

The ‘Survey Paper DDos Attack’ That’s Overwhelming Scientific Research

Originally published at <https://www.unite.ai/the-survey-paper-ddos-attack-thats-overwhelming-scientific-research/>



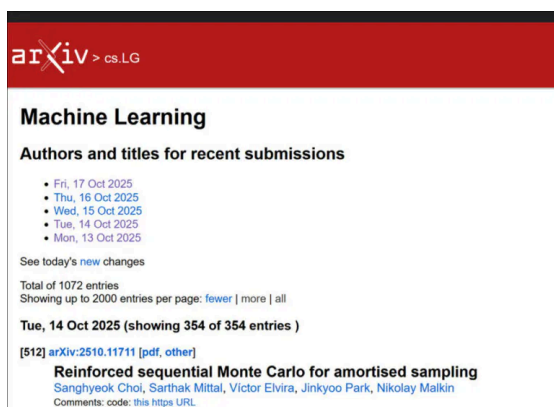
Generative AI models such as ChatGPT are now flooding academic publishing platforms with AI-generated survey papers at levels that are making the signal-to-noise ratio critical. A new study claims this flood is overwhelming researchers, distorting citations, and eroding trust in the scientific record, likening the deluge of AI-aided papers to a ‘DDoS attack’ on science itself.

(Partly) opinion Last week, for the first time in seven years of staying up-to-date with the scientific literature stream relating to AI, I had to concede defeat and admit that, at least at peak times, I must now choose being staying on top of essential new publications or having any time left to write about some of them.

The total number of entries in a very limited number of relevant categories ([Computer Vision](#), [Machine Learning](#), [Language Models](#), and a few other less-subscribed sections) stood at *significantly* over a thousand – for just one day’s submissions.

At such a volume, even skimming all the new titles and only occasionally indulging in some of the abstract summaries would have made for an unproductive day.

This was Tuesday October 7th. By contrast, in the [Machine Learning category](#), this Tuesday past (October 14th) offered a publication volume slightly less intense than the 400-odd entries for Tuesday of the prior week; it had a mere 354 entries:



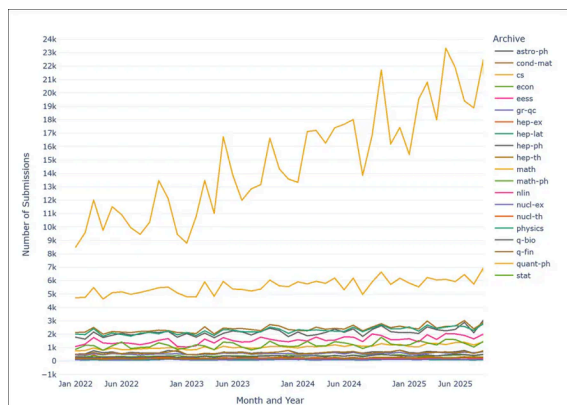
354 submissions for the Machine Learning category in one day. Source: <https://arxiv.org/>

You would have to have been reading Arxiv every day, for some years, to realize how insane these numbers are getting.

Admittedly, Tuesday is Arxiv's 'rush hour' for submissions, perhaps because it's the earliest working day that occurs away from any long weekends enjoyed by the influential people that researchers hope to reach; and the Machine Learning category is a 'catch-all' section with a lower number of unique papers (papers that are not simultaneously published in more specialized channels) than most other categories.

Nonetheless, the rise in paper submissions is already a [noted phenomena](#) in academia and in the media.

Perhaps the most shocking aspect of this escalation is how all other adjacent categories are more or less unaltered in their frequency over the past three years, whereas the Computer Science category (see if you can spot it in Arxiv's official figures below) is on a severe upward trajectory:



The rise of computer science (CS) papers over the last three years. Source: https://info.arxiv.org/about/reports/submission_category_by_year.html

Just over three years ago Arxiv's AI paper submission output was [estimated](#) to be doubling every few years; and it will be interesting to read Arxiv's own annual digest of trends at the end of 2025.

Volume at 11

The two most obvious reasons why this is happening is a) [unprecedented financial commitment](#) to generative AI is attracting massive levels of research investment in the private and academic sectors, which often collaborate; and b) the fact that AI language models systems such as ChatGPT now make submitting research papers (including papers about AI) an almost industrialized process.

However, the quality of research submissions is [not rising](#) in tandem with the volume (though AI's error-prone output tends to [make more headlines](#) in the legal sector than the academic, not least because the ramifications are more obvious there).

A zero-tolerance policy is hard to implement in this case, even if [recognizing AI-generated content](#) were easier; besides the fact that AI in itself is a manifest boon to scientific research in general, its use in research paper submissions has generally* improved the clarity of work from many non-English submitters – individuals and teams who have until now [operated at a disadvantage](#).

