

Smart enough to mislead

the functional shortcomings and ethical dilemmas of generative AI use in metadata work

Fran Frenzel  0009-0009-3890-5335

Metadata Analyst, London School of Economics and Political Science

Received: 10 June 2025 | Published: 17 June 2025

ABSTRACT

This article critically examines the applicability of generative AI in library metadata creation and cataloguing, arguing that despite growing interest and experimentation, such technologies remain fundamentally unsuited for this domain. Drawing on recent literature, surveys, and institutional case studies, the author demonstrates that generative AI tools consistently produce metadata outputs that are unreliable, inconsistent, and ethically problematic. While machine learning offers potential in specific, supervised metadata functions, generative AI's reliance on probabilistic outputs, lack of transparency, and tendency to hallucinate undermine the accuracy and reliability essential to cataloguing. The article also explores the broader ethical implications of AI adoption in libraries, including issues of bias, environmental impact, copyright concerns, and labour exploitation. The author argues that fully automated metadata creation using generative AI is neither technically viable nor ethically responsible and instead advocates for cautious, critically informed AI integration, emphasising the continued necessity of human oversight and ethical scrutiny in metadata work.

KEYWORDS generative AI; metadata creation; metadata enhancement; AI ethics

CONTACT Fran Frenzel  f.frenzel@lse.ac.uk  London School of Economics and Political Science

In early autumn 2023, *Information Technology and Libraries* published an article proclaiming that ChatGPT can “generate accurate MARC records using RDA and other standards such as the Dublin Core Metadata Element Set” (Brzustowicz, 2023, p. 2). I was intrigued and began reading, quickly followed by bewilderment and dismay about what was labelled “an accurate and effective record” (Brzustowicz, 2023, p. 2) and “comparable to the professional catalogers’ work” (Brzustowicz, 2023, p. 3), but was riddled with errors and hallucinations¹. All the article demonstrated to me was ChatGPT’s utter inadequacy for cataloguing, and I was horrified by the potential consequences of assertions, such as the ones quoted above, being taken at face value. On the positive side though, it got me interested in AI in relation to library metadata work. So, nearly 2 years on, are we any closer to AI cataloguers?

¹ The relevant mailing lists were not amused either and some responses were published in the next issue of *Information Technology and Libraries* that outline the problems very well (DeZelar-Tiedman, 2023; Amram, Malamud & Hollingsworth (2023) and Floyd (2023)).

Yes and no. In terms of AI technologies being used for cataloguing and record enhancement, yes, there are many possibilities to assist cataloguers or streamline processes. But in terms of generative AI being the AI technology to do the job, very much not. Why? Well, generative AI is singularly concerned with creating the most likely response to a prompt; whether this most likely response is factually correct or not is of no importance to it. This is in direct opposition to what one is doing and what is important when cataloguing: truthfully representing the resource, not representing it with the statistically most likely information. That is not to say that there aren't parts of records or tasks in metadata creation where generative AI's statistical approach might be useful, such as summary creation, for example.

Throughout this article I make a distinction between

- **AI**, meaning AI overall,
- **machine learning**, meaning the field of study with AI research, and
- **generative AI**, meaning a specific subtype of AI utilising generative machine learning models, deep learning and neural networks.

To understand how generative AI functions, it is useful to get a grounding on how large language models (LLMs) work. The Financial Times has published a very accessible primer on how LLMs function and the enormous difference the development of the transformer deep learning architecture has made, leading to generative pre-trained transformers (GPT, a type of large language model) which underpin generative AI tools such as OpenAI's ChatGPT, Microsoft's Copilot, Google's Gemini, Anthropic's Claude or Meta's Llama ([Murgia et al., 2023](#)). Hicks, Humphries and Slater ([2024](#)) also give a good and accessible explanation of LLMs' functionality. With this information in mind, it becomes easy to understand why LLMs produce superficially convincing-looking records that, upon closer inspection, reveal a multitude of problems.

Opportunities

Chen and Li ([2024](#)) published results of a survey conducted in early 2024 to gauge cataloguing and metadata professionals' perception of AI in relation to their roles. They found that AI does not yet play a significant role in respondents' jobs ([Chen and Li, 2024](#), p. 321), but that there is "a growing believe [sic] that AI would play an increasingly important role in their work" ([Chen and Li, 2024](#), p. 322). Among the questions asked were in which areas of respondents' work AI is currently used, and in which areas of metadata creation they think AI would be most beneficial. Translation and summary creation were mentioned most in relation to current use and also ranked highest in areas most benefiting, followed by subject headings and classmarks. I am unsure what was meant by "physical description" and "creators/contributors" in the survey, but as data transcription was not mentioned yet, my assumption is that it is covered by these two categories and/or "other" ([Chen and Li, 2024](#), pp. 322-323).

These survey results tally with my own experience in AI use for cataloguing, and where I believe AI can be of most benefit to cataloguers.

While born digital resources and digitised resources might spring to mind as the most likely candidates for AI-assisted metadata work, Lowagie (2024) has shown that there are also possibilities for physical resources.

Beyond individual record creation, machine learning, I think, has huge potential to improve metadata management tasks, though I cannot see generative AI to be useful in this respect. Data management tasks need to be transparent and produce repeatable and consistent results, all of which generative AI outputs certainly are not.

OCLC, for example, are using machine learning (not generative AI) to deduplicate WorldCat records and have removed millions of duplicate records from WorldCat with this approach (OCLC, 2025). Another great example of machine learning assisting in bulk tasks is Cornish and Scott's article in this number of Catalogue and Index (Cornish and Scott, 2025).

