

Carley, S., Porter, A.L., Youtie, J., (2019). A Multi-match Approach to the Author Uncertainty Problem. *Journal of Data and Information Science*, 4(3), 60-77.

A Multi-Match Approach to the Author Uncertainty Problem

Stephen F. Carley^a ▪ Alan L. Porter^b ▪ Jan L. Youtie^c

^aS. F. Carley (✉)

Search Technology, Inc, Norcross, GA 30092

e-mail: stephen.carley@searchtech.com

^bA. L. Porter

Search Technology, Inc, Norcross, GA 30092, USA; Georgia Institute of Technology, Atlanta, GA 30308, USA

^cJ. L. Youtie

Enterprise Innovation Institute, Georgia Institute of Technology, Atlanta, Georgia 30308 USA

Keywords

name disambiguation
author identifiers
multi-match solution

Abstract

The ability to identify the scholarship of individual authors is critical to evaluating performance. A number of factors impede this endeavor. Common and similarly spelled surnames (for a number of authors) make it difficult to isolate the scholarship of individual scholars indexed in large databases. Variations in the name spelling of individual authors further complicates matters. Common family names in scientific powerhouses like China make it problematic to distinguish between authors possessing ubiquitous, anglicized surnames (as well as the same or similar first names). The assignment of unique author identifiers has been a major step toward resolving these difficulties. We maintain, however, that in and of themselves author identifiers are not sufficient to fully address the author uncertainty problem. In this paper we build on the author identifier approach by considering commonalities in fielded data between authors containing the same surname and first initial of their first name. We illustrate our approach via three case studies.

Introduction

The ability to isolate scholarship belonging to individual authors is fundamental to assessing productivity, mobility, collaboration and scientific impact. For metrics to have meaning it is essential to build a body of scholarship unique to the subject(s) of one's study. Several undesirable outcomes can result from the inability to distinguish between individual authors. Among these, Han and colleagues (2004) note this "can affect the quality of scientific data gathering, can decrease the performance of information retrieval and web search, and can cause the incorrect identification of and credit attribution to authors." Co-authorship networks are influenced as well. Diesner and Kim (2016) show that a failure to correctly disambiguate names will "misrepresent statistical properties of co-authorship networks: It deflates the number of unique authors, number of components, average shortest paths, clustering coefficient, and assortativity, while it inflates average productivity, density, average co-author number per author, and largest component size." Negative results are not limited to the author level; they can also distort department and institution-level analyses as well.

The author uncertainty problem has become more challenging over time. The sizeable increase in scholarship from countries whose citizens share common surnames, like China (Zhou and Leydesdorff, 2006) and Korea, significantly contributes to this difficulty. The proliferation of common names (and especially anglicized Asian names) among multiple authors, variations in name spelling (for a given author's name) and prodigious scholarship (by individual authors) make name disambiguation all the more problematic. Standardized signatures shared by a myriad of authors, such as a surname followed by a first initial, further complicate the plot.

MacRoberts and MacRoberts (1989), as well as Smalheiser and Torvik (2009), provide useful overviews of the reasons why name disambiguation is as challenging as it is.

A number of solutions have been proposed for addressing the author uncertainty problem. To name a few, Han and colleagues (2004) advance “two supervised teaming approaches to disambiguate authors in the citations.” Song et al. (2007) promote a topic-based approach. Cota and colleagues (2007) advocate a hierarchical clustering technique. Among proposed solutions to date, machine learning, a field of computer science that focuses on teaching computer systems to learn, has emerged prominently. In this context, learning is synonymous with the ability to increasingly improve the performance of a given task (e.g. author name disambiguation). This improvement may occur without explicit programmer guidance.

Ferreira and colleagues (2012) conducted a survey of automatic methods for author name disambiguation to find that “[t]he majority of the surveyed methods perform disambiguation by comparing citation records using some type of similarity function.” In their analysis, a “major gap in the field is the lack of direct comparisons among the methods under the same circumstances: e.g., same collections (e.g., many methods used different versions of collections such as DBLP).” Another drawback of this approach is that when training data scarcity makes pattern detection difficult, machine learning programs typically will not produce desired results. Machine learning will be revisited later. In the interim, we now turn our attention to author identifier approaches to name disambiguation, which provide a key starting point for the method advanced in this study.