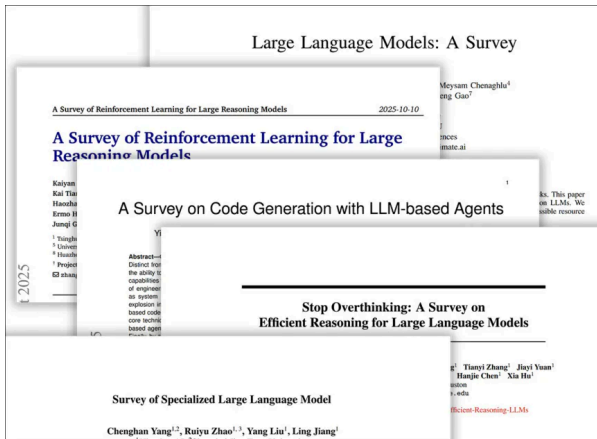
But the problem of lowering the language barrier in this way is that this also raises the sheer number of global submitters, without raising the level of human oversight that gives value to such work.

If the levels of submission continue to rise exponentially, the signal-to-noise ratio will become so ungovernable that only AI itself could possibly navigate the new floods and tributaries of AI papers; a task that it is no more suited to do than proofing its own output. Ironically, scientific research is an intensely *human* endeavor.

An Attack on Research

The cause of this reflection is an interesting [new collaboration](#) from China titled *Stop DDoS Attacking the Research Community with AI-Generated Survey Papers*.

The new position paper concentrates specifically on *survey* submissions – high-effort roundups of particular strands in research, which have traditionally both listed and contextualized, interpreting trends and making informed predictions:



A mere fraction of the vast and ever-growing body of surveys available in sections related to machine learning and AI, at arxiv.org

Since surveys curate rather than originate, they are unusually easy to automate with AI, and the authors of the new work characterize the proliferation of low-effort surveys in terms of a *security threat* to the research sector[†]:

'[The] recent surge of AI-generated surveys, especially enabled by large language models (LLMs), has transformed this traditionally labor-intensive genre into a low-effort, high-volume output. While such automation lowers entry barriers, it also introduces a critical threat: the phenomenon we term the "survey paper DDoS attack" to the research community.

'This refers to the unchecked proliferation of superficially comprehensive but often redundant, low-quality, or even hallucinated survey manuscripts, which floods preprint platforms, overwhelms researchers, and erodes trust in the scientific record.

*'[We] argue that **we must stop uploading massive amounts of AI-generated survey papers (i.e., survey paper DDoS attack) to the research community**, by instituting strong norms for AI-assisted review writing.'*

The authors assert that this unencumbered acceleration of survey production threatens to swamp the research ecosystem with polished reports that nonetheless lack critical depth, and which are likely to propagate factual errors and/or [hallucinated](#) citations.

The paper warns that without better rules or oversight, AI-generated surveys could turn into shallow copies that misrepresent which topics are important, hide meaningful analysis, and make literature reviews less trustworthy:

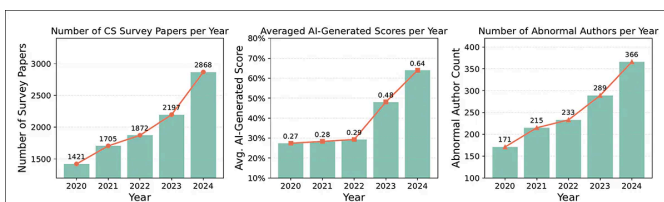
'The implications for research quality and trust are profound. First, genuine advances risk being obscured by algorithmically generated rehashes of existing work.

'Newcomers and interdisciplinary scholars may struggle to locate dependable overviews amid the noise. Moreover, errors or biases introduced by automated drafting can propagate unchecked, seeding subsequent research with faulty premises.

'In sum, the flood of non-peer-reviewed AI-generated surveys endangers both the rigor of literature reviews and the credibility of the scientific record.'

'Abnormal' Authors

The researchers of the new paper provide some interesting analyses on the evolution of survey submissions:



Left: the annual count of computer science survey papers from 2020 to 2024. Middle: average AI-generation scores for those papers over the same period. Right: number of authors flagged as abnormal (those with unusually high survey output, limited co-authorship diversity, and recurring institutional patterns) each year. All three trends show a sharp rise beginning in 2023, coinciding with the release of ChatGPT and other large-scale language models.

In the first column we see growth trends: the curve starts to steepen around 2022, just when ChatGPT emerged and large language models began to go mainstream, and follow-up models such as [Claude](#), [PaLM](#), and [Gemini](#) would keep that momentum going throughout 2023.

