

Finding Trends in Software Research

George Mathew^{ID}, Amritanshu Agrawal^{ID}, and Tim Menzies^{ID}, *Fellow, IEEE*

Abstract—Text mining methods can find large scale trends within research communities. For example, using stable Latent Dirichlet Allocation (a topic modeling algorithm) this study found 10 major topics in 35,391 SE research papers from 34 leading SE venues over the last 25 years (divided, evenly, between conferences and journals). Our study also shows how those topics have changed over recent years. Also, we note that (in the historical record) mono-focusing on a single topic can lead to fewer citations than otherwise. Further, while we find no overall gender bias in SE authorship, we note that women are under-represented in the top-most cited papers in our field. Lastly, we show a previously unreported dichotomy between software conferences and journals (so research topics that succeed at conferences might not succeed at journals, and vice versa). An important aspect of this work is that it is automatic and quickly repeatable (unlike prior SE bibliometric studies that used tediously slow and labor intensive methods). Automation is important since, like any data mining study, its conclusions are skewed by the data used in the analysis. The automatic methods of this paper make it far easier for other researchers to re-apply the analysis to new data, or if they want to use different modeling assumptions.

Index Terms—Software engineering, bibliometrics, topic modeling, text mining

1 INTRODUCTION

EACH year, SE researchers spend much effort writing and presenting papers to conferences, and journals. What can we learn about all those papers? What are the factors that prevent effective dissemination of research results? If we study patterns of acceptance in our SE papers, can we improve how we do, and report, research in software engineering?

Such introspection lets us answer important questions like:

- 1) *Hot topics*: What is the hottest topic in SE research? How is this list of “hot topics” changes over time?
- 2) *Breadth versus depth?*: Should a researcher focus on one particular research topic or venture across multiple topics?
- 3) *Gender Bias*: Is there a gender based bias in SE research?
- 4) *Where to publish: journals or conference?*: What would be the ideal venue for your latest research?

To answer such questions as these, this paper applies text mining and clustering using LDA (Latent Dirichlet Allocation) to 35,391 SE papers from the last 25 years published in 34 top-ranked conferences and journals. These venues are then clustered by their shared topics (and the topics found in this work as shown in Table 1). Using these information, we recommend answers to many questions, including those shown above.

In summary, this paper makes the following contributions:

- The authors are with the Department of Computer Science (CS), North Carolina State University (NCSSU), Raleigh, NC 27695 USA. E-mail: george.meg91@gmail.com, aagrava8@ncsu.edu, timm@ieee.org.

Manuscript received 30 March 2018; revised 1 September 2018; accepted 9 September 2018. Date of publication 14 September 2018; date of current version 18 April 2023.

(Corresponding author: Tim Menzies.)

Recommended for acceptance by A. E. Hassan.

Digital Object Identifier no. 10.1109/TSE.2018.2870388

- A fully automatic and repeatable process for discovering a topology between tens of thousands of technical papers.
- An open-source toolkit that other researchers can apply to other sets of documents: see goo.gl/Au1i53.
- A large database that integrates all our results into one, easy-to-query SQL database: see goo.gl/zSgb4i.
- The creation of a “reader” of the top venues in SE (see Table 2) that lists the most-cited papers in the identified topics (see Table 3). Note that, this “reader” is based only in the context of these 10 major topics. If the topic areas are shifted, or if a different source of data is used, or if a different measure for citation is used, the results can be quite different.
- A lightweight approach to gathering information about the entire SE community. Using that information, we can not only answer the four questions above, but many more like it.

Further to the last point, we can answer the four questions offered above. Regarding *hot topics*, in Table 1, we identify 10 major topics currently in software engineering; Also, we show how those groupings have changed over recent years.

Regarding *breadth versus depth*, we find that mono-focusing on a single topic can lead to fewer citations than otherwise.

Regarding *gender bias*, we have mixed news here. The percentage of women researchers in SE (22%) is much larger than other fields (i.e., 11% in mathematics and 13% in economics). Also, of the top 1000 most cited SE authors (out of 35,406), the percentage of published women is on par with the overall percentage of women in this field. That said, of the top 10 most cited authors, only one is female.

Regarding *where to publish*, we offer a previously unreported dichotomy between software conferences and journals. As shown in Section 5.4, SE conference publications tend to publish on different topics to SE journals and those

TABLE 1
The Top 7 Terms in the 10 Major SE Topics as Found
by This Research

Source Code	= code, source, information, tool, program, developers, patterns
Software process	= requirements, design, systems, architecture, analysis, process, development
Modeling	= model, language, specification, systems, techniques, object, uml
Program Analysis	= program, analysis, dynamic, execution, code, java, static
Metrics	= metrics, data, quality, effort, prediction, defect, analysis
Developer	= developer, project, bug, work, open, team, tools
Applications	= applications, web, systems, component, services, distributed, user
Testing	= test, testing, cases, fault, techniques, coverage, generation
Performance	= performance, time, data, algorithm, systems, problem, network, distributed
Architecture	= architecture, component, systems, design, product, reuse, evolution

Arranged top-to-bottom, most-to-least frequent. Terms in topics are sorted left-to-right most-to-least frequent (and are named using the most frequent terms).

conferences publications earn a significantly larger number of citations than journal articles (particularly in the last four years).

The rest of this paper is structured as follows. We start of with the description of the data used in this study from various sources and its consolidation in Section 2. This is followed by a description of how we used topic modeling to

find representative topics in SE. Next, Section 4 finds that the topics generated by topic modelling are reasonable. Hence, Section 5 goes on to highlight important case studies which can be used using the proposed method. The threats to the validity of this study is described in Section 6. Section 7 gives an overview of prior and contemporary studies based on bibliometrics and topic modeling in SE literature. Section 8 concludes this study by presenting the inferences of this study and how it can aid the different sections of SE research community.

Note that some of the content of this paper was presented as a short two-page paper previously published in the ICSE-17 companion [32]. Due to its small size, that document discussed very little of the details discussed here.

Additionally, this paper is a case study in the use of fully automated literature reviews for the analysis of a very large corpus of papers. It is important to stress this point since a mostly-manually analysis of a smaller number of papers may produce different conclusions. For example, our analysis uses citation counts from a service called CrossRef and this service can offer different citation counts to other tools such as Google Scholar. The reason we use CrossRef rather

TABLE 2
Corpus of Venues (Conferences and Journals) Studied in This Paper

Index	Short	Name	Type	Start	h5	Group
1	MDLS	International Conference On Model Driven Engineering Languages And Systems	Conf	2005	25	A1
	SOSYM	Software and System Modeling	Jour	2002	28	A2
2	S/W	IEEE Software	Jour	1991	34	B1
3	RE	IEEE International Requirements Engineering Conference	Conf	1993	20	C1
	REJ	Requirements Engineering Journal	Jour	1996	22	C2
4	ESEM	International Symposium on Empirical Software Engineering and Measurement	Conf	2007	22	D1
	ESE	Empirical Software Engineering	Jour	1996	32	D2
5	SMR	Journal of Software: Evolution and Process	Jour	1991	19	E1
	SQJ	Software Quality Journal	Jour	1995	24	E2
	IST	Information and Software Technology	Jour	1992	44	E3
6	ISSE	Innovations in Systems and Software Engineering	Jour	2005	12	F1
	IJSEKE	International Journal of Software Engineering and Knowledge Engineering	Jour	1991	13	F2
	NOTES	ACM SIGSOFT Software Engineering Notes	Jour	1999	21	F3
7	SSBSE	International Symposium on Search Based Software Engineering	Conf	2011	15	G1
	JSS	The Journal of Systems and Software	Jour	1991	53	G2
	SPE	Software: Practice and Experience	Jour	1991	28	G3
8	MSR	Working Conference on Mining Software Repositories	Conf	2004	34	H1
	WCRE	Working Conference on Reverse Engineering	Conf	1995	22	H2
	ICPC	IEEE International Conference on Program Comprehension	Conf	1997	23	H3
	ICSM	IEEE International Conference on Software Maintenance	Conf	1994	27	H4
	CSMR	European Conference on Software Maintenance and Re-engineering	Conf	1997	25	H5
9	ISSTA	International Symposium on Software Testing and Analysis	Conf	1989	31	I1
	ICST	IEEE International Conference on Software Testing, Verification and Validation	Conf	2008	16	I2
	STVR	Software Testing, Verification and Reliability	Jour	1992	19	I3
10	ICSE	International Conference on Software Engineering	Conf	1994	63	J1
	SANER	IEEE International Conference on Software Analysis, Evolution and Re-engineering	Conf	2014	25	J2
11	FSE	ACM SIGSOFT Symposium on the Foundations of Software Engineering	Conf	1993	41	K1
	ASE	IEEE/ACM International Conference on Automated Software Engineering	Conf	1994	31	K2
12	ASEJ	Automated Software Engineering Journal	Jour	1994	33	L1
	TSE	IEEE Transactions on Software Engineering	Jour	1991	52	L2
	TOSEM	Transactions on Software Engineering and Methodology	Jour	1992	28	L3
13	SCAM	International Working Conference on Source Code Analysis & Manipulation	Conf	2001	12	M1
	GPCE	Generative Programming and Component Engineering	Conf	2000	24	M2
	FASE	International Conference on Fundamental Approaches to Software Engineering	Conf	1998	23	M3

For a rationale of why these venues were selected, see Section 2. Note two recent changes to the above names: ICSM is now called ICMSE; and WCRE and CSMR recently fused into SANER. In this figure, the "Group" column shows venues that publish "very similar" topics (where similarity is computed via a cluster analysis shown later in this paper). The venues are selected in a diverse range of their h5 scores between 2010 and 2015. h5 is the h-index for articles published in a period of 5 complete years obtained from Google Scholar. It is the largest number "h" such that "h" articles published in a time period have at least "h" citations each.

