publications [34]. Data collection was restricted to the Science Citation Index Expanded, the Social Sciences Citation Index and the Arts & Humanities Citation Index, while only publications of the Web of Science document types ‘article’ and ‘review’ were included in the data collection.

The field of Scientometrics was delineated by selecting all publications in the following three journals: *Journal of Informetrics*, *Journal of the Association for Information Science and Technology* (including its precursor *Journal of the American Society for Information Science and Technology*), and *Scientometrics*. The field of Library & Information Science was delineated by selecting all publications in the Web of Science journal subject category Information Science & Library Science. Finally, the field of Physics was delineated by selecting all publications in the eight Physics journal subject categories in Web of Science as well as the subject category Astronomy & Astrophysics.

Graph partitioning and community detection methods. For a thorough empirical comparison, we select a large number of representative graph partitioning and community

Table 1. Statistics of citation networks of scientific publications in Web of Science. We consider three scientific fields and the entire Web of Science. See text for the definitions of the statistics and the details of the data collection procedure.

Field	Period	# Publications	# Nodes n	# Links m	Degree k	% LCC
Scientometrics	2009–2013	2,402	1,998	5,496	5.50	94.0%
Library & Infor. Sci.	1996–2013	43,741	32,628	131,989	8.09	96.7%
Physics	2004–2013	1,314,458	1,233,542	9,838,008	15.95	98.5%
All Fields	2004–2013	11,780,132	11,063,916	122,148,955	22.08	99.3%

doi:10.1371/journal.pone.0154404.t001

Table 2. Graph partitioning and community detection methods. We consider a large number of methods divided into different classes. See text for the details of methods implementation and parameters setting.

Class	Method	Description	Ref.
Spectral analysis	Grclus	k -means clustering iteration	[14]
	METIS	multi-level k -way partitioning	[13]
Map equation [36]	Infomap	information flows compression	[19]
	Hiermap	hierarchical flows compression	[20]
Modularity [37]	Louvain	greedy hierarchical optimization	[16]
	Mouvain	multi-level hierarchical optimization	[17]
	SLM	smart local moving optimization	[18]
Label propagation	LPA	label propagation algorithm	[24]
	BPA	balanced propagation algorithm	[25]
	DPA	diffusion-propagation algorithm	[26]
	HPA	hierarchical propagation algorithm	[27]
	COPRA	community overlap propagation algorithm	[28]
Statistical methods	OSLOM	order statistics local optimization method	[22]
Link clustering	Links	link similarity hierarchical clustering	[23]
Graph models	BigClam	cluster affiliation matrix factorization	[21]
	CoDA	communities through directed affiliations	[33]
Ego-networks	DEMON	democratic estimate of modular organization	[32]
Random walks	Walktrap	random walks hierarchical clustering	[29]
Cliques	SCP	sequential clique percolation	[30]
	GCE	greedy clique expansion	[31]

doi:10.1371/journal.pone.0154404.t002

detection methods [1, 35], which we refer to as clustering methods in this paper. Table 2 lists selected methods roughly divided into different classes. Due to the number of methods considered, detailed description is omitted here.

We use the source code provided by the authors of all methods in all cases except Mouvain and LPA, where we use our own implementations [18, 25]. We adopt default parameter settings of each particular algorithm. Grclus, METIS, BigClam and CoDA demand the number of clusters to be specified apriori. Thus, Grclus(S) and Grclus(L) denote the same method with the number of clusters set to $n/15$ and $n/50$, respectively, while Grclus refers to Grclus(S) on networks with $n < 10^6$ and to Grclus(L) on larger networks (similarly for METIS, BigClam and CoDA). On the other hand, Links(S) and Links(L) denote the same method with Jaccard similarity threshold [23] set to 0.1 and 0.01, respectively, whereas Links always refers to Links(S). Finally, some of the methods return overlapping clusters. For reasons of simplicity, each node in multiple clusters is assigned to the first cluster that appears in the output of the particular algorithm.

Certain otherwise prominent algorithms like Infomap can not be applied to very large networks in a time comparable with the fastest algorithms like Louvain and BPA. A straightforward solution is to first adopt some other method $\mathcal{M}$ to cut the network into smaller subgraphs and then independently apply Infomap to each of these. Let C_i be some cluster of nodes in a network, $C_i \subseteq N$, and let s_i be its size, $s_i = |C_i|$. Next, let $\mathcal{C} = \{C_i\}$ be the clustering of all the nodes in a network returned by the method $\mathcal{M}$, $\bigcup_i C_i = N$ and $C_i \cap C_j = \emptyset$, $i \neq j$. Then, for each cluster C_i with $s_i > 50$, Infomap is applied to the subgraph induced by the nodes in C_i , whereas the clustering of C_i is accepted only when it improves the log-likelihood of $\mathcal{C}$ (see Eq (5)). Several such derived methods are considered. Gracmap and Metimap refer to methods that adopt spectral algorithms Grclus and METIS for the first method $\mathcal{M}$, respectively, where

the number of clusters is set to $n/10^4$ for networks with $n < 10^6$ and to $n/(5 \cdot 10^4)$ otherwise. For comparison, we also include Louvmap and Labmap that adopt modularity optimization known as Louvain algorithm and label propagation algorithm LPA in the first step, respectively. Finally, the setting of the number of clusters in Graclus is limited to 2500. Thus, for very large networks, we use Metilus that adopts METIS for $\mathcal{M}$ and Graclus afterwards. In total, we consider 30 methods. These are the 20 methods listed in Table 2, five variations with an alternative setting of the number of clusters and five derived methods as described above.

Let $\mathcal{C} = \{C_i\}$ be the clustering returned by some method $\mathcal{M}$. $\mathcal{C}$ often includes clusters C_i that are too small or too large to be of any practical use, $s_i < s_{\text{tiny}}$ or $s_i > s_{\text{giant}}$. A straightforward solution is a two-step post-processing approach that first tries to further partition each of the giant clusters as above and then merges the tiny clusters with larger ones. We set $s_{\text{tiny}} = 15$ and $s_{\text{giant}} = 10^4$. First, for each cluster C_i with $s_i > s_{\text{giant}}$, the same clustering method $\mathcal{M}$ is applied to the subgraph induced by the nodes in C_i and the resulting clustering is accepted based on the log-likelihood of $\mathcal{C}$ as before. Note that, due to the resolution limit of community detection methods [38, 39], most will further partition cluster C_i . Next, for each cluster C_i with $s_i < s_{\text{tiny}}$, C_i is merged with a neighboring cluster that most improves or least worsens the log-likelihood of $\mathcal{C}$. While the first post-processing step can be carried out simultaneously for each of the giant clusters, the tiny clusters in the second post-processing step have to be assessed in a random order.

Graph cuts and community structure statistics. Let $\mathcal{C}$ be some clustering of network nodes as described above and let A be the network adjacency matrix, $A_{ij} = A_{ji} \in \{0, 1\}$ and $A_{ii} = 0$. To measure the structure of clustering $\mathcal{C}$, we select different representative graph cuts and community structure statistics [40]. We measure the internal connectivity of clustering $\mathcal{C}$ as the average node internal degree K [41],

$$K(\mathcal{C}) = \frac{1}{n} \sum_{ij} A_{ij} \delta(c_i, c_j), \quad (1)$$

where c_i is the cluster of node i and δ is the Kronecker delta. The external connectivity of clustering $\mathcal{C}$ is measured as the average node external degree or expansion E [41],

$$E(\mathcal{C}) = \frac{1}{n} \sum_{ij} A_{ij} (1 - \delta(c_i, c_j)). \quad (2)$$

By definition, $k = K + E$, whereas K/k is the fraction of links covered by the clustering $\mathcal{C}$. Next, the Flake function F [42] considers internal and external connectivity of clustering $\mathcal{C}$ and is defined as the fraction of nodes with larger external than internal degree,

$$F(\mathcal{C}) = \frac{\left| \left\{ i : \sum_j A_{ij} \delta(c_i, c_j) < k_i/2 \right\} \right|}{n}, \quad (3)$$

where k_i is the degree of node i . For reference with previous work, we also report the value of modularity function Q [37, 43] that compares the internal connectivity of clustering $\mathcal{C}$ to the configuration model [44], i.e. a random graph with the same degree sequence,

$$Q(\mathcal{C}) = \frac{1}{2m} \sum_{ij} \left(A_{ij} - \frac{k_i k_j}{2m} \right) \delta(c_i, c_j). \quad (4)$$

Finally, we report the posterior probability of clustering $\mathcal{C}$ or the likelihood of $\mathcal{C}$ given the network observed [45]. Assume that links in a network formed solely based on nodes' cluster membership and let θ_i be a linking probability associated with cluster C_i . Then m_i links observed

between the nodes in cluster C_i would form with probability $\theta_i^{m_i}$ and the remaining $M_i - m_i$ possible links would not form with probability $(1 - \theta_i)^{M_i - m_i}$, $M_i = s_i(s_i - 1)/2$. Let $\tilde{\theta}$ be a linking probability representing the connectivity between the clusters. Then $\tilde{m}$ links observed between the nodes in different clusters would form with probability $\tilde{\theta}^{\tilde{m}}$, $\tilde{m} = m - \sum_i m_i$, and the remaining $\tilde{M} - \tilde{m}$ possible links would not form with probability $(1 - \tilde{\theta})^{\tilde{M} - \tilde{m}}$, $\tilde{M} = n(n - 1)/2 - \sum_i M_i$. Thus, the probability that the network formed according to $\mathcal{C}$ or the likelihood of $\mathcal{C}$ is defined as

$$L(\mathcal{C}) = \tilde{\theta}^{\tilde{m}} (1 - \tilde{\theta})^{\tilde{M} - \tilde{m}} \prod_i \theta_i^{m_i} (1 - \theta_i)^{M_i - m_i}, \quad (5)$$

where $\theta_i = m_i/M_i$ and $\tilde{\theta} = \tilde{m}/\tilde{M}$ are the maximum likelihood estimators [46]. For reasons of numerical stability, we report the log-likelihood of $\mathcal{C}$ as $\log L(\mathcal{C})$.

Denote C to be a random variable corresponding to clustering $\mathcal{C}$, $P(C = C_i) = s_i/n$. The distance between two clusterings $\mathcal{C}$ and $\mathcal{D}$ is measured using the variation of information V [47] defined as

$$V(\mathcal{C}, \mathcal{D}) = H(C|D) + H(D|C), \quad (6)$$

where $H(C|D)$ and $H(D|C)$ are conditional entropies. Since $V \in [0, \log n]$, we report the normalized variation of information $V/\log n$ [48].

