

對話世界頂尖學者

Generative AI, Digital Autonomy, and the Future of Communication Research

Discussants: Dr. Hannes Cools and Dr. Justin Chun-ting Ho

Editor: Dr. Justin Chun-ting Ho

Time: June 23, 2026



Dr. Hannes Cools Dr. Justin Chun-ting
Ho

Abstract

In this dialogue, Dr. Hannes Cools examines the rise of generative AI from the perspective of journalism and communication science. Situating ChatGPT within a longer history of technological change in journalism, the conversation explores how generative AI differs from earlier technologies in its speed, scale, and impact on professional autonomy and agency. Dr. Cools reflects on the changing nature of journalistic work, the emergence of AI shame, the importance of critical AI literacy, and the methodological challenges of studying rapidly evolving technologies. The dialogue also considers digital autonomy, open-source language models, and the dependency of news organizations, universities, and

public institutions on Big Tech infrastructures. Rather than treating generative AI as either a threat or a solution, the discussion emphasizes the need for communication researchers to critically examine how these technologies are designed, adopted, normalized, and governed. It concludes by emphasizing the role of communication science in offering careful, evidence-based analysis of generative AI, with attention to how these technologies are designed, adopted, and integrated into professional practice, institutional routines, and our daily life.

Introduction to Dr. Hannes Cools

Dr. Hannes Cools is Assistant Professor of Human Factors in New Technologies at the University of Amsterdam. He is also a Marie Curie Scholar and Fellow at the Digital Democracy Centre (DDC) at the University of Southern Denmark (SDU), as well as an affiliated researcher at the Brown Institute of the Columbia Journalism School in New York City. His research focuses on generative AI, computational journalism, algorithmic recommender systems, and newsroom innovation. From 2018 to 2022, he was a PhD candidate at the Institute for Media Studies (IMS) at KU Leuven in Belgium, where he studied AI in journalism. Before entering academia, he worked as a journalist at De Morgen (DPG Media) in Belgium, covering technology and international politics. Drawing on his experience as both a journalist and a communication scholar, Dr. Cools examines how emerging technologies reshape journalistic work, media organizations, and democratic communication.

HC: Hannes Cools

JH: Justin Chun-ting Ho

JH : Let's begin with something more of a historical perspective. I know your research journey began with datafication in the newsroom. Do you see GenAI as a continuation of that trajectory, or as a qualitatively different disruption for communication and journalism?

HC : I am currently also working on a book on journalism and AI, and I was responsible for the first chapter, which historicizes ChatGPT (Thomson et al., forthcoming). In that chapter, I go all the way back to the 1840s, when the telegraph was introduced. The reason I start there is that technology has always changed journalism, but journalism has also changed technology in return.

Datafication is a more recent example of this, but it still predates ChatGPT. When I was working as a journalist myself, I had three screens in my office. One of them showed an audience metrics system that tracked how readers engaged with articles: how long they spent reading, what they clicked on, and which stories attracted attention. That system would not necessarily have existed twenty or twenty-five years ago. But by the time I was in the newsroom, it was running constantly on a third screen.

Of course, one could say that journalists still retain a great deal of autonomy. But these systems are often introduced without enough critical reflection. To some extent, they also steer editorial decisions, especially under conditions of time pressure, shrinking resources, and the expectation that journalists must do more with less.

In that sense, I do see generative AI as a continuation of earlier technological changes in journalism. It can perhaps be compared with the introduction of the internet. At the same time, the scale and acceleration feel different. When ChatGPT emerged, it quickly became the most widely downloaded applications in history in early 2023. That took many journalistic organizations by storm.

What is striking is how quickly news organizations began to respond. In the Netherlands, for example, there was already a meeting in early 2023 with editors-

in-chief from almost all major news organizations to discuss AI specifically. That is not something I remember seeing in the same way with the rise of the internet or social media. Very early on, there seemed to be a realization within the sector that journalism needed to get a grasp on what this technology was.

Here I am speaking specifically about large language models, which many in the field perceived as a potential threat to the journalistic craft. So, to put it briefly, technology has always changed journalism. But with generative AI, we are dealing with a technology whose uses in journalism are evolving extremely quickly. We now know quite a lot about specific uses of generative AI in journalism (Diakopoulos et al., 2024), but the overall scale of use remains difficult to pin down because it changes so rapidly. That is also a real challenge for communication science. Studying generative AI in journalism can feel like studying a moving target.

JH : You mentioned something very interesting: journalism, and society more broadly, has always encountered new technologies. The internet, for example, fundamentally changed journalism and everyday life. Yet I do not remember the internet being discussed in quite the same way as generative AI is today—as an almost existential threat to human beings, or at least to some essential part of what makes us human. Why do you think generative AI has produced this particular kind of anxiety? What is different about generative AI compared with earlier technologies such as the internet, social media, or audience metrics systems?

HC : I think it is very difficult to predict how the world of work will look even five years from now. We will have to wait and see. But one thing that is clearly changing is the question of autonomy and agency. That is why generative AI touches something very sensitive for many people. It is not only about whether a tool can help you do a task more efficiently. It also touches people's craft, expertise, and even pride in being human: the sense that you know how to do a particular job, that you have built up expertise, and that you have agency in how you perform your work.

Generative AI challenges that assumption. If you look at students today, many of them take these technologies for granted. They ask them almost everything, often quite uncritically. They may even refer to ChatGPT as if it were a friend: “Let’s ask Chat.” So there is already a kind of dependence emerging around large language models. Future generations may take generative AI for granted in the same way, and they may not develop the same kinds of knowledge or expertise that earlier generations had to build up. One advantage of having worked before ChatGPT is that you first had to develop your own expertise before using AI to support certain tasks. I often tell students that you should ask AI questions where you more or less already understand the answer, so that you are able to evaluate what it gives back. But for students, that is often not the case. That is where the risk of uncritical use becomes important.

So I think there are many reasons people feel threatened by this technology, but a central one is that it touches autonomy and agency. Some people are more comfortable delegating specific tasks to AI. Others resist that very strongly, especially when it concerns the core of their professional identity. Some journalists, for example, say they do not want Grammarly or a large language model to polish their writing, because they have been writing for a particular news organization for twenty years. The same may be true in science. We develop a particular way of writing for articles, papers, or other academic work, and we may not want to lose that.