Authorized licensed use limited to: NAVAL GROUP (DCNS LORIENT). Downloaded on June 25, 2025 at 14:29:07 UTC from IEEE Xplore. Restrictions apply.

TABLE 3
Most Cited Papers Within Our 10 SE Topics

Topic	Top Papers
Program Analysis	1992: Using program slicing in software maintenance; KB Gallagher, JR Lyle 2012: Genprog: A generic method for automatic software repair; C Le Goues, TV Nguyen, S Forrest, W Weimer
Software process	1992: A Software Risk Management Principles and Practices; BW Boehm 2009: Seven Process Modeling Guidelines (7PMG); J Mendling, HA Reijers, WMP van der Aalst
Metrics	1996: A validation of object-oriented design metrics as quality indicators; VR Basili, LC Briand, WL Melo 2012: A systematic literature review on fault prediction performance in software engineering; T Hall, S Beecham, D Bowes, D Gray, S Counsell
Applications	2004: Qos-aware middleware for web services composition; L Zeng, B Benatallah, AHH Ngu, M Dumas, J Kalagnanam, H Chang 2011: CloudSim: a toolkit for modeling & simulation of cloud computing; R. Calheiros, R. Ranjan, A. Beloglazov, C. De Rose, R. Buyya
Performance	1992: Spawn: a distributed computational economy; CA Waldspurger, T Hogg, BA Huberman 2010: A theoretical and empirical study of search-based testing: Local, global, and hybrid search; M Harman, P McMinn
Testing	2004: Search-based software test data generation: A survey; P McMinn 2011: An analysis and survey of the development of mutation testing; Y Jia, M Harman
Source Code	2002: CCFinder: A Multilingual Token-Based Code Clone Detection System for Large Scale Source Code; T Kamiya, S Kusumoto, K Inoue 2010: DECOR: A method for the specification and detection of code and design smells; N Moha, YG Gueheneuc, L Duchien, AF Le Meur
Architecture	2000: A Classification and Comparison Framework for Software Architecture Description Languages; N Medvidovic, RN Taylor 2010: From integration to composition: On the impact of software product lines, global development and ecosystems; J Bosch, P Bosch-Sijtsema
Modelling	1997: The model checker SPIN; GJ Holzmann 2009: The “physics” of notations: toward a scientific basis for constructing visual notations in software engineering; D Moody
Developer	2002: Two case studies of open source software development Apache and Mozilla; A Mockus, RT Fielding, JD Herbsleb 2009: Guidelines for conducting and reporting case study research in software engineering; P Runeson, M Host

The first paper in each row is the top paper since 1991 and the second paper is the top paper for the topic since 2009. For a definition of these topics, see Table 1.

than Google Scholar (or ResearchGate, or the ACM Digital Library, or a different source) is that the CrossRef Application Program Interface (API) allows ready access to 10,000+ papers. The analysis of this paper should be repeated using these other sources when they open up their APIs to large scale analysis. One reason to prefer the fully automated analysis of this paper is that such repeated analysis is possible with very little effort (and the same can not be said for the mostly manual approaches).

Further to the last point, when this paper refers to “citations”, we mean citations as reported by CrossRef.

2 DATA

For studying the trends of SE, we build a repository of 35,391 papers and 35,406 authors from 34 SE venues over a period of 25 years between 1992-2016. This time period (25 years) was chosen since it encompasses recent trends in software engineering such as the switch from waterfall to agile; platform migration from desktops to mobile; and the rise of cloud computing. Another reason to select this 25 year cut off was that we encountered increasingly more difficulty in accessing data prior to 1992; i.e., before the widespread use of the world-wide-web.

As to the venues used in this study, these were selected via a combination of on-line citation indexes (Google Scholar), feedback from the international research community (see below) and our own domain expertise:

- Initially, we selected all the non-programming language peer-reviewed conferences from the “top publication venues” list of Google Scholar Metrics (Engineering and Computer Science, Software Systems). Note that Google Scholar generates this list based on citation counts.
- Initial feedback from conference reviewers (on a rejected earlier version of this paper) made us look also at SE journals.
- To that list, using our domain knowledge, we added venues that we knew were associated with senior researchers in the field; e.g., the ESEM and SSBSE conferences.
- Subsequent feedback from an ICSE’17 companion presentation [32] about this work made us add in journals and conferences related to modeling.

This resulted in the venue list of Table 2.

For studying and analyzing those venues we construct a database of 18 conferences, 16 journals, the papers published with the metadata, authors co-authoring the papers and the citation counts from 1992-2016. Topics for SE are

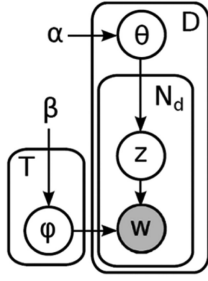


Fig. 1. LDA.

generated using the titles and abstracts of the papers published in these venues rather than the entire text of the paper. Titles and abstracts have been widely used in text based bibliometric studies [10], [19], [42] primarily due to three reasons: (a) Titles and abstracts are designed to index and summarize papers; (b) Obtaining papers is a huge challenge due to copyright violations and its limited open source access; (c) Papers contain too much text which makes it harder to summarize the content. Abstracts on the other hand are much more succinct and generate better topics.

The data was collected in five stages:

- 1) Venues are first selected manually based on top h5-index scores from Google Scholar. It should be noted that all the data collection method is automated if the source of papers from a desired venue is provided. Thus, this can be expanded to additional venues in the future.
- 2) For each venue, DOI (Document Object Identifier), authors, title, venue & year for every publication between 1992-2016 is obtained by scrapping the html page of DBLP¹. DBLP (DataBase systems & Logic Programming) computer science bibliography is an on-line reference for bibliographic information on major computer science publications. As of Jan 2017, dblp indexes over 3.4 million publications, published by more than 1.8 million authors.
- 3) For each publication, the abstract is obtained from the ACM portal via AMiner² [47] periodically on their website. Abstracts for 21,361 papers from DBLP can be obtained from the ACM dump. For rest of the papers, we use only the titles for the subsequent study.
- 4) For each publication, we obtain the corresponding citation counts using CrossRef's³ REST API. Crossref is a DOI registration agency that harness instances when the DOI was used from services like Google Scholar [2], Research Gate [3] and ACM Digital Library [1]. CrossRef was chosen for accessing the number of citations for a publication(DOI) because the DOI services mentioned earlier does not provide a public API to access the number of citations. If the publication(DOI) is registered on multiples services, the number of citations can differ based on the service accessed by CrossRef. Hence, this can lead to different results in some cases compared to complete manual analysis of this same

study using just one source. In this paper, the term "citations" refer to the citations provided by CrossRef's API (the only exception to that is in Table 2, where the 'h5' index is mentioned which uses the citations from Google Scholar).

- 5) The final stage involves acquiring the gender of each author. We used the opensource tool⁴ developed by Vasilescu et al. in their quantitative study of gender representation and online participation [53]. We used a secondary level of resolution for names that the tool could not resolve by referring the American census data⁵. Of the 35,406 authors in SE in our corpus we were able to resolve the gender for 31,997 authors.

Since the data is acquired from different sources, a great challenge lies in merging the documents. There are two major sets of merges in this data:

- *Abstracts to Papers*: A merge between a record in the ACM dump and a record in the DBLP dump is performed by comparing the title, authors, venue and the published year. If these four entries match, we update the abstract of the paper in the DBLP dump from the ACM dump. To verify this merge, we inspected 100 random records manually and found all these merges were accurate.
- *Citation Counts to Papers*: DOIs for each article can be obtained from the DBLP dump, then used to query crossref's rest API to obtain an approximate of the citation count. Of the 35,391 articles, citation counts were retrieved for 34,015 of them.

3 TOPIC MODELING

As shown in Fig. 1, the LDA topic modelling algorithm [9], [35] assumes D documents contain T topics expressed with W different words. Each document $d \in D$ of length N_d is modeled as a discrete distribution $\theta(d)$ over the set of topics. Each topic corresponds to a multinomial distribution over the words. α is the discrete prior assigned to the distribution of topics vectors(θ); and β for the distributions of words in topics(ψ).

