

On Research Quality and Impact: What Five Decades in Academia Has Taught Me

Behrooz Parhami

Department of Electrical and Computer Engineering, University of California, Santa Barbara, USA

Abstract

Research quality and impact need to be assessed for various purposes, including promotions or salary raises at universities, advancements at research labs, awarding of grants by research-funding organizations, and bestowing scientific/technical honors. Research must ultimately be at service to the society that supports it, but there is more to evaluating research than assessment of its immediate social impact. Research programs within and between various disciplines are interconnected, with the chain-link from a particular piece of research often going through multiple projects (and sometimes human generations) before it connects to a practical application. In scientific fields, research is expected to validate or challenge proposed theories of how our world works, and therein lies its impact. In technical fields, research is at the service of building better devices, systems, and processes, through the derivation of innovative designs or more accurate modeling and evaluation methods. In this paper, I present some thoughts on factors used to assess science and technology research efforts and their impact.

Keywords: Assessment, Citation, Conference, Impact, Journal, Paper, Productivity, Promotion, Scholarship.

1. Introduction

I have been involved in research quality and impact assessment for almost 50 years now. My experiences include membership/leadership on university merits-and-promotions committees (both in Iran and in the US), three years as UCSB College of Engineering's Associate Dean for Academic Personnel, extensive reviewing of research papers and grant proposals, and serving on multiple editorial boards for IEEE and other conference sponsors and journal publishers.

Here, I share some of my concerns, ideas, and experiences, as well as the dangers lurking ahead for my younger colleagues in academia and other places where research efforts and output must be evaluated.

Each discipline has its own criteria for evaluating the quality and impact of work. In way of analogy, the film industry has various awards and numerical metrics, such as ticket sales/revenues. The publishing industry also has awards and honours, as well as placement on best-sellers list and sales figures in hardback, paperback, audiobook, and other formats. Likewise, prestigious awards help assess research quality at

the highest end of the spectrum, but we need metrics that apply in a wider range.

This short article is not meant to be exhaustive. Perhaps I will expand it into a comprehensive treatise, with more details on the topics discussed and a compilation of best practices, later. For now, I will provide brief observations on five main topics: Assessing research quality; Citation metrics; Quality of conferences; Journal impact factor; and Fake Journals.

2. Assessing Research Quality

The only sure way to assess the quality and impact of a research paper, project, or program is to have other specialists in the exact same field evaluate it. Unfortunately, this is highly impractical. Often, an institution where assessment must be done for promotion or merit advancement lacks other people in the same exact field of research. Even when there is such a colleague, in-house evaluation of work is often frowned upon. Finding external reviewers who agree to do the evaluation is extremely hard, given the busy schedules of prominent researchers. Writing research grant proposals and struggling

in a publish-or-perish environment leaves very little time for community service.

With senior researchers too busy to accept refereeing or other research-assessment invitations, universities and journals are often forced to go to lower-ranked or less-qualified evaluators. The result is sometimes catastrophic or embarrassing. I often find myself cringing when reading a review or assessment that is filled with language errors and logical fallacies. The effect of low-quality reviews is amplified in highly competitive environments, leading to poor outcomes [9].

This is why it is of utmost importance to teach young researchers the proper way of evaluating a research paper (see Fig. 1, and [11]). Because in key academic hiring and promotion decisions, letters of recommendation play a key role, graduate students must be introduced to the process, with examples, as part of their training [17].

Even when qualified internal or external reviews are obtained, a phenomenon similar to grade-inflation for students can arise. Researcher X bestows lavish praise on researcher Y, with the expectation of reciprocation by Y or his/her supervisor in future. This is something I have observed directly in my various formal positions and informal advising roles, as detailed below.

If you become privy to the praise or credit given to multiple individuals for a particular contribution, you are left with the impression that each one was single-handedly responsible for the advance. Some four decades ago, I helped institute at Arya-Mehr/Sharif University of Technology a requirement that each evaluatee's submitted material include a percentage figure for every paper, indicating his/her share of the credit, with the percentages expected to add up to 100%. This approach created the need for a dialogue among the co-authors on how much credit each one should claim. It also eliminated or reduced the practice of listing a co-author who did not make any contribution to the research, as a favor to him/her. The system wasn't perfect, and it did lead to some infighting, but it was the best we could do at the time.