At the same time, much of the marketing around generative AI promises that it will make us more efficient and give us more time. Personally, I am still waiting for that extra time. Sometimes I feel more like a babysitter for these large language models, supervising them before they derail or produce nonsense. So the issue is not only efficiency. It also raises broader questions about labor: what is work, what is a job, and what will a job look like in one year, or in five years?

Another important factor is how news media report on AI. I think news organizations have a responsibility here, and they do not always take it seriously

enough. AI is often framed in either very utopian or very dystopian ways. Because of how human attention works, people tend to remember those extreme frames. I have done research on this as well, looking at which actors are mentioned in AI-related news articles. A large part of the debate is shaped by journalists themselves. Sometimes journalists interview colleagues or other media actors, rather than AI experts or people with deep technical knowledge. That shapes the public debate, and it does not always help reduce the sense of panic around the adoption and use of generative AI.

JH : You mentioned two concepts that I think are especially important here: autonomy and agency. This also connects to your new project, “Digital Autonomy that Fits.” Could you tell us more about this project? What is the rationale behind it, and what does “digital autonomy” mean in the context of journalism and generative AI?

HC : The project is carried out with two media partners in the Netherlands: De Correspondent and Follow the Money. Both are smaller media organizations within what is otherwise a highly concentrated Dutch media landscape. In the Netherlands, much of the media market is dominated by Mediahuis and DPG Media, alongside the public broadcaster and a number of smaller organizations. These smaller media organizations have realized over the years that they have built up a strong dependency on mostly American Big Tech companies. At the same time, they are now experimenting with open-source large language models, for instance to distribute their content to audiences in new formats. The project begins from that context.

The first question we ask is: what do we actually mean by digital autonomy? The term is currently very prominent in public and political debate, and it is often mentioned together with digital sovereignty. I do see differences between these two concepts, and the project tries to unpack those differences more carefully. For these two news organizations, digital autonomy is partly about finding ways to become more independent again. They are interested in building a platform that includes

LLM components and gives audiences access to news content in more personalized ways, but still under strong editorial control. The goal is not to create a filter bubble, but to explore whether personalization and editorial responsibility can coexist.

Digital autonomy has also become a major political issue in the Netherlands. Since Trump took office, many politicians, scholars, and other public figures have become more aware of how dependent European institutions and organizations are on American technology companies. In that sense, it is important to explore what alternatives exist. When we talk about large language models, most people know names such as ChatGPT or Gemini, but they may not know much about open-source alternatives.

The project has three main goals. First, we want to clarify what digital autonomy actually means. Second, we want to examine what matters when news organizations work with open-source large language models, including both technical requirements and organizational requirements. Third, we want to develop a kind of blueprint that shows what steps news organizations would need to take if they want to become more digitally autonomous.

JH : I think this is a very important point, and as you know, I am supportive of open-source approaches. At the same time, when it comes to large language models, it sometimes feels like we are simply choosing between different dependencies: do we give our data to OpenAI, Google, Meta, or another major platform? So I understand the appeal of exploring open-source alternatives. But to play devil's advocate, many so-called open-source models are still developed or heavily shaped by Big Tech companies. From that perspective, how realistic is digital autonomy through open-source AI? Do you think Europe can develop its own open-source or more publicly accountable AI infrastructure, perhaps as a matter of democratic resilience or even national security? And what would it take for European news organizations to have meaningful control over these technologies, rather than simply shifting from one dependency to another?

HC : I do not think Europe will be able to take the lead immediately, at least not this year. But there is now a growing realization in Europe that these questions matter. For example, the European debate on digital sovereignty has increasingly focused on critical infrastructure. Chips are one obvious part of that discussion, especially given the importance of semiconductor companies located in Europe. But large language models are also becoming part of that conversation.

One of the goals of our project is therefore to map which open-source models are actually available. You are right that many so-called open-source large language models are still developed by Big Tech companies. So we need to look carefully at what is really out there, who develops these models, under what conditions they can be used, and whether they actually help news organizations become more digitally autonomous. The City of Amsterdam, for example, has worked on mapping open-source large language models, including practical information such as computing power and other technical requirements. At the moment, I do not think Europe is in a very strong position to take the global lead in this area. But I do think the momentum has grown to invest more seriously in these infrastructures. In discussions about digital sovereignty, hardware is often seen as critical infrastructure, but there is also increasing awareness that large language models—and open-source or more accountable alternatives—can be part of becoming more digitally autonomous.

Another initiative we may explore is GPT-NL, a Dutch language model project in which Utrecht University is involved, together with other actors. I do not think such models will be perfect either, but they may be worth exploring as alternatives, especially in contexts where language, public values, and institutional control matter.

JH : **GenAI is often discussed through ethics: bias, transparency, and responsibility etc. But are there other communication-theory concepts that need to be reworked for the GenAI era?**

HC : We may need new frameworks to understand more complex communication chains. Generative AI makes it harder to separate human communication, human-machine

communication, and machine-to-machine communication. We may increasingly see situations where a machine communicates with a human, the human remains partly in the loop, and another machine continues the process. I also think this is not only a theoretical question, but a methodological one. We may need new methods to study generative AI, but also a return to older methods in new ways. For example, diary studies may help us understand how people's use of AI changes over time.

JH : I am especially interested in your point about methodology. Since generative AI is changing so quickly, what kinds of new methods, methodologies, or research designs do you think communication scholars need?

HC : One interesting approach is data donation. The project by Theo Araujo, for example, is useful because it allows researchers to see how people actually interact with generative AI. It is not necessarily entirely new, but it gives us a way to study real use more directly. With data donation, we can analyze prompts, outputs, topics, and patterns of use. For example, we can examine how long prompts are, how people phrase them, whether they use AI for work, news, elections, personal questions, or other purposes. This helps us move beyond self-reported use and reduces some of the problems of social desirability. But I think data donation should be combined with interviews, focus groups, or reconstructive methods.