Interestingly, the Chen and Li's survey also found that "most respondents didn't find AI tools had significant help of [sic] either quality or efficiency of their cataloging work" (Chen and Li, 2024, p. 324). I believe this is grounded in the problems with accuracy and reliability of AI outputs. If one needs to double-check everything, there is no significant time saving or efficiency in the use. In my experience, thoroughly checking and, if needed, amending a record takes the same or even more time than creating it in the first place. At LSE we investigated ExLibris' generative AI enriched Community Zone records (ExLibris, no date a; ExLibris, no date b; York and Hanegbi, 2024) some months ago and found the assigned subject headings often on the too broad side, while the quality of summaries varied significantly depending on the type of publication. We also trialled using generative AI for JEL code² assignment. In addition to not seeing a time saving due to the need to check the AI output, colleagues reported that, while they found it to be a useful assistive tool, they also felt their ability to assign codes without the use of AI and familiarity with the vocabulary as a whole declined. A sentiment that is also echoed in Chen and Li's survey results as "worry about over-reliance" (Chen and Li, 2024, p. 324)

What can machine learning and generative AI do?

Subject headings and classification are probably the areas that have been looked at the most for automation so far.

The National Library of Finland's well-known Annif³ tool (Suominen et al., 2023) has been around since 2019 (National Library of Finland, 2025) and is in use in various libraries in either fully automated workflows or human-supervised ones (Inkinen,

² The JEL (*Journal of Economic Literature*) classification is a commonly used classification scheme for scholarly works in economics. <https://www.aeaweb.org/econlit/jelCodes.php?view=jel>

³ Annif is an example of the application of machine learning and thus AI, but it does not, as far as I am aware and understand, make use of generative AI.

[Lehtinen and Suominen, 2025](#), p. 3). An Annif user survey, however, also reveals potential problems with implementing such a tool:

- The technical expertise needed. The survey shows that the most encountered problems with the tool are of a highly technical nature ([Inkinen, Lehtinen and Suominen, 2025](#), pp. 3-4), which suggests levels of technical expertise are needed for an implementation that most libraries will be unable to shoulder by themselves or at all.
- The resources needed. The implementing institutions are big ones (national libraries, university libraries) ([Inkinen, Lehtinen and Suominen, 2025](#), p. 1), which suggests resourcing that will be out of reach for most others.
- The tool not delivering the expected results/time savings ([Inkinen, Lehtinen and Suominen, 2025](#), p. 4).

The German National Library (DNB) is using Annif for its “Erschließungsmaschine” (EMa, “subject cataloguing machine”). Results published in 2021 regarding the assignment of subject headings (GND⁴ descriptors) showed a rather worrying 10% of assignments having been assessed by subject experts as “wrong” and 22% as “less useful” ([Uhlmann and Grote, 2021](#)). The DNB has since also started assigning DDC short numbers (a simplified classification system that the DNB developed ([Deutsche Nationalbibliothek, 2023](#))) automatically. The performance metrics for this indicate a very mixed picture as well, with some categories scoring pretty well, but others rather badly ([Poley et al., 2025](#), pp. 12-13). Poley et al. ([2025](#)) state that the volume of available training data is an important factor in the model’s performance, but also that it does not seem to be the only criterion. As other criteria are not mentioned, I assume the authors do not know either (black box).

Golub et al. ([2024](#)) have also used Annif to conduct research on automated Dewey Decimal Classification numbers, working with data in the Swedish union catalogue. They achieved a 66.82% accuracy rate on assigning three-digit DDC numbers by combining the results of four classification algorithms. During the research, they discovered that classifying fiction posed a problem and identifiable records for fiction were excluded from the set of records to be classified, which improved the accuracy rate to the figure mentioned above. However, fiction records remained in the training dataset and thus fiction headings are assigned. The accuracy rate might be improved by excluding fiction from the training dataset.

Chow, Kao and Li ([2024](#)) experimented with assigning Library of Congress Subject Headings using generative AI and found that generative AI only produced usable outputs in about half of their samples. They thus conclude that “while ChatGPT can access an internalized corpus of the LCSH and MARC 21 [sic] bibliographic records, the model struggles with validity, specificity, and exhaustivity in its generated subject headings” ([Chow, Kao and Li, 2024](#), p. 585) and that “in order to ensure accuracy and

⁴ The GND (Gemeinsame Normdatei, “integrated authority file”) is the standard German-language authority file and contains personal and corporate names as well as subject headings.

reliability of the cataloging process, the involvement of human catalogers remains an essential prerequisite" ([Chow, Kao and Li, 2024](#), p.586).

The Exploring Computational Description experiments by the Library of Congress (LoC) in cataloguing eBooks via generative AI also showed low quality scores getting nowhere near the goal set, but it showed some promise in extracting author names, titles and identifiers ([Weinryb-Grohsgal, Potter and Saccucci, 2024](#); [Saccucci and Potter, 2024 b](#), p. 6; [Library of Congress, no date](#)). As with the DNB activities, the LoC experiment also highlights the importance of the "quality and robustness of the training data" ([Library of Congress, no date](#)) for the success of AI-generated records. The overall conclusions of the LoC experiment, at the time of writing, remain:

- No current generative AI tool returns good enough results to run automatic cataloguing.
- "Human-in-the-loop" workflows are possible though, and should be explored.

The third phase of the experiment started in August 2024 ([Weinryb-Grohsgal, Potter and Saccucci, 2024](#)), but no results have been published yet.

The technical hurdle for the implementations and experiments above is rather high, but less tech-intensive solutions are possible as well, as Lowagie showed with the KBR's approach ([Lowagie, 2024](#)). He implemented an AI-driven solution to bottlenecks in metadata creation and retro-cataloguing backlogs by employing Power Apps to extract information from photographs of title pages, thus enabling cataloguers to concentrate on ensuring correct information rather than data entry. Lowagie also introduced creating custom application profiles with Power Apps and using them to validate records in the catalogue.