As shown in Fig. 1, the outer plate spans documents and the inner plate spans word instances in each document (so the w node denotes the observed word at the instance and the z node denotes its topic). The inference problem in LDA is to find hidden topic variables z , a vector spanning all instances of all words in the dataset. LDA is a problem of Bayesian inference. The original method used is a variational Bayes approximation of the posterior distribution [9] and alternative inference techniques use Gibbs sampling [23] and expectation propagation [33].

There are many examples of the use of LDA in SE. For example, Rocco et al. [36] used text mining and Latent Dirichlet Allocation (LDA) for traceability link recovery. Guzman and Maleej perform sentiment analysis on App Store reviews to identify fine-grained app features [24]. The features are identified using LDA and augmented with

1. <http://dblp.uni-trier.de/>

2. <https://aminer.org/citation>

3. <https://www.crossref.org/>

Authorized licensed use limited to: NAVAL GROUP (DCNS LORIENT). Downloaded on June 25, 2025 at 14:29:07 UTC from IEEE Xplore. Restrictions apply.

4. <https://git.io/vdtLp>

5. <https://www.census.gov/2010census/>

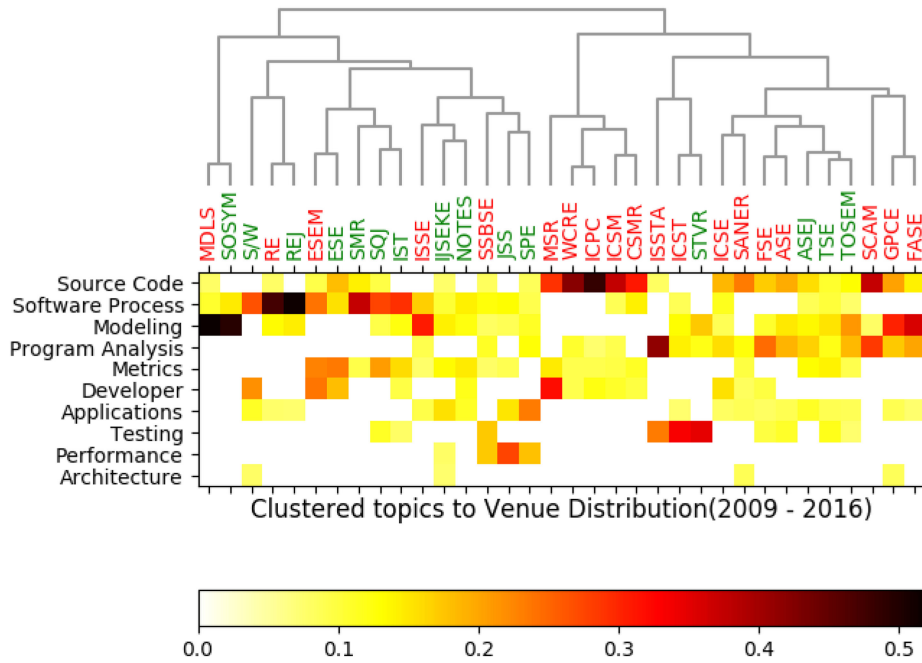


Fig. 2. Hierarchical clustering heatmap of Topics and Venues between the years 2009-2016. Along the top, **green** denotes journals and **red** denotes a conference. Along the left side, the black words are the topics of Table 1. Topics are sorted top-to-bottom, most-to-least frequent. Each cell in the heatmap depicts the frequency of a topic in a particular venue. The tree diagrams above the venues show the results of a bottom up clustering of the venues with respect to the topics. In those clusters, venues are in similar sub-trees if they have similar pattern of topic frequencies.

sentiment analysis. Thomas et al. use topic modeling in prioritizing static test cases [51].

Topic modeling is powered by three parameters; 1) k : Number of topics 2) α : Dirichlet prior on the per-document topic distributions 3) η : Dirichlet prior on the per-topic word distribution.

To set these parameters, we used *perplexity* to find the best k number of topics. Perplexity is the probability of all the words in an untrained document given the topic model, so *most* perplexity is *best*. To use perplexity, we varied topics from 2 to 50, then ran a 20 fold cross-validation study (95%/5% splits of data, clusters generated on train, perplexity assessed on test). Those runs found no significant improvement in perplexity after $k = 10$.

Next, we had to find values for $\{\alpha, \eta\}$. By default, LDA sets these to $\alpha = \frac{1}{\#topics} = 0.05$ and $\eta = 0.01$. The simplest way of tuning them would be to explore the search space of these parameters exhaustively. This approach is called grid-search and was used in tuning topic modeling [7]. The prime drawback of this approach is the time complexity and can be very slow for tuning real numbers. Alternatively, numerous researchers recommend tuning these parameters using Genetic Algorithms (GA) [29], [37], [44]. In a GA, a population of candidate solutions are mutated and recombined towards better solutions. GA is also slow when the candidate space is very large and this is true for LDA since the candidates (α and β) are real valued. Hence, inspired from the work by Fu et al. on tuning learners in SE [15], we use Differential Evolution (DE) for parameter tuning. Differential evolution randomly picks three different vectors B, C, D from a list called F (the *frontier*) for each parent vector A in F [43]. Each pick generates a new vector E (which replaces A if it scores better). E is generated as follows:

$$\forall i \in A, E_i = \begin{cases} B_i + f * (C_i - D_i) & \text{if } \mathcal{R} < cr \\ A_i & \text{otherwise} \end{cases} \quad (1)$$

where $0 \leq \mathcal{R} \leq 1$ is a random number, and f, cr are constants that represent mutation factor and crossover factor respectively (following Storn et al. [43], we use $cr = 0.3$ and $f = 0.7$). Also, one value from A_i (picked at random) is moved to E_i to ensure that E has at least one unchanged part of an existing vector. The objective of DE is to maximize the Raw Score ($\mathcal{R}_n$) which is similar to the Jaccard Similarity [16]. Agrawal et al. [4] have explained the $\mathcal{R}_n$ in much more detail. After tuning on the data, we found that optimal values in this domain are k, α & β are 11, 0.847 & 0.764 respectively.

4 SANITY CHECKS

Before delving into the results of the research it is necessary to critically evaluate the topics generated by LDA and the clustering of venues.

4.1 Are Our Topics Correct?

We turn to Fig. 2 to check the sanity of the topics. Fig. 2 shows the results of this hierarchical clustering technique on papers published between 2009 and 2016. Topics are represented on the vertical axis and venues are represented on the horizontal axis. The **green** venues represents journals and the **red** represents conferences. Each cell in the heatmap indicates the contribution of a topic in a venue. Darker colored cells indicate strong contribution of the topic towards the venue while a lighter color signifies a weaker contribution. Venues are clustered with respect to the distribution of topics using the linkage based hierarchical clustering algorithm on the *complete* scheme [34]. The linkage scheme determines the distance between sets of observations as a

function of the pairwise distances between observations. In this case we used the euclidean distance to compare the clusters with respect to the distributions of topics in each venue in a cluster. In a complete linkage scheme the maximum distance between two venues between clusters is used to group the clusters under one parent cluster.

$$\max\{\text{euclid}(a, b) : a \in \text{cluster}_A, b \in \text{cluster}_B\} \quad (2)$$

The clusters can be seen along the vertical axis of Fig. 2. Lower the height of the dendrogram, stronger the cohesion of the clusters. Several aspects of Fig. 2 suggest that our topics are “sane”. First, in that figure, we can see several examples of specialist conferences paired with the appropriate journal:

- The main modeling conference and journal (MDLS and SOSYM) are grouped together;
- The requirements engineering conference and journal are grouped together: see RE+REJ;
- The empirical software engineering conference and journals are also grouped: see ESEM+ESE;
- Testing and verification venues are paired; see ICST+STVR.

Second, the topics learned by LDA occur at the right frequencies in the right venues:

- The modeling topics appears most frequently in the modeling conference and journal (MDLAS and SOSYM);
- The software process topic appears most frequently in RE and REJ; i.e., the requirements engineering conference and journal.
- The testing topic appears most frequently in the venues devoted to testing: ISSTA, ICST, STVR;
- The metrics topic occurs most often at venues empirically assess software engineering methods: ESEM and ESE.
- The source code topic occurs most often at venues that focus most on automatic code analysis: ICSM, CSMER, MSR, WCRE, ICPC.
- Generalist top venues like ICSE, FSE, TSE and TOSEM have no outstanding bias towards any particular topic. This is a useful result since, otherwise, that would mean that supposedly “generalist venues” are actually blocking the publication of certain kinds of papers.

Third, when we examine the most-cited papers that fall into our topics, it can be seen that these papers nearly always clearly correspond to our topics names (exception: the “software process” topic, discussed below).