The middle graph shows a steep rise in submissions after 2022, coinciding with the launch of ChatGPT. One research team [found](#) that by 2024, more than 10% of scientific abstracts had been run through an LLM. A [separate report](#) from an AI detection firm put the post-ChatGPT jump at 72% for papers on arXiv that may have been written with AI help. The number of papers with high AI-generation scores also doubled in a year, from 3.6% to 6.2%.

The third, right-most graph shows a steady rise in the number of ‘abnormal’ author patterns (researchers submitting three or more surveys within a month while working with fewer than two collaborators), with a sharper rise beginning in 2022.

The authors assert that many of these survey papers may have been drafted by AI, for diverse reasons; some are written by solo authors or small groups who submit multiple surveys in a short time; many cover unrelated topics; and in some cases, the authors have no previous record in the fields they are summarizing.

Additionally, some are published under anonymous collectives with no clear institutional ties – patterns suggesting a coordinated flooding of the field with quick surveys, possibly to gain citations or improve academic profiles, rather than to make any real contribution to the literature.

Issues

Though we can’t cover all the contentions of the new paper, we ought to take a look at some of the most notable observations, as well as cast a critical eye over the authors’ suggested solutions to these issues.

Quality and Originality

The problem isn’t just volume: many AI-written surveys skip what makes a good survey useful: clear structure, deep analysis, correct and assiduous credit, and real insight. Instead, the paper suggests, AI-generated/aided surveys often read like stitched-together summaries, with none of the requisite care or curation.

The authors observe, further, that AI-written surveys often lack structure, but rather simply list papers without clear direction, skipping key sections, and failing to create context. Human-written surveys, by contrast, tend to establish proper categories, and tell a more coherent story.

Also, many potentially AI-aided surveys seem to simply copy existing topic breakdowns, sometimes straight from Wikipedia. For instance, the paper notes, multiple surveys on [Vision Transformers](#) contain common section titles and structure, belying template-driven AI output:

‘In contrast, a well-crafted human-written survey might introduce a new taxonomy, e.g., categorizing ViT by efficiency strategies. The lack of such original structure in many recent survey preprints raises concerns that they may have been generated by AI with limited human insight.’

Don’t Quote Me on That

Perhaps most publicly embarrassing, AI-written surveys often get [citations](#) wrong, missing key papers, *including* non-relevant papers, and sometimes even listing *non-existent* papers – errors that suggest the references originate from surface-level pattern matching, rather than true expertise.

The authors also point out that some recent survey papers, often from entirely different teams, share as much as 70% of their reference lists – a level of overlap so high, they argue, that it hints at a shared reliance on LLMs, which draw from the same narrow pool of source material.

Indeed, casual users of ChatGPT will know that the more obscure the topic, the fewer diverse sources there are for the model to have [generalized](#); very often, locating the model’s own limited sources on the web is more useful than trying to interact with that information via an AI which did not have adequate data in a particular domain.

A ‘Homogeneous Style’ Emerging

The authors also note that many AI-written surveys on the same topic look and sound nearly identical, as LLMs reuse phrasing and structure, especially for popular subjects, resulting in a torrent of almost identical papers that add little value, and also add significant noise to researchers seeking domain answers*:

‘When multiple authors ask an LLM to “write a literature review on X,” the model often produces very similar responses, especially for common definitions or well-known facts. [Recent research](#) has shown a sharp rise in the use of certain writing patterns linked to LLMs, suggesting many papers now share the same style.’

Your ChatGPT is Showing

The paper observes that a quick way to spot AI-written surveys is through the presence of phrases such as ‘as an AI language model’ or ‘my knowledge cutoff’, suggesting minimal or even zero curation of the output from the language models before submitting the papers (though a targeted search at the time of writing did not reveal any such tells indexed in Google Search).

The paper notes that many ‘suspect’ surveys show lower word diversity and repeated phrasing, for example, by beginning multiple paragraphs with *Furthermore*. This kind of pattern, the authors suggest, is typical of GPT-style writing, and could be a useful flag for detecting auto-generated text.

(My personal comment on this is that the strictures of online journalism often require a writer to list many items in a prose-based, non-styled form. Therefore ChatGPT and its peers are likely to have learned this bad habit from human writers who were faced with a limited number of lexical alternatives. Additionally, the authors’ conjecture shows them dabbling in the tenets of AI content-detection, which is a complex and developing field, with few enduring constants of the kind that the authors suggest)

Though the researchers go on to develop a fascinating discourse on the negative impact of AI surveys on research culture and trust, we must refer the reader to [the source paper](#) for greater depth on this.