Unique Author Identifiers

Prominent among proposed solutions to author uncertainty are author identifiers. In October of 2012, ORCID (Open Researcher and Contributor ID) was launched as an open-access database to identify individual scholars. It is self-described as “a persistent digital identifier that distinguishes you from every other researcher and, through integration in key research workflows such as manuscript and grant submission, supports automated linkages between you and your professional activities ensuring that your work is recognized.”¹ Authors are assigned an ORCID iD consisting of 16 characters. While some organizations require the adoption of an ORCID identifier, others do not, making coverage a function of organizational policy or the part of the world to which a given scholar belongs (Youtie et al., 2017).

Since 2008 the database, Web of Science (WOS), provided by Clarivate Analytics, has issued its own unique author identifier, which it calls ResearcherID. The website² for this identifier prompts authors to register with ResearcherID and link their publications with their database. If a given author has a preexisting ORCID iD he or she can link that to his or her ResearcherID (and vice versa). As of October 2017, more than 270,000 researchers have signed up for a ResearcherID³ and 9,073,149 records indexed on WOS have a ResearcherID attached.⁴ Given the fact that registration for this identifier is optional, it shouldn't come as a surprise that ResearcherID coverage is less than 100%.

¹ See <https://orcid.org>

² See <http://www.researcherid.com>

³ See <https://clarivate.com/products/researcherid>

⁴ As all ResearcherIDs begin with a letter this figure is obtained by the WOS author identifier search: “A* OR B* OR C* OR D* OR E* OR F* OR G* OR H* OR I* OR J* OR K* OR L* OR M* OR N* OR O* OR P* OR Q* OR R* OR S* OR T* OR U* OR V* OR W* OR X* OR Y* OR Z*”

The database Scopus, provided by Elsevier, issues unique author identifiers which it refers to as ‘Scopus IDs’ to every author indexed in its dataset. Scopus was launched in 2004 and claims to draw from more than 5,000 publishers to index 69 million items and 12 million author profiles.⁵ Unlike ORCID iDs and ResearcherIDs, authors do not have to create their own profile to obtain a Scopus ID. It’s noted, however, that occasionally a given scholar is unintentionally assigned more than one Scopus ID (as is the case with one of the authors on this paper). When this happens, the author has the option to request that each of his or her Scopus IDs be merged into one. In the absence of such a request, however, searching for an author with multiple Scopus IDs using just one of their identifiers will produce incomplete results. Scopus author profiles can be linked to the same author’s ORCID account (and vice versa). Gasparyan and colleagues (2017) provide a general overview of author identifier approaches.

While author identifiers are a truly valuable response to the author uncertainty problem, in and of themselves, they do not provide a comprehensive solution for several reasons. Given that not all authors have applied for, or been issued, author identifiers, a number of authors do not possess unique identifiers. Among those that do, coverage is oftentimes less than robust. An author on this paper (Porter) has a publication count of 234 on WOS (as of late 2017), but a search for his work using his ORCID iD or ResearcherID results in 40% of scholarship belonging to him.⁶ Youtie and colleagues (2017) find that 19% of all WOS documents published between 2000 and 2016 are associated with one or more ORCID iDs. Author identifiers are, without doubt, a huge step in the right direction. In and of themselves, however, they do not provide a comprehensive

⁵ See <https://www.elsevier.com/solutions/scopus/content>

⁶ The same author has been issued no fewer than eight Scopus IDs, making a search for his scholarship on SCOPUS problematic as well.

response to the author uncertainty problem at this time. In this study we seek a more holistic approach.