Another low technical hurdle experiment was undertaken by Taniguchi ([2024](#)), who used the illustrative sources of information in *Maxwell's Handbook for RDA* to generate records using ChatGPT. The conclusion here is also that the generative AI produced records with significant errors and "struggled with complex bibliographic patterns and nuanced cataloging rules", but could conceivably be used as an assistive tool for human cataloguers ([Taniguchi, 2024](#), p. 544).

All these examples show that AI technologies can be leveraged to assist in cataloguing and metadata maintenance, but, apart from Taniguchi ([2024](#)), they are also all very far from the "prompt in chat to ingestible record" scenario that started this piece. I fully agree with Moulaison-Sandy and Coble ([2024](#)) in their assessment that "the perception that [AI] is able to solve specialized problems in cataloging easily, with the click of a button, if only the right prompt is created, is problematic to perpetuate" but also that "now is the time to look to the future and to be creative, but with a sense of the full understanding of the limitations and affordances." ([Moulaison-Sandy and Coble, 2024](#), p. 382)

This sentiment is further echoed in a survey the Program for Cooperative Cataloging (PCC) ran in March 2024 to gauge current AI activities and what their impact on cataloguing and metadata work is. The first theme emerging from it is:

“The need to clearly communicate to library administrators and the broader cataloguing community that AI is not an easy fix or money saver. AI and ML [machine learning] technologies take time and careful consideration in order to be implemented effectively and must be done in concert with cataloging and metadata experts.” ([Program for Cooperative Cataloguing, 2024](#), p. 2)

The newest generation of generative AI models are no longer pure LLMs but Large Multimodal Models (LMMs) able to handle not just text in- and output, but other media such as images, audio and video as well ([Wu et al., 2023](#)). Large Reasoning Models (LRMs) are developed with better reasoning and fact-checking abilities to improve performance ([Mollick, 2025 a](#)). On the other hand, Apple just released a paper that reckons this is all just an “illusion of thinking” and “that frontier [large reasoning models] face a complete accuracy collapse beyond certain complexities” ([Shojaee et al., 2025](#), p.1). Furthermore, there is evidence that newer models hallucinate more, and the developers and researchers do not understand why ([OpenAI, 2025](#), p. 4; [Chowdhury et al., 2025](#); [Metz and Weise, 2025](#)).

What could be promising though, is a combination of generative AI, computer vision tools, good old database queries (not everything needs to be generated new, sometimes just finding a good, existing record and verifying it is all that’s needed) and incorporation of local documentation. The use of the latter two, I recently learned, actually has a name and framework: Retrieval Augmented Generation (RAG). RAG retrieves information from external sources (e.g. a knowledge base, database, etc.) and uses this to augment the LLM response. By doing this, the LLM can return more accurate and contextually relevant responses ([Google, no date](#)).

So, yes, there are opportunities to employ AI technologies (and generative AI can be part of this) to create or enhance metadata, but can it be done with the needed reliability and accuracy, or can entire records be created? Absolutely not – at least not for the time being.

With generative AI bullshitting⁵ and hallucinations seemingly not going anywhere, through them being an inherent part of it⁶, I cannot see generative AI-driven solutions for metadata creation and enhancement being able to operate without oversight; we need the human in the loop to check outputs.

Time and emerging technologies may well change this view.

⁵ Strictly in the Frankfurtian sense outlined by [Hicks, Humphries and Slater, 2024](#).

⁶ “Despite our best efforts, they will always hallucinate” Amr Awadallah, formerly of Google and now CEO of a startup building AI tools for businesses, told the New York Times ([Metz and Weise, 2025](#)).

Should we just because we can?

Now that we have covered the technical possibilities, let's have a look at the ethical side of things. Berkowitz (2025) argues that libraries tend to choose quick adoption of emergent technologies (AI use in this case) "for the sake of being perceived as cutting-edge early adopters" over "a deeply methodical intent", i.e. they tend to opt for FOMO over slow-mo (Berkowitz, 2025, p. 52). He calls for libraries to champion AI ethics and concentrate on ethical scrutiny and developing policies and ethical frameworks for AI use rather than quick adoption.

AI bias

Bias can enter generative AI outputs in various ways:

- It can be present in the training data, for example, because the data is not representative, omits or obscures information. There can also be problems with inconsistent training data labelling. However, even data that is otherwise sound will still reflect structural and historical biases.
- Secondly, the training and inference algorithms may display bias or amplify biases in the training data.
- Further biases can be introduced through the evaluation of a model and the used benchmark dataset(s).
- Finally, models may be used in scenarios they were not intended for and thus produce biased, harmful outputs (Gallegos et al., 2024, p. 1107).

For a much deeper dive into AI biases, their evaluation, and techniques for bias mitigation, I refer you to Resnik (2025) and Gallegos et al. (2024).

What are the consequences of an agent with harmful biases creating metadata? Well, metadata that furthers and perpetuates those, of course. A lot of work has been done in libraries to overcome harmful language in records as well as to bring materials by and about marginalised groups out of their space of marginalisation and othering. By using generative AI to help us in creating and managing metadata, are we negating at least some of this work? As Corrado (2021, p. 402) asks: "how will [AI] satisfy the ethical concerns related to representation and identity in metadata?" As humans we can of course use our judgement and awareness of our own biases to ensure we do not perpetuate inequities or "hide" content behind overly broad or othering headings or by omission. I agree with Corrado that "it is yet to be seen how artificial intelligence will deal with this fluid space. Unless librarians and other advocates push for this, the answer may very well be that it won't." (Corrado, 2021, pp. 402-403)

Given the importance of representation in training data the German National Library (Poley et al., 2025) and Library of Congress (Library of Congress, no date) have found in their respective subject indexing and metadata creation experiments as well as the struggles with specificity Chow, Kao and Li (2024, p. 585) found, I have doubts that AI-

generated subject headings can adequately represent material on subjects underrepresented in the training data or on emerging subjects. Poley et al. in fact recognise that this is the case, stating that “subject areas where automatic subject cataloguing does not work, or does not work well, must first be intellectually indexed in order to generate training data to improve possible machine models.” ([Poley et al., 2025](#), p. 25)

AI’s black box nature further obscures things, making it very difficult indeed to both identify and address biases in outputs. How can one intervene in a “thinking” process whose workings are not understood even by the people developing them ([OpenAI, 2025](#); [Metz and Weise, 2025](#))?