Clustering robustness plots $R(\mathcal{M}, \alpha)$ [48] estimate the robustness of clustering $\mathcal{C}$ or the respective clustering method $\mathcal{M}$ under random perturbations of network links. R is defined as the distance between $\mathcal{C}$ and $\mathcal{C}_\alpha$,

$$R(\mathcal{M}, \alpha) = V(\mathcal{C}, \alpha) = V(\mathcal{C}, \mathcal{C}_\alpha), \quad (7)$$

where $\mathcal{C}_\alpha$ is obtained by $\mathcal{M}$ after randomly rewiring α links in the network.

Bibliometric clustering criteria. Let $\mathcal{C}$ be some clustering of network nodes as described above. To measure the utility of clustering $\mathcal{C}$, we select different bibliometric clustering criteria. We report the average cluster size S and the fraction of covered links K/k already introduced above. Next, we define the orders of magnitude covered by cluster sizes O as

$$O(\mathcal{C}) = \log_{10} \frac{s_L}{s_S}, \quad (8)$$

where s_L is the size of the largest cluster and s_S is the size of the smallest. Note that twice the value of s_S , which is negligible, has the same effect on O as twice the value of s_L , which is substantial. We thus report 5-percentile effective orders O_5 defined as

$$O_5(\mathcal{C}) = \log_{10} \frac{s_L}{s_5}, \quad (9)$$

where s_5 is the size of the smallest remaining cluster after removing the 5% smallest clusters. To measure the diameter of clusters in $\mathcal{C}$, we compute the 90-percentile effective cluster diameter D_{90} [49], i.e. the average number of hops to reach 90% of all the nodes within a cluster. The value of D_{90} is estimated from 1000 randomly selected seed nodes. Finally, the robustness of clustering $\mathcal{C}$ [48] or equivalently the uncertainty U of the respective clustering method $\mathcal{M}$ is defined as the distance between the clusterings $\mathcal{C}_1$ and $\mathcal{C}_2$ obtained by two consecutive

realizations of $\mathcal{M}$ (see Eq (6)),

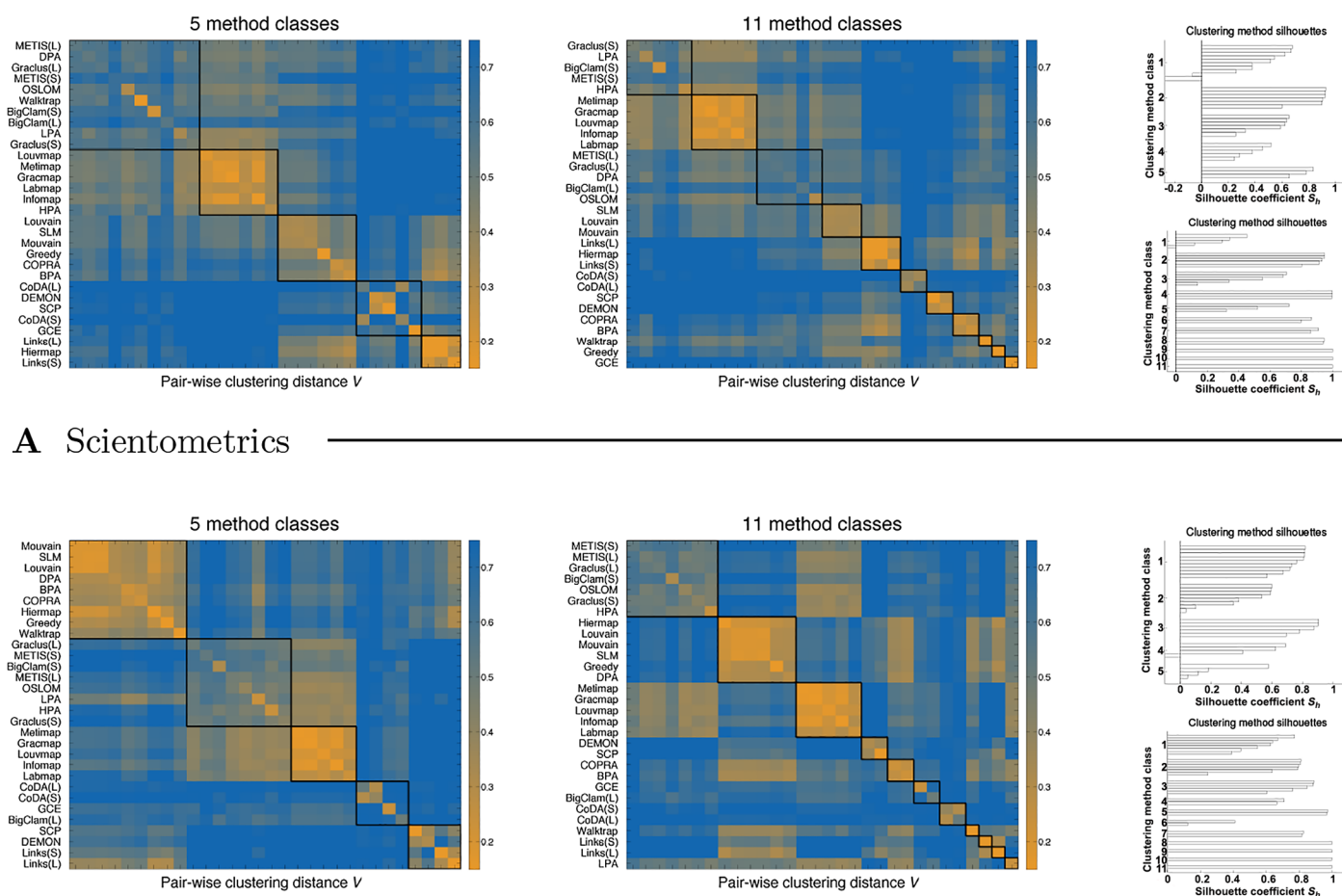
$$U(\mathcal{M}) = V(\mathcal{C}_1, \mathcal{C}_2). \quad (10)$$

All values, plots and diagrams reported in Results are averages over 100 realizations for Scientometrics, 10 realizations for Library & Information Science, two realizations for Physics and a single realization for All Fields.

Results

We start by directly comparing the clusterings obtained by all 30 clustering methods described in Methods to derive a manageable set of representatives. Next, we analyze structural and bibliometric statistics of the clusterings obtained by representative methods, and perform an expert-based assessment of the clusterings. Last, we analyze also the large-scale behavior of the most prominent methods.

Pair-wise clustering comparison. Fig 1 shows heatmaps of the pair-wise distances between the clusterings returned by the considered methods (see Eq (6)). The methods are applied to



B Library & Information Science

Fig 1. Pair-wise distances between the clusterings obtained by the considered methods. Panel A shows the heatmaps of clustering distances for the Scientometrics citation network, where the methods are clustered into 5 and 11 classes (left- and right-hand side, respectively). Note that this merely implies the ordering of the rows/columns. Insets on the right show the method silhouette coefficients. Panel B shows the same for the Library & Information Science citation network. See [Methods](#) for the definition of the clustering distance and text for the details of the method clustering procedure.

doi:10.1371/journal.pone.0154404.g001

two citation networks representing the fields of Scientometrics and Library & Information Science (see Table 1). To gain insight into different classes of methods, we apply the k -means data clustering algorithm [50] to the rows/columns of the heatmaps with the number of classes set to 5 and 11 (left- and right-hand side of Fig 1, respectively). The classes of methods are shown in the order of decreasing size and the methods within each class are listed in the order of decreasing silhouette coefficients S_h [51]. $S_h(\mathcal{M})$ of some method $\mathcal{M}$ is defined as a normalized difference between the lowest average inter-class dissimilarity and the average intra-class dissimilarity, for which we adopt the standard cosine similarity.

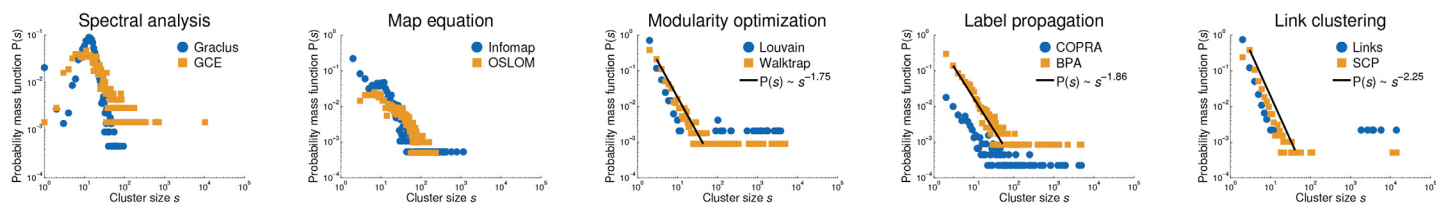
We observe compact classes of methods, most notably pronounced for the larger network (see right-hand side of Fig 1, panel B). Namely, the largest three classes represent spectral and statistical methods (e.g. Graclus, METIS and OSLOM), modularity optimization (e.g. Louvain and SLM) and map equation algorithms (e.g. Gracmap, Metimap and Infomap). Other smaller classes correspond to label propagation algorithms (e.g. LPA, BPA and COPRA), random walks (e.g. Walktrap), link clustering (i.e. Links), methods based on cliques (i.e. GCE and SCP) and other methods. Thus, despite the large number of methods considered, these can be divided into only a handful of truly different classes, but the differences between the classes can be rather substantial. In the following we limit the analysis to the 15 class representatives explicitly stated above, although the actual subset of methods considered depends on the size of the network analyzed.

Structural clustering analysis. Past literature often reported a power-law form $s^{-\gamma}$ of the cluster size distribution $P(s)$ [15, 52], to the extent that $s^{-\gamma}$ is also incorporated into the standard network benchmarks for testing clustering methods [53, 54]. Nevertheless, this may be merely an artifact of the power-law degree distribution $P(k) \sim k^{-\gamma}$ observed in real-world networks [55], while recent work on principled clustering methods sheds further doubts on the power-law form of $P(s)$ [56].