To prevent the embarrassing situation when each co-author claims full credit for the ideas in a paper (confidentiality of reviews and assessments thwarts the exposure of such fraud), many journals are now requiring a statement within the paper about the role played by and

contributions of each co-author. Here is an example statement: "W conceived the project. X and Y designed the experiments. W and Z performed the experiments. X and Y analysed the data. W wrote the article's first draft and revised it after receiving input from all the others."

This scheme is significantly more detailed but much harder to quantify. Asking for a single percentage figure makes the assessment task easier, but, like all single numerical indicators, quantifying the contribution with a lone fractional number has its drawbacks.

Aiming to characterize the contribution of each co-author, whether qualitatively or quantitatively, does not address the research quality question. As a substitute for directly assessing the quality of a piece of research, one may use indirect metrics. Citations and the reputation of publication venues are possible surrogate measures. In the next three sections, I discuss the use of citations, quality of conferences, and journal reputations as indirect measures of the quality of published research. I then tackle the relatively new phenomenon of fake journals providing easy/fast publication venues, while pretending to have high quality and genuine peer review.

Given wide variations in refereeing quality, even for respected journals and conferences, any aggregate measure of quality or impact will be imprecise. Yet, in the absence of better alternatives, such aggregate measures are widely used. Research assessment procedures are like tax laws, in that the introduction of excessive complexity in an attempt to ensure fairness invites misunderstanding and abuse. So, simplification isn't always bad!

3. Citation Metrics

One indicator of the impact of a paper is how many times it is cited by other researchers. I say "by other researchers," because self-citations are often discounted in citation analysis. Number of citations must be used with care in evaluating research impact. Here are some of the main caveats:

- Citation frequencies vary across (sub)disciplines
- Books tend to garner more citations than papers
- Survey/tutorial articles often get more citations
- Specialized papers tend to get fewer citations
- Hot/fashionable topics can get undeserved citations
- Famous researchers tend to be cited more
- More-recent publications have fewer citations

To the general considerations above one must add the possibility of fraud and abuse. Researchers can inflate their citation counts by scheming to provide mutual or circular citations.

Studies showing citation metrics to be non-robust, in the sense of being influenced by a number of peripheral, non-scholarly factors, are abundant. I will cite two such studies very briefly as cautionary tales for the use of citation metrics. In one study of two leading citation-indexing systems, Scopus and Web of Science, a negative correlation was observed between the number of hyphens appearing in a paper's title and its citation count [20]. Another study concluded that male authors loading their abstracts with words like 'novel,' 'unique' and 'excellent,' tend to generate more citations by peers [7], exacerbating an already-significant bias against citing women authors [14].

A. Theoretical/Conceptual Soundness • Is the theory, if any, behind the research logically applied and thoroughly justified? • Do the authors correctly interpret and appropriately synthesize relevant prior work? • Are the hypotheses, if any, derived from the theory to be tested, clearly stated? B. Methodological Soundness • Is the research design (sample, methods, measures) appropriate for the problem? • Are appropriate analytical techniques applied to the data collected? • Are the results of the study correctly interpreted? • Are conclusions and/or implications properly derived from the research findings? C. Contribution • Does the article advance knowledge within/about the discipline? • Are the findings and their implications noteworthy, rather than routine drudgery? • Is the article of interest to many people, or at least one segment/group, in the field? D. Communication • Is the article clearly written (simple, direct language, free from embellishments)? • Is the article laid out in a logical format, with proper groupings of related ideas? • Are the major points easily grasped, with clear delineation of what's new?

Fig. 1 . Guidelines for article evaluation [8]

Despite the inevitable drawbacks, citation count does provide a good metric at the macro level. An article with 100 citations almost certainly has higher impact than one with 5 citations. The problem arises when one tries to make very fine distinctions between papers with 10 vs. 15 citations, say. As to where to get citation data, Google Scholar has emerged in recent years as a comprehensive and accessible source. Web of Science and SciVerse Scopus are other broad-based databases. Discipline-specific databases, such as SciFinder Scholar (chemistry), PsychInfo, and PubMed also exist. Besides use for assessing the quality of a particular paper or piece of research, citation counts can be used to assess the impact of a researcher, a research group, or an organization, such as a university or research center.

