

FOCAS: Penalising friendly citations to improve author ranking

Jorge Silva

CRACS/INESC TEC & University of Porto
Porto, Portugal
jorge.m.silva@inesctec.pt

Pedro Ribeiro

CRACS/INESC TEC & University of Porto
Porto, Portugal
pribeiro@dcc.fc.up.pt

David Aparício

CRACS/INESC TEC & University of Porto
Porto, Portugal
daparicio@dcc.fc.up.pt

Fernando Silva

CRACS/INESC TEC & University of Porto
Porto, Portugal
fmsilva@dcc.fc.up.pt

ABSTRACT

Scientific impact is commonly associated with the number of citations received. However, an author can easily boost his own citation count by (i) publishing articles that cite his own previous work (*self-citations*), (ii) having co-authors citing his work (*co-author citations*), or (iii) exchanging citations with authors from other research groups (*reciprocated citations*). Even though these *friendly citations* inflate an author's perceived scientific impact, author ranking algorithms do not normally address them. They, at most, remove self-citations. Here we present Friends-Only Citations AnalySer (FOCAS), a method that identifies friendly citations and reduces their negative effect in author ranking algorithms. FOCAS combines the author citation network with the co-authorship network in order to measure author proximity and penalises citations between *friendly authors*. FOCAS is general and can be regarded as an independent module applied while running (any) PageRank-like author ranking algorithm. FOCAS can be tuned to use three different criteria, namely authors' distance, citation frequency, and citation recency, or combinations of these. We evaluate and compare FOCAS against eight state-of-the-art author ranking algorithms. We compare their rankings with a ground-truth of best paper awards. We test our hypothesis on a citation and co-authorship network comprised of seven Information Retrieval top-conferences. We observed that FOCAS improved author rankings by 25% on average and, in one case, leads to a gain of 46%.

CCS CONCEPTS

• **Human-centered computing** → **Social network analysis**; • **Computing methodologies** → **Ranking**;

KEYWORDS

Author ranking, self-citations, friendly citations, citation networks, co-authorship networks

ACM Reference Format:

Jorge Silva, David Aparício, Pedro Ribeiro, and Fernando Silva. 2020. FOCAS: Penalising friendly citations to improve author ranking. In *The 35th ACM/SIGAPP Symposium on Applied Computing (SAC '20)*, March 30–April 3, 2020, Brno, Czech Republic. ACM, New York, NY, USA, 9 pages. <https://doi.org/10.1145/3341105.3373991>

1 INTRODUCTION

Deciding where (or to whom) to allocate research funding is a problem that affects all scientists directly. This is typically done by attempting to assess the impact of a scientist, that is, to determine *how much of his research work has contributed to advance his scientific field*. The impact of scientists is also commonly used to pick scientific committees, attribute research grants, or decide faculty promotions. These processes are not fully automated and are traditionally done by peers. However, bibliometrics can be of help since they provide an unbiased estimator of scientific impact. For example, the h-index [5] counts the number of publications that a scientist (or author) has with more than h citations (e.g., an author has h-index = 7 if he has 7 papers with at least 7 citations). Many variations of the h-index have been proposed [8, 10] but the h-index remains widely used.

Another common approach to evaluate an author's impact is to use graph metrics on citation networks. Computing graph metrics is computationally more expensive than calculating bibliometrics, but has some advantages, namely (i) they give credit for indirect citations (i.e., if A cites B, and B cites C, C receives part of the credit of the citation of A to B), and (ii) they measure the author's impact at a group scale, that is the impact of each author depends on the impact of the authors that cite him. PageRank [9] is the most widely used graph algorithm to measure author's impact, and many variations have been proposed specifically for author ranking [2, 3, 11, 14, 16, 19]. One of PageRank's major algorithmic ideas is that nodes are not all equal, i.e., in its original context of hyperlinks, *it is good that any webpage points at yours but it is better that important webpages point at yours*. This idea naturally extends to author citation networks, meaning that *it is good to be cited by any author but it is better to be cited by important authors*.

Regardless of the metric used to evaluate scientific impact (e.g., bibliometrics or graph metrics), citations are important and several works study how an author can increase his number of citations. Undoubtedly the quality of the author's work is correlated with his number of citations [17]. However, other factors such as the author's co-authorship network [12] and his social behaviour [4, 15]

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

SAC '20, March 30–April 3, 2020, Brno, Czech Republic

© 2020 Association for Computing Machinery.

ACM ISBN 978-1-4503-6866-7/20/03...\$15.00

<https://doi.org/10.1145/3341105.3373991>

also have a big effect on his citation count. In this work, we focus on the aspect of social behaviour which studies the abuse of certain citation patterns to boost someone's (or their own) citation count, thus increasing the author's (perceived) scientific impact. There are three types of patterns often used to boost citations, namely (i) self-citations, when an author cites his own work, (ii) reciprocated citations, when authors or groups of authors interchangeably cite each other, and (iii) co-author citations, when authors cite works of their co-authors. Note that there may be nothing inherently malicious in using these citation patterns since in certain cases it makes sense for an author to cite his previous work, cite other people that have cited him, or cite the work of his co-authors, as long as the publications are in the same research line. However, by abusing these practices, authors can unreasonably boost their number of citations and consequently their (perceived) scientific impact, as shown in several studies [1, 4, 13, 15]. Thus, it is important to provide the scientific community with tools to mitigate the effect of citation boosting in author ranking algorithms.

In this paper, we merge self-citations and co-author citations into a group named *friendly citations* and we propose Friends-Only Citation AnalySer (FOCAS). FOCAS is a penalty estimation algorithm that analyses the co-authorship and citation networks in order to decrease the effect of friendly citations in author ranking algorithms. We propose three different criteria used to capture friendly citations (i.e., authors' distance, co-authorship frequency, and co-authorship recency). FOCAS does not rank authors by itself but it is designed to be easily integrated with any existing PageRank-based author ranking algorithms. Since we did not find a single superior state-of-the-art (STOA) author ranking method, as their performance greatly depends on what the user considers to be relevant criteria to rank the authors by, we decided to create a general tool. Thus, we augment four STOA author ranking methods SCEAS [14], SARA [11], RLPR [2], and five variations of OTARIOS [16], with FOCAS. We run all methods on a (citation and co-authorship) network comprised of seven Information Retrieval top-conferences and evaluate how well their produced rankings match a ground-truth ranking built using best paper awards information. We observe that FOCAS improves STOA author ranking algorithms by 25% on average (i.e., across author ranking algorithms) and 46% at best in our experiments (namely for one variation of OTARIOS). These results suggest that addressing the negative effect of friendly citations can improve author ranking algorithms.

Our paper is organised as follows. Section 2 presents related work. In Section 3 we describe our method, FOCAS. We detail our experiments and show results in Section 4. Finally, we discuss our main conclusions and propose future work in Section 5.