4.2 Are 10 Topics Enough?

After instrumenting the internals of LDA, we can report that the 10 topics in Table 1 covers over 95% of the papers. While increasing the number of topics post 10 reported no large change in perplexity, those occur at diminishingly low frequencies. As evidence of this, recall that the rows of Fig. 2 are sorted top-to-bottom most-to-least frequent. Note that the bottom few rows are mostly white (i.e., occur at very low frequency) while the upper rows are much darker (i.e., occur at much higher

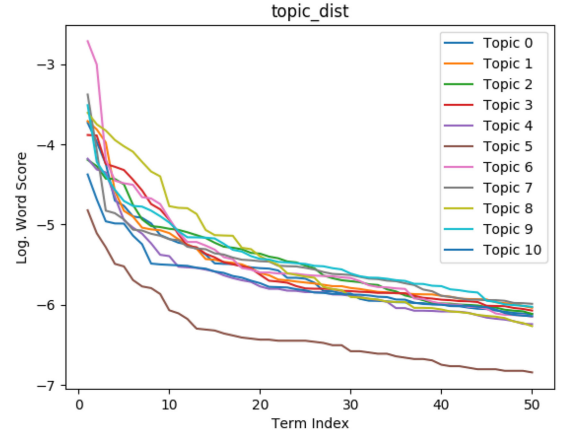


Fig. 3. Log Score of terms in each topic.

frequency). That is, if we reported *more* than 10 topics then the 11th, 12th etc would occur at very low frequencies. Thus all the less frequent topics are grouped into a single topic called Miscellaneous.

We note that this study is not the only one to conclude that SE can be reduced to 10 topics. Other researchers [10], [11], [20] also report that 90% of the topics can be approximated by about a dozen topics.

4.3 Are Our Topics Correctly Labelled?

Another question to ask is whether or not the topics of Table 1 have the correct labels. For example, we have assigned the label “Program analysis” to the list of terms “program, analysis, dynamic, execution, code, java, static”. More generally, we have labelled all our topics using first one or two words (exception: “software process”, which is discussed below). Is that valid?

We can verify this question two different ways:

- *Mathematically:* Fig. 3 shows the LDA score for each term in our topics. The x -axis orders the terms in same order as the right-hand-column of Table 1. The y -axis of that plot logarithmic; i.e., there is very little information associated with the latter words. It can be seen that terms from one topic (Miscellaneous) has relatively lower scores compared to other topics. This further justifies the grouping of all the lesser topics into a single topic.
- *Empirically:* Table 3 enlists the top papers associated with each topics. The abstracts of these papers contain the terms constituting the topics.

Overall, the evidence to hand suggests that generated labels from the first few terms is valid. That said, one topic was particularly challenging to label. For papers in the “software process” topic, we found a wide variety of research directions. For example, the title of two papers listed in this topic in Table 3 includes “software risk management” and “process modeling guidelines”. Other top-cited papers in this group discussed methods for configuration and dynamic analysis of software tools. As we read more papers from this topic, a common theme emerged; i.e., how to decide between possible parts from a set of risk, process, or product options.

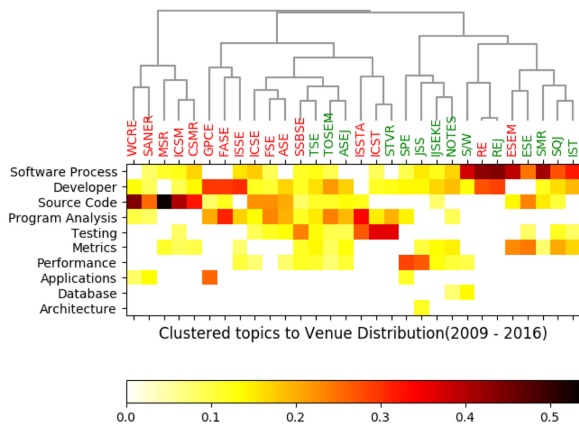


Fig. 4. Are our topics stable? Here, we show a heat map of topics in venues (2009-2016) after removing four *MDLS*, *SOSYM*, *SCAM*, *ICPC*. Note that even with those removals, we find nearly the same patterns as Fig. 2.

4.4 Do Topics Adapt to Change in Venues?

Finally we check how the topics vary when venues are added. It should be noted that no study can cover all papers from all venues. Hence, it is important to understand how the addition of more venues might change the observed structures. For our work, we kept adding new venues until our topics stabilizes. To demonstrate that stability, we show here the structures seen *before* and *after* removing the last venues added to this study:

- *MDLS* and *SOSYM*: These two venues focused extensively on modeling. It should be noted that modeling was the third most frequent topic in the literature during 2009-2016 and was primarily due to these two venues.
- *SCAM*: A venue which heavily focuses on Source Code and Program Analysis.

- *ICPC*: Another venue that focuses heavily on Source Code.

Once these venues were removed, we performed LDA using the same hyper-parameters. The resultant clustered heatmap is shown in Fig. 4. We can observe that

- 1) The topic “Modeling” is now not present in the set of topics anymore. Rather a weak cluster called “Database” is formed with the remnants.
- 2) “Software process” is now the most popular topic since none of the venues removed were contributing much towards this topic.
- 3) “Source Code” has gone down the popularity since two venues contributing towards it were removed.
- 4) “Program Analysis” stays put in the rankings but “Developer” goes above it since one venue contributing towards “Program Analysis” was removed.
- 5) The sanity clusters we identified in Section 4.2 stays intact implying that the clustering technique is still stable.

Based on these results, we can see that a) on large changes to venues contributing towards topics, the new topics are formed and the existing topics almost stays intact; b) on small changes to venues contributing towards topics, the topics either become more or less prominent but new topics are not created. Thus, we can see that the topics produced by this approach is relatively adaptable.

5 DISCUSSION

5.1 What are the Hottest Topics in SE Research?

Figs. 5 and 6 show how the topics change in conferences and journals over the last 25 years of software engineering. These are *stacked diagrams* where, for each year, the *more popular* topic appears *nearer the bottom* of each stack. Arrows are added to highlight major trends in our corpus:

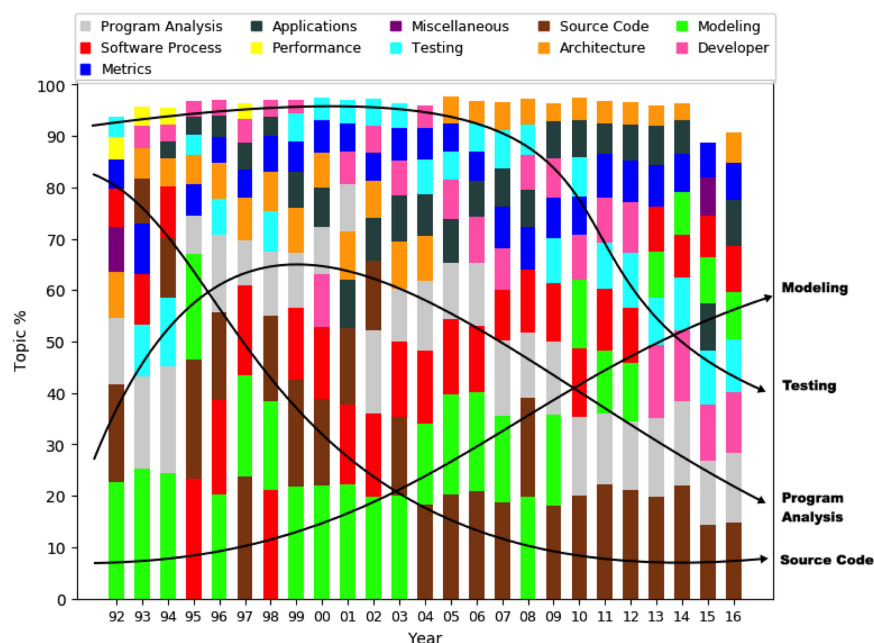


Fig. 5. Changes in **conference** topics, 1992-2016. Arrows are used to highlight the major trends. This is a *stacked diagram* where, for each year, the *most popular* topic appears at the *bottom* of each stack. That is, as topics become *more* frequent, they fall towards the *bottom* of the plot. Conversely, as topics become *less* frequent, then rise towards the *top*.

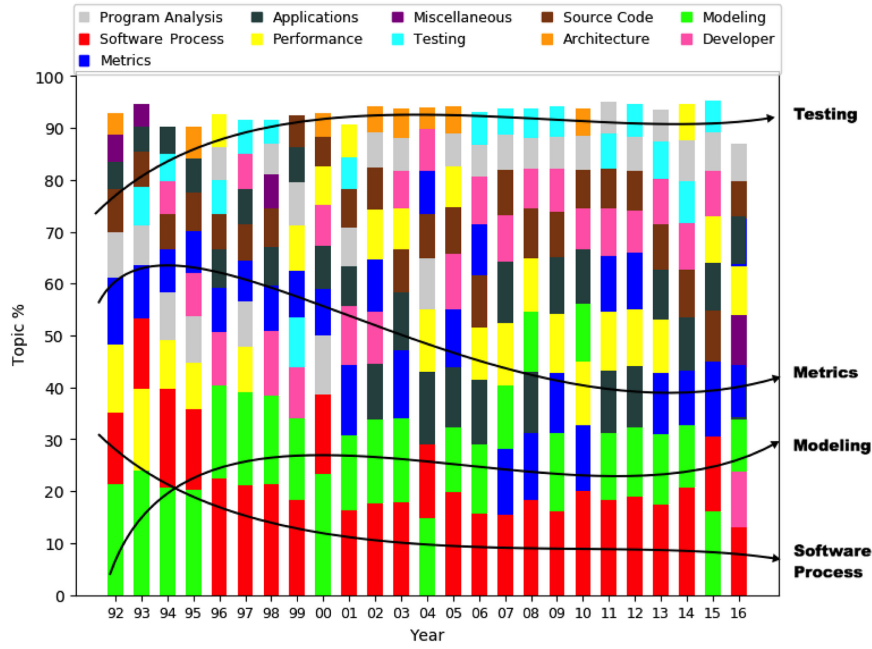


Fig. 6. Changes in **journal** topics, 1992-2016. Same format as Fig. 5. Note that some popular conference topics (e.g., source code, program analysis, testing) are currently *not* popular journal topics.