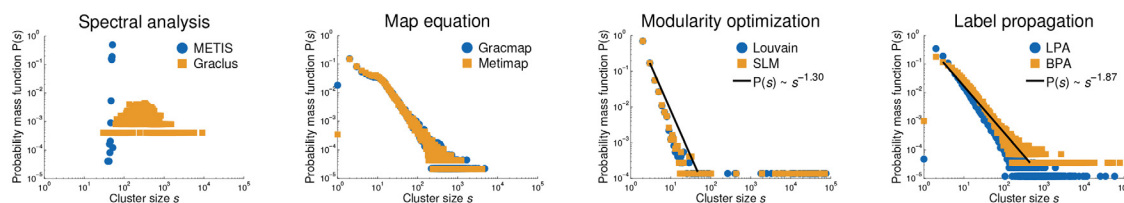
Fig 2 shows the distributions $P(s)$ of the clusterings returned by representative methods applied to the Library & Information Science and Physics citation networks (see Table 1). The methods are paired according to a similar shape of $P(s)$, where each pair is named by its most “famous” representative. Statistical methods are thus reported under map equation, while methods based on cliques appear under spectral analysis and link clustering. Notice that the validity of the power-law claim $P(s) \sim s^{-\gamma}$ clearly depends on the particular method considered. For instance, there is evidently a peak in the distributions of spectral methods with a lack of heavy tail (see left-hand side of Fig 2, panel A). Furthermore, in the case of map equation and statistical methods, the power-law form $s^{-\gamma}$ is violated for small and moderate s . On the other hand, the distributions for modularity optimization, label propagation and link clustering seem to follow the power-law scaling over several orders (see right-hand side of Fig 2, panel A) with the power-law exponent γ increasing from left to right. In the extreme case, link clustering produces a few very large clusters covering most of the nodes in the network, while the size distribution of the remaining ones follows a power-law. The observed differences between the clustering methods are even more striking on a larger network (see Fig 2, panel B).

Table 3 shows structural statistics of the clusterings obtained by representative methods applied to the Library & Information Science citation network. Most methods return a little less than 2000 clusters with some notable exceptions. Modularity optimization method Louvain, and also the methods based on dynamical processes (e.g. Walktrap and BPA), return a much smaller number clusters. On the other hand, link clustering and some other methods (e.g. COPRA) return a much larger number of clusters.

Table 3 further shows the average internal degree of the nodes in the clusters K and the average external degree or expansion E (see Eqs (1) and (2)). Although most methods achieve $K \gg E$, there are some important differences between the methods. The Flake function



A Library & Information Science



B Physics

Fig 2. Size distributions of the clusterings obtained by representative methods. Panels A and B show cluster size distributions $P(s)$ for the Library & Information Science and Physics citation networks, respectively. Wherever plausible, the power-laws $s^{-\gamma}$ are fitted to the tails of the distributions by maximum likelihood estimation, $\gamma = 1 + n(\sum_i \log s_i / s_{min})$ for $s_{min} > 1$.

doi:10.1371/journal.pone.0154404.g002

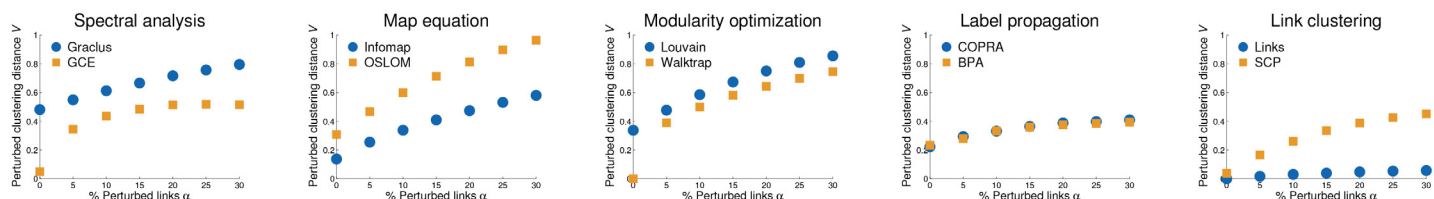
F measures the fraction of nodes with larger external than internal cluster degree (see Eq (3)). Notice that the values of F reflect the differences in the cluster size distributions $P(s)$ observed in Fig 2. Modularity optimization and other methods that return clusterings with a power-law distribution $P(s) \sim s^{-\gamma}$ can, due to a number of very large clusters, effectively cover many of the links in the network, giving low F (e.g. Louvain, Walktrap and BPA). On the contrary, spectral methods with a rather homogeneous distribution $P(s)$ must inevitably cut a large number of links between the clusters, thus giving very high F (e.g. Graculus). As in Fig 2, the middle ground between these two regimes is represented by map equation and statistical methods (e.g. Infomap and OSLOM).

Mainly for reference with previous work, Table 3 shows the values of modularity Q (see Eq (4)). Expectedly, the modularity optimization method Louvain gives the highest Q . Table 3 also

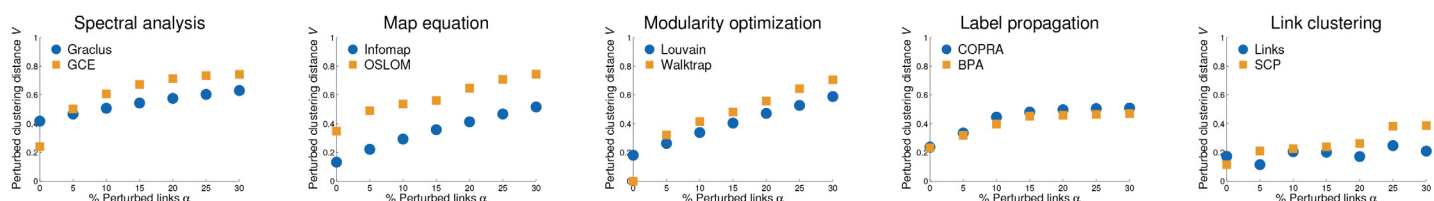
Table 3. Structural statistics of the clusterings obtained by representative methods. The methods are applied to the Library & Information Science citation network. See Methods for the definitions of the statistics and text for the interpretation.

Method	# Clusters	Degree K	Expansion E	Flake F	Modularity Q	Likelihood log L
Louvain	488.2	6.81	1.28	3.3%	0.734	-978498.8
GCE	682.0	4.06	4.03	28.9%	0.431	-997346.0
BPA	1001.9	7.00	1.09	3.0%	0.664	-975063.7
Walktrap	1127.0	6.47	1.62	7.0%	0.686	-968783.9
Infomap	1871.2	5.00	3.09	19.3%	0.602	-836963.9
OSLOM	1914.2	3.79	4.30	36.9%	0.453	-932170.7
SCP	1969.0	4.92	3.17	37.2%	0.217	-1103053.0
Graculus	2175.0	2.36	5.73	52.4%	0.290	-1003511.5
Links	2933.1	6.39	1.70	20.0%	0.093	-1173310.5
COPRA	3825.5	6.83	1.26	15.1%	0.645	-993909.5

doi:10.1371/journal.pone.0154404.t003



A Scientometrics



B Library & Information Science

Fig 3. Robustness of the clusterings obtained by representative methods. Panels A and B show clustering robustness plots $V(\alpha)$ for the Scientometrics and Library & Information Science citation networks, respectively. These show the distances between the clusterings obtained after randomly rewiring α links. See [Methods](#) for the definitions of clustering distance and robustness.

doi:10.1371/journal.pone.0154404.g003

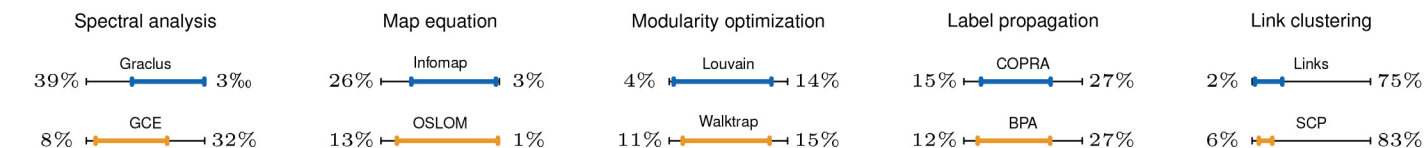
reports the log-likelihood $\log L$ of the clusterings given the network observed (see [Eq \(5\)](#)). The most likely clustering is obtained by Infomap, yet it should be stressed that the map equation is actually a likelihood criterion.

[Fig 3](#) shows the robustness plots $V(\alpha)$ of the clusterings returned by representative methods for the Scientometrics and Library & Information Science citation networks (see [Eq \(7\)](#)). The plots measure the distances between the clusterings obtained by the same method after randomly rewiring α links in the network. Although initially introduced as a measure of network community structure [\[48\]](#), we here adopt the same approach to measure the robustness of different clusterings.

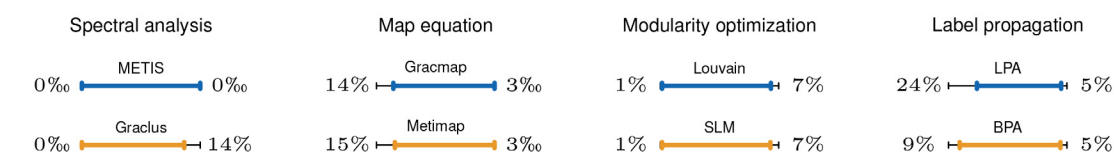
The methods in [Fig 3](#) are paired as in [Fig 2](#). Since many of them are nondeterministic, most of the plots do not start in the origin. The clusterings obtained by spectral and statistical methods (e.g. Graclus and OSLOM) prove to be the least robust with high values of V even for small α (see left-hand side of [Fig 3](#)). Map equation algorithm Infomap, and modularity optimization on the larger network (see middle of [Fig 3](#), panel B), seem to give stable clusterings with gradually increasing V over all α . Label propagation methods and link clustering appear very robust at first sight with surprisingly low V even for very large α (see right-hand side of [Fig 3](#)). For instance, the clustering returned by Links stays almost unchanged even after rewiring 30% of the links in the network. Nevertheless, this is a consequence of the existence of a few very large clusters that occupy the majority of the nodes in the network (see [Figs 2](#) and [4](#)) and change very little compared to the clusterings returned by other methods.

Bibliometric clustering analysis. The above structural analysis of the clusterings of citation networks would most likely be of interest to network scientists, but might provide limited value to the bibliometric community. In the following, we therefore analyze the clusterings also from an alternative perspective.