2 RELATED WORK

There are two main types of metrics used to measure scientific impact: bibliometrics and graph metrics. Next, we focus on these metrics and their advantages (see 1).

There are two groups of graph metrics used for author ranking: paper-level and author-level [18]. Paper-level graph metrics are measured in paper citation networks (i.e., *paper p' cites paper p''*) and they first estimate the credit of papers and then distribute their credit to their respective authors. On the other hand, author-level

graph metrics directly estimate the credit of authors using an author citation network (i.e., *author a' cites author a''*). Regardless of being paper-level or author-level, most graph metrics for author ranking are based on PageRank [9].

PageRank has two main steps: node initialisation and score diffusion. During node initialisation, PageRank assigns initial scores to the nodes based on *a priori* information about them. The simplest possibility is to assign the same score to all nodes; this is done by some author ranking methods [2, 14]. Other methods assign initial scores based on statistics of the author's publications [3, 6, 11, 16] (e.g., number of publications) or attributes from the citation network [16] (e.g., venue prestige of authors' publications). The score diffusion step is an iterative process that diffuses the initial scores through the network using the weights associated with the edges. This process is mostly universal across author ranking methods, however they differ in the way they calculate the edges' weights. They can favour authors that (i) publish with fewer co-authors [2, 11, 14, 16], (ii) receive more recent citations [3, 6, 16], or (iii) receive citations from more prestigious venues [6, 16].

Despite several studies showing that the abuse of certain citation patterns (i.e., self, reciprocated, and co-author citations) leads to undeserved (perceived) scientific impact [4, 15, 17], author ranking algorithms do not address their negative effect in ranking estimation. At most, some remove self-citations from the networks in a pre-processing step [11, 16].

We should point out the existence of the c-index [1], an h-index based bibliometric that counts the number of citations that an author has received from a distance bigger than c in his co-authorship network. However, the c-index does not analyse the citation network, thus it has the same drawbacks of other bibliometrics.

In this paper we propose FOCAS, a method that simultaneously analyses the co-authorship and citation networks and estimates penalties for all citations based on the co-authorship relation between citing and cited authors. FOCAS aims to reduce the negative effect of the abuse of *friendly citations* on author ranking algorithms. Due to the closer similarities between author-level and co-authorship networks (i.e., the nodes represent authors in both) we build FOCAS to be easily integrated with Author-Level Author Ranking (ALAR) algorithms.

3 METHODOLOGY

FOCAS is a citation penalty estimator for author-level citation networks that aims to reduce the effect of friendly citations in author ranking. FOCAS, by itself, does not rank authors; instead it estimates a penalty $p \in [0, 1]$ for every citation. Thus, ALAR algorithms can easily incorporate this information in their ranking estimations. Furthermore, FOCAS is a flexible algorithm that estimates different penalties depending on user-defined criteria.

While ALAR algorithms differ in the criteria used to rank authors, they are similarly divided into two main steps. On the first one a vector R is defined as the initial score for all the authors (*score initialisation*) and on the second one a vector S , containing the scores of the authors, is continuously updated using the edges' weights until convergence (*score diffusion*)¹, thus leading to the authors final ranking.

¹Note that at the first iteration: $R = S$

Table 1: Notation table.

Notation	Description
$a' \rightarrow a''$	author a' cites author a''
$(a' \rightarrow a'')_y$	year of the citation $a' \rightarrow a''$
$(a' \rightarrow a'')_w$	weight of the citation $a' \rightarrow a''$
$(a' \rightarrow a'')_p$	penalty of the citation $a' \rightarrow a''$
$a' \leftrightarrow a''$	author a' and a'' are co-authors on a publication
$(a' \leftrightarrow a'')_y$	year of the collaboration $a' \leftrightarrow a''$
$C_{(a' \leftrightarrow a'')}$	set of collaborations between a' and a''
$\Delta(C_{(a' \leftrightarrow a'')}, y)$	year of the most recent collaboration between authors a' and a'' prior to year y
$\Phi(C_{(a' \leftrightarrow a'')}, y)$	number of collaborations between authors a' and a'' prior to year y
$p(a' \leftrightarrow a'')$	path between a' and a'' in the co-authorship network
$P(a' \leftrightarrow a'')$	all paths between a' and a'' in the co-authorship network
R	authors' initial score
S	authors' estimated score
$S(a' \rightarrow a'')$	score of citation $a' \rightarrow a''$
$R_{a'}$	a' initial score
$S_{a'}$	a' estimated score

We present our method in the following sections. First we describe how friendly citations are penalised. Then, we detail how an initial version (named FOCAS-naive) applies the penalties before score initialisation. Finally, we put forward FOCAS, which iteratively applies the penalties during the score diffusion step.

3.1 Penalising friendly citations

We use citation and co-authorship networks to calculate penalties for friendly citations. Table 1 presents the notation used throughout this document. The citation network has authors as nodes and citations as edges, i.e., *author a' cites author a''* . Each citation $(a' \rightarrow a'')$ is made in year $(a' \rightarrow a'')_y$ and has weight $(a' \rightarrow a'')_w$. The weight is computed by an ALAR algorithm (e.g., RLPR) and measures the impact of the citation (i.e., citations with higher weights have higher impact on the ranking of the cited author). The co-authorship network has authors as nodes and collaborations as

edges, i.e., *author a' co-authors an article with author a''* . Each collaboration $(a' \leftrightarrow a'')$ is published in year $(a' \leftrightarrow a'')_y$. Note that authors a' and a'' may collaborate multiple times, thus we define $C_{(a' \leftrightarrow a'')}$ as the set of collaborations between them. Additionally, $\Delta(C_{(a' \leftrightarrow a'')}, y)$ is the year of their most recent collaboration prior to year y and $\Phi(C_{(a' \leftrightarrow a'')}, y)$ is the number of times they collaborated prior to year y .

