

Can AI Develop a Nose for News?

Originally published at <https://www.unite.ai/can-ai-develop-a-nose-for-news/>



AI is getting better at writing news stories, but not getting much better at identifying them.

Opinion In the five years since I [last took a look](#) at AI's ability to find a hot news story, the landscape has changed considerably, with increased levels of AI-driven automation accompanied by the inevitable growing pains and [controversies](#).

Recently, a [WSJ report](#) on a prolific, AI-aided *Fortune* contributor presented the journalist of the future as emancipated from scutwork such as transliterating press releases, leaving them scope to write the features and do the digging that only larger publications usually have the budget for.

But what we hear about much less often is AI's ability to *spot* a news story.

Noise Abatement

In the 2021 piece, I concentrated on writers covering the research beat, since that's where I spend most of my time; and perhaps the biggest effect that the new AI revolution has had on that is that it created an [ungovernable blizzard](#) of AI-powered research paper submissions, raising the signal-to-noise ratio so high that even covering the Arxiv AI-related domains comprehensively is now beyond the efforts of a single person.

Surely this is where AI excels – at iterating through vast tranches of data that humans can't resolve, to find 'outliers' (which we will get to shortly) in seconds that would have taken people days, if they could have done it at all.

Why, then, is AI still so bad at identifying a hot news story from the thousands, even tens of thousands, of daily contenders?

Backward-Looking AI

This massive proliferation of AI-generated content is happening far beyond the academic sector that I discussed previously. Late last year it was estimated that half of *all new writing* on the web [is now 'written by AI'](#), with even greater acceleration of this trend presumed to be coming. Therefore the noise is deafening *everywhere*, not just in academia.

Though there has been some progress in AI/algorithmic identification of a 'hot' story in the last few years, these systems tend to concentrate on stratified and predictably-organized data feeds, meaning that they can operate only in a rather brittle context.

In this regard, Stanford post-doctoral researcher and former New York Times journalist [Alexander Spangher](#) has made several forays into defining 'newsworthiness' in terms that can be applied to machine learning processes and statistical analysis; and has produced evidence of automated lead-generation in corpora [such as court filings](#), state bills, and city council meetings, as well as [general public documents](#) – the kind of schema-driven output that *Fortune*'s prolific AI-powered scribe can turn into 6-7 news pieces a day:

Δ Word Distributions for Newsworthy vs. Non-Newsworthy Text							
Policy Text				Meeting Speech		Public Comment	
authorizing	-0.41	housing	0.35	supervisor	1.98	budget	0.4
county	-0.3	health	0.31	think	0.89	philippines	0.16
grant	-0.26	board	0.3	know	0.82	solar	0.15
lawsuit	-0.25	ordinance	0.29	want	0.78	medical	0.15
bonds	-0.23	covid	0.28	people	0.76	covid	0.14
settlement	-0.22	department	0.23	like	0.58	caltrain	0.14
contract	-0.21	cannabis	0.22	need	0.43	rooms	0.13
expend	-0.19	election	0.21	president	0.37	amendments	0.12

The ‘heat’ of word distributions gleaned from corpora of public documents. In this case, we can see that ‘authorizing’ has a high score, perhaps because it represents decision, change and novelty. [Source](#)

However, the problem with approaches such as the Spangher-led 2023 [offering](#) *Tracking the Newsworthiness of Public Documents*, is that in typical AI fashion, they center on *observed trends in the data*. In other words, they observe things that made good news before, and go on to look for more of the same.

In the real world, unexpected sources nearly always turn out to be a ‘one hit wonder’; and for how obscure they were, no-one could have predicted their sudden prominence. Then, having been fruitful once, and in spite of occasional attempts to capitalize on fleeting fame/notoriety, they will usually [never produce anything useful again](#).

Sign of the Times

Therefore, since monitoring this kind of one-and-done’ news source is usually just going to add more noise to the general blizzard, could AI instead not identify the *signifiers* of a source that will one day become fruitful? If one could find out what type of source may eventually yield news, one could focus on its *characteristics* rather than its context, or its methods.

By that logic, one could deduce from the Edward Snowden revelations of the 2010s that anyone who recently left the employ of the CIA (or [a similar organization](#)) would be worth following as a potential source of a future scoop.

However, there are no RSS feeds or APIs that are likely to be able to automate this kind of ongoing monitoring, since LinkedIn and many other once-open sources of data are [retrenching](#) in the face of rapacious and [scofflaw](#) AI web-scrappers. Even if there were, frequency would be an issue, because you can’t poll an API or a site every five seconds; aside from the resource-cost, IP-banning responses from the platforms would make this an unsustainable activity.

Further, there is clearly a ‘human dimension’ to such disclosures that is hard to automate.