For an individual researcher, one can order his/her publications in descending order of the number of citations, with the most-cited work appearing first. Consider the following hypothetical list for Researcher X, with a letter identifying a publication followed by the number of its citations:

X: A 150; B 122; C 70; D 56; E 26; F 8; G 6; H 3; ...

This particular researcher has a few high-impact papers (those with 100+ citations) but the number of citations drops after a handful of items at the top. Now, consider Researcher Y with the following citations profile:

Y: P 45; Q 41; R 36; S 33; T 30; U 28; V 25; W 23; ...

Even though the highest citation number is smaller for Researcher Y, the overall impact is intuitively higher, as it is spread over a larger number of publications. This broader impact is sometimes measured by h-index, defined as the top-cited h publications having at least h citations each (Fig. 2). For Researcher X above, the h-index is 6, because the first 6 publications have 6 or more citations (it isn't 7, because the top 7 publications do not have 7 or more citations each). The h-index for Researcher Y is at least 8; we don't know for sure, because the top 8 most-cited publications have 23 or more citations each, so if the pattern continues with 20, 18, ... citations, the h-index can be much larger.

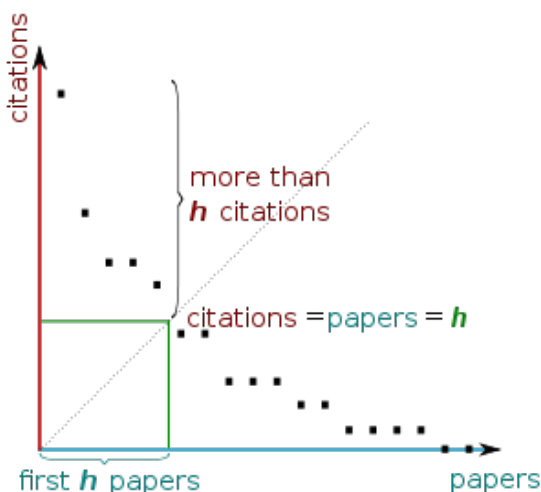


Fig. 2. Graphical illustration of h -index. (Wikipedia)

Google Scholar compiles both a general h-index and a "recent h-index" for publications over the past 5 years. The latter number is often much lower, in view of both the limited time span included and the smaller elapsed time since publication, but it indicates the impact of a researcher's recent work. For a senior researcher, a large h-index may represent citations of his/her early work published decades ago. A large recent-h-index indicates ongoing excellence, because it points to the continued impact of a researcher's work.

A third metric supplied by Google Scholar is a researcher's i10-index, the number of publications with at least 10 citations. Ten is a somewhat arbitrary threshold for distinguishing low-impact from high-impact publications. The vast majority of papers published garner citations in the single digits, so, if a researcher has an i10 index of 50, say, it represents a healthy impact over a substantial number of publications.

To summarize, a Google Scholar h-index of 25, say, shows significant impact, as it indicates that 25 of the researcher's publications have garnered 25 or more citations. The indices just discussed can be defined for groups of researchers in a department, university, and so on, but such uses are less common.

4. Quality of Conferences

In many scientific and technical disciplines, there are two main venues for publishing research results: Conferences (a catch-all term that also covers congresses, conventions, symposia, and workshops) and journals. Distinguishing and appropriately weighing a researcher's publications in conferences and journals has been a source of tension and problems in some disciplines.

In long-established scientific disciplines, the two kinds of publication venues are complementary, not competing. Conferences are viewed as venues for reporting breaking or in-progress research, which will then be published in an appropriate journal, once it is refined as a result of additional research and peer review. In such disciplines, conferences may have little or no peer review, and are viewed primarily as places to network and rub shoulders with leaders of one's discipline. It's not unusual to have in such networking-focused conferences, often co-located with job fairs, publisher book/periodical booths, and industrial expositions, hundreds of presentations and/or posters, in several parallel tracks, with each participant choosing to attend talks in at most a couple of the tracks best matched to his/her research interests.