[Table 4](#) shows bibliometric statistics of the clusterings obtained by representative methods applied to the Library & Information Science citation network. The average cluster sizes S can



A Library & Information Science



B Physics

Fig 4. Degeneracy of the clusterings obtained by representative methods. Panels A and B show clustering degeneracy diagrams D for the Library & Information Science and Physics citation networks, respectively. These display the non-degenerate ranges of the clusterings, while the percentages show the fraction of nodes in tiny clusters $\sum_{s_i < s_{tiny}} s_i/n$ and in the largest cluster s_L/n (left- and right-hand side, respectively). See text for the definition of clustering degeneracy.

doi:10.1371/journal.pone.0154404.g004

be interpreted as the number of clusters in Table 3. For most methods, $S \approx 15$. Modularity optimization method Louvain gives almost five times larger clusters on average, while link clustering and some other methods (e.g. COPRA) return much smaller clusters with $S \approx 10$. Table 4 further shows 5-percentile effective orders O_5 that measure the orders of magnitude covered by cluster sizes s (see Eq (9)). For many practical applications, the clusters ideally should span no more than a single order of magnitude giving $O_5 \approx 1$. This turns out to be an illusive goal as $O_5 \gg 1$ for all methods except the spectral ones (e.g. Gracius), which one can observe also in Fig 2. Next, the 90-percentile effective diameter D_{90} measures the average number of hops to reach most of the nodes in a cluster (see Methods). Most methods return clusterings with small D_{90} consistent with the small-world network structure [57]. On the other hand, $D_{90} > 10$ for methods based on cliques (i.e. GCE and SCP) and link clustering, indicating the existence of some very large clusters, which is rather inconvenient in practice.

Table 4. Bibliometric statistics of the clusterings obtained by representative methods. The methods are applied to the Library & Information Science citation network. See Methods for the definitions of the statistics and text for the interpretation.

Method	Size S	Orders O_5	Diameter D_{90}	Coverage K/k	Uncertainty U	Complexity T
Louvain	66.7	3.33	9.13	84.5%	0.194	0.6 sec
GCE	47.8	3.32	11.99	50.1%	0.241	26.5 sec
BPA	32.0	3.61	7.28	86.2%	0.213	3.3 sec
Walktrap	29.0	3.39	7.80	79.9%	0.000	34.9 sec
Infomap	17.3	2.68	4.32	61.5%	0.133	9.6 sec
SCP	16.6	4.15	23.12	60.8%	0.021	1.4 sec
OSLOM	16.0	2.61	4.82	45.9%	0.364	94.9 sec
Gracius	15.0	1.13	3.38	29.2%	0.417	6.4 sec
Links	10.1	4.31	11.09	78.0%	0.048	10.0 sec
COPRA	8.8	3.97	6.91	84.9%	0.217	27.0 sec

doi:10.1371/journal.pone.0154404.t004

[Table 4](#) also shows the fractions of the links covered by different clusterings K/k (see [Methods](#)). Notice substantial diversity between the methods, which can again be interpreted in terms of different cluster size distributions $P(s)$ (see [Fig 2](#)). The methods that return clusterings with a power law $P(s) \sim s^{-\gamma}$, namely modularity optimization (e.g. Louvain), link clustering and methods based on dynamical processes (e.g. Walktrap, COPRA and BPA), can effectively cover over 80% of the links in the network. However, spectral and statistical methods (e.g. Graclus and OSLOM) that are characterized by a rather homogeneous $P(s)$ give K/k as low as 30%. The middle ground is again represented by the map equation algorithm Infomap with K/k around 60%.

The uncertainty U measures the stability of a method or equivalently the distance between the clusterings obtained by two consecutive realizations of the same method (see [Eq \(10\)](#)). Note that $U = V(0)$ in [Fig 3](#). [Table 4](#) shows the uncertainties of representative clustering methods. Spectral and statistical methods (e.g. Graclus and OSLOM) are substantially less stable than the rest with $U \approx 0.4$. Due to the existence of a few very large clusters already discussed above, link clustering and some other methods (i.e. Walktrap and SCP) appear very robust with $U \approx 0$. For the rest, $U \approx 0.2$.

The method complexity T in [Table 4](#) is measured as the execution time on a 2.3 GHz Intel Core i7 processor with a sufficient amount of memory. The fastest methods are those based on modularity optimization (i.e. Louvain), label propagation (e.g. BPA) and also spectral analysis (e.g. Graclus). Notice that the map equation algorithm Infomap takes only about ten seconds on the Library & Information Science citation network. Although this does not seem much, the network is relatively small. In fact, the algorithm takes almost three hours on the Physics citation network (results not shown) and would probably take several days to cluster the All Fields citation network (see [Table 1](#)).

[Fig 4](#) shows the degeneracy diagrams D of the clusterings returned by representative methods on the Library & Information Science and Physics citation networks. These display the non-degenerate or effective ranges of the clusterings that span the fraction of nodes not covered by tiny clusters with $s < s_{\text{tiny}}$, $s_{\text{tiny}} = 15$, or the largest or giant cluster. Hence, the degeneracy diagram D is defined as a range $(\sum_{s_i < s_{\text{tiny}}} s_i/n, 1 - s_L/n)$, where s_L is the size of the largest cluster. In the best-case scenario, the ranges in [Fig 4](#) would span from left to right. Any deviation from right or left signifies the existence of at least one very large cluster or many tiny clusters, respectively.

The methods in [Fig 4](#) are paired as in [Fig 2](#). The map equation algorithm Infomap and spectral and statistical methods (e.g. Graclus and OSLOM) return clusterings without a giant cluster spanning a large fraction of the nodes (see left-hand side of [Fig 4](#), panel A). However, these can include many tiny clusters. On the other hand, modularity optimization and label propagation methods (e.g. Louvain and BPA) return clusterings with at least one very large cluster (see right-hand side of [Fig 4](#), panel A). Even more, in the case of link clustering and some other methods (e.g. SCP), the giant cluster contains almost all the nodes in the network. Although the existence of a giant cluster and tiny clusters is not clearly visible in the case of a larger network (see [Fig 4](#), panel B), we stress that even a slight deviation from right or left is already substantial.

Expert-based clustering assessment. An expert-based assessment was performed on the clusterings obtained by representative methods on the Library & Information Science citation network. Within this network, the assessment focused on clusters covering topics or research areas in the field of scientometrics. Scientometrics can be seen as a subfield of the broader field of library and information science. The assessment was performed jointly by the second and the third author (NJvE and LW), who both have an extensive expertise in the field of scientometrics. A detailed investigation and comparison of the different clusterings was done with the

help of the CitNetExplorer software tool for visualizing and analyzing citation networks of publications [58].

We start by comparing the obtained clusterings based on the resolution they provide. A clustering consisting of a small number of clusters, with each cluster including a relatively large number of publications, has a low resolution. On the other hand, a clustering consisting of a large number of clusters, each including only a small number of publications, has a high resolution.

There are a number of clusterings for which we consider the resolution to be too high. This is the case for spectral methods Graclus(S), Graclus(L), METIS(S) and METIS(L). In these clusterings, topics that we would expect to be represented by a single cluster were instead represented by multiple clusters, each covering a subset of the publications dealing with a topic. For instance, the clustering returned by Graclus(L) includes four clusters that all cover part of the literature on the topic of the h-index, a very prominent topic in the field of scientometrics. Of these four clusters, there is one that clearly has its own focus. This cluster includes publications studying the mathematical properties of the h-index. Having a separate cluster for these publications is probably defensible. However, the other three clusters all seem to cover very similar publications, and therefore we see no justification for the fact that these publications are distributed over three clusters rather than all being assigned to the same cluster.

Other clusterings have a resolution that is too low for a meaningful analysis of the scientometric literature. The clusterings for which this is the case are obtained by BPA and Walktrap. One of the clusters created by BPA for instance consists of 3,808 publications and essentially covers the entire scientometric literature. This cluster seems to properly delineate the scientometric literature from the rest of the library and information science literature. Hence, if one's purpose is to identify subfields within the field of library and information science, then BPA may provide good results. However, in our case, we are interested in identifying topics rather than entire subfields, and for this purpose the results provided by BPA are not helpful.

The clusterings with a resolution that matches reasonably well with the idea of identifying topics within the subfield of scientometrics are obtained by the statistical method OSLOM and the map equation algorithms Infomap and Metimap. In addition to the clustering methods presented in Methods, we here consider also a variant of the Louvain modularity optimization method with a resolution parameter [59] that one can tune to customize the clustering resolution [18]. Setting the resolution parameter to 10 gives the most suitable resolution here, which we denote Louvain(10). We next analyze OSLOM, Infomap, Metimap and Louvain(10) in more detail.

The clustering obtained by OSLOM has a relatively high resolution. It includes only three clusters with more than 100 scientometric publications, which means that most scientometric publications are assigned to small clusters. As a consequence, some topics that we would expect to be represented by a single cluster are in fact distributed over multiple clusters. Important examples are the topic of webometrics and the topic of patents. These topics are each distributed over two clusters of approximately equal size, which we consider an unsatisfactory result. A more general problem of OSLOM is that we observe a relatively large number of publications that are assigned to a cluster where they do not seem to belong. For instance, there is a cluster covering the topic of the analysis and visualization of bibliometric networks, but this cluster includes a significant number of publications dealing with other topics, such as the topic of indicators for citation analysis.

Louvain(10) clustering is characterized by a somewhat unusual cluster size distribution. Compared with other clusterings, it includes a relatively large number of clusters with more than 100 publications and a relatively small number of clusters with a number of publications between 10 and 100. As a consequence, there are a number of larger scientometric clusters for

which there is no similar cluster in other clusterings, for instance obtained by Metimap or Infomap. A detailed examination of these clusters indicates that they do not cover easily recognizable topics. Publications included in these clusters usually do have something in common. For instance, there are clusters in which many publications relate to a specific country or a specific geographical region, such as China or Africa. However, our overall impression is that the clusters are of a somewhat heterogeneous nature and that it would have been better if the publications in the clusters had been distributed over a number of smaller clusters. The presence of these heterogeneous clusters is a significant weakness of Louvain(10).

The clusterings that we are most satisfied with are obtained by Metimap and Infomap. In Table 5, we present for each of these clusterings a list of all scientometric clusters with at least 50 publications. For each cluster, we report the number of publications included in the cluster or equivalently the cluster size s and we provide an indication of the topic that is represented by the cluster. Fig 5 compares the Metimap and Infomap clusterings by showing the overlap of scientometric clusters using an alluvial diagram.

Metimap and Infomap both offer a reasonable perspective on the main topics in the field of scientometrics. As can be seen in Table 5, the clustering returned by Metimap has a somewhat higher resolution than that of Infomap and consequently some topics that are covered by a

Table 5. Statistics of the clusterings obtained by the map equation methods Metimap and Infomap.