Solutions?

The paper’s solution is fascinating, radical, and at the same time strangely unoriginal: that the utility of survey papers should be replaced by a *Dynamic Live Survey* – by interpretation, a kind of hybrid between a Wiki and a GitHub page, constantly fed with new data from LLM and other AI systems, but with commits being made only by humans, so that AI cannot essentially ‘auto-publish’ updates.

The proposed system would share the versioning and branching of GitHub, essentially turning an information resource into a constantly-updating list similar to the [‘awesome’](#) strain of curated lists at GitHub:

‘Under this framework, a community member first establishes a survey topic wiki by specifying the scope, key research questions, and seminal references, which thereby sets a clear thematic boundary and initial structure.

‘Thereafter, an LLM-based ingestion agent continuously monitors preprint archives, conference proceedings, and benchmark leaderboards. It automatically extracts abstracts, figures, and key performance metrics; synthesizes concise summaries of new results; updates the citation graph to reflect inter-paper relationships; and flags emerging research trends for further review.

‘By design, these automated updates occur within hours of publication, ensuring that the repository remains at the cutting edge.

‘Human contributors then step in to provide the interpretive depth that machines alone cannot offer. They refine evolving taxonomies to capture subtle methodological distinctions, coordinate conflicting interpretations of algorithmic innovations across different subfields, and provide deeper critical comparisons to the document.’

The Book of Changes

The authors expound enthusiastically and at length on this proposal, and essentially justify it with something that is very true: high-effort human-written surveys on volatile topics around AI *age so quickly that they are hardly worth writing*; and the paper notes that a three-month turnaround on a new survey paper will likely mean that it will be out of date (or even *severely* out of date) by the time its scheduled publication day arrives:

‘Year after year, communities are flooded with repetitive or superficial overviews that quickly lose relevance, leaving practitioners and newcomers alike struggling to distinguish signal from noise. The traditional publication cycle (i.e., draft, submit, review, and publish) can span several months, by which time critical breakthroughs may already have shifted the landscape.

‘Moreover, the increasing volume of static surveys adds to cognitive overload, as readers must sift through numerous overlapping documents to find substantive insights.’

Unfortunately, the paper’s solution shares many of the worst and most derided qualities of Discord: most especially that it would be a constantly shifting and changing resource.

Since any part of a Dynamic Live Survey might disappear or be amended at any time, it would be impossible to use as a citable, stable source; except, perhaps, by linking to a ‘previous commit’, in much the same way archive.is and the Wayback Machine, among other archive sites, provide linkable snapshots of web-page content, frozen at a particular time. But what resources would such a commit need, and could it be relied on to stay live over time?

Additionally, a platform/Wiki with constantly changing definitions and content would be challenging to index, either by traditional search engines or LLMs.

Perhaps the weakest part of the proposed system is the idea that real people should oversee the commits from the LLM agents; as ever, real people are expensive. What is being proposed is something between a museum and a library – both of which are going to need meatware provisioning proportionate to the volume of data and number of topics covered.

If ‘use real people’ is the only answer to an AI development problem, it’s fair to say that the problem remains open and unsolved.

Conclusion

At the moment the short shelf-life of survey papers on AI is annoying; if the current trend towards high-scale automated writing and submission continues, as envisaged in the new paper, the signal-to-noise ratio will become chronic, and the literature ungovernable.

In such a situation, it would be even harder than it is now for minor, sub-FAANG voices to be heard in the storm of submissions, and major market leaders would likely gain even greater prominence.

Besides live surveys, the new paper proposes that authors not only be constricted to self-declare when AI is used in any part of a submission, but also that AI-aided sections be explicitly labeled within a paper (perhaps with a sidecar JSON file...?).

Since this is an onerous prospect, the paper alternately suggests what I can only characterize as an ‘AI ghetto’ – a distinct section in the submission which is set aside for AI contributions.

In short, the new work has, at least in my opinion, no realistic answers to offer; but the authors have performed a useful service in framing the challenges ahead.

The paper **Stop DDoS Attacking the Research Community with AI-Generated Survey Papers** can be found at <https://arxiv.org/abs/2510.09686>, and is written by six authors across departments at Shanghai Jiao Tong University.

* [Not all feel that this is the case.](#)

† Authors’ emphasis, not mine. Also, where applicable, my conversion of authors’ inline citations into hyperlinks.

First published Friday, October 17, 2025

Writer on machine learning, domain specialist in human image synthesis. Former head of research content at Metaphysic.ai, until its dissolution into DNEG’s Brahma.ai.

Portfolio site: martinanderson.ai

Contact: martin@martinanderson.ai



Discover More