Computer science and engineering is one of the few exceptions to the scheme described above. There are quite a few computing conferences with rigorous peer-review processes (at least in theory) and very low acceptance rates. It is common for a paper published in the proceedings of such a conference to never appear in a journal. Proceedings of such conferences are indexed and are widely available, much like top-tier journals.

I have written elsewhere on the dangers of relying on conference publications as archival records of research. First, low acceptance rate is not synonymous with high quality of accepted papers [9]. Faced with impossible deadlines, referees place an inordinate amount of emphasis on the authors' reputation rather than the quality of the specific paper being considered. Over time, such conferences tend to develop a

closed circle of contributors and committees, making it difficult for newcomers to make headway in contributing papers or participating in running the conference. Another side-effect of tight deadlines is authors not getting a chance to rebut referee evaluations or to apply revisions. With very few exceptions, conferences render an accept/reject decision unilaterally, without a revision round. The typical selection process entails the referees being asked to rate submissions numerically, with the numbers averaged to form a rank-ordered list. Then, papers at the very top are accepted and those at the very bottom rejected automatically, without further discussion. In the case of journals, there is a dialog, sometimes in multiple rounds, between authors and anonymous reviewers, which leads to quality improvement.

We see in Fig. 3 the possible harmful impact of low acceptance rates in the face of error-prone referee evaluations, which is inevitable with tight conference deadlines. When $k = 3$ referees evaluate each paper and each one has an error rate of 32%, say, we see from the heavy black dot in Fig. 3 that, with an acceptance rate of 10%, an equal number of good and bad papers will be accepted. Increasing the acceptance rate to 20%, say, actually improves the fraction of accepted papers that are good! For precise definitions of good and bad papers, please refer to the paper by Parhami (2016).

An analogy may help clarify the notions above without reading the original paper by Parhami (2016). Suppose you want to select people afflicted with disease D to attend an informational gathering. Let's say 10% of the population has disease D . Having disease D is the analog of a paper being in the top 10%, ideally chosen for conference presentation. Refereeing is like an imperfect diagnostic test with 10%, say, false positives and false negatives. There are 200 papers submitted. A perfect paper selection process identifies the 20 best papers. A selection process with 10% false positives/negatives will pick $0.9 \times 20 = 18$ of the good papers along with $0.10 \times 180 = 18$ of the bad ones. Since all 36 papers passing this phase are generally assessed very positively by referees, whatever secondary criteria you use to cut down the

number from 36 to 20, for a 10% acceptance rate, is likely to choose as many bad papers as good papers.

The concepts above are depicted graphically in Fig. 4. The population in our example consists of 200 submitted papers, represented by the orange blob. The top 10% of submissions (or fraction of the sick population) is the region on the right, demarcated by the vertical line segment. The test (refereeing) is imprecise, with false positives (bad papers selected for presentation) and false negatives (good papers rejected). A small false-positive fraction leads to the possibility of many bad papers being selected for presentation.

The move within computer science and engineering (CSE) to the conference-centered publication model has been justified by the fast-moving nature of the field that makes short-turnaround publication highly desirable. It was noted by proponents of using conferences as the primary publication venue that some journals of the field had turnaround times measured in years, not weeks or even months, making archival publications stale or perhaps superseded by newer results.

In retrospect, using conferences as the primary publication venues in CSE was the wrong solution to the very real problem of slow turnaround for research publications. A better path would have been to make journals more responsive to the demands of a fast-moving field. It is noteworthy that some of the most-prestigious science journals, such as *Science* and *Nature*, have turnaround times of days to weeks. Fortunately, a movement toward the goal of reducing turnaround times in CSE journals seems to be afoot. More incentives for referees to submit timely reviews are needed.