Observation studies are also important, especially in workplaces. We need to understand how these tools are embedded into people's daily workflows: in journalism, universities, Microsoft Word, Outlook, Google, and other platforms. Eye-tracking could also be useful, because it can show what people actually look at in AI-generated outputs, how long they look at specific parts, and whether they search for trust cues. This kind of granular research can also inform policy. If we understand how people actually use and evaluate AI systems, we can better think through regulation, literacy, transparency, and responsible use.

We also need older methods, such as diary studies, observation, and longitudinal research, used in new ways. Because the technology is changing so quickly, it is

important to capture snapshots over time—perhaps yearly—to see how adoption and use evolve. Ultimately, I think the key is triangulation: combining old and new methods to understand not only whether people use generative AI, but how they use it, how that use changes, and what it means for communication, journalism, and society.

JH : That is very interesting, because many of the methods you mentioned—interviews, focus groups, diary studies, and observation—are in some sense “old” methods. In computational communication research, there has often been a tendency to move toward increasingly abstract methods: from keywords, to embeddings, to complex models that even researchers may find difficult to interpret.

HC : But with generative AI, I feel there is also a movement back to the basics: toward methods that allow us to understand what people are actually doing, why they are doing it, and how they make sense of these tools in everyday life. In that sense, generative AI is a very new technology, but it also makes many older questions and methods newly relevant again.

JH : ChatGPT became widely visible at the end of 2022, and 2023 felt like the first year of the generative AI era. Now, a few years later, the field has produced a lot of research, discussion, and methodological reflection on generative AI. Do you think communication science is now better prepared to study and respond to generative AI? Or are we still trying to catch up with a technology that continues to move faster than our theories, methods, and institutions?

HC : I think there is still a lot of AI shame. In our focus groups, we saw very different attitudes: some people refuse to use AI altogether, while others seek answers from AI for almost everything as if it were an oracle (Ho et al., 2026). That makes open discussion difficult. People do not always tell colleagues when they use AI, especially for writing, because it may not look good or may invite judgment. In that sense, AI use still feels personal, and sometimes almost taboo.

So I would not say the field is fully ready for generative AI. But we are having many more conversations than before. One positive effect of ChatGPT is that it forces the field to talk seriously about AI, ethics, responsibility, and research practice. That makes me hopeful. There should not be such a thing as AI shame. We should be open about how we use these tools, while continuing to debate what counts as ethical and responsible use. So, no, we are not necessarily ready. But we are better prepared than we were in 2023. We are now studying it, discussing it, and beginning to understand how people actually use generative AI and what it may mean for the field, organizations, and the public.

JH : I totally agree with the AI shame. In one of my own presentations, I asked the audience whether they had used generative AI, and almost everyone raised their hand. So the adoption rate seems extremely high, but we rarely see AI disclosure in formal academic contexts. I think this creates a difficult situation. Once you disclose that you use AI, you make yourself vulnerable to criticism or judgment from others. It takes courage to be open about it. But I also feel that the field is beginning to realize that this silence itself is a problem.

Let me move to the final part of our conversation: the role of communication researchers. Throughout our discussion, we have talked about journalists, students, and other users who are adopting generative AI, sometimes without fully understanding how these systems work or how to use them critically.

With that in mind, what do you think are the most important skills for aspiring communication researchers in the era of generative AI? What should young scholars learn in order to study these technologies critically, use them responsibly, and still develop their own expertise?

HC : Critical thinking will remain an essential skill for communication researchers. That may sound obvious, but I think it is more important than ever in the era of generative AI. I was recently asked whether writing will still be an important skill in five or

ten years. I think writing as a skill may change. Research papers may also change in format and perhaps become more interactive. But what worries me is that generative AI removes a lot of friction from the research process.

I feel lucky that I knew academic work before ChatGPT. We had to write our theses ourselves, struggle with concepts, get stuck, and find our own way through. That difficulty was not always pleasant, but it was valuable. Sometimes students and researchers need time to sit with a concept, a research question, or a method. With generative AI, the tool is always there asking, “How can I help you?” And once you get stuck, it becomes very tempting to ask ChatGPT immediately. That is not necessarily bad. But sometimes it is important to resist that impulse and think through the problem yourself. A life without friction is not necessarily a desirable life, especially for learning and research.

This connects to autonomy. Future researchers need to know how to use AI while still remaining in control of the information they produce, the code they run, and the arguments they make. If we ask a model to write a text and simply copy and paste it without reading it carefully, what are we giving away? What are we losing? At that point, the text is no longer fully ours; it becomes a surrogate for what our own thinking and writing could have produced. Some skills may change or move into the background, including certain forms of writing. These tools are here to stay. As Shoshana Zuboff (2019) has argued, once a technology can automate particular tasks, there is often strong pressure for those tasks to be automated. In that sense, what can be automated will be automated. But we pay a price for that automation.

So I hope future communication researchers still experience some friction. They should learn to think critically, question the design and interests behind AI systems, and reflect on the ethical consequences of using them. At the same time, I am hopeful. Younger generations are growing up with these tools, but they may also become more critical of them. As scholars and teachers, we have a clear responsibility to help them understand both the possibilities and the risks.

JH : I really like your answer, especially because it is not new—and I mean that in a positive sense. Critical thinking is something academia has emphasized for a very long time, perhaps since the Greek philosophers. It speaks to something fundamental about being a researcher, and perhaps about being human: no matter what tools we use, how fast we work, or how much technology changes, we still need the capacity to think critically, to preserve autonomy, and to reflect on what we are doing.

At the same time, the fear about new technology is also very old. When the internet became popular, people worried that it would make us dumber. When Google became central to knowledge-seeking, people worried that we would stop remembering things. And even in Plato's (2009) *Phaedrus*, Socrates warns that writing could weaken memory and create only the appearance of wisdom.

HC : I do not think generative AI will necessarily make us dumb. What worries me more is the possibility of a frictionless life, where we no longer have to struggle with ideas, decisions, or forms of expression. That is a more philosophical concern. This is why discussion is so important. By talking openly about how we use these technologies—and by getting rid of AI shame—we can better calibrate what we want to make of them. We should not assume that everyone already understands generative AI just because ChatGPT has been around for several years. Some people still do not know what these systems are or how they work. This is true among older people, younger people, journalists, and even within universities.