- *Downward* arrows denote topics of *increasing* popularity.;
- *Upward* arrows show topics of *decreasing* popularity.

One clear trend in both figures is that some topics are far more popular than others. For example, at the conferences (Fig. 5):

- Source code, program analysis, and testing have become very popular topics in the last decade.
- Performance and security concerns have all disappeared in our sampled venues. Performance appears, a little, in journals but is far less popular than many other topics.
- Modeling is a topic occurring with decreasing frequency in our sample.

This is *not* to say that performance, modeling and security research has “failed” or that no one works on these topics

anymore. Rather, it means that those communities have moved out of mainstream SE to establish their own niche communities.

As to the low occurrence of performance and the absence of any security terms in Table 1, this is one aspect of these results that concerns us. While we hear much talk about security at the venues listed in Table 2, Figs. 5 and 6, security research is not a popular topic in mainstream software engineering. Given the current international dependence on software, and all the security violations reported due to software errors⁶, it is surprising and somewhat alarming that security is not more prominent in our community. This is a clear area where we would recommend rapid action; e.g.,

- Editors of software journals might consider: increasing the number of special issues devoted to software security;
- Organizers of SE conferences might consider changing the focus of their upcoming events.

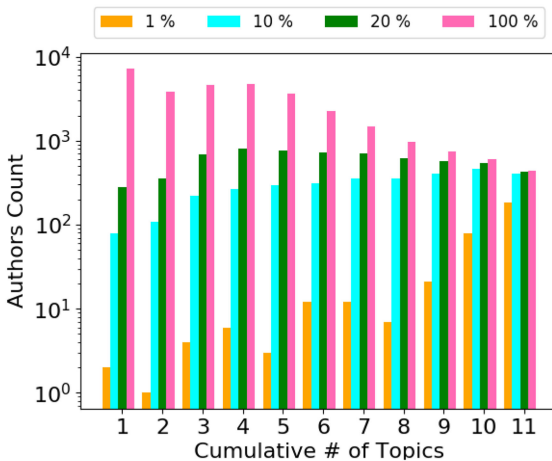


Fig. 7. Represents number of authors (on log scale) exploring different number of topics for the top 20%, 10%, and 1% authors based on citation counts.

5.2 Breadth Versus Depth?

What patterns of research are most rewarded in our field? For greater recognition, should a researcher focus on multiple topics or stay dedicated to just a few? We use Fig. 7 to answer this question. This figure compares all authors to those with 1%, 10%, and 20% of the top ranked authors based on their total number of citations. The x -axis represents the cumulative number of topics covered (ranging from 1-10). The y -axis shows the number of authors working on x_i topics where x_i is a value on the x -axis. The patterns in the figure are quite clear

- More authors from the corpus focus on fewer number of topics.

6. See <https://catless.ncl.ac.uk/Risks/search?query=software+security>

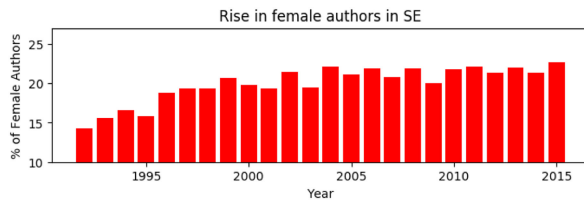


Fig. 8. % of female authors from 1991-2015. From the figure, we can see that the SE community saw a great rise in female authors until the late 90s. Since then, the percentage of women in SE has remained around 20-23%

- This trend reverses when it comes to the top authors. We can see that in the top 1% and 10% of authors (and to a lesser degree in the top 20%), more authors focus on fewer number of topics.
- If an author focuses on all the topics, almost always the author would be in the top 20% of the SE research community.
- Very few authors (3 to be precise) in the top 1% authors focus specifically on a single topic.

From this, we offer the following advice. The field of SE is very broad and constantly changing. Some flexibility in research goals tend to be most rewarding for researchers who explore multiple research directions.

5.3 Gender Bias?

Recent research has raised the spectre of gender bias in research publications [40], [49], [55]. In our field, this is a major concern. For example, a recent study by Roberts et al. show that conference reviewing methods have marked an impact on the percentage of women with accepted SE papers [40]. Specifically, they showed that the acceptance of papers coauthored by women increased after the adoption of the double blind technique.

Accordingly, this section checks if our data can comment on gender bias in SE research. The conjecture of this section is that if the percentage of women authors is $W_0\%$, then the percentage of women in top publications should also be $W_1 \approx W_0\%$. Otherwise, that would be evidence that while this field is accepting to women authors, our field is also denying those women senior status in our community.

Based on our survey of 35,406 authors in our corpus, the ratio of women in SE was increasing till 2004 and since then it has stayed around 22%. This can be inferred from Fig. 8 which is a bar chart representing percentage of women in SE each year. Thus, for this study we can say that $W_0 = 22\%$

Fig. 9 shows a line chart with the percentage of women W_1 in the top 10 to 1000 most-cited papers in SE (red solid line). The chart also shows the expected value of the percentage of women W_0 (blue dashed line). These figures gives us mixed results

- On the upside 23% of the authors in the top 100+ authors are women.
- But on the downside only 15% authors in the top 20 are women and it drops down to 10% in the top 10.

Although there is an evident bias towards the male community when it comes to the upper echelon of SE, on a larger scale there is evidence of increasing participation in

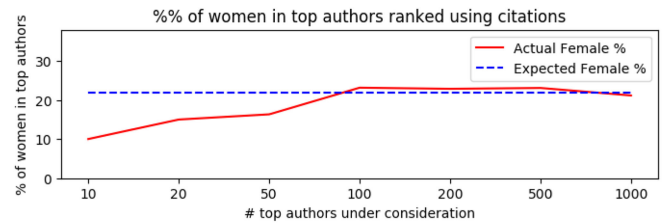


Fig. 9. % of women in top 10 to 1000 authors in SE using the total number of citations. Red solid line represents the actual value (W_1). While the blue dotted line represents the expected value (W_0) which is the percentage of women in SE.

SE research. This is observed in the rise of female authors in SE from less than 8% in 1991 to close to 22% in 2015 and over 23% of the top 100+ authors in SE being female. When compared to other traditional areas of research, the SE field is much ahead (mathematics - 11.2% and economics - 13.1%). But there is no denying that the SE community is well behind other scientific fields like veterinary medicine (35.1%) and cognitive science (38.3%) where more women are collaborating in research [55].

This section has explored one specific bias within the SE community. Apart from the above, there may well be other kinds of biases that systematically and inappropriately disadvantage stratifications within our community. One role for the tools discussed in this paper allows us to introspect and detect those other kinds of biases.

5.4 Where to Publish: Journals or Conference?

Figs. 5 and 6 report sharp differences between the kinds of topics published at SE journals and conferences. For example:

- In journals, metrics, modeling and software process appear quite often. But in conferences they appear either rarely or (in the case of modeling) remarkably decreasing frequency in the last decade.
- In conferences, source code, program analysis, and testing appear to have become very prominent in the last decade but these topics barely appear in journals.

Hence, we say that SE journals and SE conferences clearly accept different kinds of papers. Is this a problem? Perhaps so. Using our data, we can compute the number of citations per year for papers published at different venues. Consider Table 4 which shows the median (Med) and inter-quartile range (IQR) for the average number of citations per year for conferences and journals. Table 4 shows that the median of average cites per year (in the future for the papers published in a given year) has been steadily increasing since 2000. When same data from Table 4 is looked at but this time broken down into conferences and journals:

- A red background indicates when conferences are receiving more citations per year;
- A green background indicates when journals papers are receiving more citations per year;
- Otherwise, there is no statistically difference between the two distributions.

For this analysis, “more citations” means that the distributions are statistically significantly different (as judged via the 95% confident bootstrap procedure recommended by Efron & Tibshirani [13, p220-223]) and that difference is not trivially small (as judged by the A12 test endorsed by Arcuri et al. at ICSE ’11 [6]).

Note that until very recently (2009), journal papers always achieved statistically more cites per year than conference papers. However, since 2012, conference papers now receive more cites per year. These statistics justify Fernandes’ 2014 study [14] where he states that since 2008, there has been a significant increase (almost double) in the number of conference publications over journal articles.