Copyright

Copyright legislation is woefully behind AI development, and many questions are unanswered regarding copyright ownership of AI-generated content as well as copyright infringements in training AI models. Most generative AI training data is scraped from publicly available internet resources, but it also includes material from platforms that contain copyrighted material (e.g. the recent outcry over the Library Genesis dataset). Alex Reisner’s article on the subject is a scary read indeed ([Reisner, 2025](#)). He details that both Meta and OpenAI argue that their use of copyrighted material for training without a license falls under “fair use”. I believe the courts have yet to make a judgment on this, but the US Copyright Office certainly begs to differ ([United States Copyright Office, 2025](#); [Constantino, 2025](#)).

Ball ([2025](#)) adds that “To add insult to injury, academic publishers are now beginning to license access to their content to AI companies, some without providing academics the opportunity to opt out. This forces complicity on academics, turning their intellectual contributions into commodities for AI profit without their consent and with no remuneration for them or their institutions.” See also Battersby ([2024](#)) and Eaton ([2024](#)) on this subject.

If libraries engage in AI use, I think they should think hard about these issues and whether it is ethical to condone such practices by using tools built on them.

Environment

“There is still much we don’t know about the environmental impact of AI but some of the data we do have is concerning”. This is the statement of the Chief Digital Officer of the United Nations Environment Programme in a news piece summarising report findings ([UN Environment Programme, 2024](#)). Hugging Face, which provides a platform for sharing machine learning models and datasets, also acknowledges that “the nature and extent of AI’s effects are under-documented, ranging from its

embodied and enabled emissions to rebound effects due to its increased usage” ([Luccioni, Trevelin and Mitchell, 2024](#))⁷.

Generative AI companies are not forthcoming with accurate and complete data on the environmental footprint of their products, and available figures rely on lab-based research, such as that carried out by Luccioni and Strubell⁸ as well as “limited company reports; and data released by local governments” ([Crawford, 2024](#)). Annual reports of big tech companies show that they are not meeting their sustainability targets ([Barker, 2025](#)).

The negative environmental impact of AI splits between its electricity and water consumption, resources needed to manufacture equipment and the regurgitation of it as electronic waste ([UN Environment Programme, 2024](#)). It includes not just the training of models, but also their usage.

The training and operation of AI demands vast quantities of computational power, and the data centres housing the servers that do the work need electricity and water for cooling. Loads of figures are floating about on the energy consumption of generative AI data centres:

- Mollick ([2025 b](#)) states that there are now AI models in use that consume as much computing power for training as it takes to “[run] a modern smartphone for 634,000 years or the Apollo Guidance Computer that took humans to the moon for 79 trillion years”.⁹
- Zewe ([2025](#)) states that the energy demand of data centres in North America is estimated to have increased “from 2,688 megawatts at the end of 2022 to 5,341 megawatts at the end of 2023”, an increase “partly driven by the demands of generative AI”. He also mentions that the global data centre electricity consumption reached 460 terawatts in 2022, which makes it the “11th largest electricity consumer in the world, between the nations of Saudi Arabia (371 terawatts) and France (463 terawatts)” and that it is expected to be closer to 1,050 terawatts by 2026.
- The UN Environment Programme ([2024](#)) gives the example of Ireland, which hosts many data centres, stating that the International Energy Agency estimates that “the rise of AI could see data centres account for nearly 35 per cent of the country’s energy use by 2026”.
- Taking the research by Strubell, Ganesh and McCallum ([2020](#)) as the basis, Luccioni, Trevelin and Mitchell ([2024](#)) state that “training [an LLM with] 213 million parameters was responsible for ... [the] equivalent to the lifetime emissions of five cars, including fuel”. For comparison, OpenAI’s most recent

⁷ Luccioni is Hugging Face’s Climate Lead, Mitchell its Chief Ethics Scientist and Trevelin its Legal Counsel.

⁸ The reference is to [Strubell, Ganesh and McCallum, 2020](#), [Luccioni, Viguier and Ligozat, 2023](#) and [Luccioni, Jernite and Strubell, 2024](#).

⁹ In the correct unit of measurement that’s 10²⁶ FLOPS (Floating point operations per second).

model, GPT-4o, allegedly¹⁰ uses 200 billion parameters ([Ben Abacha et al., 2025](#)). Zewe (2025) states that the electricity needed for training a model like GPT-3 was estimated to consume the equivalent of 120 average U.S. homes yearly energy consumption.

Training a model is just one part of it though: energy is consumed every time the model is used, and, with the rapid development of new models, training needs to be repeated for them frequently ([Zewe, 2025](#)).

In terms of model usage, Luccioni, Jernite and Strubell (2024) found clear differences between modalities, with image-based tasks and generation of new content using the most energy. A report prepared by Goldman Sachs found “that a ChatGPT search consumes around 6x-10x the power as a traditional Google search” ([Goldman Sachs, 2024](#), p. 13). This report also shows quite frightening predictions for power use by AI. As Luccioni, Trevelin and Mitchell state, “the growing energy demand for AI is significantly outpacing the increase in renewable energies – entailing substantial new [greenhouse gas] emissions and squeezing an already tight renewable energy market.” ([Luccioni, Trevelin and Mitchell, 2024](#))

Next on the list is water consumption: water is needed to cool the servers in the data centres, and it needs to be cooled so it can absorb the heat from the machines. Additionally, neither salt nor grey water can be used for this process as this damages the cooling systems. Approximately 30-40% of the electricity consumed by data centres is used for water cooling ([Luccioni, Trevelin and Mitchell, 2024](#)). The amount of water needed depends largely on the size of the data centre. The biggest, hyperscale data centres are reported to use 2.1 million litres of water a day, while smaller ones are reported to use 68,000 litres a day ([Zhang, 2024](#)).