5. Journal Impact Factor

The impact-factor (IF) or journal-IF (JIF) metric was proposed by ISI founder Eugene Garfield [1][4] to measure the quality of scientific journals and the research published therein. In 1964, Garfield's Institute for Scientific Information (ISI) published the first Science Citation Index (SCI) spanning 2200 journals. Eight years later, Garfield released his first set of journal impact factors in another article [5]

A particular journal's IF for year y , calculated and reported in year $y + 1$, is defined as total citations in year y to papers the journal published in years $y - 1$ and $y - 2$, divided by the total number of papers published in those two years. The impact factor of the prestigious journal *Nature* for 2017 based on 2017 citations to 2016 and 2015 papers is computed as: $(32,389 + 41,701)/(880 + 902) = 41.6$ [18]. This is an

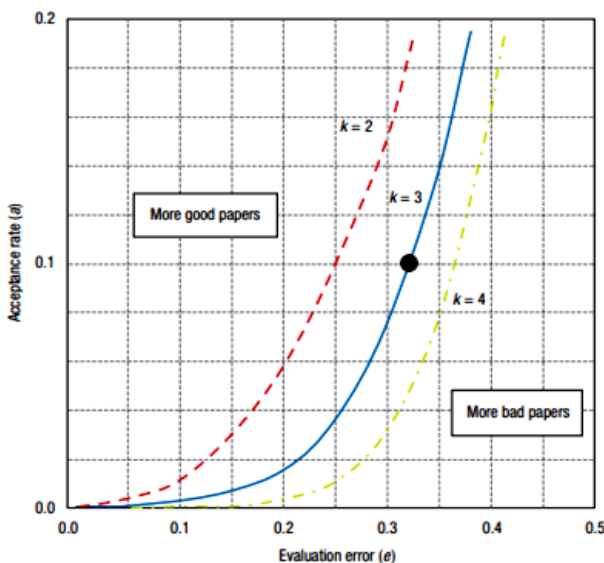


Fig 3. Graphical illustration of the possible harmful effect of low acceptance rates. [9]

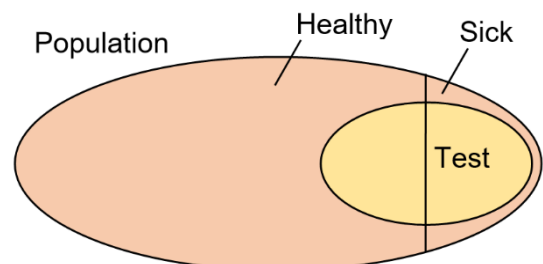


Fig. 4. Illustration of test specificity vs. sensitivity

impressive number, as JIFs of specialized journals tend to be fairly small numbers, sometimes even falling below 1.

Definition and use of JIF is a brilliant idea, but like everything else in the world, it is subject to abuse and manipulation by crooks (think of Web optimization to get an artificially high ranking on Google searches), especially since the academic progress and promotions of researchers, and thus their salaries, have come to depend on it. Even setting aside the problem of a paper being widely cited because it is intensely criticized (rather than valued), citing a paper as a favor to its authors or the publication venue isn't uncommon. Some journal editors publish papers that are likely to garner many citations early in the volume/year to give them more time to collect citations before the following year's JIF calculation.

As an editor for multiple technical journals over the past few decades, I have noticed the tendency of some reviewers to insist that their work be cited as a condition of acceptance for publication, even if their work is only marginally relevant to the paper under review. Some journal editors have been known to encourage authors to cite papers in recent issues of their journal in order to increase its impact factor.

Like any single numerical metric, JIF is flawed, but it can be valuable if used with care. In my teachings on parallel computation, I often remind students that numerical measures of performance (peak FLOPS or sustained FLOPS, say) are okay for quickly comparing supercomputers, but they should be augmented by other measures to go beyond a superficial judgment. In the case of journals, JIF can be augmented by qualitative criteria such as inclusion in various tiers of journals for a particular scientific or technical field.

An even more approximate version of this approach is to divide journals into two piles: reputable and questionable. Inclusion in established indexing databases, such as ISI, is sometimes taken to be a sign of a journal's repute [16]. Using inclusion in ISI as a binary measure of journal quality is particularly common in many Third-World countries with no established research traditions of their own.

High-quality journals not only are selective in accepting papers, they are also vigilant in monitoring reactions to published work, so as to discover errors or research misconduct. Detected errors necessitate issuing corrections in a subsequent issue. Scientific misconduct, when established through a fair and thorough review after an initial accusation, leads to retraction of the work and disciplinary action against the authors. The retraction index of a journal is defined as: $\text{Articles retracted} \times 1000 / \text{Articles published}$.