We now describe how citations penalties are computed (Algorithm 1). We iterate over all citations $(a' \rightarrow a'')$ in the citation network N_c (line 1). Initially, we assign no penalty to the citation (line 2), i.e., if authors a' and a'' never co-authored a paper together, and there is not a path between a' and a'' , the citation has no penalty. Otherwise, we find set $Q \subseteq P(a' \leftrightarrow a'')$ which contains all paths between a' and a'' in the co-authorship network N_a constrained by $(a' \rightarrow a'')_y$, i.e., only co-authorships previous to the citation are used to calculate penalties (line 3). This means that citations from author a' to author a'' can have different penalties in different years (e.g., two authors have no penalty in year y because they never collaborated; then, if they co-author a paper in year y , they will have a penalty in year $y + 1$). For efficiency purposes, we only consider paths with distance $p(a' \leftrightarrow a'')_d \leq 3$. Our co-authorship network is a small-world network, which is typical for collaboration networks between scientists [7]. Thus, even for small distances, we find paths between many authors that are not co-authors. We then calculate the penalty p_q for all $p(a' \leftrightarrow a'') \in Q$ using criteria θ (e.g., frequency) and save the largest penalty found as the final penalty $(a' \rightarrow a'')_p$ for the citations between a' and a'' in year $(a' \rightarrow a'')_y$ (lines 4-9); finally, we add the penalty to the list of penalties P_c . Note that we calculate penalties for direct co-authors and the total penalty is the product of the penalties in the path (e.g., if a is a co-author of b with penalty 0.5, and b is a co-author of c with penalty 0.7, the penalty of path a to c is $0.5 * 0.7$) (lines 6-7). If multiple paths exist between a' and a'' , the final penalty between a' and a'' is the maximum penalty found (e.g., if we found three paths with penalties 0.5, $0.7 * 0.3$, and 0.2, the final penalty is 0.5) (lines 8-9).

We use three different criteria to compute penalties, i.e., distance, frequency, and recency (Equations 1, 2, and 3, respectively). The three criteria capture different properties of collaborations and give higher penalties to citations (i) between authors that are close in the co-authorship network (D-FOCAS), (ii) between more frequent collaborators (F-FOCAS), (iii) or between more recent collaborators (R-FOCAS). Distance applies a penalty of $d = 0.75$ to co-authors and 0 otherwise². Frequency and recency use a decay parameter λ to regulate the function's slope. In our experiments we set $\lambda = 4$.

$$D(a' \rightarrow a'') = d = 0.75 \quad (1)$$

$$F(a' \rightarrow a'') = 1 - e^{\left(\frac{\Phi(C_{(a' \leftrightarrow a'')}, (a' \rightarrow a'')_y)}{\lambda} \right)^{-1}} \quad (2)$$

$$R(a' \rightarrow a'') = e^{\left(\frac{(a' \rightarrow a'')_y - \Delta(C_{(a' \leftrightarrow a'')}, (a' \rightarrow a'')_y)}{\lambda} \right)^{-1}} \quad (3)$$

Note that it is possible to combine the different criteria, e.g., combine frequency with recency, thus decreasing the weight of the

²0.75 is approximately the highest expected penalty for R-FOCAS and F-FOCAS.

Algorithm 1 Penalty estimation.

Input: Co-authorship network N_a , citation network N_c , criteria θ .

Output: Penalties $P_c = \{(a' \rightarrow a'')_p : \forall (a' \rightarrow a'') \in N_c\}$.

```

1: for  $a' \rightarrow a'' \in N_c$  do
2:    $(a' \rightarrow a'')_p = 0$ 
3:    $Q = \text{getCoAuthorPaths}(a', a'', (a' \rightarrow a'')_y, N_a)$ 
4:   for  $p(a' \leftrightarrow a'') \in Q$  do
5:      $p_q = 1$ 
6:     for  $a^i \leftrightarrow a^j \in p(a' \leftrightarrow a'')$  do
7:        $p_q = p_q \times \text{calculatePenalty}(a^i \leftrightarrow a^j, \theta)$ 
8:     if  $p_q > (a' \rightarrow a'')_p$  then
9:        $(a' \rightarrow a'')_p = p_q$ 
10:   $P_c = P_c \cup (a' \rightarrow a'')_p$ 
11: return  $P_c$ 

```

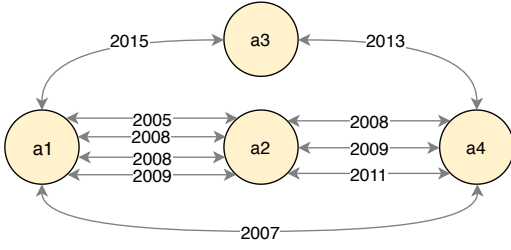


Figure 1: Example of a co-authorship network.

citations between authors that co-authored many articles recently; our nomenclature for that variation is FR-FOCAS (others are DF-FOCAS, DR-FOCAS, and DFR-FOCAS, for a total seven variations). We should note that FOCAS handles self-citations as a special case, i.e., independently of the criteria used, the penalty $(a' \rightarrow a'')_p = 1$ when $a' = a''$. Thus, self-citations have weight $(a' \rightarrow a'')_w = 0$ and are removed from the citation network.

We show an example in Table 2 of how penalties are calculated using different criteria for a given citation of the co-authorship network from Figure 1. This example highlights how different the penalties are when using different criteria. In this case, we are considering different paths between the same citing and cited author. However, if we consider a case where the only path available is the direct one (i.e., $a1 \leftrightarrow a2$ in the example), then the penalty applied to the citation varies from a maximum of 0.75 (D-FOCAS) to a minimum of 0.10 (R-FOCAS). Thus, one must carefully decide which criteria to use.

3.2 FOCAS-naive

We now describe FOCAS-naive, an initial version of FOCAS which applies penalties, calculated as described in Section 3.1, before score initialisation. Thus, FOCAS-naive can be used as a pre-processing step of ALAR algorithms.

FOCAS-naive (Algorithm 2) iterates over all citations ($a' \rightarrow a''$) in the citation network N_c (line 1) and calculates the new citation weight $(a' \rightarrow a'')'_w$ based on the original weight $(a' \rightarrow a'')_w$ and the penalty $(a' \rightarrow a'')_p$ (line 2). The citation network with new citation weights is then used by the ALAR algorithms during the score diffusion step; thus, they will obtain different author rankings.

In each iteration of the score diffusion step of ALAR algorithms, every citing author divides his score (from the previous iteration) and distributes it among his cited authors according to his citations weights. Therefore, cited authors with higher weights receive more score (e.g., $(a1 \rightarrow a2)_w = 0.6$ and $(a1 \rightarrow a3)_w = 0.3$, if the score from the previous iteration $S_{a1} = 0.8$, $a2$ receives $0.8 \times \frac{0.6}{0.6+0.3} \approx 0.53$ and $a3$ receives $0.8 \times \frac{0.3}{0.6+0.3} \approx 0.27$). FOCAS-naive decreases the weight of friendly citations. Consequently, whenever a citation $a1 \rightarrow a2$ is penalised, the score received by other authors cited by $a1$ is increased (e.g., from the previous example, if $(a1 \rightarrow a2)_p = 0.5$ and $(a1 \rightarrow a2)_w = 0$, then the new weight $(a1 \rightarrow a2)_{w'} = 0.6 \times (1 - 0.5) \approx 0.3$ which is the same as $(a1 \rightarrow a3)_{w'}$, thus both $a2$ and $a3$ increase their scores by 0.4 from $a1$'s citation). Therefore, FOCAS-naive not only decreases the score/impact of authors with

Table 2: Penalties using three different criteria for citation $a1 \rightarrow a4$ in 2016 from the co-authorship network of Figure 1. Penalties for co-authors (i.e., direct connections) are calculated using Equations 1, 2, and 3. Penalties for indirect connections are the product of penalties of the co-authors chain (e.g. $(a1 \rightarrow a2 \rightarrow a4)_p = (a1 \rightarrow a2)_p \times (a2 \rightarrow a4)_p$). η is the number of collaborations between two co-authors, δ is the difference in years between the citation and the most recent collaboration of two co-authors (e.g., 2016 - 2009). Bold values indicate the path from $a1$ to $a2$ with the highest penalty for the respective criteria.