Crucially, data centres are often located in areas with already limited water supply and exacerbate problems in those areas ([Barratt and Gambarini, 2025](#)).

Finally, there is the extraction of the raw materials needed to build servers and other data centre equipment as well as the waste they eventually become. The mining for the metals needed has its own environmental problems, and some are so-called “conflict minerals”, which means “that they are mined or traded in areas of conflict, and contribute towards perpetuating human rights abuses and armed conflict” ([Luccioni, Trevelin and Mitchell, 2024](#)).

Wang et al. (2024) predict that by 2030 generative AI could add up to 5 million metric tons of electronic waste to the global total. Given, that is a relatively small proportion of the global total, but, as experts warn, a significant one ([Crownhart, 2024](#)). Electronics often contain hazardous or toxic materials such as lead, mercury and chromium, and, if not disposed of responsibly, these can harm the environment. Another problem is the waste of valuable metals such as copper, gold and rare earth

¹⁰ It seems that parameter counts are not readily published information. The cited paper seems to be the source of the 200 billion figure floating about for GPT-4o.

elements, when electronic waste is not recycled ([Crownhart, 2024](#)). According to the 2024 Global E-Waste Monitor ([Baldé et al., 2024](#), p. 9), only 22.3% of electronic waste is formally collected and recycled in an environmentally sound manner.

I am fully aware than I am neglecting to mention the positive environmental impacts of AI, for example through helping investigating and addressing environmental problems ([UN Environment Programme, 2024](#)) as well as the steps that are being taken by authorities and cooperations to mitigate negative impacts ([Luccioni, Trevelin and Mitchell, 2024](#); [UN Environment Programme, 2024](#); [Barker, 2024](#); [Ren and Wierman, 2024](#)). Given the missing of sustainability goals by the big tech corporations ([Barker, 2025](#)) and in light of humanity's current track record in taking care of our planet, I find it not very believable that we can mitigate such a huge projected increase in resource consumption and emissions successfully. Hence, I believe it is important to showcase the massive negative environmental impact of AI clearly.

Labour exploitation and inequity

In addition to the issues regarding exploitation and inequity with mining mentioned by Luccioni, Trevelin and Mitchell ([2024](#)), there is also the problem of environmental impacts being very unfairly distributed across the planet and some regions and communities being disproportionately affected, for example, through air pollution from local fossil fuel consumption ([Ren and Wierman, 2024](#)). A particular example raised by Ren and Wierman ([2024](#)) is Google's data centre in Finland operating on 97% carbon-free energy as opposed to its ones in Asia, which only use 4%-18% carbon-free energy.

Ball ([2025](#)) also raises issues around exploitation and inequities:

"The extraction of vast amounts of data without informed consent, perpetuates a system of surveillance and control that undermines democratic principles and disproportionately affects vulnerable populations. The reliance on low-paid workers in the Global South to perform data labelling and content moderation tasks further exacerbates global inequalities, exposing these individuals to exploitative practices and precarious working conditions."

Barriers

Apart from the ethical considerations, there are also other barriers to AI use for metadata work.

AI literacy

The previously mentioned survey by Chen and Li revealed a lack of adequate training and support in relation to AI use ([Chen and Li, 2024](#), pp. 321-322), which may go some way in explaining respondents' concerns about "misunderstandings about the capabilities and limitations of AI in cataloging, which may lead to unrealistic

expectations or disappointment with the results” ([Chen and Li, 2024](#), p. 324) as well as “reservations about rushing AI integration without considering potential consequences” ([Chen and Li, 2024](#), p. 326)

This clearly shows that more needs to be done to improve AI literacy and understanding, not only for metadata professionals, but also for managers and decision makers.

Resourcing

PCC’s report on strategic planning for AI and machine learning highlights the issue of library resourcing being prohibitive for investigating, experimenting or even implementing potential AI-driven workflows:

“General concern about a lack of resources in order to investigate and implement AI. Many institutions are involved with system migrations, training for Official RDA and/or linked data, or are generally under-resourced or too small to realistically spend time working with AI.” ([Program for Cooperative Cataloguing, 2024](#), p. 2)

Few libraries around the world have the resourcing, technical expertise and equipment at their disposal to spend time on experimenting with a technology that, in order to deliver usable results, needs a deep understanding of machine learning techniques and algorithms as well as the ability to set up tools, fine-tune them to their respective needs and maintain them.

Planning, building, testing and implementing a machine learning solution such as the ones outlined earlier takes a long time, years even.

It is probably also worth saying that what works for library A does not automatically also work for library B. When it comes to metadata, we all have our local practices and idiosyncrasies to account for. Models trained on someone else’s data might not do very well with one’s own.

It would be very nice to see solutions come out of the community rather than AI-assisted cataloguing becoming yet another area where libraries need to rely on vendor-provided solutions ([Moulaison-Sandy and Coble, 2024](#), p. 381).

The Black Box

Generative AI’s “black box” nature is also a concern. Something is considered a black box when input and outputs can be seen, but how the inputs are turned into the outputs, i.e. the internal workings, are mysterious and cannot be seen ([Kosinski, no date; Bagchi, 2023](#)). Additionally, even when algorithms are known, for deep learning (which generative AI is based on), the learning process itself creates connections and patterns that mean even the creators of these processes cannot understand how they

actually work ([Kosinski, no date](#); [Metz and Weise, 2025](#); [OpenAI, 2025](#)). This means that even open-source models using deep learning are essentially black boxes.