The methods are applied to the Library & Information Science citation network and the largest scientometric clusters with $s \geq 50$ are shown. See Fig 5 for a comparison of the clusterings and text for the interpretation.

Method	Topic	Size s
Metimap	Citation analysis: h-index	262
	Webometrics	256
	Collaboration	224
	Bibliometric networks (1) + Interdisciplinarity	163
	Patents + Nanotechnology	137
	Bibliographic databases	115
	Citation analysis: Advanced indicators	107
	Social sciences and humanities	95
	Citation analysis: Journal impact factor	87
	Bibliometric networks (2)	69
	Citation analysis: Foundations	59
	Citation distributions and citation dynamics	56
	Peer review	56
Infomap	Citation analysis: h-index + Bibliographic databases	358
	Collaboration	308
	Bibliometric networks	254
	Webometrics	250
	Citation analysis: Advanced indicators & Journal impact factor	220
	Patents + Nanotechnology	216
	Social sciences and humanities	104
	Country-specific case studies	87
	Citation analysis: Foundations	85
	Peer review	67
	Gender differences	59
	Interdisciplinarity	59
	University rankings	57
	Citation distributions and citation dynamics	56

doi:10.1371/journal.pone.0154404.t005

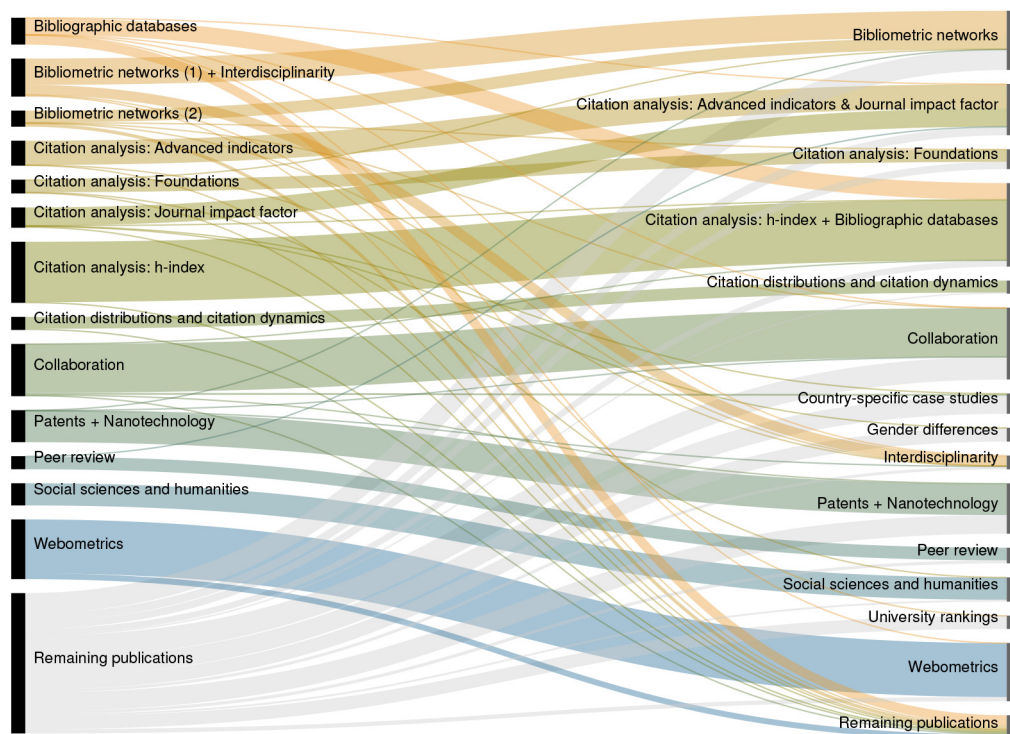


Fig 5. Alluvial diagram of the clusterings obtained by the map equation methods Metimap and Infomap. The diagram shows the overlap between the largest scientometric clusters returned by Metimap and Infomap on the Library & Information Science citation network (left and right, respectively). 'Remaining publications' are included in one of the clusters in the Metimap (Infomap) clustering but not included in any of the clusters in the Infomap (Metimap) clustering. See [Table 5](#) for details of the clusterings.

doi:10.1371/journal.pone.0154404.g005

single cluster in the case of Infomap are distributed over multiple clusters in the case of Metimap. We have a slight preference for Infomap over Metimap because the way in which topics are distributed over multiple clusters in the case of Metimap does not always seem fully satisfactory to us. For instance, we prefer to have a single cluster covering the topic of bibliometric networks instead of the two clusters that are provided by Metimap. However, we emphasize that the differences between the two clusterings are small and that we have only a weak preference for Infomap. Furthermore, even though Metimap and Infomap gave the best clusterings obtained in our study, it should be mentioned that these clusterings sometimes suffer from questionable assignments of publications to clusters. This is a problem especially for smaller clusters. In the case of clusters with fewer than 100 publications, we often observe that a significant share of the publications assigned to a cluster (e.g. about 25% of the publications) are only weakly related to the main topic of the cluster.

In the case of the clusterings obtained by Metimap and Infomap, we also investigated the effect of applying our post-processing approach (see [Methods](#)). Due to the relatively small size of the Library & Information Science citation network, the effect of the post-processing approach on the main clusters obtained in the Metimap and Infomap clusterings is small. The number of publications that are reassigned from small clusters to larger clusters, i.e. clusters with at least 50 publications, is very limited. Given the small effect of the post-processing approach, no significant influence on the quality of the clusters could be observed.

Large-scale clustering analysis. In the following, we analyze the large-scale behavior of different clustering methods. We limit the analysis to the Louvain modularity optimization method, the map equation algorithm Metimap, the label propagation algorithm BPA and the

Table 6. Bibliometric statistics of the clusterings obtained by selected methods. The methods are applied to Physics citation network and bibliometric statistics of the clusterings with and without post-processing are shown. See [Methods](#) for the definitions of statistics and the details of clustering post-processing approach.

Method	Size S	Orders O_5	Diameter D_{90}	Coverage K/k	Uncertainty U	Complexity T
Louvain	169.5	4.62	9.88	88.3%	0.172	89.8 sec
Metilus	50.0	2.29	4.53	37.5%	0.330	140.7 sec
BPA	43.5	4.58	5.36	76.7%	0.212	276.0 sec
Metimap	26.5	3.28	3.68	58.8%	0.122	459.5 sec
Louvain+post.	147.5	3.70	6.92	73.1%	0.238	134.9 sec
Metilus+post.	51.3	2.23	4.69	37.4%	0.331	144.7 sec
BPA+post.	72.6	4.56	5.39	74.9%	0.217	340.8 sec
Metimap+post.	44.1	3.29	4.28	59.0%	0.148	500.3 sec

doi:10.1371/journal.pone.0154404.t006

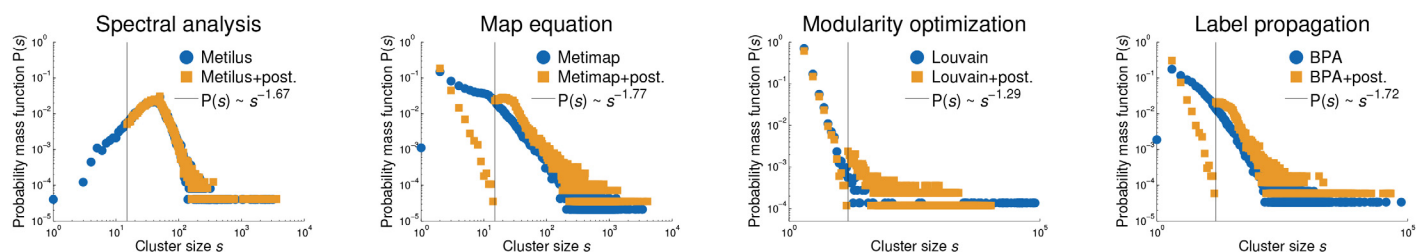
spectral analysis approach Metilus. These were selected since they can cluster the All Fields citation network in about an hour.

[Table 6](#) shows bibliometric statistics of the clusterings obtained by the selected methods applied to the Physics citation network (see [Table 1](#)). Compared to the clusterings obtained for the Library & Information Science network in [Table 4](#), one can observe a notable increase in the average cluster size S and the effective orders of magnitude O_5 . The clusterings thus include at least some much larger clusters. Yet, the effective diameter D_{90} and the clustering coverage K/k remain comparable. The clusterings returned by modularity optimization and label propagation methods (i.e. Louvain and BPA) again cover around 80% of the links, while the spectral method Metimap gives K/k below 40%. Finally, despite a substantial increase in the network size, the method uncertainty U stays about the same, while the complexity T obviously increases.

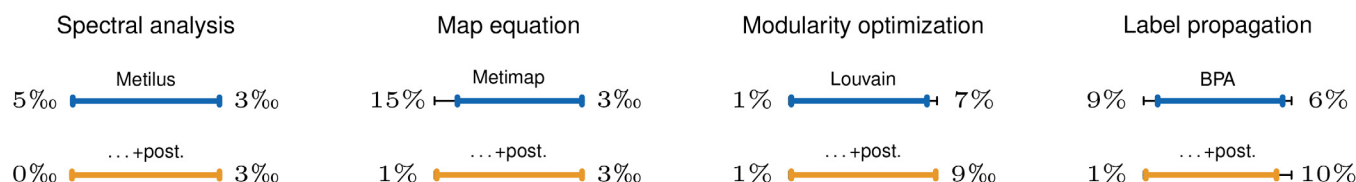
[Table 6](#) also shows the effect of the clustering post-processing approach presented in Methods that first tries to further partition the largest clusters with $s > s_{\text{giant}}$ and then merges the tiny clusters with larger ones for $s < s_{\text{tiny}}$, $s_{\text{tiny}} = 15$ and $s_{\text{giant}} = 10^4$. In the case of the map equation, label propagation and spectral methods (i.e. Metimap, Metilus and BPA), the post-processing approach has no apparent effect on the largest clusters. Due to the merging of tiny clusters, the average cluster size S increases, while all the remaining statistics remain roughly the same (see [Table 6](#)). On the other hand, the post-processing manages to further partition the largest clusters returned by the modularity optimization method Louvain. This decreases the cluster size S , and also the effective orders O_5 and the effective diameter D_{90} . However, the clustering coverage K/k decreases as well, while the method uncertainty U increases (see [Table 6](#)).