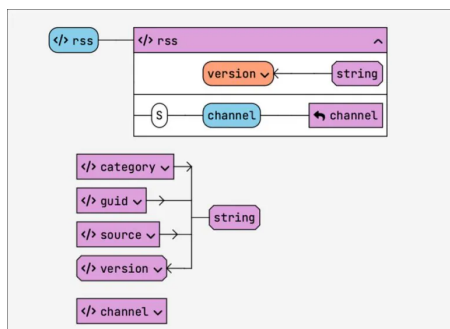
News-gathering with the personal touch: capture from a disc release of the 1976 Alan J. Pakula movie ‘All The President’s Men’, featuring the informant coming out of the shadows. [Source](#)

Also, in the real world, it is awfully hard to identify the *defining characteristics* of a future news source. It probably isn’t ‘people who left the CIA recently’, and it certainly isn’t defined by a protocol: platforms such as X or GitHub output far too much signal in themselves, and even narrowing down on search terms or post categories makes little difference – only if you’re involved in the problem, and engaged with the community (or repo, etc.) are you really likely to recognize the significance of a development.

Even a term [such as ‘security alert’](#) can’t contextualize the true severity or newsworthiness of an incident, since references of that kind are thrown around daily, by the thousand, in such communities – yet have no broad news value; and even if one restricts that kind of monitoring to the English language alone, the potential variations in idiom, along with the use of oblique language, would make it very hard to parse an ‘in the wild’ post into a true news alert.

The Narrow Path

The current crop of AI-powered newsworthiness-detection systems depend on formalized data structures (such as JSON output, from an API), or else on informal data structures that AI-developed algorithms *have a chance of parsing* into a structured schema (such as press releases from a particular organization):



A parsed RSS/XML feed, revealing the rigid hierarchy of data containers. [Source](#)

Clearly, approaches of this kind are well-suited for programmatic output, such as the mundane work that the aforementioned Fortune reporter declares AI has freed him from, including weather, stocks, and sports scores reporting, as well as routine press releases from municipal and other government organizations.

While it is possible to attach 'human-alerting' triggers to statistical feeds like weather (sudden storms), stocks (sudden plunges) and sports (unexpected victories/losses, with some prep-work), again, human attention would still be needed even for very stratified government releases, in order to gauge newsworthiness.

Though terms such as 'death', 'unexpected illness', 'leak', and 'accident' can all help drill down to newsworthy events, they only address 'routine' eventualities, and also can't account for alternative language (or languages).

Return of the Elite Writers?

In recent years, [data-driven journalism](#) has become [an ascendant plank](#) in news reporting, with editorial departments no longer limited to sweetheart 'scoop' deals granting them early release on special reports and white papers from major publishers; instead, they can crunch the numbers themselves.

However, this is no free lunch; as the evident value of parsing public data with AI in this way has grown, a rent-seeking/AI-blocking response has followed – or even anticipated – demand, driving the data-hungry major AI players [into stealth tactics](#).

The added friction of the New Retrenchment arguably restores a certain amount of power from 'citizen journalists' back to legacy media – or at least, well-funded news organizations that have the bandwidth to absorb the extra manual work required in gathering, refining and evaluating data, in an era where publishers and domains are increasingly restricting casual access.

So, in a way, perhaps in the spirit of the time, the practical manifestation of AI in journalism, in terms of the way that major players and markets have responded to AI-based innovation and adoption, may actually be taking us back in time: de-democratizing the means of news production, and adding roadblocks to meaningful data-driven newsworthiness-evaluation systems.

Common Instincts

These strictures clearly lead us back to 'gut feeling' as an inevitable component in evaluating the newsworthiness of a story.

Naturally, this is comforting for those who are engaged professionally in this aspect; but complacency would be a mistake, since this instinct can, to a certain extent, be distilled and operationalized in a very general way that does not depend on studying the obsessions or hobby-horses of any one individual or organization: in a 2022 [study](#), researchers from Northwestern University used crowd-sourced evaluations of potentially newsworthy stories to train a predictive model, specifically concerned with the newsworthiness of newly-published Arxiv research papers:

Survey Questions

Now, please answer each of the following questions, and provide the appropriate explanations:

1. Please type in the fourth letter of the abstract of the paper in the text box, so we know that you are paying attention:

2. Please rate the potential impact of the research described on society, such as to individuals, organizations, politics, or the economy.

1 Strongly Negative 2 Negative 3 Neither Negative nor Positive 4 Positive 5 Strongly Positive

Please explain and justify your rating in 2-3 sentences:

3. Please rate how strongly you agree or disagree with the following statement:

"The research described has the potential to impact a significant number of people."