This is problematic as it is hard to trust an output if it is not transparent how it was arrived at, and impossible to validate its path through the model. Even if the output is correct, maybe the model arrived at it for the wrong reasons (the “Clever Hans effect”). Due to the lack of understanding of the internal workings, adjusting a model that makes wrong decisions or produces bad outputs is very difficult ([Kosinski, no date](#); [Blouin, 2023](#)).

A black box model can hide security vulnerabilities. If one doesn’t know how something works, one cannot tell if it has been modified in malicious ways ([Kosinski, no date](#); [Bagchi, 2023](#)).

Black box models might also exacerbate algorithmic bias and lead to bad, maybe even outright harmful and illegal outcomes. While biases will be present in outputs if they are present in the training data, assessing if a bias exists and finding what its cause is, is especially hard in black box models ([Kosinski, no date](#); [Blouin, 2023](#)).

Lastly, Kosinski ([no date](#)) mentions trouble assessing whether one is compliant with regulations regarding the use of sensitive data in AI tools, such as for example the Artificial Intelligence Act of the European Union¹¹.

Researchers are working on improving insights into model workings, but sufficient transparency does not seem to be on the horizon ([Kosinski, no date](#))

Conclusion

All of the above may read like I oppose AI use in metadata work, but this is not the case. I am not a technophobe, and I truly believe that there are opportunities to improve record quality and to assist cataloguers and metadata managers in their work. Maybe my attitude can be best described as that of a “curator” as defined by Rosser and Hanegan ([2024](#)). In my exploration of the subject over the last couple of years I felt that, while the limitations of generative AI are mentioned and the conclusions generally align with my own here, not enough space was given to critical exploration of the technical possibilities and ethical dilemmas associated with generative AI use and this article is merely an attempt to more explicitly state the limitations and issues.

In summary, regarding the question of generative AI being able to catalogue: no, it absolutely cannot, and I believe it will not. Can machine learning catalogue? Well, maybe, but not yet. Can machine learning assist cataloguers in their work? Yes, absolutely, but the human in the loop remains a non-negotiable necessity!

¹¹ See <https://artificialintelligenceact.eu/>

I cited Moulaison-Sandy and Coble (2024) before in this article and do so again here, as they brilliantly sum up the matter:

“[N]ow is the time to look to the future and to be creative, but with a sense of the full understanding of the limitations and affordances. Yes, finding new ways in which AI can support the work of librarians, especially technical services librarians like catalogers, will be critical to future success” (Moulaison-Sandy and Coble, 2024, p. 383)

On the ethical side and the question of whether we should implement AI, I think Berkowitz (2025) has a point: we need more ethical scrutiny, policies and frameworks. For small-scale experiments and implementations, this might not be as crucial, but certainly for the adoption of AI tools, e.g. via vendor products, that rely on mainstream tools such as ChatGPT, Copilot, etc., we need to think properly about all implications and whether they outweigh the usefulness of the tool.

References

- Amram, T., Malamud, R. G., and Hollingsworth, C. (2023) Response to "From ChatGPT to CatGPT". *Information Technology and Libraries*, 42(4). Available at: <https://doi.org/10.5860/ital.v42i4.16983>
- Bagchi, S. (2023) What is a black box? A computer scientist explains what it means when the inner workings of AIs are hidden. *The Conversation*, 22 May. Available at: <https://theconversation.com/what-is-a-black-box-a-computer-scientist-explains-what-it-means-when-the-inner-workings-of-ais-are-hidden-203888> [Accessed: 24 May 2025]
- Baldé, C. P., Kuehr, R., Yamamoto, T., McDonald, R., D'Angelo, E., Althaf, S., Bel, G., Deubzer, O., Fernandez-Cubillo, E., Forti, V., Gray, V., Heart, S., Honda, S., Iattoni, G., Khetriwal, D. S., Luda di Cortemiglia, V., Lobuntsova, Y., Nnorom, I., Pralat, N. and Wagner, M. (2024) *Global E-waste Monitor 2024*. Geneva/Bonn: International Telecommunication Union and United Nations Institute for Training and Research. Available at: https://ewastemonitor.info/wp-content/uploads/2024/12/GEM_2024_EN_11_NOV-web.pdf [Accessed: 28 May 2025]
- Ball, Caroline (2025) The Unethical Underbelly of AI: A Call for Universities to Take a Stand. *UKSG News*, 586. Available at: <https://www.uksg.org/newsletter/uksg-enews-586/enews-586-editorial/> [Accessed 21 May 2025]
- Barker, C. (2024) Artificial intelligence and the environment: Taking a responsible approach. *JISC Artificial Intelligence*, 18 September. Available at: <https://nationalcentreforai.jiscinvolve.org/wp/2024/09/18/artificial-intelligence-and-the-environment-taking-a-responsible-approach/> [Accessed: 24 May 2025]
- Barker, C. (2025) Artificial intelligence and the environment: The current landscape. *JISC Artificial Intelligence*, 28 March. Available at: <https://nationalcentreforai.jiscinvolve.org/wp/2025/03/28/artificial-intelligence-and-the-environment-the-current-landscape/> [Accessed: 24 May 2025]
- Barratt, L. and Gambarini, C. (2025) Revealed: Big tech's new datacentres will take water from the world's driest areas. *The Guardian* (Online), 9 April. Available at: <https://www.theguardian.com/environment/2025/apr/09/big-tech-datacentres-water> [Accessed: 9 April 2025]