Having a paper retracted is the ultimate punishment for a researcher, so such actions provide strong disincentives for data-falsification or other types of fraud in research reporting. On the other hand, the strong correlation of journal quality, as reflected in its impact factor, and its retraction index (Fig. 5) is an indication that authors might be willing to take greater risks in having work published in high-impact journals.

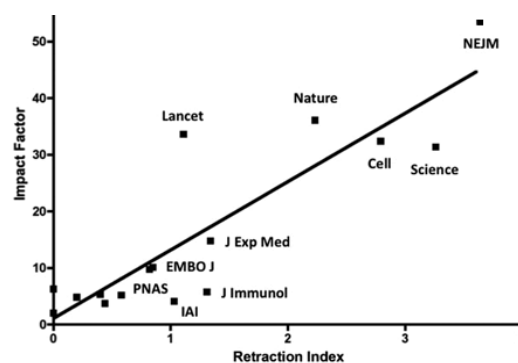


Fig. 5. Correlation between journal impact factor and retraction index [3]

An interesting argument exists for using post-publication reviews in lieu of pre-publication refereeing. Let researchers publish as many papers as they want, now that on-line self-publication is essentially cost-free. Such a publication would be deemed worthless if it does not garner positive peer reviews after it is made visible in cyberspace. This approach is theoretically sound, but figuring out how to avoid "commentary pollution" or spam, of the types one sees on other on-line forums, is nontrivial. Citations are basically a kind of binary or yes/no post-publication reviews.

6. Fake Journals

Financing scientific and technical journals has always been a thorny issue. Relying on volunteer (unpaid) work by editors and reviewers, scientific societies have managed over the years to publish journals of fairly high quality. But the task has been getting more and more difficult. Right now, finding multiple reviewers for a submitted paper is an extremely challenging task, given researchers' busy schedules and the proliferation of journals and conferences.

When journals were primarily in hard-copy format, research libraries had shelf after shelf of past issues of each journal, bound neatly into hard-copy tomes. Proliferation of journals made subscribing to every pertinent journal all but impossible, so libraries became very selective, relying on collaborations and inter-library loans to supply their patrons with needed references.

Now, as journals move to on-line or electronic format (with hard copies printed and distributed as a relic for those who find the transition difficult), another problem has surfaced. Who pays for the cost of maintaining journal archives, so that they remain accessible in perpetuity? The cost is non-trivial and, as collections grow, organized and accessible archiving is a challenge. For journals published by both professional societies and private publishers, one worries that the archives may disappear, should the entity maintaining them become inactive or go bankrupt.

Open-access publishing has come into existence, in part to solve the problem above. Old publication models entailed cost-free publishing for authors, but access to research results required payment (e.g., via subscription to journal or membership in the professional society). Open-access

publishing flips this model. The author, or his/her institution, pays the publication and archiving costs, with the work then being freely accessible on-line. Open-access papers tend to be read by more people, thereby potentially increasing their citation counts.

The move to open-access publishing is a major shift, and mostly a positive one with respect to accessibility of research results, even though there are still wrinkles to be ironed out. The downside is also significant. If each author pays \$1000, say, to publish his/her paper, the journal publisher is much less motivated to rigorously evaluate each paper and publish only 15% of the submissions, say. So, the journal's acceptance rate begins creeping up to increase revenue. And this is just for the well-meaning publishers and editors.

Another drawback of open-access publishing is that it is biased toward better-funded researchers and financially well-endowed institutions, given the sometimes-exorbitant submission and/or publication fees. Many open-access journals offer need-based discounts, but the offers are ad-hoc and not consistently extended. Despite its drawbacks, open-access publishing does hold the key to future dissemination of research results.

Over the past decade, many thousands of journals have been launched by predatory publishers [19] to cash in on the open-access publishing trend. Many of these fake journals have their headquarters in Third-World countries, where labor is much cheaper. The journal (often part of a collection that includes dozens or even hundreds of titles) typically has a nameless managing editor who solicits papers by providing various incentives to potential authors, such as stroking their egos or offering fee discounts.