So critical thinking also requires training. But that training cannot be a one-time checklist or a certificate that says, "This is ethical AI use for the next five years." Responsible use is contextual, and it needs to be practiced continuously. That also makes generative AI an important object of study. We need to examine AI literacy across different groups and contexts, and identify not only risks and challenges, but also opportunities. As a field, we often focus on what is going wrong, and a lot is going wrong. But we should not become blind to what might be valuable or productive.

At the same time, communication science should be more proactive about transparency, especially in academic practice. Right now, there is still a taboo around AI use, and some people fear that disclosing it could hurt their chances of publication. We are still making sense of how to approach and study these technologies.

But I think this is also a valuable problem for the field. It forces us to reflect on fundamental questions: what methods do we need, what theories do we use, and how should we study a technology that is itself changing how communication works?

JH : I like how our discussion began with the threats and problems created by generative AI, but ended up with how it also creates opportunities. Because it is unsettling so many existing practices, assumptions, and routines, it also opens up many new things to study. There are many research questions waiting for us to tackle. That may be a good thing for us as scholars, and perhaps also for society. Whether these changes are ultimately for better or worse, something important is happening, and communication science has a role to play in making sense of it. Before we finish, do you have any final words to add?

HC : One final thought is that communication science can also serve as a counterweight to Big Tech. We have a responsibility not simply to accept the marketing and public relations narratives of these companies, but to study what is actually happening. In that sense, our responsibility has become even more important. Almost like journalists, we need to hold powerful actors to account through our research, using rigorous and vetted methodologies. Communication science can have a kind of watchdog function. We are independent, and we should be able to ask the research questions we think are important, choose the methods that are appropriate, and remain critical of these technologies.

That role will only become more important as we try to make sense of generative AI. It will also inform how we use these technologies ourselves. I do not see our personal use of AI and our research on AI as entirely separate. How we deal with these technologies personally can shape our research questions and perspectives,

and we should not forget that positionality. So I think being a counterweight and taking up that responsibility will be very important for the field going forward. We should not feel too threatened by these technologies. One of the things I have learned over the years is that it is not only up to Big Tech to determine what our future will look like.

JH : I think this is a very profound message for all of us to remember. Big Tech often provides services that are extremely useful, convenient, and sometimes even free. But these companies are not charities. They offer these services for a reason, and they have their own interests, incentives, and agendas. Generative AI also does not exist in a vacuum. It is shaped by business models, infrastructures, and political contexts. As communication researchers, we need to remember that background. We should study not only what these technologies can do, but also who builds them, who benefits from them, and what kinds of dependencies they create.

AI Disclaimer:

The original interview was conducted in English. ChatGPT (GPT-5.5; accessed on 24 June 2026) was used to assist with the Chinese translation of the manuscript using the following prompt: “Translate this article into Taiwan Chinese, using academic writing style: {English manuscript}”. The editor reviewed, edited, and verified the final translation, and assumes full responsibility for the accuracy and integrity of the content presented in this publication.

References

- Diakopoulos, N., Cools, H., Li, C., Helberger, N., Kung, E., & Rinehart, A. (2024). *Generative AI in journalism: The evolution of newswork and ethics in a generative information ecosystem*. Associated Press. <https://doi.org/10.13140/RG.2.2.31540.05765>
- Ho, J. C., Cools, H., Strycharz, J., Möller, A. M., Van Berlo, Z. M. C., Stolwijk, S., Trilling, D., & Araujo, T. (2026). *Guidance over guidelines? Unpacking the uses and concerns of generative AI in Communication Science*. SocArXiv. https://doi.org/10.31235/osf.io/dncuf_v1
- Plato (2009). *Phaedrus*. Oxford University Press.
- Thomson, T. J., Thomas, Ryan. J., & Cools, H. (Forthcoming). *Generative AI and journalism: From production to perception*. Routledge.
- Zuboff, S. (2019). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. Profile Books.

對話世界頂尖學者

生成式人工智慧、數位自主性與傳播研究的未來 Generative AI, Digital Autonomy, and the Future of Communication Research

Discussants: Dr. Hannes Cools、Dr. Justin Chun-ting Ho

Editor: Dr. Justin Chun-ting Ho

Time: June 23, 2026



Dr. Hannes Cools Dr. Justin Chun-ting
Ho

摘要

在本對話中，Hannes Cools 博士從新聞學與傳播科學的視角，檢視生成式 AI 的興起。對話將 ChatGPT 置於新聞業技術變遷的長歷史脈絡中，探討生成式 AI 如何在速度、規模，以及對專業自主性與能動性的影響上，有別於早期技術。Cools 博士反思新聞工作的性質如何改變、AI 羞恥感的出現、批判性 AI 素養的重要性，以及研究快速演變技術所面臨的方法論挑戰。本對話也討論數位自主性、開源語言模型，以及新聞組織、大學與公共機構對大型科技公司基礎設施的依賴。與其將生成式 AI 視為威脅或解方，本討論強調傳播研究者有必要批判性地檢視這些技術如何被設計、採用、

常態化與治理。最後，本對話強調傳播科學在生成式 AI 研究中所扮演的角色：提供審慎且以證據為基礎的分析，並關注這些技術如何被設計、採用，以及如何整合進專業實作、制度性慣例與日常生活之中。

Dr. Hannes Cools 博士介紹

Hannes Cools 博士現任阿姆斯特丹大學 Human Factors in New Technologies 助理教授。他同時也是瑪麗·居里學者人才基金學者，並於南丹麥大學數位民主中心擔任研究員，亦為紐約哥倫比亞大學新聞學院 Brown Institute 的兼任研究員。他的研究聚焦於生成式 AI、運算新聞學、演算法推薦系統與新聞室創新。2018 至 2022 年間，他曾於比利時魯汶大學媒體研究所攻讀博士，研究新聞業中的 AI。在進入學術界之前，他曾任職於比利時 De Morgen (DPG Media) 擔任記者，報導科技與國際政治。憑藉其作為記者與傳播學者的雙重經驗，Cools 博士檢視新興技術如何重塑新聞工作、媒體組織與民主傳播。