[Fig 6](#) shows the impact of the post-processing approach on the cluster size distributions $P(s)$ and the clustering degeneracy diagrams D . All distributions $P(s)$ remain conceptually the same, with the difference that most tiny clusters have been merged with larger ones (see [Fig 6](#), panel A). Notice that a small number of tiny clusters with $s < 15$ remain, which correspond to disconnected components that could obviously not be merged with other clusters (see [Table 1](#) for the size of LCC). Still, the degeneracy diagrams D show that post-processing effectively removes tiny clusters, and also the giant cluster in the case of the modularity optimization method Louvain, but fails to further partition the giant cluster in the case of the label propagation algorithm BPA (see right-hand side of [Fig 6](#), panel B).

Last, we apply the selected methods to the All Fields citation network (see [Table 1](#)). [Table 7](#) shows different statistics of the obtained clusterings. Compared to those obtained for the Physics citation network in [Table 6](#), we can again observe an increase in the average cluster size S



A Size distributions



B Degeneracy diagrams

Fig 6. Size distributions and degeneracy of the clusterings obtained by the selected methods. The methods with and without post-processing are applied to the Physics citation network, while the panels **A** and **B** show cluster size distributions $P(s)$ and clustering degeneracy diagrams D , respectively. Vertical lines in panel **A** represent the threshold size $s_{tiny} = 15$. See text for the definition of clustering degeneracy and Methods for the details of the clustering post-processing approach.

doi:10.1371/journal.pone.0154404.g006

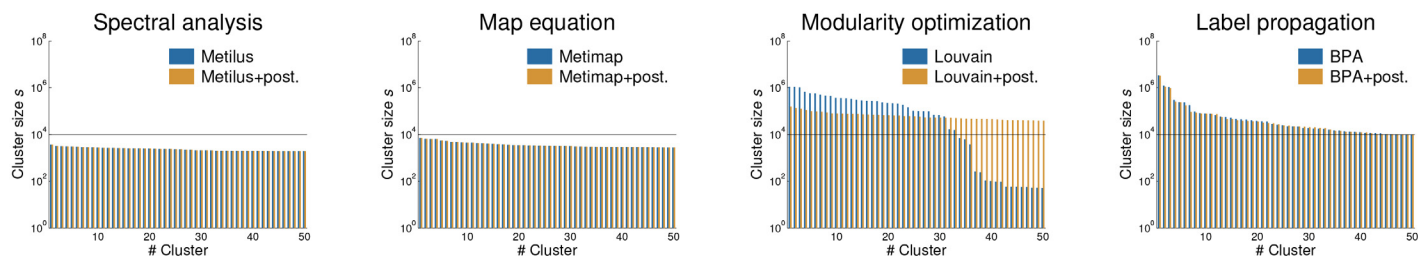
and the effective orders O_5 . Thus the size of the largest clusters further increases. Yet, as before, the clustering coverage K/k of different methods remains roughly the same, while the differences between the methods can also clearly be observed in the average internal degree K . [Table 7](#) also shows the statistics of the clusterings after the post-processing approach, which has exactly the same effect on the clusterings as in [Table 6](#). Notice also that the post-processing does not substantially increase the running time of the methods.

To better understand the nature of different clusterings and the effects of the post-processing approach, [Fig 7](#) shows the sizes s and coverage K/k of the largest 50 clusters returned by the selected methods (see [Methods](#)). The coverage K/k of an individual cluster is defined as the

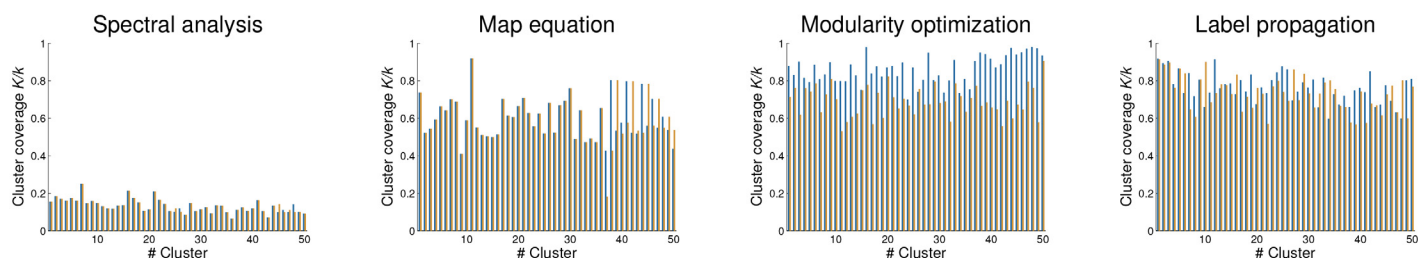
Table 7. Statistics of the clusterings obtained by the selected methods. The methods are applied to the All Fields citation network and different statistics of the clusterings with and without post-processing are shown. See [Methods](#) for the definitions of the statistics and the details of the clustering post-processing approach.

Method	Size S	Orders O_5	Degree K	Coverage K/k	Flake F	Complexity T
Louvain	334.4	5.74	18.53	83.9%	5.3%	52.1 min
BPA	105.4	6.22	18.50	83.8%	7.2%	66.2 min
Metilus	50.0	2.33	5.91	26.8%	68.9%	30.0 min
Metimap	33.2	3.55	10.30	46.6%	45.0%	94.2 min
Louvain+post.	320.9	4.88	15.20	68.8%	17.1%	78.9 min
BPA+post.	167.1	6.20	18.04	81.7%	9.0%	114.3 min
Metilus+post.	51.5	2.24	5.92	26.8%	68.9%	34.3 min
Metimap+post.	58.9	3.55	10.33	46.8%	44.5%	98.9 min

doi:10.1371/journal.pone.0154404.t007



A Cluster sizes



B Cluster coverage

Fig 7. Sizes and coverage of the largest clusters obtained by the selected methods. The methods with and without post-processing are applied to the All Fields citation network, while the panels **A** and **B** show the sizes s and coverage K/k of the largest 50 clusters, respectively. Horizontal lines in panel **A** represent the threshold size $s_{giant} = 10^4$. See text for the definition of cluster coverage.

doi:10.1371/journal.pone.0154404.g007

average internal degree of the nodes in the cluster divided by the total degree of these nodes. As already lengthily discussed above, the spectral analysis approach Metilus returns clusters with very low $K/k \approx 15\%$ (see left-hand side of Fig 7, panel **B**), while the modularity optimization and label propagation methods (i.e. Louvain and BPA) give clusters with very high $K/k \approx 80\%$ (see right-hand side of Fig 7, panel **B**). For the map equation algorithm Metimap, $K/k \approx 60\%$. One can also observe that, in the case of the label propagation algorithm BPA, the post-processing approach fails to further partition the largest clusters with $s > s_{giant}$, where s_{giant} is represented by horizontal lines in Fig 7, panel **A**. On the contrary, the post-processing does partition the largest clusters in the case of the modularity optimization method Louvain. However, the results are far from satisfactory. Each cluster with $s > s_{giant}$ is indeed split into smaller clusters, but the number of such clusters thus actually increases (see middle of Fig 7, panel **A**).

Discussion

Which methods for graph partitioning and community detection perform best for the purpose of grouping scientific publications into clusters? In this paper, we have carried out an extensive analysis comparing the performance of a large number of methods. The methods have been applied to a number of networks of publications connected by direct citation relations. We have studied the statistical properties of the results provided by the different methods, and we have also performed an expert-based assessment of the results.

From a bibliometric point of view, a good clustering of publications ideally should have a number of properties. First of all, although it is natural to expect that there will be larger and

smaller clusters, it is inconvenient for practical purposes if there are very large differences in the size of clusters. As a rule of thumb, we ideally would like the difference in size between the largest and the smallest clusters to be no more than an order of magnitude. Second, if it turns out to be inevitable that some publications end up in very small clusters, for instance because these publications have almost no citation relations with other publications, then at least we would prefer the number of publications assigned to these insignificant clusters to be as limited as possible. Third, we would like the results of a clustering method to be reasonably stable. Many methods include a random element, in which case different runs of a method may yield different results. However, running the same method multiple times should not affect the results too much, and the results should also be reasonably robust to small changes in a citation network of publications. Fourth, the computing time of a clustering method should not be excessive. This is especially important when one aims to apply a method to networks consisting of large numbers of publications and citation relations. Finally, and perhaps most importantly, the results produced by a clustering method should make intuitive sense. Experts should be able to recognize the scientific topics represented by clusters of publications.

Our analysis shows that most clustering methods yield results with large differences in the size of clusters. The larger clusters are typically several orders of magnitude larger than the smaller clusters. Sometimes more than half of the publications in a citation network are all assigned to the same cluster. This was for instance observed for the results obtained from the Links and SCP methods in the Library & Information Science citation network. The only methods that yield clusters of more or less similar size are the spectral methods (e.g. Graclus). These methods produce results that are characterized by a much more uniform cluster size distribution. Depending on the cluster size distribution and also on the resolution of a clustering, there can be large differences in the share of all citation relations that are covered by clusters. Coverage for instance ranges from less than 30% to more than 85% in the Library & Information Science citation network. Clustering methods also often assign a significant share of the publications in a citation network to very small clusters. In the Library & Information Science citation network, the Graclus and Infomap methods for instance assign more than 25% of the publications to clusters consisting of fewer than 15 publications. The stability or robustness of the results obtained from a clustering method also partly depends on the size of the clusters produced by the method. Not surprisingly, methods that produce one or more very large clusters tend to yield relatively robust results. Furthermore, in the Library & Information Science citation network, spectral and statistical methods (e.g. Graclus and OSLOM) produce results with a relatively low robustness, while Infomap and modularity optimization yield quite robust results.

In terms of computing time, there are substantial differences between the various methods. For instance, clustering the publications in the Library & Information Science citation network takes more than 100 times longer for the slowest method than for the fastest method. Modularity optimization methods (e.g. Louvain), label propagation (e.g. BPA), and spectral analysis methods (e.g. Graclus) perform best in terms of computing time. Other methods require a more significant amount of computing time, making them less suitable for applications on large citation networks.

Turning now to the expert-based assessment of the results produced by different clustering methods for the scientometrics subfield within the Library & Information Science citation network, we find that the Infomap and Metimap (i.e. Infomap combined with spectral method METIS) methods give the most satisfactory results, with a slight preference for the Infomap results over the results obtained from Metimap. Other methods, such as OSLOM and Louvain, provide less satisfactory results.