Fake journals are very generous in their lavish praise of your research, whether or not you deserve it. They routinely send researchers invitations to submit papers, guest-edit special issues, and join their often-obscure editorial boards. They sometimes promise publication within days, such as soliciting papers for the June issue of a particular journal in late May.

As they say, it takes two to tango. Academics, who may not make it onto the editorial boards of prestigious journals, jump at the opportunity to be listed on multiple such boards and padding their CVs with publications and editorial-board memberships. A hallmark of fake journals is their pretentious and often very-general titles, which tend to include everything. I have made up the following title for comic effect, but real titles aren't much better: *Global Journal of Innovative and Advanced Engineering, Science, Management, and Social Research*.

Various lists of predatory journals [13] and publishers [2] have been compiled. Many such lists qualify their compilation by including the word "Potential," both to legally protect themselves and because any long list is bound to contain some inaccuracies, which are corrected over time.

Alongside carefully-compiled lists to consult, researchers should be thought simple rules of thumb for spotting predatory

A. Unusual fees and/or conditions

The journal charges a submission fee or editing fee (in addition to publication fee), insists on keeping the copyright, or imposes other unreasonable conditions

B. Small or as-yet nonexistent editorial board

Legitimate editorial boards are always prominently displayed in print and on-line, citing contact information and affiliations

C. Myriads of new journals launched by a publisher

The journal names often begin with or contain the same set of words, usually including pretentious or scientifically-imprecise terms

D. Delayed publication of overdue issues

Lack of success in attracting a sufficient number of submissions to publish issues regularly and on time is a sign that authors do not deem the journal worthy

E. Web site that isn't professionally designed

A clean, user-friendly Web site is a key requirement for any serious journal, particularly if it is open-access

F. Claimed academic affiliations do not pan out

The journal or its Web site contains the name of a prestigious organization, which does not match the editorial members' affiliations or journal's location

G. Grammatical errors in paper titles or abstracts

Such errors are indicative of the editors' lack of familiarity with the subject matter or they signal very light or even non-existent refereeing

H. Content does not match the journal title or scope

Content that isn't relevant to the journal's title or stated scope is a sure sign of desperation in filling the journal's issues and maintaining a healthy cash flow

Fig. 6. Some of the common signs of predatory publishers and journals [10]

journals and publishers. Eight tell-tale signs are listed in Fig. 6 for easy reference [10]. See also [6].

As the numbers of legitimate and predatory journals rise, it becomes increasingly difficult to tell them apart. Whereas fake journals publish almost anything to make money, some legitimate scholarly journals may not fare much better. In the latter case, the incentive often isn't financial, but promoting viewpoints or ideologies. There are examples of satirical [12] or computer-generated articles [15] having been published in peer-reviewed journals, somehow falling through the cracks because of extremely busy or perhaps careless referees.

7. Conclusion

In this paper, I have shared some thoughts and useful tips on the challenges and perils of evaluating research quality and impact. My equivocations should not be interpreted as a stance that such evaluations are futile. Rather, they should be viewed as pointing out the need for greater caution and for being mindful that the process is error-prone. We can't quite put error bars around research evaluation results, but just the awareness that the outcomes aren't precise goes a long way toward applying due diligence in administering fair evaluations.

The existence of errors in assessing research quality and impact necessitates that we maintain flexibility and not blindly follow rigid guidelines, no matter how carefully they are formulated and how much detail they provide. Bean-counting

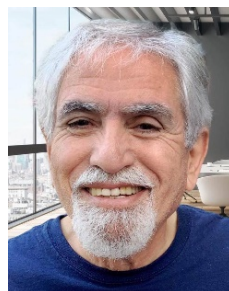
is still bean-counting, even if the beans are sorted by size and color!

Multidisciplinary research fields provide particularly thorny problems in assessing quality and impact. Combining knowledge and methods from multiple fields is inherently difficult and worthy of being rewarded. On the other hand, I have seen instances of trivial results and methods from one field being touted as ingenious or earth-shattering to those working in other fields, who are not equipped to evaluate the claims. In other words, multidisciplinary work should be state-of-the-art in each of the disciplines involved.

I view this article as a start. It can be extended and expanded in many different directions. More details are needed if this work is to be used for educating the next generation of researchers and to serve as a resource for university promotions and merit-advancement committees. I am working in this direction and would appreciate any comments, pointers, or critiques that might help improve the depth and coverage.