The Use of Author Identifiers in Conjunction with Other Fielded Data

How best then to address author uncertainty? The method we propose uses author identifiers as a first step in the name consolidation process. Our approach is deductive in nature – we begin with a large dataset consisting of a given author’s true positives along with a significant amount of noise. The starting point is a WOS search for a given author’s surname, followed by a comma, followed by the first initial of his or her first name—e.g. a WOS search for author John Doe would assume the form: ‘Doe, J’. This search results in all authors with last name ‘Doe’ and first initial ‘J’. Some of these authors – e.g. Doe, Jane – are clearly not the John Doe we seek. Other author names pose a challenge, however. For example, in our search results we see the following: (i) Doe, J, (ii) Doe, JE, (iii) Doe, J E, (iv) Doe, John and (v) Doe, John E. It isn’t immediately apparent which of these refers to the John Doe we seek. By way of comparison, we note that one of the authors on this paper has more than 200 records indexed on WOS, and, among these, has the following variations in name spelling: (i) PORTER, AL, (ii) Porter, Alan L, (iii) Porter, Alan, (iv) Porter, A and (v) Porter, A L. While we know with certainty that each of the preceding five names refers to our colleague and co-author, the same conclusion is more difficult in the case of a John Doe who we’ve never met or with whom we have not interacted.

Author identifiers offer assistance for the conundrum we face in some instances, but not others. In the case of John Doe, for instance, 13% of the search results for Author=‘Doe, J’ are assigned an ORCID iD. More to the point, just one of the five preceding names that might possibly refer to the John Doe we’re interested in is assigned an ORCID iD in WOS search results (Doe,

John). It isn't currently possible to know, on the basis of ORCID identifiers alone, whether 'Doe, John' and 'Doe, J' refer to one and the same person. Given coverage issues, we must rely on additional approaches for isolating the scholarship of the John Doe we seek.

The approach advanced in this study is based on commonalities among author data in search results. We cast a broad initial net—i.e. a WOS search for a given author's last name, followed by a comma, followed by the first initial of his or her first name (e.g. 'Doe, J'). Results for this search will typically contain all of the scholarship legitimately belonging to this author in the given database (i.e. all of his or her true positives) ,along with a large amount of noise, or scholarship not belonging to this author (i.e., a large number of false positives). From this corpus we proceed to iteratively weed out false positives and retain true positives. Author identifiers provide a good starting point – if 'Doe, J' and 'Doe, John' share the same author identifier, that is sufficient for us to conclude these are the same individual. We find email addresses similarly adequate—i.e. if two author names which share the same surname and same first initial have the same email address in common, we conclude these authors are the same person. As previously noted, however, author identifier data isn't always available. The same holds true for email addresses as well. When this occurs, other fields are used to address the author uncertainty problem.

Commonalities among author data other than unique identifiers and email addresses is less conclusive for name consolidation purposes. For example, if 'Doe, John' and 'Doe, J' have an affiliation in common, do we conclude that these names belong the same person? They may or may not; affiliations have employed two or more faculty members sharing the same last and first initial. Similarly, it's conceivable that two individuals with the same last and first initial publish in

the same journal, publish with the same co-authors, and/or cite the same references. Should we then ignore commonalities among these fields and conclude they're too imprecise for name consolidation purposes? It is our position that such commonalities are indeed valuable for addressing the author uncertainty problem, but more so when used in combination. We illustrate this in the case studies that follow, but first outline the basic mechanics of the script on which our procedure is based.

When analyzing a modest number of records, manual inspection has been shown to be an effective technique (Iversen et al., 2007), but when dealing with “large-scale applications... it is necessary to automate the disambiguation process as much as possible, to keep the approach feasible and easy to maintain over time, as more and more data becomes available.” (D'Angelo et al., 2011). Our approach is somewhat of a hybrid, relying initially on author identifiers, then commonalities among fielded data other than author identifiers, and finally manual inspection. To achieve name consolidation independent of author identifier matches, we have developed a procedure that is used with bibliometric software called VantagePoint.⁷ While the application of our technique isn't dependent on VantagePoint, this is the software we find efficient in the following analysis. The script we developed to implement this procedure is made available at the VantagePoint Institute (VPI).⁸ It begins by prompting the user for a surname and a first initial of first name (for any given author of interest). It then asks the user to select a WOS field on which to consolidate author names. After this the user is prompted to point to the name of the authors field, and finally asked to identify a specific author name (referred to by the script as the primary

⁷ see www.thevantagepoint.com

⁸ see <http://vpinstitute.org/wordpress>

author) within this field whom the user knows to be a true positive (a convenient technique for this is to point to an author name associated with one of the records that has the author's ORCID iD or email address attached to it). The script proceeds to identify and combine all author names sharing the primary author's surname and first initial of his or her last name who share commonalities in the WOS field on which the user was prompted to consolidate author names. Initial dataset size is significantly reduced, and after the procedure is finished, the user is left with a much smaller (and more manageable) dataset to manually inspect.