HC: Hannes Cools

JH: Justin Chun-ting Ho

JH： 讓我們先從歷史視角談起。我知道您的研究歷程始於新聞室中的資料化。您認為生成式 AI 是這條軌跡的延續，還是對傳播與新聞業而言具有質變意義的顛覆？

HC： 我目前正在撰寫一本關於新聞業與 AI 的書，並負責第一章，該章旨在將 ChatGPT 歷史化（Thomson et al., forthcoming）。在那一章中，我一路追溯到 1840 年代電報被引入之時。我之所以從那裡開始，是因為技術一直都在改變新聞業，但新聞業也反過來改變了技術。

資料化是較近的一個例子，但它仍早於 ChatGPT。當我自己擔任記者時，我的辦公桌上有三個螢幕。其中一個顯示的是受眾指標系統，用來追蹤讀者如何與文章互動：他們閱讀多久、點擊了什麼，以及哪些報導吸引注意。這樣的系統在二十或二十五年前未必存在。但當我在新聞室工作時，它已經持續地運作在第三個螢幕上。

當然，人們可以說，記者仍然保有相當程度的自主性。然而，這些系統往往是在缺乏充分批判性反思的情況下被引入。在某種程度上，它們也會引導編輯決策，尤其是在時間壓力、資源縮減，以及記者被期待以更少資源完成更多工作的條件下。

從這個意義上說，我確實將生成式 AI 視為新聞業早期技術變遷的延續。它或許可以與網際網路的引入相提並論。與此同時，其規模與加速度又顯得不同。當 ChatGPT 出現時，它在 2023 年初迅速成為史上下載最廣泛的應用程式之一。這讓許多新聞組織措手不及。

令人印象深刻的是，新聞組織開始回應的速度相當快。例如在荷蘭，早在 2023 年初，幾乎所有主要新聞組織的總編輯便已召開會議，專門討論 AI。這並不是我記憶中網際網路或社群媒體興起時曾以相同方式出現的情況。很早期，新聞業內部似乎就已意識到，新聞界需要理解並掌握這項技術究竟是什麼。

這裡我特別指的是大型語言模型，許多業界人士將其視為對新聞技藝的潛在威脅。因此，簡要而言，技術始終都在改變新聞業。但在生成式 AI 的情況下，我們面對的是一種其新聞用途正在極快速演變的技術。如今，我們已經對

生成式 AI 在新聞業中的特定用途有相當多了解 (Diakopoulos et al., 2024)，但其整體使用規模仍難以精確掌握，因為它變化得太快。這對傳播科學而言也是一項真實挑戰。研究新聞業中的生成式 AI，有時就像是在研究一個不斷移動的目標。

JH： 您剛才提到一個非常有趣的觀點：新聞業，乃至更廣泛的社會，向來都會遭遇新技術。例如，網際網路根本性地改變了新聞業與日常生活。然而，我不記得網際網路曾以今天生成式 AI 被討論的方式來討論——也就是被視為對人類，或至少對某種構成人之為人的本質面向，幾乎具有存在性威脅。您認為為什麼生成式 AI 會引發這種特定形式的焦慮？相較於網際網路、社群媒體或受眾指標系統等早期技術，生成式 AI 有何不同？

JH： 我認為，要預測五年後的工作世界會是什麼樣貌，是非常困難的。我們必須拭目以待。但有一件事明顯正在改變，那就是自主性與能動性的問題。這也是為什麼生成式 AI 觸及了許多人非常敏感的部分。它不只是關於某個工具是否能幫助你更有效率地完成任務；它也觸及人們的技藝、專業知識，甚至作為人的驕傲：也就是你知道如何完成某項工作、你已經累積了專業能力，並且你對如何執行自己的工作具有能動性的那種感受。

生成式 AI 挑戰了這個假設。若你觀察今天的學生，許多人已將這些技術視為理所當然。他們幾乎什麼都問它們，而且往往相當不加批判。他們甚至可能像稱呼朋友一樣稱呼 ChatGPT：「我們來問 Chat。」因此，圍繞大型語言模型，已經出現某種依賴。未來世代也可能以同樣方式將生成式 AI 視為理所當然，而他們也許不會發展出早期世代必須累積的相同知識或專業能力。在 ChatGPT 出現之前工作過的一個優勢是，你必須先發展自己的專業知識，然後再使用 AI 來支援特定任務。我經常告訴學生，你應該向 AI 提問那些你大致已經理解答案的問題，如此你才有能力評估它回覆的內容。但對學生而言，這往往並非事實。這正是非批判性使用的風險變得重要之處。

因此，我認為人們之所以感到受到這項技術威脅，有許多原因，但其中一個核心原因是它觸及自主性與能動性。有些人較能接受將特定任務委派給 AI；

另一些人則強烈抗拒，尤其當這涉及其專業認同的核心時。例如，有些記者表示，他們不希望 Grammarly 或大型語言模型潤飾自己的寫作，因為他們已經為特定新聞組織寫作二十年了。在科學領域可能也是如此。我們會為文章、論文或其他學術工作發展出特定的寫作方式，而我們未必想失去它。

與此同時，許多圍繞生成式 AI 的行銷話語都承諾，它會讓我們更有效率，並給我們更多時間。就我個人而言，我還在等待那個「多出來的時間」。有時候，我反而覺得自己更像是在照顧這些大型語言模型的保母，在它們失控或產生胡言亂語之前監督它們。因此，問題不只是效率。它也引發了更廣泛的勞動問題：什麼是工作？什麼是職業？一年後，或五年後，一份工作會是什麼樣子？

另一個重要因素是新聞媒體如何報導 AI。我認為新聞組織在這方面負有責任，但它們並不總是足夠認真看待這項責任。AI 經常被以非常烏托邦或反烏托邦的方式框架。由於人類注意力運作的方式，人們往往會記住這些極端框架。我也曾針對這一點進行研究，分析 AI 相關新聞文章中提及哪些行動者。相當大一部分的討論是由記者自身形塑的。有時記者訪問的是同業或其他媒體行動者，而不是 AI 專家或具有深厚技術知識的人。這會形塑公共討論，也不總是有助於降低圍繞生成式 AI 採用與使用的恐慌感。