Path	Distance (D)	Frequency (F)	Recency (R)
$a1 \leftrightarrow a2$	0.75	$(\eta = 4) 0.63$	$(\delta = 7) 0.17$
$a2 \leftrightarrow a4$	0.75	$(\eta = 3) 0.53$	$(\delta = 5) 0.29$
$a1 \leftrightarrow a3$	0.75	$(\eta = 1) 0.22$	$(\delta = 1) 0.78$
$a3 \leftrightarrow a4$	0.75	$(\eta = 1) 0.22$	$(\delta = 3) 0.47$
$a1 \leftrightarrow a2$	0.75	$(\eta = 1) 0.22$	$(\delta = 9) 0.10$
$a1 \leftrightarrow a2 \leftrightarrow a4$	$0.75 \times 0.75 = 0.56$	$0.63 \times 0.53 = 0.34$	$0.17 \times 0.29 = 0.05$
$a1 \leftrightarrow a3 \leftrightarrow a4$	$0.75 \times 0.75 = 0.56$	$0.22 \times 0.22 = 0.05$	$0.78 \times 0.47 = 0.37$

(many) friendly citations but it also increases the impact of authors without (many) friendly citations.

This idea fits our goal, but FOCAS-naive fails to penalise some kinds of friendly citations. Let us consider that $a1$ only cites $a2$ and $a3$ with respectively citation weights w_2 and w_3 ; if the same (or a similar) penalty is calculated for both citations, then the authors still receive (nearly) the same score from $a1$ as if their citations were not penalised. Furthermore, in cases where $a1$ only cites $a2$ (once or many times), $a2$'s receives a score from $a1$ that is independent of the penalty assigned to the citations. To overcome this limitation, we propose FOCAS in the next section.

3.3 FOCAS

We now describe FOCAS, an improved version of FOCAS-naive, which applies penalties during the score diffusion step of ALAR methods. Therefore, FOCAS is integrated during runtime with ALAR methods.

FOCAS (Algorithm 3)³ calculates penalties $(a' \rightarrow a'')_p$ for every citation as described in Section 3.1. Initially, the ALAR method calculates vector R which contains the authors' initial score (line 1; we do not specify parameters for the function since they depend on the ALAR method). Then, the authors' scores are initialised with the initial scores (line 2) and the score diffusion step begins (line

³For reference, lines 5 and 8-11 in Algorithm 3 are specific to FOCAS while the remaining lines are general to ALAR methods.

Algorithm 2 FOCAS-naive.

Input: Citation network N_c , citation penalties

$P_c = \{(a' \rightarrow a'')_p : \forall (a' \rightarrow a'') \in N_c\}$.

Output: Citation network N_c with redefined weights.

```

1: for  $a' \rightarrow a'' \in N_c$  do
2:    $(a' \rightarrow a'')_{w'} = (a' \rightarrow a'')_w \times (1 - (a' \rightarrow a'')_p)$ ;
3: return  $N_c$ 
```

3). At each iteration, all authors' scores S and their total penalised score ($total_penalty$) are initialised 0 (lines 4 and 5, respectively). Then, the ALAR method iterates over all citations $a' \rightarrow a''$ in the citation network N_c (line 6). For every citation, the ALAR algorithm calculates the score $S(a' \rightarrow a'')$ given from author a' to a'' based on the properties of the citation $a' \rightarrow a''$ and the previous iteration score $S'_{a'}$ (line 7; we do not specify the calculation of this step since it is ALAR dependent). Then, the ALAR method adds the score $S(a' \rightarrow a'')$ to the cited author score $S''_{a''}$ and moves on to the next citation (this would imply skipping lines 8-11 in Algorithm 3). When FOCAS is integrated with the ALAR method, there is a new step, before adding the score to $S''_{a''}$, where friendly citations are penalised. First, FOCAS removes a portion $(a' \rightarrow a'')_p \in [0, 1]$ from $S(a' \rightarrow a'')$ (line 8). In order to maintain the scores stable (i.e., the sum of all the authors' scores equal to 1) and to guarantee that the ALAR method will eventually converge, FOCAS cannot simply remove scores. Thus it stores the total penalised score ($total_penalty$) to reallocate it at a later stage (line 9). After iterating over all citations, FOCAS iterates over all authors (line 10) and gives a portion of the total penalty $total_penalty$ to each one according to their initial score (line 11). Finally, the normal process of ALAR algorithms resumes, i.e., (i) checking if the stopping criteria is met (i.e., if the scores have converged, calculated by comparing if the scores S are *too similar* to the scores from the previous iteration S') (lines 12 and 13) and (ii) updating the scores for the next iteration, if necessary (line 14). At the end of the process, FOCAS obtains the authors' scores calculated by an ALAR method and penalising friendly citations.

Algorithm 3 FOCAS.

Input: Citation network N_c , citation penalties

$P_c = \{(a' \rightarrow a'')_p : \forall (a' \rightarrow a'') \in N_c\}$.

Output: Author scores S .

```

1:  $R = ALAR\_Initialisation()$ 
2:  $S' = R$ 
3: while True do
4:    $S = \{S'_a = 0 : \forall a' \in R\}$ 
5:    $total\_penalty = 0.0$ 
6:   for  $a' \rightarrow a'' \in N_c$  do
7:      $S(a' \rightarrow a'') = ALAR\_Score(a' \rightarrow a'', S'_{a'})$ 
8:      $S_{a''} += S(a' \rightarrow a'') \times (1 - (a' \rightarrow a'')_p)$ 
9:      $total\_penalty += S(a' \rightarrow a'') \times (a' \rightarrow a'')_p$ 
10:  for  $a' \in S$  do
11:     $S_{a'} = total\_penalty \times R_{a'}$ 
12:  if converged( $S, S'$ ) then
13:    break
14:   $S' = S$ 
15: return  $S$ 
```

3.4 FOCAS-naive versus FOCAS

FOCAS decreases the score/impact given through friendly citations, consequently, FOCAS decreases the score of authors that have many friendly citations. However, by reducing the score of an author (and in order to keep the ranking system stable) ALAR methods automatically benefit some other authors (i.e., scores cannot be

removed, so they redistributed to other authors). FOCAS-naive and FOCAS differ on how they redistribute the penalised score to other authors. FOCAS-naive benefits authors that are not being cited by friendly citations, but that the citing author is using friendly citations to cite other authors (e.g., a_2 has a friendly citation from a_1 and a_3 has a regular citation from a_1 , FOCAS-naive penalising the citation $a_1 \rightarrow a_2$ results in higher score given from a_1 to a_3). On the other hand, FOCAS redistributes the penalised scores through the authors according to the score initialisation of authors (i.e., authors with higher score initialisation receive a larger part of the penalised score) which is calculated by the ALAR method.