There are many actions journal editors could undertake to mitigate this trend. For example, a recent trend in SE conferences are the presentation of “Journal-First” papers that have already appeared in journals. Since they are being seen by a larger audience, such journal-first papers might receive more citations per year. Also, once a journal becomes known for a Journal-first program, then that could increase the number of submissions to that journal.

6 THREATS TO VALIDITY

This paper is less prone to *tuning bias* and *order bias* than other software analytics papers. As discussed earlier in Section 3, we use differential evolution to find our tunings.

The main threat to validity of our results is *sampling bias*. The study can have three kinds of sampling bias

- This study only accessed the accepted papers *but not the rejected ones*. Perhaps if both accepted and rejected papers are studied, then some patterns other than the ones reported here might be discovered. That said, sampling those rejected papers is problematic. We can access 100% of accepted papers via on-line means. Sampling rejected papers imply asking researchers to supply their rejected papers. The experience in the research community with such surveys is that only a small percent of respondents reply [41] in which case we would have another sampling bias amongst the population of rejected papers. Additionally, most researchers alter their rejected papers in-line with an alternate conference/journal and make a new submission to the venue. At the time of this writing, we do not know how to resolve this issue.
- The paper uses 34 top venues considering their online citation indices and feedback from the software engineering community (see Section 2). Thus, there will always be a venue missing from such a study and can be considered. This can raise questions on the choice of the 10 topics which model SE. But, the scale of a study to consider all the publications even remotely related to SE is extremely large. Hence in Section 4.2 we show that the topics remain similar when a reduced set of venues are used. Thus, this bias can be overcome to a great extent if a diverse set of venues are considered.
- An external tool, CrossRef (see Section 2), is used to harness citations for the documents due to the

Authorized licensed use limited to: NAVAL GROUP (DCNS LORIENT). Downloaded on June 25, 2025 at 14:29:07 UTC from IEEE Xplore. Restrictions apply.

TABLE 4
Median (Med) and Inter Quartile Range (IQR=Inter-Quartile Range= (75-25)th Range) of Average Cites Per Year for Articles Published in Conferences and Journals Between 1992-2015

Year	Conference		Journal	
	Med	IQR	Median	IQR
1992	0.08	0.20	0.12	0.60
1993	0.08	0.33	0.13	0.46
1994	0.09	0.35	0.13	0.43
1995	0.27	0.64	0.14	0.45
1996	0.10	0.33	0.05	0.48
1997	0.10	0.30	0.10	0.50
1998	0.26	0.63	0.11	0.63
1999	0.06	0.28	0.06	0.50
2000	0.12	0.59	0.06	0.47
2001	0.13	0.50	0.25	1.00
2002	0.20	0.60	0.27	1.13
2003	0.29	0.71	0.36	1.14
2004	0.23	0.69	0.38	1.23
2005	0.33	0.75	0.25	0.92
2006	0.18	0.82	0.45	1.36
2007	0.40	1.00	0.50	1.50
2008	0.44	1.11	0.78	1.78
2009	0.38	1.00	0.63	1.75
2010	0.57	1.43	0.71	1.71
2011	0.67	1.50	0.67	1.83
2012	0.80	2.00	0.80	1.60
2013	0.75	2.50	0.75	1.75
2014	1.00	2.00	0.67	2.00
2015	1.00	2.00	0.50	2.00

Conference Journal

Column one colors denote the years when either Conferences or Journals received a statistically significantly larger number of citations per year.

absence of a public API for services like Google Scholar, ACM-DL and ResearchGate. Thus, this tool can provide different citations in case a document is registered on multiple services. Hence, performing a manual analysis using just one single service can recommend a different paper for the top papers listed in Table 3. That being said, performing manual analysis for over 34,000 papers where the citation counts are always increasing, escalates the time and resources of such a study.

7 RELATED WORK

To the best of our knowledge, this paper is the largest study of the SE literature yet accomplished (where “large” is measured in terms of number of papers and number of years).

Unlike prior work [20], [22], [38], [56] our analysis is fully automatic and repeatable. The same is not true for other studies. For example, the Ren and Taylor [38] method of ranking scholars and institutions incorporate manual weights assigned by the authors. Those weights were dependent on expert knowledge and has to be updated every year. Our automation means that this analysis is quickly repeatable whenever new papers are published.

7.1 Bibliometrics Based Studies

Multiple bibliometric based studies have explored patterns in SE venues over the last decades. Some authors have explored only conferences [46], [54] or journals [25] independently while some authors have explored a mixture of both [10], [14], [38].

Early bibliometric studies were performed by Ren & Taylor in 2007 where they assess both academic and industrial research institutions, along with their scholars to identify the top ranking organizations and individuals [38]. They provide an automatic and versatile framework using electronic bibliographic data to support such rankings which produces comparable results as those from manual processes. This method although saves labor for evaluators and allow for more flexible policy choices, the method does not provide a spectrum of the topics and the publication trends in SE.

Later in 2008, Cai & Card [10] analyze 691 papers from 14 leading journals and conferences in SE. They observe that 89% of papers in SE focus on 20% of the subjects in SE, including software/program verification, testing and debugging, and design tools and techniques. We repeat this study over a wider spread of publications, venues and authors (see Section 2) and the results conform with their observations.

In 2012, Hamadicharef performed a scientometrics study on IEEE Transactions of Software Engineering (TSE) for three decades between 1980-2010 [25]. He analyzes five different questions

- Number of publications: A quasi-linear growth in number of publications each year.
- Authorship Trends: It is observed that, a large number of articles in the 80s had only one or two co-authors but since the turn of the century it seemed to increase more papers having 5, 6 or 7 co-authors.
- Citations: The most cited TSE paper in this period had 811 cites and on an average a paper was cited 22.22 times with a median of 9 cites. On the other hand 13.38% papers were never cited. A larger study over 34 venues was repeated in this paper comparing the citation trends between conferences and journals (See Section 5.4).
- Geographic Trends: 46 different countries contribute to TSE between this period. 57% of the contributions to TSE come from USA, close to 21% come from Europe and less than 1% publications come from China.
- # of References: The average number of references per article increases 12.9 to 42.9 between 1980 to 2010.

A much larger study using 70,000 articles in SE was conducted in 2014 by Fernandes [14]. He observes that the number of new articles in SE doubles on an average every decade and since 2008 conferences publish almost twice as many papers as journals every year. Though the paper fails to address how the citation trends have been varying between conferences and journals and if it attributes towards the increased number of publications in conferences over journals.

More recently, Garousi and Fernandes [18] et al. performed a study based on citations to identify the top cited paper in SE. This study was based on two metrics: a) total number of citations and b) average annual number of citations to identify the top papers. The authors also go to the extent of characterizing the overall citation landscape in Software Engineering hoping that this method will encourage further discussions in the SE community towards further analysis and formal characterization of the highly-cited SE papers.

Geographical based bibliometric studies on Turkish [17] and Canadian [21] SE communities were performed by Garousi et al. to study the citation landscape in the respective countries. They identify a lack of diversity in the general SE spectrum, for example, limited focus on requirements engineering, software maintenance and evolution, and architecture. They also identify a low involvement from the industry in SE. Since these studies were localized to a certain country, it explored lesser number of papers

Citation based studies have also evolved into adoption of measures such as h-index and g-index to study the success of an SE researcher. Hummel et al. in 2013 analyzed the expressiveness of modern citation analysis approaches like h-index and g-index by analyzing the work of almost 700 researchers in SE [28]. They concluded that on an average h-index for a top author is around 60 and g-index is about 130. The authors used citations to rank authors and some researchers are apprehensive to this definition of success [12], [14].

Vasilescu et al. studied the health of SE conferences with respect to community stability, openness to new authors, inbreeding, representatives of the PC with respect to the authors' community, availability of PC candidates and scientific prestige [54]. They analyzed conference health using the R project for statistical computing to visualize and statistically analyze the data to detect patterns and trends [48]. They make numerous observations in this study

- Wide-scoped conferences receive more submissions, smaller PCs and higher review load compared to narrow scoped conferences.
- Conferences considered in the study are dynamic and have greater author turnover compared to its previous edition.
- Conferences like ASE, FASE and GPCE are very open to new authors while conferences like ICSE are becoming increasingly less open.
- Lesser the PC turnover, greater the proportion of papers accepted among PC papers.
- Narrow-scoped conferences have more representative PCs than wide-scoped ones.
- Not surprisingly, the higher the scientific impact of a conference, the more submissions it attracts and tend to have lower acceptance rates. The authors work is very detailed and gives a detailed summary of SE conferences.

Although this paper is very extensive, the authors do not explain the topics associated with venues (Sections 4.3 and 5.1) or how conferences and journals are similar/different from each other (Section 5.4).