Our analysis seems to provide most support for the use of Infomap and related methods such as Metimap to cluster the publications in a citation network. Infomap has the best

performance in our expert-based assessment, and it yields quite robust results. Compared with some of the other methods, Infomap has a relatively high computing time, but this can be overcome by using Metimap in larger citation networks. The price that we pay for the good performance of Infomap seems to be the assignment of a relatively large number of publications to small clusters. Paying this price seems necessary to obtain high-quality clustering results. In large citation networks, a post-processing procedure can be applied to minimize the number of small clusters, but the effect of the use of such a procedure on the quality of the clustering results is not clear.

The promising results obtained for Infomap are in line with earlier findings reported in the network science literature [60]. Although Infomap has been introduced in the bibliometric literature [61] and has been applied to citation networks in a number of studies [19, 20, 62, 63], the method has not yet gained a widespread popularity in the bibliometric community, where researchers seem to prefer the use of modularity-based methods. Our findings suggest that the bibliometric community could benefit from exploring the use of other clustering methods in addition to modularity-based methods. Infomap seems to be of particular interest. Future studies should reveal whether Infomap indeed consistently performs well in applications to citation networks.

Limitations of the analysis. It is important to emphasize that our results should be interpreted cautiously because of a number of limitations of our analysis. One obvious limitation is that, despite the large number of clustering methods included in our analysis, we did not exhaustively cover all methods proposed in the literature. The selection of the methods included in our analysis was made based on the popularity of a method and to some degree also on our familiarity with a method. In addition, the availability of source code played a role as well. Many methods discussed in the literature are not included in our analysis. In particular, methods that produce overlapping clusters [64, 65] or clusters at multiple levels of resolution [66, 67] are not covered. Also, we for instance do not cover some recently developed principled methods based on statistical inference [56].

A second limitation is that each clustering method was applied using the default parameter settings. We did not try to optimize the parameter values of the different methods. So the performance of some methods may have been better if we had used optimized parameter values for these methods. Some methods for instance have a parameter that can be used to fine-tune the level of granularity of the clustering results. One could use such a parameter to try to obtain results at similar levels of granularity for different methods, and in that way a more accurate comparison between different methods may be possible. We did not explore this possibility in our analysis, but we do consider this an interesting direction for future research. We note that the clustering method proposed by two of us in an earlier paper [10] requires a careful choice of parameter values. For this reason, this method was not included in our present analysis.

A third limitation is our exclusive focus on undirected and unweighted networks of direct citation relations between publications. We did not consider the possibility of taking into account the direction of a citation relation, and we did not test the effect of assigning weights to citation relations [10]. We also did not study the use of indirect citation relations between publications, in particular co-citation and bibliographic coupling relations.

Finally, we should emphasize the limitations of our expert-based assessment of the clustering results obtained for the scientometrics subfield within the Library & Information Science citation network. The expert-based assessment was carried out at a high level of detail by two experts with an extensive expertise in the field of scientometrics. Nevertheless, any expert-based assessment will necessarily be of a subjective nature, and different experts therefore may not always reach the same conclusions. Moreover, experts typically have a deep understanding of the literature only in a relatively small area of science. This for instance explains why in our

expert-based assessment we could not cover the entire field of library and information science but only the subfield of scientometrics. Unfortunately, it is difficult to say to what extent conclusions reached for such a relatively small area of science can be expected to generalize to other areas. For this reason, the findings of our expert-based assessment should be interpreted with some caution.

Acknowledgments

We thank numerous authors for kindly providing the source code of their methods. This work has been supported in part by the Slovenian Research Agency Program No. P2-0359.

Author Contributions

Conceived and designed the experiments: LŠ NJvE LW. Performed the experiments: LŠ. Analyzed the data: LŠ NJvE LW. Contributed reagents/materials/analysis tools: NJvE. Wrote the paper: LŠ LW.

References

1. Fortunato S. Community detection in graphs. *Phys Rep*. 2010; 486(3–5):75–174. doi: [10.1016/j.physrep.2009.11.002](https://doi.org/10.1016/j.physrep.2009.11.002)
2. Boyack KW, Newman D, Duhon RJ, Klavans R, Patek M, Biberstine JR, et al. Clustering more than two million biomedical publications: Comparing the accuracies of nine text-based similarity approaches. *PLoS ONE*. 2011; 6(3):e18029. doi: [10.1371/journal.pone.0018029](https://doi.org/10.1371/journal.pone.0018029) PMID: [21437291](https://pubmed.ncbi.nlm.nih.gov/21437291/)
3. Janssens F, Leta J, Glänzel W, De Moor B. Towards mapping library and information science. *Inform Process Manag*. 2006; 42(6):1614–1642. doi: [10.1016/j.ipm.2006.03.025](https://doi.org/10.1016/j.ipm.2006.03.025)
4. Boyack KW, Klavans R. Co-citation analysis, bibliographic coupling, and direct citation: Which citation approach represents the research front most accurately? *J Am Soc Inf Sci Tec*. 2010; 61(12):2389–2404. doi: [10.1002/asi.21419](https://doi.org/10.1002/asi.21419)
5. Jarneving B. Bibliographic coupling and its application to research-front and other core documents. *J Infometr*. 2007; 1(4):287–307. doi: [10.1016/j.joi.2007.07.004](https://doi.org/10.1016/j.joi.2007.07.004)
6. Small H, Griffith BC. The structure of scientific literatures I: Identifying and graphing specialties. *Sci Stud*. 1974; 4(1):17–40. doi: [10.1177/030631277400400102](https://doi.org/10.1177/030631277400400102)
7. Janssens F, Glänzel W, De Moor B. A hybrid mapping of information science. *Scientometrics*. 2008; 75(3):607–631. doi: [10.1007/s11192-007-2002-7](https://doi.org/10.1007/s11192-007-2002-7)
8. Small H. Update on science mapping: Creating large document spaces. *Scientometrics*. 1997; 38(2):275–293. doi: [10.1007/BF02457414](https://doi.org/10.1007/BF02457414)
9. Waltman L, Van Eck NJ, Noyons ECM. A unified approach to mapping and clustering of bibliometric networks. *J Infometr*. 2010; 4(4):629–635. doi: [10.1016/j.joi.2010.07.002](https://doi.org/10.1016/j.joi.2010.07.002)
10. Waltman L, Van Eck NJ. A new methodology for constructing a publication-level classification system of science. *J Am Soc Inf Sci Tec*. 2012; 63(12):2378–2392. doi: [10.1002/asi.22748](https://doi.org/10.1002/asi.22748)
11. Boyack KW, Klavans R. Including cited non-source items in a large-scale map of science: What difference does it make? *J Infometr*. 2014; 8(3):569–580. doi: [10.1016/j.joi.2014.04.001](https://doi.org/10.1016/j.joi.2014.04.001)
12. Klavans R, Boyack KW. Which type of citation analysis generates the most accurate taxonomy of scientific and technical knowledge? e-print arXiv:151105078v2. 2016;p. 1–26.
13. Karypis G, Kumar V. A fast and high quality multilevel scheme for partitioning irregular graphs. *SIAM J Sci Comput*. 1998; 20(1):359–392. doi: [10.1137/S1064827595287997](https://doi.org/10.1137/S1064827595287997)
14. Dhillon IS, Guan Y, Kulis B. Weighted graph cuts without eigenvectors: A multilevel approach. *IEEE T Pattern Anal*. 2007; 29(11):1944–1957. doi: [10.1109/TPAMI.2007.1115](https://doi.org/10.1109/TPAMI.2007.1115)
15. Clauset A, Newman MEJ, Moore C. Finding community structure in very large networks. *Phys Rev E*. 2004; 70(6):066111. doi: [10.1103/PhysRevE.70.066111](https://doi.org/10.1103/PhysRevE.70.066111)
16. Blondel VD, Guillaume JL, Lambiotte R, Lefebvre E. Fast unfolding of communities in large networks. *J Stat Mech*. 2008;P10008. doi: [10.1088/1742-5468/2008/10/P10008](https://doi.org/10.1088/1742-5468/2008/10/P10008)
17. Rotta R, Noack A. Multilevel local search algorithms for modularity clustering. *ACM J Exp Algorithmics*. 2011; 16:2.3.