4 RESULTS AND DISCUSSION

In this section we study a real-world (co-authorship and citation) network and we integrate FOCAS-naive and FOCAS with existing ALAR methods. Our aim is to show (i) that friendly citations are frequent in real-world datasets and (ii) that both FOCAS-naive and FOCAS improve the authors rankings of ALAR methods.

4.1 Experimental setup

In order to create a test scenario, we build a citation and co-authorship network using the publications extracted from the DBLP dataset⁴ for 7 top-tier conferences (KDD, CIKM, PODS, SIGMOD, VLDB, WWW, SIGIR) in the area of Information Systems from Computer Science. There are a total of 28,266 different authors in these publications (which corresponds to the number of nodes in each network). Furthermore, there are 5.77M citations (edges in citation network) and 0.15M collaborations (edges in the co-authorship network). We create a ground-truth author ranking based on the best paper awards given by the 7 conferences⁵. We counted each awarded publication as a unit of merit which is equally divided by the authors of the publications. As a result, we are assuming that authors that have won more best paper awards with fewer co-authors should be ranked higher. The idea of using best paper awards as a human judged ground-truth was already used by previous ALAR algorithms [11, 16].

In our experiments, rankings produced by ALAR methods (or ALAR + FOCAS methods) are compared to the ground-truth ranking; i.e., methods that produce rankings similar to the ground-truth ranking are considered better. We use the Normalized Discounted Cumulative Gain (NDCG) to compare two rankings. NDCG penalises methods that rank more impactful authors below less impactful ones (where the impact of an author is defined by the ground-truth ranking). The value of NDCG ranges between 0 and 1, i.e., 1 represents a perfect match between two ranks. NDCG only considers the number of incorrect ranking placements for the top n authors from the ground truth. In our tests, we consider the following values of n : 5, 10, 20, 50 and 100; however, for space concerns, we only show the average NDCG obtained for these values.

We augment eight different ALAR algorithms with FOCAS-naive and FOCAS: RLPR [2], SCEAS [14], SARA [11] and five variants of OTARIOS [16]. RPLR is a PageRank algorithm used in the context of author ranking. SCEAS adapts PageRank to achieve faster convergence. SARA is a PageRank algorithm but the score initialisation

⁴<https://aminer.org/citation>

⁵Awards information obtained from: https://jeffhuang.com/best_paper_awards.html

Table 3: Distribution of the co-authorship distance of the citations. $L-X$ represents the level of distance with $L-0$ corresponding to auto-citations and $L-N$ corresponding to 4 or more. Network represents the citations for all the authors while T represents the ones incoming to authors with best paper awards. $T@N$ represents the top N authors with the most awards.

	# Cits	L-0	L-1	L-2	L-3	L-N
Network	5.77M	2.08%	16.63%	64.45%	8.93%	7.92%
T@397	0.84M	1.91%	7.56%	12.98%	21.83%	55.73%
T@100	0.33M	2.04%	7.17%	13.27%	22.11%	55.42%
T@50	0.21M	2.23%	7.60%	13.16%	22.18%	54.82%
T@10	0.05M	2.94%	8.68%	11.15%	20.51%	56.71%
T@1	0.01M	5.69%	17.47%	15.42%	37.27%	24.15%

is based on author’s productivity. OTARIOS is also a PageRank algorithm but with multiple user-defined parameters that estimate different citation weights and score initialisation. OTARIOS variants are numbered from 1 to 5 according to their respective criteria ($L + A + AW$), ($L + AVW + AW$), ($AP + L + AW$), ($AP + A + AW$) and ($AV + VW + AW$)⁶. We run each ALAR on the citation network and calculate the average NDCG as our baselines to beat. The goal of our experiment is not to determine which is the best ALAR algorithm, instead we focus on measuring the improvement of the produced rankings (i.e., the average NDCG) after adding FOCAS-naive and FOCAS to the ranking process. To measure the gain of the methods we use the following equation:

$$gain = \frac{NDCG_{FOCAS} - NDCG_{ALAR}}{\min(NDCG_{FOCAS}, NDCG_{ALAR})} * 100\% \quad (4)$$

4.2 Motivation

In order to show the frequency of friendly citations in real citation networks and how they diverge for different groups of authors, we measure the co-authorship distance between citing and cited author in the citation network (Table 3). We observe that $> 92\%$ of the citations are friendly citations and that most of them (i.e., $> 64\%$) have a co-authorship distance of 2. We filter the citations to compare the differences between the friendly citations received in general (i.e., the whole network) and the most prestigious authors (i.e., the ones with at least one best paper award). We observe that the distribution is similar within these groups of best authors across different levels of prestigious authors (i.e., T@397, T@100, T@50, T@10) and that the distribution is very different from the whole network. For the awarded authors, $> 55\%$ of their citations have a co-authorship distance higher than 3; the case of citations coming towards the most awarded author (T@1) is the exception, i.e., in this case his citations are on average closer to his co-authors when compared to other awarded authors, but they still are much farther away when compared to the whole network. We should note that the small-world network effect [7] is a justification for the high number of friendly citations in the network. However, it does not explain the different distribution between the whole network and

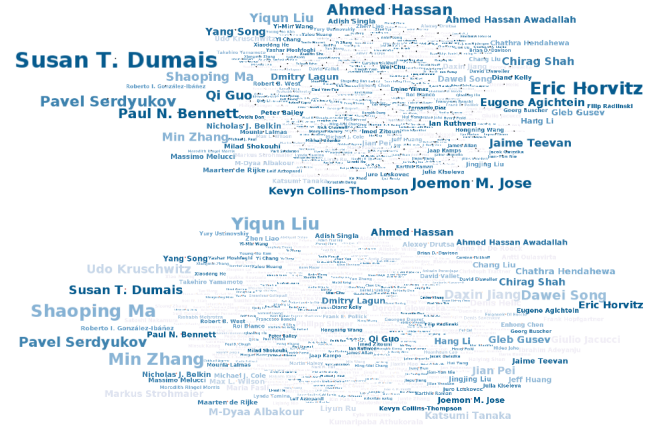


Figure 2: Ego-networks of the citations received by Ryen White (the best author according to the ground-truth) without any penalties (top figure) and with D-FOCAS penalty applied to the citation weights (bottom figure). Larger author names indicate that they have higher weights in Ryan White’s citation network. Additionally, darker colours indicate that the author is close to Ryen White in his co-authorship network.

the awarded authors, since we are just filtering citations and not recalculating their co-authorship distances.