7.2 Topic Modeling

Another class of related work are *Topic Modeling* based studies which has been used in various spheres of Software Engineering. According to a survey reported by Sun et al. [45], topic modeling is applied in various SE tasks, including source code comprehension, feature location, software defects prediction, developer recommendation, traceability link recovery, re-factoring, software history comprehension, software testing and social software engineering. There are works in requirements engineering where it was necessary

to analyze the text and come up with the important topics [8], [31], [50]. People have used topic modeling in prioritizing test cases, and identifying the importance of test cases [26], [57], [58]. Increasingly, it has also become very important to have automated tools to do SLR [52]. We found these papers [5], [30], [39] who have used clustering algorithms (topic modeling) to do SLR.

Outside of SE, in the general computer science (CS) literature, a 2013 paper by Hoonlor et al. highlighted the prominent trends in CS [27]. This paper identified trends, bursty topics, and interesting inter-relationships between the American National Science Foundation (NSF) awards and CS publications, finding, for example, that if an uncommonly high frequency of a specific topic is observed in publications, the funding for this topic is usually increased. The authors adopted a Term Frequency Inverse Document Frequency (TFIDF) based approach to identify trends and topics. A similar approach can be performed in SE considering how closely CS is related to SE.

Garousi and Mantyla recently have adopted a Topic Modeling and Word Clustering based approach to identify topics and trends in SE [19] similar to the research by Hoonlor et al. [27]. Although their method is very novel and in line with the current state of the art, but they used only the titles of the papers for modeling topics. This might lead to inaccurate topics as titles are generally not very descriptive of the field the paper is trying to explore. This issue is addressed in the current work where we use the abstracts of the paper which gives more context while building models.

In 2016, Datta et al. [11] used Latent Dirichlet Allocation to model 19000 papers from 15 SE publication venues over 35 years into 80 topics and study the half life of these topics. They coin the term “Relative Half Life” which is defined as the period between which the “importance” of the topic reduces to half. They further define two measures of importance based on the citation half life and publication half life. The choice of 80 topics is based on lowest log likelihood and although very novel but the authors do not shed light on the individual topic and the terms associated with it. Note that we do not recommend applying their kind of analysis since it lacks automatic methods for selecting the control parameters for LDA whereas our method is automatic.

8 CONCLUSION

Here, text mining methods were applied to 35,391 documents written in the last 25 years from 34 top-ranked SE venues. These venues were divided, nearly evenly, between conferences and journals. An important aspect of this analysis is that it is fully automated. Such automation allows for the rapid confirmation and updates of these results, whenever new data comes to hand. To achieve that automation, we used a topic modeling technique called LDA (augmented with an automatic assistant for tuning the control parameters called differential evolution).

Sometimes we are asked what is the point of work like this? “Researchers,” say some, “should be free to explore whatever issues they like, without interference from some

burdensome supervision body telling them what they should, or should not conduct particular kinds of research”. While this is a valid point, we would say that this actually endorses the need for the research in this paper. We fully accept and strongly endorse the principle that researchers should be able to explore software engineering, however their whims guide them. But if the structure of SE venues is inhibiting, then that structure should change. This is an important point since, as discussed above, there are some troubling patterns within the SE literature:

- There exists different sets of topics that tend to be accepted to SE conferences or journals; researchers exploring some topics contribute more towards certain venues.
- We show in Section 5.4 that a recent trend where SE conference papers are receiving significantly larger citations per year than journal papers. For academics whose career progress gets reviewed (e.g., during the tenure process; or if ever those academics are applying for new jobs), it is important to know what kinds of venues adversely affect citation counts.
- We also highlighted gender bias issues that need to be explored more in future work.

We further suggest that, this method can also be used for strategic and tactical purposes. Organizers of SE venues could use them as a long-term planning aid for improving and rationalizing how they service our research community. Also, individual researchers could also use these results to make short-term publication plans. Note that our analysis is 100% automatic, thus making it readily repeatable and easily be updated.

It must further be stressed that, a) the topics reported here should not be “set in stone” as the only topics worthy of study in SE; b) our recommendations to the SE community is purely based on the analysis of the data used in our study. Rather our goal is to report that (i) text mining methods can detect large scale trends within our community; (ii) those topics change with time; so (iii) it is important to have automatic agents that can update our understanding of our community whenever new data arrives.

REFERENCES

- [1] The ACM digital library. Accessed: Jul. 30, 2018. [Online]. Available: <https://dl.acm.org>
- [2] Google scholar. Accessed: Jul. 30, 2018. [Online]. Available: <http://scholar.google.com>
- [3] Researchgate | share and discover research. Accessed: Jul. 30, 2018. [Online]. Available: <https://www.researchgate.net/>
- [4] A. Agrawal, W. Fu, and T. Menzies, “What is wrong with topic modeling?(and how to fix it using search-based software engineering),” *Inf. Softw. Technol.*, vol. 98, pp. 74–88, 2018.
- [5] A. Al-Zubidy and J. C. Carver, “Review of systematic literature review tools,” Dept. Comput. Sci., Univ. Alabama, Tech. Rep. SERG-2014-03, 2014.
- [6] A. Arcuri and L. Briand, “A practical guide for using statistical tests to assess randomized algorithms in software engineering,” in *Proc. Int. Conf. Softw. Eng.*, 2011, pp. 1–10.
- [7] A. Asuncion, M. Welling, P. Smyth, and Y. W. Teh, “On smoothing and inference for topic models,” in *Proc. 25th Conf. Uncertainty Artif. Intell.*, 2009, pp. 27–34.
- [8] H. U. Asuncion, A. U. Asuncion, and R. N. Taylor, “Software traceability with topic modeling,” in *Proc. IEEE/ACM 32nd Int. Conf. Softw. Eng.*, 2010, pp. 95–104.