- Battersby, M. (2024) Academic authors 'shocked' after Taylor & Francis sells access to their research to Microsoft AI. *The Bookseller*, 19 July. Available at: <https://www.thebookseller.com/news/academic-authors-shocked-after-taylor--francis-sells-access-to-their-research-to-microsoft-ai> [Accessed: 24 May 2025]
- Ben Abacha, A., Yim, W.-W., Fu, Y., Sun, Z., Yetisgen, M., Xia, F. and Lin, T. (2025) *MEDEC: A Benchmark for Medical Error Detection and Correction in Clinical Notes*. Preprint. Available at: <https://doi.org/10.48550/arXiv.2412.19260>
- Berkowitz, A. E. (2025) "Slow-MO or FOMO": AI conversations at library conferences. *Public Services Quarterly*, 21(1), pp. 51-70. Available at: <https://doi.org/10.1080/15228959.2024.2442657>
- Blouin, L. (2023) AI's mysterious 'black box' problem, explained. *University of Michigan-Dearborn News*, 6 March. Available at: <https://umdearborn.edu/news/ais-mysterious-black-box-problem-explained> [Accessed: 24 May 2025]
- Brzustowicz, R. (2023) From ChatGPT to CatGPT: The Implications of Artificial Intelligence on Library Cataloging. *Information Technology and Libraries*, 42(3). Available at: <https://doi.org/10.5860/ital.v42i3.16295>
- Chen, S. and Li, M. (2024) AI for Cataloging and Metadata Creation: Perspectives and Future Opportunities from Cataloging and Metadata Professionals. *Technical Services Quarterly*, 41(4), pp. 317-332. Available at: <https://doi.org/10.1080/07317131.2024.2394919>
- Chow, E. H. C., Kao, T. J. and Li, X. (2024) An Experiment with the Use of ChatGPT for LCSH Subject Assignment on Electronic Theses and Dissertations. *Cataloging & Classification Quarterly*, 62(5), pp. 574-588. Available at: <https://doi.org/10.1080/01639374.2024.2394516>
- Chowdhury, N., Johnson, D., Huang, V., Steinhardt, J. and Schwettmann, S. (2025) Investigating truthfulness in a pre-release o3 model. Available at: <https://transluce.org/investigating-o3-truthfulness> [Accessed 24 May 2025]
- Constantino, Tor (2025) U.S. Copyright Office Shocks Big Tech With AI Fair Use Rebuke. *Forbes*, 29 May. Available at: <https://www.forbes.com/sites/torconstantino/2025/05/29/us-copyright-office-shocks-big-tech-with-ai-fair-use-rebuke/> [Accessed 31 May 2025]
- Cornish, B. and Scott, B. (2025) Identify, Obtain, Explore: using NLP to link article and journal records in the NHM library catalogue. *Catalogue & Index*, 211, pp. 20-28. Available at: <https://journals.cilip.org.uk/catalogue-and-index/article/view/749>
- Corrado, E. M. (2021) Artificial Intelligence: The Possibilities for Metadata Creation. *Technical Services Quarterly*, 38(4), pp. 395-405. Available at: <https://doi.org/10.1080/07317131.2021.1973797>
- Crawford, K. (2024) Generative AI's environmental costs are soaring - and mostly secret. *Nature*, 626, p. 693. Available at: <https://doi.org/10.1038/d41586-024-00478-x>
- Crownhart, C. (2024) AI will add to the e-waste problem. Here's what we can do about it. *MIT Technology Review*, 28 October. Available at: <https://www.technologyreview.com/2024/10/28/1106316/ai-e-waste/> [Accessed: 24 May 2025]
- Deutsche Nationalbibliothek (2023) *Launch of cataloguing machine EMa*. Available at: https://jahresbericht.dnb.de/Webs/jahresbericht/EN/2022/Hoehepunkte/Erschliessungsmaschine/erschliessungsmaschine_node.html [Accessed: 24 May 2025]
- DeZelar-Tiedman, C. (2023) Response to "From ChatGPT to CatGPT". *Information Technology and Libraries*, 42(4). Available at: <https://doi.org/10.5860/ital.v42i4.16991>