Case Study #1: Alan Porter (Georgia Tech and Search Technology)

Alan Porter has served as Professor Emeritus at the Georgia Institute of Technology and Director of R&D at Search Technology since 2002. As of late 2017 he's been published 234 times on WOS. We know his exact publication count after asking him to confirm publications belonging to him from a WOS search for Author=Porter, A (which resulted in 3,617 records). According to the U.S. Census Bureau, the surname Porter is the 159th most common in the United States. His publication count, along with the fact there are a number of additional authors with the same name as himself indexed on WOS, make him a good case study for present purposes. A search for Professor Porter's work on WOS solely on the basis of his ORCID iD (or his ResearcherID) yields only 40% of his scholarship in WOS, which points to the drawback of relying exclusively on author identifiers for name consolidation or bibliometric purposes (at least for some authors—it's noted that other authors have full and complete ORCID coverage, but without asking them to personally verify results, it is difficult to know what their coverage is).

If we apply the match technique advanced in this paper to the Source field (in the 3,617 records from a search for Author=Porter, A) our dataset moves from 3,617 records to 2,377 (a

34% reduction). Proceeding in similar fashion for other match fields in our dataset we have the following:

MATCH FIELD	ORIGINAL DATASET SIZE	REDUCED DATASET SIZE	PERCENT REDUCTION
Source	3,617	2,377	34%
Co-authors	3,617	2,465	32%
Title	3,617	2,826	22%
ISSN	3,617	2,486	31%
Publication Year	3,617	2,905	20%
Affiliation	3,617	NA	NA
Cited References	3,617	NA	NA
Email	3,617	NA	NA
All of the above	3,617	1,750	52%

Table 1: Match Results for Alan Porter

Fielded data for Affiliation, Cited References and Email Address wasn't available, but among match fields that were available, we see that Source, Co-authors and ISSN all reduced the initial dataset by more than 30%. If we apply all of the above match fields to the initial dataset it is reduced by more than 50% and 233 out of Alan Porter's 234 true positives are retained (the one true positive that was lost occurred during the co-authorship round of reduction).

From our first case study, note that the more fielded data that is available and the higher the coverage for the same the better. Matching on the basis of email address would have been ideal, but that data wasn't available. Moreover, data for Cited References, a common match field for name disambiguation purposes, also wasn't available. Among the match fields that were available, Source, Co-authors and ISSN performed comparably (in terms of dataset reduction). We also note that while matching on the basis of individual fields reduces the initial dataset by as much as 34%, matching on the basis of multiple fields reduces the initial dataset by more than half (i.e. the more match fields used the better). The new dataset still contains a substantial

degree of noise, but the user is now approaching a much more manageable dataset to manually reduce.

Case Study #2: Zhong Lin Wang (Georgia Tech)

Zhong Lin Wang serves as the Hightower Chair in Materials Science and Engineering at the Georgia Institute of Technology. He is also a Regents' Professor. From 2009 to the present he has been published 700 times on WOS. We know his exact publication count after asking him to confirm publications from a list of WOS search results. According to the Chinese Ministry of Public Security, Wang is the most common surname in mainland China, making Professor Wang a unique challenge for name disambiguation purposes. Applying our match technique to a WOS search for him results in the following:

MATCH FIELD	ORIGINAL DATASET SIZE	REDUCED DATASET SIZE	PERCENT REDUCTION
Source	4,810	3,560	26%
Affiliation	4,810	4,147	14%
Web of Science Category	4,810	4,173	13%
Co-authors	4,810	4,175	13%
Title	4,810	3,794	21%
ISSN	4,810	3,555	26%
Publication Year	4,810	4,349	10%
Cited References	4,810	4,347	10%
Email	4,810	3,349	30%
All of the above	4,810	2,894	40%

Table 2: Match Results for Zhong Lin Wang

Fielded data was in greater supply for this case study. As might be expected, the email address field performed the best of all match fields in terms of list reduction results. Source and ISSN tied for second place. As was the case with Professor Porter, the more match fields used, the greater the total reduction. Applying our script for each of the above fields results in 40% list reduction,

while retaining 695 (99%) true positives. The retention of all true positives is particularly difficult for scholarship as voluminous as Professor Wang's given that the likelihood of misspelled, incomplete and/or mis-categorized records grows substantially.