JH： 您提到兩個我認為在此特別重要的概念：自主性與能動性。這也連結到您的新計畫「合宜的數位自主性」。您能否進一步介紹這個計畫？它背後的理據是什麼？在新聞業與生成式 AI 的脈絡中，「數位自主性」意味著什麼？

HC： 這個計畫與荷蘭兩個媒體夥伴合作進行：De Correspondent 與 Follow the Money。兩者都是較小型的媒體組織，而它們所在的荷蘭媒體環境則高度集中。在荷蘭，多數媒體市場由 Mediahuis 與 DPG Media 主導，此外還有公共廣播機構與若干較小型組織。這些較小型的媒體組織多年來已意識到，它們對主要來自美國的大型科技公司形成了強烈依賴。與此同時，它們現在也正在實驗開源大型語言模型，例如用來以新的格式將內容傳遞給受眾。這個計畫正是從這樣的脈絡出發。

我們提出的第一個問題是：我們實際上所謂的數位自主性究竟是什麼？這

個詞目前在公共與政治討論中非常突出，而且經常與數位主權一起被提及。我確實認為這兩個概念之間存在差異，而本計畫試圖更仔細地拆解這些差異。對這兩個新聞組織而言，數位自主性部分意味著尋找重新變得更獨立的方法。它們有興趣建立一個包含大型語言模型元件的平台，讓受眾能以更個人化的方式接觸新聞內容，但同時仍置於強而有力的編輯控制之下。其目標不是製造資訊同溫層，而是探索個人化與編輯責任是否能夠共存。

數位自主性在荷蘭也已經成為一個重要的政治議題。自川普上任以來，許多政治人物、學者與其他公共人物更加意識到歐洲機構與組織對美國科技公司的依賴程度。從這個意義上說，探索有哪些替代方案非常重要。談到大型語言模型時，多數人知道 ChatGPT 或 Gemini 等名稱，但他們可能不太了解開源替代方案。

這個計畫有三個主要目標。第一，我們希望釐清數位自主性的實際意涵。第二，我們希望檢視新聞組織在使用開源大型語言模型時，哪些因素是重要的，包括技術需求與組織需求。第三，我們希望發展一種藍圖，說明若新聞組織希望變得更具數位自主性，需要採取哪些步驟。

JH： 我認為這是一個非常重要的觀點。您也許知道，我個人非常支持開源。與此同時，當涉及大型語言模型時，有時感覺我們只是從不同依賴之間做選擇：我們究竟要把資料交給 OpenAI、Google、Meta，還是另一個大型平台？因此，我理解探索開源替代方案的吸引力。但若扮演魔鬼代言人，許多所謂的開源模型仍然是由大型科技公司開發，或受到它們的深刻形塑。從這個角度來看，透過開源 AI 實現數位自主性有多現實？您認為歐洲是否能發展自己的開源或更具公共問責性的 AI 基礎設施，或許作為民主韌性、甚至國家安全的一部分？而歐洲新聞組織若要對這些技術擁有有意義的控制權，而不只是從一種依賴轉向另一種依賴，需要具備什麼條件？

HC： 我不認為歐洲能夠立即取得領先地位，至少今年不會。但歐洲現在越來越意識到這些問題的重要性。例如，歐洲關於數位主權的討論已越來越聚焦於關鍵基礎設施。晶片顯然是這項討論中的一部分，尤其考量到位於歐洲的半導體公司的重要性。但大型語言模型也逐漸成為這場討論的一部分。

因此，我們計畫的目標之一，就是盤點目前實際可用的開源模型。你說得沒錯，許多所謂的開源大型語言模型仍由大型科技公司開發。因此，我們需要仔細考察實際上有哪一些模型、由誰開發、在什麼條件下可以使用，以及它們是否真正有助於新聞組織變得更具數位自主性。舉例而言，阿姆斯特丹市政府曾致力於盤點開源大型語言模型，包括運算能力與其他技術需求等實用資訊。目前，我不認為歐洲在這個領域處於能夠引領全球的強勢位置。但我確實認為，投資這些基礎設施的動能已經增強。在數位主權的討論中，硬體經常被視為關鍵基礎設施；但人們也越來越意識到，大型語言模型——以及開源或更具問責性的替代方案——也可以成為實現更高數位自主性的一部分。

另一個我們可能會探索的倡議是 GPT-NL，這是一個荷蘭語言模型計畫，烏特勒支大學以及其他行動者均參與其中。我不認為這類模型會是完美的，但它們或許值得作為替代方案加以探索，尤其是在語言、公共價值與制度控制具有重要性的脈絡中。

JH： 生成式 AI 經常透過倫理面向被討論：偏誤、透明度與責任等等。但是否還有其他傳播理論概念需要為生成式 AI 時代重新調整？

HC： 我們或許需要新的框架，來理解更複雜的傳播鏈。生成式 AI 使人類傳播、人機傳播與機器對機器傳播之間的界線更難區分。我們可能越來越常看到這樣的情境：一台機器與人類溝通，人類仍部分處於迴路之中，而另一台機器接續這個過程。我也認為，這不只是理論問題，也是方法論問題。我們可能需要新的方法來研究生成式 AI，但也需要以新的方式回到較舊的方法。例如，日記研究可以幫助我們理解人們的 AI 使用如何隨時間改變。

JH： 我對您關於方法論的觀點特別感興趣。既然生成式 AI 變化如此快速，您認為傳播學者需要哪些新方法、方法論或研究設計？

HC： 一個有趣的取徑是資料捐贈。例如 Theo Araujo 的計畫便很有用，因為它讓研究者得以觀察人們實際上如何與生成式 AI 互動。這不一定是全然新的方法，但它提供了一種更直接研究真實使用的方式。透過資料捐贈，我們可以分析提示詞、輸出、主題與使用模式。例如，我們可以檢視提示詞有多長、人們如何措辭、他們是否將 AI 用於工作、新聞、選舉、個人問題或其他目的。這有助