References

- [1] Baldwin, M. "Origins of the Journal Impact Factor": <https://physicstoday.scitation.org/doi/10.1063/PT.5.9082/full/>, 2017.
- [2] Beall's List "Potential Predatory Scholarly Open-Access Publishers": <https://beallslist.net/>, 2020.
- [3] F. C. Fang, and A. Casadevall, "Retracted Science and the Retraction Index," *Infection and Immunology*, vol. 10, pp. 3855-3859, 2011.
- [4] E. Garfield, "Citation Indexes for Science", *Science*, vol. 122, no. 3159, pp. 108-111, 1995.
- [5] Garfield, E. "Citation Analysis as a Tool in Journal Evaluation," *Science*, vol. 178, no. 4060, pp. 471-479, 1972.
- [6] Glasson, V. 2020. "6 Ways to Spot a Predatory Journal": <https://www.rxcms.com/blog/6-ways-spot-predatory-journal/>
- [7] Lerchenmueller, M. J., O. Sorenson, and A. B. Jena, "Gender Differences in How Scientists Present the Importance of Their Research: Observational Study," *BMJ*, vol. 367, l6573, pp. 1-8, 2019.
- [8] Ng, I. C. L. (undated. "Guidelines on How to Evaluate a Research Article," On-line document. https://www.academia.edu/1744631/Guidelines_on_how_to_evaluate_a_research_article
- [9] Parhami, B., "Low Acceptance Rate of Conference Papers Considered Harmful," *IEEE Computer*, vol. 49, no. 4, pp. 70-73, 2016.
- [10] Prater, C. "Eight Ways to Identify Questionable Open Access Journals," American Journal Experts Web site, <https://www.aje.com/arc/8-ways-identify-questionable-open-access-journal/>, 2020.
- [11] Smith, A. J. "The Task of the Referee," *IEEE Computer*, vol. 23, no. 4, pp. 65-71, 1990.
- [12] Sokal, A. D., *The Sokal Hoax: The Sham that Shook the Academy*, University of Nebraska Press, 271 pp, 2000.
- [13] Stop Predatory Journals 2020. "List of Predatory Journals": <https://predatoryjournals.com/journals/>
- [14] Sugimoto, C. R. et al., "Global Gender Disparities in Science," *Nature*, vol. 504, pp. 211-213, 2013.
- [15] Van Noorden, R., "Publishers Withdraw More than 120 Gibberish Papers," *Nature*, News & Comment, 2014.
- [16] Web of Science, "History of Citation Indexing": <https://clarivate.com/webofsciencegroup/essays/history-of-citation-indexing/>, 2020.
- [17] WikiHow "How to Write a Letter of Recommendation": <https://www.wikihow.com/Write-a-Letter-of-Recommendation>, 2019.
- [18] Wikipedia, "Impact Factor": https://en.wikipedia.org/wiki/Impact_factor, 2020a.
- [19] Wikipedia, "Predatory Publishing": https://en.wikipedia.org/wiki/Predatory_publishing, 2020b.
- [20] Zhou, Z. Q., Tse, T. H., and Witheridge, M. "Metamorphic Robustness Testing: Exposing Hidden Defects in Citation Statistics and Journal Impact Factors," *IEEE Trans. Software Engineering*, early access, pp. 1-22, May 2019.



Behrooz Parhami is Professor of Electrical and Computer Engineering at University of California, Santa Barbara, and former Associate Dean for Academic Personnel, UCSB College of Engineering. His research interests include computer arithmetic, parallel processing, and dependable computing, areas in which he has authored graduate-level textbooks. He

has served on the editorial boards of multiple journals, including current Associate Editor positions with *IEEE Trans. Computers* and *IEEE Trans. Sustainable Computing*.

Email: parhami@ece.ucsb.edu

Paper Handling Data:

Submitted: 06-09-2020

Received in revised form: 10-03-2020

Accepted: 10-04-2020

Corresponding author: Behrooz Parhami

Affiliation of the corresponding author: Department of Electrical and Computer Engineering, University of California, Santa Barbara, CA 93106-9560, USA