Our exploration suggests that friendly citations are frequent in the network. Additionally, authors do not receive the same amount of friendly citations, and authors that (on average) receive the most friendly citations are placed (on average) lower in a human judged ranking. These facts highlight the need of penalising friendly citations in order to obtain more reliable author rankings.

Next we give an example of how friendly citations penalties affect the citation network. For this purpose we chose to single-out Ryen White (the best author according to our ground-truth) and we created two ego-networks of the citations he receives (Figure 2). The first ego-network uses a traditional approach [2, 11, 14] to calculate citation weights. The second ego-network uses the same approach to calculate citation weights and also applies the D-FOCAS penalties (i.e., penalties based on co-authorship distance). There are 1614 different authors citing Ryen White in our dataset. To ease visualisation, for both ego-networks, we removed citing authors whose citation weight (towards Ryen White) is lower than the average weight in the ego-network. The ego-network without D-FOCAS (top figure) and the ego network with D-FOCAS (bottom figure) have an average weight of 1.27 and 0.86, respectively, and there are 392 and 447 authors with weights above the average weight (i.e., visible in the figure), respectively. The effect of D-FOCAS is very noticeable by looking at the difference between the average weight in both ego-networks. D-FOCAS decreases the average weight by 0.41, which is a total loss of $\approx 32\%$ on the citation weights. We also observe that, without any penalties, there are only a few citing authors that contribute the most to the citation weight (i.e., in-weight) received by Ryen White (indicated by larger authors names). Furthermore, these authors are close to Ryen White

⁶Please refer to the OTARIOS paper for more details about each OTARIOS variant

in the co-authorship network (indicated by darker colours in the authors' names). Overall, the in-weight received by Ryen White is heavily based on a small set of authors which are close to him in the co-authorship network. After applying the D-FOCAS penalties, we now observe an increase on the number of authors that contribute the most to Ryen White's citation weight (i.e., there are more authors with larger names) and these authors have different co-authorship distances to Ryen White. Overall, the new weights are more evenly distributed by citing authors and less dependent on friendly co-author relations when no penalties are applied. As an example, for the case without penalties, citing authors Susan T. Dumais and Eric Horvitz were the ones with the highest contribution to Ryen White's in-weights with 33.73 and 26.67 respectively. After applying D-FOCAS penalties, their contribution decreased to 8.65 and 6.92 and they are now the 6th and 10th authors that contribute the most to Ryen White's in-weights.

There is a high correlation between the citation weights received by authors and their score in ALAR algorithms. Frequently co-authors (or authors close in the co-authorship network) are the ones that contribute the most for an author received citation weight. Although it is normal for an author to cite his co-authors, abusing this practice can lead to undeserved (perceived) scientific impact. In these cases, FOCAS tries to bring fairness to the author ranking process by making the authors' score calculation based on a more evenly distributed citation weight network (i.e., authors do not benefit as much from their co-authors).

We should point that although Ryen White citation weights are penalised in our example, that does not mean that his citing authors abuse the co-author citation pattern. As we will see in more detail in the next section, due to the nature of FOCAS penalty estimation, the common case in the citation network is that authors lose citation weight and consequently score. However, this does not indicate that their final ranking position decreases.

4.3 FOCAS' impact on author ranking

Here we combine FOCAS-naive and FOCAS with eight different ALAR methods and measure their improvements when compared to the original ALAR method.

Table 4 shows our baselines, i.e., the average NDCG obtained for the rankings produced by each ALAR method. The results show that OTARIOS₄ (0.267), OTARIOS₃ (0.265) and SCEAS (0.261) are the best methods, while RLPR (0.176) and SARA (0.160) are the worst. Removing self-citations is a common practice used by ALAR methods; thus, we measure the gain in NDCG obtained by removing self-citations from the citation network when compared against the baselines. OTARIOS₅ is the algorithm that benefits the most from removing self-citations (11% gain). On the other hand, there are three ALAR methods that have negative gain: OTARIOS₁ (-3.6%), SARA (-2.3%) and SCEAS (-2.1%). Furthermore, we observe that removing self-citations only has a gain of (1.4%) on average.

Table 5 shows the gains of combining ALAR methods with FOCAS-naive using different penalty criteria. There are three ALAR methods (RLPR, SCEAS, and SARA) that systematically have negatives gains regardless of the penalty criteria used by FOCAS-naive. We should point out that these methods all share the same method to calculate citation weights in the network. In the worst case,

Table 4: Results of the average NDGC @(5,10,20,50,100) for the STOA methods.

	Baseline	No self-citations
RLPR	0.176	0.4%
SCEAS	0.261	-2.1%
SARA	0.160	-2.3%
OTARIOS ₁	0.213	-3.6%
OTARIOS ₂	0.212	0.2%
OTARIOS ₃	0.265	4.9%
OTARIOS ₄	0.267	2.7%
OTARIOS ₅	0.238	11.0%
<i>Average gain</i>		1.4%

Table 5: Gain on the average NDCG obtained by the ALAR algorithms after combining them with FOCAS-NAIVE using 7 different criteria. Bold value per row represents the criterion with the most gain.

	D	F	R	DF	DR	FR	DFR
RLPR	-8.3%	2.3%	-7.8%	2.3%	-3.9%	0.5%	0.5%
SCEAS	-9.9%	-2.4%	-9.2%	-2.3%	-7.7%	-2.4%	-2.1%
SARA	-9.7%	-5.9%	-13.3%	-6.0%	-6.9%	-3.2%	-3.2%
OTARIOS ₁	36.1%	0.5%	5.2%	-0.3%	-2.5%	-3.6%	-3.6%
OTARIOS ₂	34.0%	-0.2%	1.9%	0.2%	-2.5%	0.0%	-0.2%
OTARIOS ₃	22.6%	6.7%	2.3%	6.9%	4.9%	5.7%	5.7%
OTARIOS ₄	16.0%	-0.1%	-1.9%	0.0%	4.5%	2.9%	2.7%
OTARIOS ₅	24.8%	20.1%	12.1%	29.9%	14.1%	12.0%	12.0%
<i>Average gain</i>	13.2%	2.5%	-1.3%	3.8%	0.0%	1.5%	1.5%

R-FOCAS-naive has a gain of -13.3% for SARA. On the other hand, the five OTARIOS variants consistently improve their rankings with FOCAS-naive. However, for four out of the five variants, the gains are only significantly high for the distance criteria (D). Overall, OTARIOS₅ is the only method that has significantly high gains (> 10%) for all criteria. We also observe that distance is the only criteria that has significantly high gains (13.2% average) and recency is the worst one (-1.3% average). The remaining criteria have gains comparable to simply removing self-citations.