於我們超越自我報告式使用，並減少社會期許偏誤所帶來的一些問題。不過，我認為資料捐贈應與訪談、焦點團體或重構式方法結合。

觀察研究也很重要，尤其是在工作場域中。我們需要理解這些工具如何嵌入人們的日常工作流程：在新聞業、大學、Microsoft Word、Outlook、Google 以及其他平台之中。眼動追蹤也可能有用，因為它能顯示人們實際上在 AI 生成的輸出中看了什麼、在特定部分停留多久，以及他們是否尋找信任線索。這類細緻研究也能為政策提供參考。如果我們理解人們實際上如何使用並評估 AI 系統，就能更好地思考監管、素養、透明度與負責任使用。

我們也需要以新的方式運用較舊的方法，例如日記研究、觀察與縱貫研究。由於技術變化如此快速，重要的是在不同時間點捕捉截面——或許每年一次——以觀察採用與使用如何演變。最終，我認為關鍵在於三角檢證：結合新舊方法，以理解人們不只是是否使用生成式 AI，而是如何使用、這種使用如何改變，以及它對傳播、新聞業與社會意味著什麼。

JH： 這非常有趣，因為您提到的許多方法——訪談、焦點團體、日記研究與觀察——在某種意義上都是「舊」方法。在運算傳播研究中，往往有一種傾向，是走向越來越抽象的方法：從關鍵字，到詞嵌入，再到連研究者自己都可能覺得難以詮釋的複雜模型。

HC： 但在生成式 AI 的情況下，我也感覺到一種回到基礎的趨勢：回到那些能讓我們理解人們實際在做什麼、為什麼這樣做，以及他們如何在日常生活中理解這些工具的方法。從這個意義上說，生成式 AI 是一種非常新的技術，但它也讓許多較舊的問題與方法重新變得重要。

JH： ChatGPT 在 2022 年底變得廣為人知，而 2023 年像是生成式 AI 時代的第一年。如今，幾年過去，這個領域已經產出大量關於生成式 AI 的研究、討論與方法論反思。您認為傳播科學現在是否已更準備好研究並回應生成式 AI？還是我們仍在試圖追趕一項持續比我們的理論、方法與制度移動得更快的技術？

HC： 我認為 AI 羞恥感仍然相當普遍。在我們的焦點團體中，我們看到非常不同的態度：有些人完全拒絕使用 AI，而另一些人則幾乎什麼都向 AI 尋求答案，彷彿它是先知一般（Ho et al., 2026）。這使得開放討論變得困難。人們並不總是

會告訴同事自己使用 AI，尤其是在寫作方面，因為這可能看起來不太好，或招致他人的評價。從這個意義上說，AI 使用仍然顯得很個人，有時甚至近乎禁忌。

因此，我不會說這個領域已經完全準備好面對生成式 AI。但我們確實比以前有更多對話。ChatGPT 的一個正面效果是，它迫使這個領域嚴肅談論 AI、倫理、責任與研究實作。這讓我感到樂觀。AI 羞恥感不應該存在。我們應該開放地討論自己如何使用這些工具，同時持續辯論什麼才算是合乎倫理且負責任的使用。所以，不，我們未必已經準備好了。但我們比 2023 年時準備得更充分。現在，我們正在研究它、討論它，並開始理解人們實際上如何使用生成式 AI，以及它對這個領域、組織與公眾可能意味著什麼。

JH： 我完全同意 AI 羞恥感這一點。在我自己的某次演講中，我問聽眾是否曾使用過生成式 AI，幾乎所有人都舉手。如此看來，AI 採用率似乎非常高，但我們很少在正式學術脈絡中要求 AI 使用揭露。我認為這造成了一個困難局面。一旦你揭露自己使用 AI，就讓自己容易受到他人的批評或評價。要對此保持開放是需要勇氣的。但我也感覺到，這個領域開始意識到，這種沉默本身就是一個問題。

讓我進入我們對話的最後一部分：傳播研究者的角色。整個討論中，我們談到了記者、學生與其他使用者如何採用生成式 AI，有時並未充分理解這些系統如何運作，也未必知道如何批判性地使用它們。在這樣的脈絡下，您認為在生成式 AI 時代，有志成為傳播研究者的人最重要的能力是什麼？年輕學者應該學習什麼，才能批判性地研究這些技術、負責任地使用它們，同時仍發展出自己的專業能力？

HC： 批判性思考仍將是傳播研究者的一項核心能力。這聽起來或許顯而易見，但我認為在生成式 AI 時代，它比以往任何時候都更重要。我最近被問到，五年或十年後寫作是否仍會是一項重要能力。我認為，寫作作為一項能力可能會改變。研究論文的形式也可能改變，或許會變得更具互動性。但令我擔憂的是，生成式 AI 從研究過程中移除了許多摩擦力。

我很慶幸自己在 ChatGPT 出現之前就已經了解學術工作。我們必須自己寫論文，與概念搏鬥，經歷卡關，並自己找出前進的道路。那種困難並不總是令人愉快，但它是有價值的。有時學生與研究者需要時間與一個概念、一個研究問題或一種方法共處。生成式 AI 出現後，工具總是在那裡問你：「我可以怎麼幫你？」而一旦你卡住，就很容易立刻去問 ChatGPT。這不一定是壞事。但有時候，抵抗這種衝動、自己思考問題是很重要的。沒有摩擦的人生未必是值得嚮往的人生，尤其對學習與研究而言更是如此。

這也連結到自主性。未來的研究者需要知道如何使用 AI，同時仍然掌控自己所產出的資訊、執行的程式碼，以及提出的論證。如果我們要求模型寫一段文字，然後未經仔細閱讀便直接複製貼上，我們究竟交出了什麼？又失去了什麼？在那個時刻，文本便不再完全屬於我們；它成為我們自身思考與寫作本可產生之成果的替代品。有些能力可能會改變，或退居背景，包括某些形式的寫作。這些工具將會持續存在。正如 Shoshana Zuboff (2019) 所主張的，一旦某項技術能夠自動化特定任務，就常會出現強大的壓力，促使這些任務被自動化。從這個意義上說，凡是能被自動化的，都將被自動化。但我們也會為這種自動化付出代價。