Case Study #3: Haesun Park (Georgia Tech)

Haesun Park serves as Professor of Computational Science and Engineering at Georgia Tech. She is an Institute of Electrical and Electronics Engineers Fellow, as well as a Society for Industrial and Applied Mathematics Fellow. Professor Park has co-authored more than 100 articles in peer-reviewed journals and conferences. Given that Park is the 3rd most common surname in Korea,⁹ disambiguating scholarship associated with her name also poses a challenge. Professor Park's publications are identified from her personal website.¹⁰

A WOS CORE Collection search for AUTHOR=Park, H from 2010 to the present yields 23,298 results. We select this timeframe because the number of results grows considerably if no time constraints are used. Applying the match technique, Professor Park's WOS scholarship results in the following:

MATCH FIELD	ORIGINAL DATASET SIZE	REDUCED DATASET SIZE	PERCENT REDUCTION
Co-authors	23,298	8,527	63%
Source	23,298	14,174	39%
Affiliation (Organization Only)	23,298	20,867	10%
Title	23,298	4,978	79%
ISSN	23,298	14,762	37%
1st Author	23,298	14,791	37%
ORCID iD	23,298	3,288	86%
Researcher ID	23,298	2,903	88%
Web of Science Category	23,298	21,906	6%
Publication Year	23,298	23,094	1%
Country	23,298	23,044	1%

⁹ See https://en.wiktionary.org/wiki/Appendix:Korean_surnames

¹⁰ See <https://www.cc.gatech.edu/~hpark/>

MATCH FIELD	ORIGINAL DATASET SIZE	REDUCED DATASET SIZE	PERCENT REDUCTION
All of the above	23,298	2,319	90%

Table 3: Match Results for Haesun Park

Interestingly enough our procedure works particularly well when applied to the scholarship of Professor Park, achieving a total list reduction of 90% (while retaining all true positives). A number of explanations might account for this. Professor Park had excellent coverage for the match fields that appear in Table 3 (her mean coverage for these fields is 88%). Having a high coverage in more specific match fields (e.g. author identifier, email, co-authors, etc.) will produce more precise results and greater list reduction for the procedure we propose. In addition, our third case study had a more modest number of publications indexed on WOS. The sheer volume of research output by our first two case authors makes their list reduction more problematic. We note that volume of output on WOS is sensitive to a number of factors; one of which is the subject area in which the scholar under consideration publishes.

Results Compared

Comparing results from the preceding case studies yields the following:

	Professor Porter	Professor Wang	Professor Park
Top 3 list reduction match fields	Source (34%), Co-authors (32%) and ISSN (31%)	Email (30%), Source (26%), ISSN (26%)	Researcher ID (88%), ORCID (86%), Title (79%)
Original dataset size	3,617	4,810	23,298
Total list reduction	52%	40%	90%
Retained all true positives	no	no	yes

Table 4: Comparison of Disambiguation Results for the Three Cases

From Table 4 we note that Professor Porter and Professor Wang share Source and ISSN in their three most effective match fields, while the most effective match fields for Professor Park were

author identifiers and Title. Interestingly enough a very large number of titles (274) were used in multiple records, which might be expected when a large number of document types are present (Professor Park had 16 to be exact). We find that match field effectiveness is largely a function of coverage. Comparing original dataset size, the timeframe analyzed for each case study is not the same, nor are the subject areas in which they publish. For reasons mentioned previously our procedure is more effective when applied to our third case study, both in terms of list reduction and 100% retention of true positives (case studies #1 and #2 had a 99% retention of true positives, so not bad).