Table 6 shows the gains of combining ALAR algorithms with FOCAS using different penalty criteria. We observe that seven out of the eight ALAR methods (i.e., all except SCEAS) have positive gains for FOCAS regardless of the criteria used to calculate penalties. Six of the ALAR methods have gains $\geq 30\%$ for some criteria. OTARIOS₂ has the highest gain of 46.4%. SCEAS is the only method that does not have positive gains for any criteria and has the lowest gain of -8.7% for D-FOCAS and R-FOCAS. Overall, all FOCAS criteria present significantly high gains across ALAR algorithms. D-FOCAS has the highest average gain with 25.4% and FR-FOCAS has the lowest average gain with 3.8%.

Comparing FOCAS-naive against FOCAS, our results indicate that FOCAS is considerably better than FOCAS-naive. FOCAS improves more different ALAR algorithms with different criteria than FOCAS-naive, and FOCAS also improves them more than FOCAS-naive (i.e., higher average gains). Furthermore, our results show

Table 6: Gain on the average NDCG obtained by the ALAR algorithms after combining them with FOCAS using 7 different criteria. Bold value per row represents the criterion with the most gain.

	D	F	R	DF	DR	FR	DFR
RLPR	11.3%	3.1%	4.5%	3.2%	7.8%	2.1%	1.2%
SCEAS	-8.7%	-3.3%	-8.7%	-4.1%	-2.3%	-3.0%	-2.0%
SARA	30.0%	5.5%	16.3%	2.9%	5.6%	2.3%	-0.9%
OTARIOS ₁	32.0%	39.3%	-2.1%	39.0%	2.2%	2.1%	3.3%
OTARIOS ₂	46.4%	39.0%	2.7%	38.6%	2.4%	0.7%	3.4%
OTARIOS ₃	32.5%	8.8%	26.9%	8.7%	13.1%	8.8%	7.1%
OTARIOS ₄	23.3%	21.6%	11.4%	21.7%	16.0%	2.0%	3.2%
OTARIOS ₅	36.6%	32.0%	17.8%	31.7%	28.8%	15.5%	26.0%
Average gain	25.4%	18.6%	8.6%	17.7%	9.2%	3.8%	5.2%

that the gains of FOCAS-naive are highly dependent on the process that estimates citation weights. RLPR, SCEAS, and SARA all share the same strategy to calculate citation weights and they all have very similar gains for all the criteria when FOCAS-naive is used. On the other hand, we observe that FOCAS gains are dependent on the quality of the score initialisation. RLPR and SCEAS are the only methods that use an uniform score initialisation strategy and, as a result, they are the ones with the smallest gains when FOCAS is used. We should also point out that SCEAS is the only method that does not have positive gains for neither FOCAS nor FOCAS-naive. This is not surprising because SCEAS converges in fewer iterations, meaning that the effect of the penalties (which grows as the iterative process of PageRank continues) are not noticeable.

Regarding the best criteria to penalise friendly citations, we observe that measuring the co-authorship distance between citing and cited authors and/or the frequency of their collaborations (i.e., D-FOCAS, F-FOCAS, and DF-FOCAS) yields the highest gains. Measuring how recent a collaboration is prior to a citation and its combinations with other base criteria (R, DR, FR and DFR) also yields positive gains; however they are much smaller.

In order to demonstrate the effect of FOCAS when measuring author impact, we compare the rankings and scores of authors on the OTARIOS₂ variant before and after applying the D-FOCAS penalties⁷. For brevity, we restrict our analysis to the top-10 authors from the ground-truth (Table 7). We must first highlight the difficulty of ALAR methods in producing rankings similar to the rankings created using human judgement (in this case, using best paper awards). Only one of the top-10 authors of the ground truth is placed in the top-10 of the ranking produced by OTARIOS₂ (with or without D-FOCAS) and six of the top-10 authors are placed outside the top-350 predicted authors. Regarding the differences in authors' scores after applying D-FOCAS, we observe that the top-10 authors lose 14% of their score on average. Ryen W. White loses the most score with a gain of -44% and Edo Liberty was the only one that presents a positive gain of 8%. The loss of score for the top-10 authors is not surprising; if we consider that due to their impact they are more likely to be cited (not only in quantity but also by different authors) and that due to the small-world effect of the

⁷We chose OTARIOS₂ with D-FOCAS because, overall, this is the combination that produces the best results, with an average NDCG of 0.351.

Table 7: Impact of FOCAS with criterion distance (D) on the OTARIOS₃ baseline on the top 10 most awarded authors. Author names are sorted from the most awarded author to the lowest awarded one. BR: Baseline Rank, PR: Penalty Rank, RI: Rank Improvement, BS: Baseline Score, DFSG: D-FOCAS Score Gain and # CIT: number of citations received. The number of citations only considers citations received from publications from the 7 conferences of our dataset.

Author	BR	PR	RI	10 ⁻² × BS	DFSG	# CIT
Ryen W. White	31	28	+3	0.225	-44%	5749
Pedro M. Domingos	24	14	+10	0.246	-13%	9202
Marcelo Arenas	381	372	+9	0.044	-20%	2602
Leonid Libkin	607	483	+124	0.029	5%	1433
Gerhard Weikum	29	18	+11	0.228	-17%	11566
Georg Gottlob	628	601	+27	0.029	-7%	2329
Edo Liberty	751	598	+153	0.025	8%	244
Ian Ruthven	531	675	-144	0.033	-35%	681
Jan Van den Bussche	2347	2192	+155	0.008	13%	554
Thorsten Joachims	2	1	+1	0.619	-30%	10984

co-authorship network (i.e., there is a small distance between most pair of authors) they are more likely to receive a friendly citation, it is expected that their score is negatively affected. Our results also show that despite the fact that the top-10 lose score after applying D-FOCAS penalties, these authors actually improve their ranking position on an average of 35 places. Jan Van den Bussche has the highest ranking improvement, jumping 155 positions, while Ian Ruthven is the only author whose rank position decreases, going down 144 positions. We should point out that the variations on ranking positions after applying FOCAS are more volatile for authors ranked lower because the difference between their scores and authors at the same level are smaller (i.e., a small change of the score is required to change ranking position).