- [9] D. M. Blei, A. Y. Ng, and M. I. Jordan, "Latent dirichlet allocation," *J. Mach. Learn. Res.*, vol. 3, no. Jan, pp. 993–1022, 2003.
- [10] K.-Y. Cai and D. Card, "An analysis of research topics in software engineering–2006," *J. Syst. Softw.*, vol. 81, no. 6, pp. 1051–1058, 2008.
- [11] S. Datta, S. Sarkar, and A. Sajeve, "How long will this live? Discovering the lifespans of software engineering ideas," *IEEE Trans. Big Data*, vol. 2, no. 2, pp. 124–137, Jan. 2016.
- [12] Y. Ding, E. Yan, A. Frazho, and J. Caverlee, "Pagerank for ranking authors in co-citation networks," *J. Amer. Soc. Inf. Sci. Technol.*, vol. 60, no. 11, pp. 2229–2243, 2009.
- [13] B. Efron and R. J. Tibshirani, *An Introduction to the Bootstrap*. London, U.K.: Chapman and Hall, 1993.
- [14] J. M. Fernandes, "Authorship trends in software engineering," *Scientometrics*, vol. 101, no. 1, pp. 257–271, 2014.
- [15] W. Fu, T. Menzies, and X. Shen, "Tuning for software analytics: Is it really necessary?," *Inf. Softw. Technol.*, vol. 76, pp. 135–146, 2016.
- [16] L. V. Galvis Carre no and K. Winbladh, "Analysis of user comments: An approach for software requirements evolution," in *Proc. IEEE Int. Conf. Softw. Eng.*, 2013, pp. 582–591.
- [17] V. Garousi, "A bibliometric analysis of the turkish software engineering research community," *Scientometrics*, vol. 105, no. 1, pp. 23–49, 2015.
- [18] V. Garousi and J. M. Fernandes, "Highly-cited papers in software engineering: The top-100," *Inf. Softw. Technol.*, vol. 71, pp. 108–128, 2016.
- [19] V. Garousi and M. V. Mäntylä, "Citations, research topics and active countries in software engineering: A bibliometrics study," *Comput. Sci. Rev.*, vol. 19, pp. 56–77, 2016.
- [20] V. Garousi and G. Ruhe, "A bibliometric/geographic assessment of 40 years of software engineering research (1969–2009)," *Int. J. Softw. Eng. Knowl. Eng.*, vol. 23, no. 09, pp. 1343–1366, 2013.
- [21] V. Garousi and T. Varma, "A bibliometric assessment of canadian software engineering scholars and institutions (1996–2006)," *Comput. Inf. Sci.*, vol. 3, no. 2, 2010, Art. no. 19.
- [22] R. L. Glass and T. Y. Chen, "An assessment of systems and software engineering scholars and institutions (1999–2003)," *J. Syst. Softw.*, vol. 76, no. 1, pp. 91–97, 2005.
- [23] T. L. Griffiths and M. Steyvers, "Finding scientific topics," *Proc. Nat. Acad. Sci. USA*, vol. 101, no. suppl 1, pp. 5228–5235, 2004.
- [24] E. Guzman and W. Maalej, "How do users like this feature? A fine grained sentiment analysis of app reviews," in *Proc. IEEE 22nd Int. Requirements Eng. Conf.*, 2014, pp. 153–162.
- [25] B. Hamadicharef, "Scientometric study of the ieee transactions on software engineering 1980–2010," in *Proc. 2nd Int. Congr. Comput. Appl. Comput. Sci.*, 2012, pp. 101–106.
- [26] H. Hemmati, Z. Fang, and M. V. Mantyla, "Prioritizing manual test cases in traditional and rapid release environments," in *Proc. IEEE 8th Int. Conf. Softw. Testing, Verification Validation*, 2015, pp. 1–10.
- [27] A. Hoonlor, B. K. Szymanski, and M. J. Zaki, "Trends in computer science research," *Commun. ACM*, vol. 56, no. 10, pp. 74–83, 2013.
- [28] O. Hummel, A. Gerhart, and B. Schäfer, "Analyzing citation frequencies of leading software engineering scholars," *Comput. Inf. Sci.*, vol. 6, no. 1, 2013, Art. no. 1.
- [29] S. Lohar, S. Amornborvornwong, A. Zisman, and J. Cleland-Huang, "Improving trace accuracy through data-driven configuration and composition of tracing features," in *Proc. 9th Joint Meeting Found. Softw. Eng.*, 2013, pp. 378–388.
- [30] C. Marshall and P. Brereton, "Tools to support systematic literature reviews in software engineering: A mapping study," in *Proc. IEEE Symp. Empirical Softw. Eng. Meas.*, 2013, pp. 296–299.
- [31] A. K. Massey, J. Eisenstein, A. I. Antón, and P. P. Swire, "Automated text mining for requirements analysis of policy documents," in *Proc. IEEE 21st Int. Requirements Eng. Conf.*, 2013, pp. 4–13.
- [32] G. Mathew, A. Agrawal, and T. Menzies, "Trends in topics at se conferences (1993–2013)," in *Proc. IEEE 39th Int. Conf. Softw. Eng.*, 2017, pp. 397–398.
- [33] T. Minka and J. Lafferty, "Expectation-propagation for the generative aspect model," in *Proc. 18th Conf. Uncertainty Artif. Intell.*, 2002, pp. 352–359.
- [34] D. Müllner, "Modern hierarchical, agglomerative clustering algorithms," 2011, *arXiv:1109.2378*.
- [35] S. I. Nikolenko, S. Koltcov, and O. Koltsova, "Topic modelling for qualitative studies," *J. Inf. Sci.*, vol. 43, pp. 88–102, 2017.
- [36] R. Oliveto, M. Gethers, D. Poshyvanyk, and A. De Lucia, "On the equivalence of information retrieval methods for automated traceability link recovery," in *Proc. IEEE Int. Conf. Prog. Comprehension*, 2010, pp. 68–71.
- [37] A. Panichella, B. Dit, R. Oliveto, M. Di Penta, D. Poshyvanyk, and A. De Lucia, "How to effectively use topic models for software engineering tasks? An approach based on genetic algorithms," in *Proc. IEEE Int. Conf. Softw. Eng.*, 2013, pp. 522–531.
- [38] J. Ren and R. N. Taylor, "Automatic and versatile publications ranking for research institutions and scholars," *Commun. ACM*, vol. 50, no. 6, pp. 81–85, 2007.
- [39] A. Restificar and S. Ananiadou, "Inferring appropriate eligibility criteria in clinical trial protocols without labeled data," in *Proc. ACM 6th Int. Workshop Data Text Mining Biomed. Inform.*, 2012, pp. 21–28.
- [40] S. G. Roberts and T. Verhoef, "Double-blind reviewing at evolang 11 reveals gender bias," *J. Lang. Evol.*, vol. 1, no. 2, pp. 163–167, 2016.
- [41] G. Robles, J. M. González-Barahona, C. Cervigón, A. Capiluppi, and D. Izquierdo-Cortázar, "Estimating development effort in free/open source software projects by mining software repositories: A case study of openstack," in *Proc. 11th Work. Conf. Mining Softw. Repositories*, 2014, pp. 222–231.
- [42] S. Sarkar, R. Lakdawala, and S. Datta, "Predicting the impact of software engineering topics: An empirical study," in *Proc. 26th Int. Conf. World Wide Web Companion*, 2017, pp. 1251–1257.
- [43] R. Storn and K. Price, "Differential evolution—a simple and efficient heuristic for global optimization over continuous spaces," *J. Glob. Optim.*, vol. 11, no. 4, pp. 341–359, 1997.
- [44] X. Sun, B. Li, H. Leung, B. Li, and Y. Li, "MSR4SM: Using topic models to effectively mining software repositories for software maintenance tasks," *Inf. Softw. Technol.*, vol. 66, pp. 1–12, 2015.
- [45] X. Sun, X. Liu, B. Li, Y. Duan, H. Yang, and J. Hu, "Exploring topic models in software engineering data analysis: A survey," in *Proc. IEEE/ACIS 17th Int. Conf. Softw. Eng., Artif. Intell., Netw. Parallel Distrib. Comput.*, 2016, pp. 357–362.
- [46] T. Systä, M. Harsu, and K. Koskimies, "Inbreeding in software engineering conferences," 2012. [Online]. Available: <https://citeseerx.ist.psu.edu/document?repid=rep1&type=pdf&doi=5a99289403465289a2a5d9dee42599daf35a7aa4>
- [47] J. Tang, J. Zhang, L. Yao, J. Li, L. Zhang, and Z. Su, "Arnetminer: Extraction and mining of academic social networks," in *Proc. Int. Conf. Knowl. Discov. Data Mining*, 2008, pp. 990–998.
- [48] R Core Team, "R: A language and environment for statistical computing," *R Found. Stat. Comput.*, 2018. [Online]. Available: <https://www.R-project.org/>
- [49] J. Terrell et al., "Gender differences and bias in open source: Pull request acceptance of women versus men," *PeerJ Comput. Sci.*, vol. 3, 2017, Art. no. e111.
- [50] S. W. Thomas, B. Adams, A. E. Hassan, and D. Blostein, "Studying software evolution using topic models," *Sci. Comput. Program.*, vol. 80, pp. 457–479, 2014.
- [51] S. W. Thomas, H. Hemmati, A. E. Hassan, and D. Blostein, "Static test case prioritization using topic models," *Empirical Softw. Eng.*, vol. 19, no. 1, pp. 182–212, 2014.
- [52] G. Tsafnat, P. Glasziou, M. K. Choong, A. Dunn, F. Galgani, and E. Coiera, "Systematic review automation technologies," *Systematic Rev.*, vol. 3, no. 1, 2014, Art. no. 1.
- [53] B. Vasilescu, A. Capiluppi, and A. Serebrenik, "Gender, representation and online participation: A quantitative study," *Interacting Comput.*, vol. 26, no. 5, pp. 488–511, 2013.
- [54] B. Vasilescu, A. Serebrenik, T. Mens, M. G. van den Brand, and E. Pek, "How healthy are software engineering conferences?," *Sci. Comput. Program.*, vol. 89, pp. 251–272, 2014.
- [55] J. West and J. Jacquet, *Women as Academic Authors, 1665–2010*. Washington, D.C., USA: The Chronicle of Higher Education, 2012.
- [56] C. Wohlin, "An analysis of the most cited articles in software engineering journals–2000," *Inf. Softw. Technol.*, vol. 49, no. 1, pp. 2–11, 2007.
- [57] G. Yang and B. Lee, "Predicting bug severity by utilizing topic model and bug report meta-field," *KIISE Trans. Comput. Pract.*, vol. 21, no. 9, pp. 616–621, 2015.
- [58] Y. Zhang, M. Harman, Y. Jia, and F. Sarro, "Inferring test models from kate's bug reports using multi-objective search," in *Search-Based Software Engineering*. Berlin, Germany: Springer, 2015, pp. 301–307.



George Mathew is currently working toward the third year PhD degree with the Department of Computer Science, North Carolina State University. With industrial experience with Facebook, CrowdChat and Microsoft, He explores automated program repair, multi-objective optimization and applied machine learning algorithms in software engineering. For more information on his research and interests, For more information, please visit <http://bigfatnoob.us>



Amritanshu Agrawal is currently working toward the fourth year PhD degree with the Department of Computer Science, North Carolina State University, Raleigh, NC. His research involves defining better, and simpler, quality improvement operators for software analytics. He aims at providing better next generation AI software to software engineers by making them literate with AI. He has industrial experiences with IBM, Lucidworks and LexisNexis. For more, For more information, please visit <http://www.amritanshu.us>.



Tim Menzies (Fellow, IEEE) received the PhD degree from UNSW, in 1995. He is a full professor in CS with NC State University, where he explores SE, data mining, AI, search-based SE, programming languages and open access science. He is the author of more than 250 referred publications and was co-founder of the ROSE festivals and PROMISE conferences devoted to reproducible experiments in SE (<http://tiny.cc/seacraft>). He does, or has, served, as associated editor of many journals: *IEEE Transactions on Software Engineering*, *ACM Transactions on Software Engineering and Methodology*, *Empirical Software Engineering*, *Journal of the American Society of Echocardiography*, *Information Software Technology*, *IEEE Software*, and *Software Quality Journal*. For more, For more information, please visit <http://menzies.us>.

▷ **For more information on this or any other computing topic, please visit our Digital Library at www.computer.org/csdl.**