Discussion

In this study we present a practical and replicable strategy for addressing the author uncertainty problem. Put another way, this paper uses a “forward stagewise” approach by adding variables until a good combination for disambiguation is achieved. An alternative would be to use a “backward stagewise” approach that applies all the variables, then removes those with less discriminatory power. We did not use the latter approach because we found the former to yield more parsimonious combinations of fields, but future research could apply the latter and compare the approaches. While machine learning is considered authoritative by many, we don’t see it as practical or replicable. The procedure advanced in this paper is both practical, replicable and relatively user friendly. It might be categorized somewhere between ORCID and machine learning. Machine learning approaches typically look for commonalities among citation data, which isn’t always available, structured or easy to work with. The procedure we advance is intended to be applied across numerous fields in a dataset of interest (e.g. emails, co-authors, affiliations, etc.), resulting in multiple rounds of reduction. Results indicate that effective match

fields include author identifiers, emails, source titles, co-authors and ISSNs. While the script we present is not likely to result in a dataset consisting solely of true positives (at least for more common surnames), it does significantly reduce manual effort on the user's part. Results of this procedure are sensitive to a number of factors:

(i) *How common the surname under consideration is*

The more common the surname under consideration the more challenging it will be to identify a specific author of interest. While Professor Porter's initial dataset was reduced by 52%, Professor's Wang dataset reduced by 40%. This finding is unsurprising given that Wang is a relatively more common surname than is Porter. While we're satisfied that a dataset reduction of 40% was achieved for the most common surname in the most populous country, we note that more common names are likely to involve greater manual effort after our procedure is applied.

(ii) *Match field availability and coverage*

In the case of Professor Porter, fielded data were not available for Affiliation, Cited References and Email Address. Had these data been available, it is likely that the reduction of his initial dataset would have been significantly more than 52%. Moreover, among fielded data that were available, coverage was sparse at times, rendering certain fields unusable for present purposes. We recommend a coverage threshold of 80% or greater to obtain valid results.

(iii) *The subject area the author of interest is associated with*

Professor Wang publishes in subject areas typically associated with a high number of co-authors who enjoy greater publication counts than social science (and some other) disciplines. The user ought to bear in mind that, as was the case with common surnames, subject areas with higher publication counts are likely to involve more manual effort after our procedure is run. The

user ought to keep in mind as well that when evaluating a reduced dataset, if ‘Doe, John’ and ‘Doe, J’ publish in two very different subject areas they are unlikely to be the same person.

(iv) The ethnicity of the author of interest

A search for Asian names is likely to be very different than a search for Western names. The transliteration of differently spelled Asian names into the same English name, notable increase in the number of active scientists, and the sharing of a modest number of family names by a large number of scholars, all add to a number of factors that make Asian name disambiguation a true challenge (Tang and Walsh, 2010). When profiling the work of an Asian scholar using our procedure, it is advisable to use as many and as precise match fields as possible. Other notable challenges not directly addressed here include multiple spellings in many Middle Eastern names (e.g., Mohammed) and more complicated Hispanic names (multiple parts, variations in ordering).

(v) The country of the author of interest

Our procedure uses author identifiers as a starting point. As is noted by Youtie and colleagues (2017), author identifier coverage varies significantly by country. Generally speaking, coverage tends to be stronger in Europe and weaker in Asia. When analyzing research from a specific country or set of countries using the procedure advanced here, the user should be cognizant of accompanying author identifier coverage. If author identifier coverage is poor or nonexistent for a given author, other match fields can be used, but we only advocate using them if they are reasonably precise and have adequate coverage (i.e., 80% or higher).

(vi) The number and type of match fields used

As mentioned, we take the position that matching on the basis of either author identifier or email address is adequately conclusive. Matching on the basis of other fields is less straightforward, however. As a general principle we find matches made on the basis of three or more of the following to be adequate: co-authors, affiliations, journals, titles and ISSNs. For fields less precise than these (e.g. subject areas, publication years, cited authors, cited journals, keywords, etc.), we advocate a higher threshold. After matches are made, the procedure advanced in this paper will produce match results for the user's consideration. If names A and B share the same family name and first initial of first name, but publish in very different subject areas (e.g. Music and Physics), we caution that these are unlikely to be the same individual. If, however, names A and B share the same family name and first initial, as well as commonality among one or more key match fields (i.e. co-authors, affiliations, journals, titles or ISSNs), further investigation is encouraged. After all relevant matches are made, the user is advised to manually inspect results.