18. Waltman L, Van Eck NJ. A smart local moving algorithm for large-scale modularity-based community detection. *Eur Phys J B*. 2013; 86(11):471. doi: [10.1140/epjb/e2013-40829-0](https://doi.org/10.1140/epjb/e2013-40829-0)
19. Rosvall M, Bergstrom CT. Maps of random walks on complex networks reveal community structure. *P Natl Acad Sci USA*. 2008; 105(4):1118–1123. doi: [10.1073/pnas.0706851105](https://doi.org/10.1073/pnas.0706851105)
20. Rosvall M, Bergstrom CT. Multilevel compression of random walks on networks reveals hierarchical organization in large integrated systems. *PLoS ONE*. 2011; 6(4):e18209. doi: [10.1371/journal.pone.0018209](https://doi.org/10.1371/journal.pone.0018209) PMID: [21494658](https://pubmed.ncbi.nlm.nih.gov/21494658/)
21. Yang J, Leskovec J. Overlapping community detection at scale: A nonnegative matrix factorization approach. In: *Proceedings of the ACM International Conference on Web Search and Data Mining*. Rome, Italy; 2013. p. 587–596.
22. Lancichinetti A, Radicchi F, Ramasco JJ, Fortunato S. Finding statistically significant communities in networks. *PLoS ONE*. 2011; 6(4):e18961. doi: [10.1371/journal.pone.0018961](https://doi.org/10.1371/journal.pone.0018961) PMID: [21559480](https://pubmed.ncbi.nlm.nih.gov/21559480/)
23. Ahn YY, Bagrow JP, Lehmann S. Link communities reveal multiscale complexity in networks. *Nature*. 2010; 466(7307):761–764. doi: [10.1038/nature09182](https://doi.org/10.1038/nature09182) PMID: [20562860](https://pubmed.ncbi.nlm.nih.gov/20562860/)
24. Raghavan UN, Albert R, Kumara S. Near linear time algorithm to detect community structures in large-scale networks. *Phys Rev E*. 2007; 76(3):036106. doi: [10.1103/PhysRevE.76.036106](https://doi.org/10.1103/PhysRevE.76.036106)
25. Šubelj L, Bajec M. Robust network community detection using balanced propagation. *Eur Phys J B*. 2011; 81(3):353–362. doi: [10.1140/epjb/e2011-10979-2](https://doi.org/10.1140/epjb/e2011-10979-2)
26. Šubelj L, Bajec M. Unfolding communities in large complex networks: Combining defensive and offensive label propagation for core extraction. *Phys Rev E*. 2011; 83(3):036103. doi: [10.1103/PhysRevE.83.036103](https://doi.org/10.1103/PhysRevE.83.036103)
27. Šubelj L, Bajec M. Group detection in complex networks: An algorithm and comparison of the state of the art. *Physica A*. 2014; 397:144–156. doi: [10.1016/j.physa.2013.12.003](https://doi.org/10.1016/j.physa.2013.12.003)
28. Gregory S. Finding overlapping communities in networks by label propagation. *New J Phys*. 2010; 12(10):103018. doi: [10.1088/1367-2630/12/10/103018](https://doi.org/10.1088/1367-2630/12/10/103018)
29. Pons P, Latapy M. Computing communities in large networks using random walks. In: *Proceedings of the International Symposium on Computer and Information Sciences*. Istanbul, Turkey; 2005. p. 284–293.
30. Kumpula JM, Kivelä M, Kaski K, Saramäki J. Sequential algorithm for fast clique percolation. *Phys Rev E*. 2008; 78(2):026109. doi: [10.1103/PhysRevE.78.026109](https://doi.org/10.1103/PhysRevE.78.026109)
31. Lee C, Reid F, McDaid A, Hurley N. Detecting highly overlapping community structure by greedy clique expansion. In: *Proceedings of the ACM SIGKDD Workshop on Social Network Mining and Analysis*. Washington, DC, USA; 2010. p. 33–42.
32. Coscia M, Rossetti G, Giannotti F, Pedreschi D. DEMON: A local-first discovery method for overlapping communities. In: *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. Beijing, China; 2012. p. 615–623.
33. Yang J, McAuley J, Leskovec J. Detecting cohesive and 2-mode communities in directed and undirected networks. In: *Proceedings of the ACM International Conference on Web Search and Data Mining*. New York, NY, USA; 2014. p. 323–332.
34. Olensky M, Schmidt M, Van Eck NJ. Evaluation of the citation matching algorithms of CWTS and iFQ in comparison to Web of Science. *J Assoc Inf Sci Tec*. 2015; In press. doi: [10.1002/asi.23590](https://doi.org/10.1002/asi.23590)
35. Hric D, Darst RK, Fortunato S. Community detection in networks: Structural communities versus ground truth. *Phys Rev E*. 2014; 90(6):062805. doi: [10.1103/PhysRevE.90.062805](https://doi.org/10.1103/PhysRevE.90.062805)
36. Rosvall M, Bergstrom CT. An information-theoretic framework for resolving community structure in complex networks. *P Natl Acad Sci USA*. 2007; 104(18):7327–7331. doi: [10.1073/pnas.0611034104](https://doi.org/10.1073/pnas.0611034104)
37. Newman MEJ, Girvan M. Finding and evaluating community structure in networks. *Phys Rev E*. 2004; 69(2):026113. doi: [10.1103/PhysRevE.69.026113](https://doi.org/10.1103/PhysRevE.69.026113)
38. Fortunato S, Barthelemy M. Resolution limit in community detection. *P Natl Acad Sci USA*. 2007; 104(1):36–41. doi: [10.1073/pnas.0605965104](https://doi.org/10.1073/pnas.0605965104)
39. Traag VA, Van Dooren P, Nesterov Y. Narrow scope for resolution-limit-free community detection. *Phys Rev E*. 2011; 84(1):016114. doi: [10.1103/PhysRevE.84.016114](https://doi.org/10.1103/PhysRevE.84.016114)
40. Yang J, Leskovec J. Defining and evaluating network communities based on ground-truth. In: *Proceedings of the ACM SIGKDD Workshop on Mining Data Semantics*. Beijing, China; 2012. p. 1–10.
41. Radicchi F, Castellano C, Cecconi F, Loreto V, Parisi D. Defining and identifying communities in networks. *P Natl Acad Sci USA*. 2004; 101(9):2658–2663. doi: [10.1073/pnas.0400054101](https://doi.org/10.1073/pnas.0400054101)
42. Flake GW, Lawrence S, Giles CL. Efficient identification of web communities. In: *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. Boston, MA, USA; 2000. p. 150–160.

43. Newman MEJ. Modularity and community structure in networks. *P Natl Acad Sci USA*. 2006; 103(23):8577–8582. doi: [10.1073/pnas.0601602103](https://doi.org/10.1073/pnas.0601602103)
44. Newman MEJ, Strogatz SH, Watts DJ. Random graphs with arbitrary degree distributions and their applications. *Phys Rev E*. 2001; 64(2):026118. doi: [10.1103/PhysRevE.64.026118](https://doi.org/10.1103/PhysRevE.64.026118)
45. Peixoto TP. The entropy of stochastic blockmodel ensembles. *Phys Rev E*. 2012; 85(5):056122. doi: [10.1103/PhysRevE.85.056122](https://doi.org/10.1103/PhysRevE.85.056122)
46. Casella G, Berger RL. *Statistical Inference*. Belmont: Duxbury Press; 1990.
47. Meila M. Comparing clusterings: An information based distance. *J Multivariate Anal*. 2007; 98(5):873–895. doi: [10.1016/j.jmva.2006.11.013](https://doi.org/10.1016/j.jmva.2006.11.013)
48. Karrer B, Levina E, Newman MEJ. Robustness of community structure in networks. *Phys Rev E*. 2008; 77(4):046119. doi: [10.1103/PhysRevE.77.046119](https://doi.org/10.1103/PhysRevE.77.046119)
49. Leskovec J, Kleinberg J, Faloutsos C. Graph evolution: Densification and shrinking diameters. *ACM Trans Knowl Discov Data*. 2007; 1(1):1–41. doi: [10.1145/1217299.1217301](https://doi.org/10.1145/1217299.1217301)
50. MacQueen JB. Some methods for classification and analysis of multivariate observations. In: *Proceedings of Berkeley Symposium on Mathematical Statistics and Probability*. Berkeley, CA, USA; 1967. p. 281–297.
51. Rousseeuw PJ. Silhouettes: A graphical aid to the interpretation and validation of cluster analysis. *J Comput Appl Math*. 1987; 20:53–65. doi: [10.1016/0377-0427\(87\)90125-7](https://doi.org/10.1016/0377-0427(87)90125-7)
52. Newman MEJ. Fast algorithm for detecting community structure in networks. *Phys Rev E*. 2004; 69(6):066133. doi: [10.1103/PhysRevE.69.066133](https://doi.org/10.1103/PhysRevE.69.066133)
53. Lancichinetti A, Fortunato S, Radicchi F. Benchmark graphs for testing community detection algorithms. *Phys Rev E*. 2008; 78(4):046110. doi: [10.1103/PhysRevE.78.046110](https://doi.org/10.1103/PhysRevE.78.046110)
54. Lancichinetti A, Fortunato S. Benchmarks for testing community detection algorithms on directed and weighted graphs with overlapping communities. *Phys Rev E*. 2009; 80(1):016118. doi: [10.1103/PhysRevE.80.016118](https://doi.org/10.1103/PhysRevE.80.016118)
55. Karrer B, Newman MEJ. Stochastic blockmodels and community structure in networks. *Phys Rev E*. 2011; 83(1):016107. doi: [10.1103/PhysRevE.83.016107](https://doi.org/10.1103/PhysRevE.83.016107)
56. Peixoto TP. Model selection and hypothesis testing for large-scale network models with overlapping groups. *Phys Rev X*. 2015; 5(1):011033.
57. Watts DJ, Strogatz SH. Collective dynamics of 'small-world' networks. *Nature*. 1998; 393(6684):440–442. doi: [10.1038/30918](https://doi.org/10.1038/30918) PMID: [9623998](https://pubmed.ncbi.nlm.nih.gov/9623998/)
58. Van Eck NJ, Waltman L. CitNetExplorer: A new software tool for analyzing and visualizing citation networks. *J Infometr*. 2014; 8(4):802–823. doi: [10.1016/j.joi.2014.07.006](https://doi.org/10.1016/j.joi.2014.07.006)
59. Reichardt J, Bornholdt S. Statistical mechanics of community detection. *Phys Rev E*. 2006; 74(1):016110. doi: [10.1103/PhysRevE.74.016110](https://doi.org/10.1103/PhysRevE.74.016110)
60. Lancichinetti A, Fortunato S. Community detection algorithms: A comparative analysis. *Phys Rev E*. 2009; 80(5):056117. doi: [10.1103/PhysRevE.80.056117](https://doi.org/10.1103/PhysRevE.80.056117)
61. Bohlin L, Edler D, Lancichinetti A, Rosvall M. Community detection and visualization of networks with the map equation framework. In: Ding Y, Rousseau R, Wolfram D, editors. *Measuring Scholarly Impact*. Switzerland: Springer International Publishing; 2014. p. 3–34.
62. Rosvall M, Bergstrom CT. Mapping change in large networks. *PLoS ONE*. 2010; 5(1):e8694. doi: [10.1371/journal.pone.0008694](https://doi.org/10.1371/journal.pone.0008694) PMID: [2011700](https://pubmed.ncbi.nlm.nih.gov/2011700/)
63. Mirshahvalad A, Beauchesne OH, Archambault É, Rosvall M. Resampling effects on significance analysis of network clustering and ranking. *PLoS ONE*. 2013; 8(1):e53943. doi: [10.1371/journal.pone.0053943](https://doi.org/10.1371/journal.pone.0053943) PMID: [23372677](https://pubmed.ncbi.nlm.nih.gov/23372677/)
64. Ball B, Karrer B, Newman MEJ. Efficient and principled method for detecting communities in networks. *Phys Rev E*. 2011; 84(3):036103. doi: [10.1103/PhysRevE.84.036103](https://doi.org/10.1103/PhysRevE.84.036103)
65. Gopalan PK, Blei DM. Efficient discovery of overlapping communities in massive networks. *P Natl Acad Sci USA*. 2013; 110(36):14534–14539. doi: [10.1073/pnas.1221839110](https://doi.org/10.1073/pnas.1221839110)
66. Ronhovde P, Nussinov Z. Multiresolution community detection for megascale networks by information-based replica correlations. *Phys Rev E*. 2009; 80(1):016109. doi: [10.1103/PhysRevE.80.016109](https://doi.org/10.1103/PhysRevE.80.016109)
67. Traag VA, Krings G, Van Dooren P. Significant scales in community structure. *Sci Rep*. 2013; 3:2930. doi: [10.1038/srep02930](https://doi.org/10.1038/srep02930) PMID: [24121597](https://pubmed.ncbi.nlm.nih.gov/24121597/)