1 Strongly Disagree 2 Disagree 3 Neither Agree nor Disagree 4 Agree 5 Strongly Agree

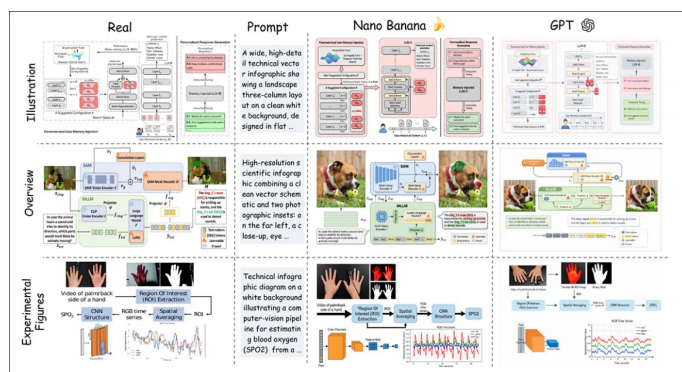
Survey questions given to study participants in order to obtain training data for a 'newsworthiness-prediction' AI model. [Source](#)

The system ranks candidates fairly well, with about 80% of its top ten picks also judged newsworthy by experts. However, agreement with experts proved only moderate, with the results missing factors such as framing, or audience fit.

The system is predicated on the principles outlined in the [2020 paper](#) *Computational News Discovery: Towards Design Considerations for Editorial Orientation Algorithms in Journalism*. As with most similar projects, this work tackles science journalism rather than abstract news-gathering – perhaps because the scientific literature tends towards templated output that could potentially be parsed into trainable and interpretable data points.

Well, as I observed back in 2021, this *would* be the case, except that research scientists frequently abuse the conventions of research paper submission to hide or downplay unimpressive results, or even outright failure.

Even more of a challenge is the [great difficulty](#) that AI systems have in interpreting figures and tables in scientific papers, to the point where this pursuit has, lately, [become](#) an [active strand](#) in the literature:



From the paper 'SciFigDetect: A Benchmark for AI-Generated Scientific Figure Detection', showing real scientific figures, their generation prompts, and synthetic counterparts produced by Nano Banana and GPT across three categories: illustration, overview, and experimental figures. [Source](#)

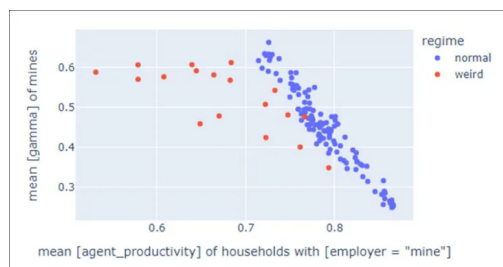
It's often the case that a graph or table will contain results which the main body of the paper will either report with selective bias, or else where it will outright ignore any negative consequences implicit in the table/graph's results. Therefore this roadblock in AI-driven science journalism is not a minor one.

More tellingly, the fact that a paper is derivative, or only a minor advance (if any) on the state-of-the-art, is often buried in an almost impenetrable citation (i.e., you would need to search for the term, locate a readable PDF copy, and understand the extent of the prior art, before comprehending the lack of originality or novelty in the *new* work).

Alone Again, Naturally

The crowd-sourced method outlined above suggests some possible agreement between common consensus on potential news stories, and professional evaluation of the same. But without context, only the broadest strokes of newsworthiness can apparently be determined.

The very strength of AI lies in its ability, depending on configuration, to isolate [outliers](#) – either for the purpose of [discarding them](#) as a curve-blowing and non-meaningful exception to a trends in a dataset, or (more relevant for news-gathering) to identify meaningful and valuable uncommon instances and occurrences:



Outliers (in red) in a scatter plot. [Source](#)

On the principle that lightning rarely strikes twice, nearly all hit news stories are outliers. In cases where they issue from an active and volatile domain, such as an ongoing war, that domain can be vigilantly scanned with a high probability of newsworthy stories emerging – but at the cost of massive contention, since common attention is likely also focused on the domain.

Many newsworthy scientific leads are, by definition, [not the center of the language distribution](#). They are rare combinations of methods, surprising negative results or anomalous replications. If the model's competence degrades disproportionately on such low-frequency groupings, then the very region where an editorial 'nose' needs to be sharp, becomes the region where the model is least reliable.

Trust Issues

In seeking new stories, journalists balance multiple constraints, including time, access, credibility, audience, and organizational priorities), leading to non-obvious choices. A 2022 [literature review](#) from Denmark characterized journalists as balancing multiple concerns, acutely aware that sources may have agendas or be misinformed; and often bypassing direct checking in favor of indirect trust cues when operating under pressure.

These same 'trust issues' would be a developmental hurdle in any definitive AI-driven newsworthiness-identification system, since engagement with such a platform requires the user to trust that any algorithmically-discarded articles are indeed not worthy of the writer's time.

Extensive beta-testing and retraining or [fine-tuning](#), with human oversight picking up strays and stragglers, could eventually improve the reliability of such an approach; but a shift in national or global culture – such as surprising changes in the political landscape, or the outbreak of war/s – could inevitably upend all the base priorities of such a finely-calibrated system, leaving the AI-dependent writer to rebuild their necessary 'internal domain model' almost from scratch.

First published Monday, April 20, 2026.

Amended Thursday, April 23, 2026 14:13:25, to substitute 'Fortune' for 'WSJ' in 'The Narrow Path', para 2 (thanks to Mark Riley of mathison.ai for pointing it out).

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