因此，我希望未來的傳播研究者仍然經歷某些摩擦力。他們應該學會批判性思考，質疑 AI 系統背後的設計與利益，並反思使用它們的倫理後果。同時，我也抱持希望。年輕世代正與這些工具一同成長，但他們也可能對這些工具變得更加批判。作為學者與教師，我們有明確的責任，幫助他們理解這些工具的可能性與風險。

JH： 我非常喜歡您的回答，尤其是因為它並不新——我是以正面意義來說。批判性思考是學術界長久以來一直強調的事，或許自希臘哲學家以來便如此。它觸及作為研究者的某種根本性，也或許觸及作為人的某種根本性：無論我們使用什麼工具、工作速度多快，或技術如何變遷，我們仍然需要批判性思考、保有自主性，並反思自己正在做什麼的能力。

與此同時，對新技術的恐懼也非常古老。當網際網路開始普及時，人們擔心它會讓我們變笨。當 Google 成為知識搜尋的核心時，人們擔心我們會停止

記憶事物。甚至在柏拉圖（2009）的《斐德羅篇》中，蘇格拉底也曾警告，書寫可能削弱記憶，並只創造出智慧的表象。

HC： 我不認為生成式 AI 必然會讓我們變笨。我更擔憂的是一種無摩擦生活的可能性，在那樣的生活中，我們不再需要與想法、決策或表達形式搏鬥。這是一種更哲學性的憂慮。這也是為什麼討論如此重要。透過開放談論我們如何使用這些技術——並消除 AI 羞恥感——我們能更好地校準自己希望如何看待並運用它們。我們不應該假設，僅僅因為 ChatGPT 已經存在數年，每個人就都已經理解生成式 AI。有些人仍然不知道這些系統是什麼，或它們如何運作。這在年長者、年輕人、記者，甚至大學內部都是真實存在的情況。

因此，批判性思考也需要訓練。但這種訓練不能只是一份一次性的檢核表，或一張聲稱「這就是未來五年倫理 AI 使用方式」的證書。負責任的使用是脈絡性的，而且需要持續實踐。這也使生成式 AI 成為重要的研究對象。我們需要檢視不同群體與脈絡中的 AI 素養，並辨識風險與挑戰之外的機會。作為一個領域，我們經常聚焦於哪些事情出了問題，而確實有很多事情正在出問題。但我們也不應因此對可能具有價值或生產性的面向視而不見。

與此同時，傳播科學應該更主動地處理透明度問題，尤其是在學術實作中。目前，圍繞 AI 使用仍存在禁忌，有些人擔心揭露 AI 使用可能損害其發表機會。我們仍在理解應如何接近並研究這些技術。

但我認為，這對本領域而言也是一個有價值的問題。它迫使我們反思一些根本問題：我們需要什麼方法？我們使用哪些理論？我們應如何研究一項本身正在改變傳播運作方式的技術？

JH： 我喜歡我們的討論從生成式 AI 所造成的威脅與問題開始，但最後卻走向它如何也創造了機會。正因為它擾動了許多既有的實作、假設與慣例，它也開啟了許多新的研究可能。有許多研究問題正等待我們處理。這對我們作為學者而言或許是好事，對社會而言也可能如此。無論這些變化最終是好是壞，某些重要的事情正在發生，而傳播科學有責任協助理解它。在我們結束之前，您還有什麼最後想補充的話嗎？

HC： 最後我想說的是，傳播科學也可以作為大型科技公司的制衡力量。我們有責任

不只是接受這些公司的行銷與公關敘事，而是研究實際上正在發生什麼。從這個意義上說，我們的責任變得更加重要。幾乎像記者一樣，我們需要透過研究，以嚴謹且經過檢證的方法論，要求權力行動者負責。傳播科學可以具有某種監督功能。我們是獨立的，也應該能夠提出我們認為重要的研究問題，選擇適切的方法，並對這些技術保持批判。

隨著我們試圖理解生成式 AI，這個角色只會變得更加重要。它也會影響我們自己如何使用這些技術。我不認為我們個人對 AI 的使用，以及我們對 AI 的研究，是完全分離的。我們個人如何處理這些技術，可能形塑我們的研究問題與視角，而我們不應忘記這種位置性。因此，我認為成為一種制衡力量，並承擔這份責任，對本領域未來而言將非常重要。我們不應該對這些技術感到過度威脅。這些年來我學到的一件事是，決定我們未來樣貌的，並不只是大型科技公司。

JH： 我認為這是一個非常深刻、值得我們所有人記住的訊息。大型科技公司經常提供極其有用、便利，有時甚至免費的服務。但這些公司並不是慈善機構。它們提供這些服務有其原因，也有自己的利益、誘因與議程。生成式 AI 也並非存在於真空之中。它受到商業模式、基礎設施與政治脈絡的形塑。作為傳播研究者，我們需要記得這個背景。我們不僅應研究這些技術能做什麼，也應研究誰建構它們、誰從中受益，以及它們創造了哪些形式的依賴。

AI 使用聲明：

原始訪談以英文進行。本文使用 ChatGPT（GPT-5.5；於 2026 年 6 月 24 日使用）協助將稿件翻譯為台灣中文，使用的提示語如下：「Translate this article into Taiwan Chinese, using academic writing style: {English manuscript}」。編輯已審閱、修訂並確認最終譯文，並對本出版內容之準確性與完整性負全部責任。

參考文獻

- Diakopoulos, N., Cools, H., Li, C., Helberger, N., Kung, E., & Rinehart, A. (2024). *Generative AI in journalism: The evolution of newswork and ethics in a generative information ecosystem*. Associated Press. <https://doi.org/10.13140/RG.2.2.31540.05765>
- Ho, J. C., Cools, H., Strycharz, J., Möller, A. M., Van Berlo, Z. M. C., Stolwijk, S., Trilling, D., & Araujo, T. (2026). *Guidance over guidelines? Unpacking the uses and concerns of generative AI in Communication Science*. SocArXiv. https://doi.org/10.31235/osf.io/dncuf_v1
- Plato (2009). *Phaedrus*. Oxford University Press.
- Thomson, T. J., Thomas, Ryan. J., & Cools, H. (Forthcoming). *Generative AI and journalism: From production to perception*. Routledge.
- Zuboff, S. (2019). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. Profile Books.