4.4 So, authors shouldn't collaborate?

FOCAS only penalises self-citations and co-author citations. As a result, one might be misled into thinking that methods like FOCAS encourage authors to avoid collaborations since having co-authors makes you closer to everyone else in the community, and thus it will result in higher penalties for your citations. Part of this assumption is correct since FOCAS is less likely to penalise authors with less co-authors. However, it does not necessarily mean that these authors are going to obtain higher rankings. There are studies that have shown that collaboration is key to achieve career success in research [12]. As a result, an author that does not collaborate is more likely to obtain fewer citations which would result in a lower author ranking compared to having collaborations and having some of his citations penalised. FOCAS does not aim to discourage collaborations, instead it aims to identify and mitigate the abuse of friendly citations patterns on author ranking.

5 CONCLUSIONS AND FUTURE WORK

We found that ALAR methods did not address the problem of friendly citations, i.e., when authors boost their own scientific impact by citing themselves or exchanging citations with co-authors.

We thus put forward FOCAS, a method that penalises friendly citations based on (a) authors' distance, (b) citation frequency, and (c) citation recency. We assessed if FOCAS improved author ranking methods on a (citation and co-authorship) network comprised of seven Information Retrieval top-conferences. In our experiments, we verified that FOCAS improved state of the art author ranking methods by 25% on average and 46% at best. The most important criteria to improve rankings was the distance between authors, highlighting the importance of graph-based methods since traditional bibliometrics can not capture this information. Our experiments also suggested that the frequency of the citations seems to be more important than the recency of the citations. Another relevant result obtained in our study was that the traditional approach of removing self-citations has minimal gains ($\approx 1\%$ on average, and 11% at most for one of the tested methods). This latter result highlights why current state of the art is lacking and the importance of our approach.

With respect to future work there are a few possible directions. One would be to improve penalty estimation at two different levels: authors and citations. For authors we plan to study co-authorship and citation networks in order to estimate thresholds for the acceptable amount of citations received from co-authors and only penalise authors with an excessive amount of co-author citations compared to the normal patterns of the network. With respect to citations we aim to compare the context of two publications (i.e., how the content of a citing publications matches the content of the cited publication) and assign higher penalties to publications which context is much different from each other. Another topic of research that we wish to tackle in the near future is adding penalties for reciprocated citations. Several studies have shown that the abuse of reciprocated citation between groups of authors leads to undeserved scientific merit and again, the problem is not handled by author ranking algorithms. Furthermore, we plan to study different decay parameters for the R-FOCAS and F-FOCAS equations and how they differ (or not) in different networks. For instance, in some areas the expected gap between collaborations might be longer than in other areas, so a smaller decay should be used.

6 ACKNOWLEDGEMENTS

This work is financed by National Funds through the Portuguese funding agency, FCT - Fundação para a Ciência e a Tecnologia, within project UIDB/50014/2020. Jorge Silva is also supported by a FCT/MAP-i PhD grant PD/BD/128157/2016.

REFERENCES

- [1] Maria Bras-Amorós, Josep Domingo-Ferrer, and Vicenç Torra. 2011. A bibliometric index based on the collaboration distance between cited and citing authors. *Journal of Informetrics* 5, 2 (2011), 248–264.
- [2] Ying Ding. 2011. Applying weighted PageRank to author citation networks. *Journal of the American Society for Information Science and Technology* 62, 2 (2011), 236–245.
- [3] Marcel Dunański and Willem Visser. 2012. Comparing paper ranking algorithms. In *Proceedings of the South African Institute for Computer Scientists and Information Technologists Conference*. ACM, 21–30.
- [4] James Fowler and Dag Aksnes. 2007. Does self-citation pay? *Scientometrics* 72, 3 (2007), 427–437.
- [5] Jorge E Hirsch. 2005. An index to quantify an individual's scientific research output. *Proceedings of the National academy of Sciences* 102, 46 (2005), 16569–16572.
- [6] Won-Seok Hwang, Soo-Min Chae, Sang-Wook Kim, and Gyun Woo. 2010. Yet another paper ranking algorithm advocating recent publications. In *Proceedings of the 19th international conference on World wide web*. ACM, 1117–1118.
- [7] Mark EJ Newman. 2001. The structure of scientific collaboration networks. *Proceedings of the national academy of sciences* 98, 2 (2001), 404–409.
- [8] Eric Oberesch and Sven Groppe. 2017. The mf-index: A Citation-Based Multiple Factor Index to Evaluate and Compare the Output of Scientists. *OJW* 4, 1 (2017), 1–32.
- [9] Lawrence Page, Sergey Brin, Rajeev Motwani, and Terry Winograd. 1999. *The PageRank citation ranking: Bringing order to the web*. Technical Report. Stanford InfoLab.
- [10] Alex Post, Adam Y Li, Jennifer B Dai, Akbar Y Maniya, Syed Haider, Stanislaw Sobotka, and Tanvir F Choudhri. 2018. c-index and Subindices of the h-index: New Variants of the h-index to Account for Variations in Author Contribution. *Cureus* 10, 5 (2018).
- [11] Filippo Radicchi, Santo Fortunato, Benjamin Markines, and Alessandro Vespignani. 2009. Diffusion of scientific credits and the ranking of scientists. *Physical Review E* 80, 5 (2009), 056103.
- [12] Emre Sarigöl, René Pfitzner, Ingo Scholtes, Antonios Garas, and Frank Schweitzer. 2014. Predicting scientific success based on coauthorship networks. *EPJ Data Science* 3, 1 (2014), 9.
- [13] Marco Seeber, Mattia Cattaneo, Michele Meoli, and Paolo Malighetti. 2017. Self-citations as strategic response to the use of metrics for career decisions. *Research Policy* (2017).
- [14] Antonis Sidiropoulos and Yannis Manolopoulos. 2006. Generalized comparison of graph-based ranking algorithms for publications and authors. *Journal of Systems and Software* 79, 12 (2006), 1679–1700.
- [15] Antonis Sidiropoulos and Yannis Manolopoulos. 2019. Reciprocity and impact in academic careers. *EPJ Data Science* 8, 20 (2019).
- [16] Jorge Silva, David Aparicio, and Fernando Silva. 2018. OTARIOS: OpTimizing Author Ranking with Insiders/Outsiders Subnetworks. In *International Conference on Complex Networks and their Applications*. Springer, 143–154.
- [17] Roberta Sinatra, Dashun Wang, Pierre Develle, Chaoming Song, and Albert-László Barabási. 2016. Quantifying the evolution of individual scientific impact. *Science* 354, 6312 (2016), aaf5239.
- [18] Hao Wang, Hua-Wei Shen, and Xue-Qi Cheng. 2016. Scientific credit diffusion: Researcher level or paper level? *Scientometrics* 109, 2 (2016), 827–837.
- [19] Jevin D West, Michael C Jensen, Ralph J Dandrea, Gregory J Gordon, and Carl T Bergstrom. 2013. Author-level Eigenfactor metrics: Evaluating the influence of authors, institutions, and countries within the social science research network community. *Journal of the American Society for Information Science and Technology* 64, 4 (2013), 787–801.