REFERENCES

- Cota, R. G., Gonçalves, M. A., & Laender, A. H. F. (2007). A heuristic-based hierarchical clustering method for author name disambiguation in digital libraries. In *Proceedings of the XXII Brazilian symposium on databases* (pp. 20–34). João Pessoa, Paraíba, Brazil.
- D'Angelo, C. A., Giuffrida, C., & Abramo, G. (2011). A heuristic approach to author name disambiguation in bibliometrics databases for large-scale research assessments. *Journal of the American Society for Information Science & Technology*, 62(2), 257-269. doi:10.1002/asi.21460
- Diesner, J., & Kim, J. (2016). Distortive Effects of Initial-Based Name Disambiguation on Measurements of Large-Scale Co-authorship Networks. *Journal of the Association for Information Science and Technology*, 67 (6), 1446-1461.
- Ferreira, A., Goncalves, M.A., Laender, A.H. F. (2012). A Brief Survey of Automatic Methods for Author Name Disambiguation. *Sigmod Record*, 41 (2), 15-26.
- Gasparyan A.Y., Nurmashev B., Yessirkepov M., Endovitskiy D.A., Voronov A.A., Kitas G.D. (2017). Researcher and Author Profiles: Opportunities, Advantages, and Limitations. *J Korean Med Sci.*; 32 (11): 1749-1756. <https://doi.org/10.3346/jkms.2017.32.11.1749>
- Han, H., Giles, C. L., Zha, H., Li, C., & Tsioutsoulouklis, K. (2004). Two supervised learning approaches for name disambiguation in author citations. In *Proceedings of the 4th ACM/IEEE-CS joint conference on digital libraries* (pp. 296–305). Tuscon, USA.
- Haesun Park (2006). Georgia Tech. Retrieved December 14, 2017 from <https://www.cc.gatech.edu/~hpark/>
- Iversen, E.J., Gulbrandsen, M., & Klitkou, A. (2007). A baseline for the impact of academic patenting legislation in Norway. *Scientometrics*, 70 (2), 393–414.

- Macroberts, M. H., & Macroberts, B. R. (1989). Problems of citation analysis: A critical review. *Journal of the American Society for Information Science*, 40, 342–349.
- ORCID Connecting Research and Researchers (2010). ORCID. Retrieved December 13, 2017 from <https://orcid.org>
- ResearcherID (2008). Thompson Reuters. Retrieved November 6, 2017 from www.researcherid.com
- ResearcherID (2011). Clarivate Analytics. Retrieved November 13, 2017 from <https://clarivate.com/products/researcherid>
- Smalheiser, N. R., & Torvik, V. I. (2009). Author name disambiguation. In B. Cronin (Ed.), *Annual review of information science and technology* (Vol. 43). Maryland, USA: American Society for Information Science and Technology (ASIST).
- Song, Y., Huang, J., Councill, I. G., Li, J., & Giles, C. L. (2007). Efficient topic-based unsupervised name disambiguation. In *Proceedings of the 7th ACM/IEEE joint conference on digital libraries* (pp. 342–351). Vancouver, BC, Canada.
- Tang, L., & Walsh, J.P. (2010). Bibliometric fingerprints: name disambiguation based on approximate structure equivalence of cognitive maps, *Scientometrics* DOI 10.1007/s11192-010-0196-6.
- Who uses Scopus (2015). Elsevier. Retrieved November 15, 2017 from <https://www.elsevier.com/solutions/scopus/content>
- Wiktionary (2007). Wikimedia Foundation. Retrieved December 19, 2017 from https://en.wiktionary.org/wiki/Appendix:Korean_surnames

Youtie, J., Carley, S., Porter, A.L. et al. *Scientometrics* (2017) 113: 437.

<https://doi.org/10.1007/s11192-017-2473-0>

Zhou, P. & L. Leydesdorff (2006). The emergence of China as a leading nation in science. *Research Policy*, 35 (1), 83-104.