- Eaton, L. (2024) Research Insights #12: Copyrights and Academia: Scholarly authors are not going to be happy... *AI+Edu=Simplified*, 23 July. Available at: <https://aiedusimplified.substack.com/p/research-insights-12-copyrights-and> [Accessed: 24 May 2025]
- ExLibris (no date a) *AI Bibliographic Records Enrichment*. Available at: https://knowledge.exlibrisgroup.com/Content/Knowledge_Articles/Alma/Knowledge_Articles/AI_Bibliographic_Records_Enrichment [Accessed: 31 May 2025]
- ExLibris (no date b) *AI Metadata Enrichment for Libraries*. Available at: https://knowledge.exlibrisgroup.com/Alma/Product_Materials/010Roadmap/AI_Metadata_Enrichment_for_Libraries [Accessed: 31 May 2025]
- Floyd, D. (2023) Response to "From ChatGPT to CatGPT". *Information Technology and Libraries*, 42(4). Available at: <https://doi.org/10.5860/ital.v42i4.16995>
- Gallegos, I. O., Rossi, R. A., Barrow, J., Tanjim, M. M., Kim, S., Deroncourt, F., Yu, T., Zhang, R. and Ahmed, N. K. (2024) Bias and Fairness in Large Language Models: A Survey. *Computational Linguistics*, 50(3), pp. 1097-1179. Available at: https://doi.org/10.1162/coli_a_00524
- Goldman Sachs (2024) *AI, data centers and the coming US power demand surge*. Available at: <https://www.goldmansachs.com/pdfs/insights/pages/generational-growth-ai-data-centers-and-the-coming-us-power-surge/report.pdf> [Accessed: 24 May 2025]
- Golub, K., Suominen, O., Mohammed, A.T., Aagaard, H. and Osterman, O. (2024) Automated Dewey Decimal Classification of Swedish library metadata using Annif software. *Journal of Documentation*, 80(5), pp. 1057-1079. Available at: <https://doi.org/10.1108/JD-01-2022-0026>
- Google (no date) *What is Retrieval-Augmented Generation (RAG)?* Available at: <https://cloud.google.com/use-cases/retrieval-augmented-generation> [Accessed: 03 June 2025]
- Hicks, M. T., Humphries, J. and Slater, J. (2024) ChatGPT is bullshit. *Ethics and Information Technology*, 26(2). Available at: <https://doi.org/10.1007/s10676-024-09775-5>
- Inkinen, J., Lehtinen, M. and Suominen, O. (2025) *Annif Users Survey: Understanding Usage and Challenges*. National Library of Finland. Available at: <https://urn.fi/URN:ISBN:978-952-84-1301-1> [Accessed: 24 May 2025]
- Kosinski, M. (no date) *What is black box artificial intelligence (AI)?* Available at: <https://www.ibm.com/think/topics/black-box-ai> [Accessed: 31 May 2025]
- Library of Congress (no date) *Exploring Computational Description*. Available at: <https://labs.loc.gov/work/experiments/ECD> [Accessed: 24 May 2025]
- Lowagie, H. (2024) Harnessing Power Apps and AI for Automated Cataloguing: Innovations in Bibliographic Record Creation. *Catalogue & Index*, 209. Available at: <https://journals.cilip.org.uk/catalogue-and-index/article/view/697> [Accessed: 01 May 2025]
- Luccioni, A. S., Jernite, Y. and Strubell, E. (2024) Power Hungry Processing: Watts driving the cost of AI deployment?, *FACt '24: The 2024 ACM Conference on Fairness, Accountability, and Transparency*, Rio de Janeiro, 3-6 June. Available at: <https://doi.org/10.1145/3630106.3658542>
- Luccioni, A. S., Viguier, S. and Ligozat, A. L. (2023) Estimating the carbon footprint of BLOOM, a 176B parameter language model. *Journal of Machine Learning Research*, 24(253), pp. 1-15. Available at: <https://jmlr.org/papers/v24/23-0069.html> [Accessed: 31 May 2025]
- Luccioni, S., Trevelin, B. and Mitchell, M. (2024) The Environmental Impacts of AI - Policy Primer. *Hugging Face Blog*, 3 September. Available at: <https://doi.org/10.57967/hf/3004>

- Metz, C. and Weise, K. (2025) AI is getting more powerful, but its hallucinations are getting worse. *New York Times* (Online), 5 May. Available at: <https://www.nytimes.com/2025/05/05/technology/ai-hallucinations-chatgpt-google.html> [Accessed: 31 May 2025]
- Mollick, E. (2025 a) The End of Search, The Beginning of Research: The first narrow agents are here. *One Useful Thing*, 3 February. Available at: <https://www.oneusefulthing.org/p/the-end-of-search-the-beginning-of> [Accessed 6 April 2025]
- Mollick, E. (2025 b) A new generation of AIs: Claude 3.7 and Grok 3. *One Useful Thing*, 24 February. Available at: <https://www.oneusefulthing.org/p/a-new-generation-of-ais-claude-37> [Accessed: 6 April 2025]
- Moulaison-Sandy, H. and Coble, Z. (2024) Leveraging AI in Cataloging: What Works, and Why? *Technical Services Quarterly*, 41(4), pp. 375-383. Available at: <https://doi.org/10.1080/07317131.2024.2394912>
- Murgia, M., Clark, D., Learner, S., de la Torre Arenas, I., Joiner, S., Hemingway, E. and Hawkins, O. (2023) Generative AI exists because of the transformer. *Financial Times*, 12 September. Available at: <https://ig.ft.com/generative-ai/> [Accessed: 24 May 2025]
- National Library of Finland (2025) *Annif*. Available at: <https://annif.org/> [Accessed: 24 May 2025]
- OCLC (2025) *Implementing AI to further scale and accelerate WorldCat de-duplication*. Available at: <https://www.oclc.org/en/news/announcements/2025/ai-worldcat-deduplication.html> [Accessed: 24 May 2025]
- OpenAI (2025) *OpenAI o3 and o4-mini System Card*. Available at: <https://cdn.openai.com/pdf/2221c875-02dc-4789-800b-e7758f3722c1/o3-and-o4-mini-system-card.pdf> [Accessed: 1 June 2025]
- Poley, C., Uhlmann, S., Busse, F., Jacobs, J.-H., Kähler, M., Nagelschmidt, M. and Schumacher, M. (2025) Automatic Subject Cataloguing at the German National Library. *Liber Quarterly*, 35. Available at: <https://doi.org/10.5337/lq.19422>
- Program for Cooperative Cataloguing (2024) *PCC Task Group on Strategic Planning for AI and Machine Learning: Final Report Transmittal & Tracking Sheet*. Available at: <https://www.loc.gov/aba/pcc/taskgroup/TG-Strategic-Planning-AI-final-report.pdf> [Accessed: 21 May 2025]
- Resnik, P. (2025) Large Language Models Are Biased Because They Are Large Language Models. To be published in *Computational Linguistics* [Peer-reviewed accepted version]. Available at: https://doi.org/10.1162/coli_a_00558
- Reisner, A. (2025) The Unbelievable Scale of AI's Pirated-Books Problem. *The Atlantic* (Online), 20 March. Available at: <https://www.theatlantic.com/technology/archive/2025/03/libgen-meta-openai/682093/> [Accessed: 28 May 2025]
- Ren, S. and Wierman, A. (2024) The Uneven Distribution of AI's Environmental Impacts. *Harvard Business Review*, 15 July. Available at: <https://hbr.org/2024/07/the-uneven-distribution-of-ais-environmental-impacts> [Accessed: 24 May 2025]
- Rosser, C. and Hanegan, M. (2024) Cyborgs and Centaurs, Prophets and Priests: Anywhere Left for Curators? *Atla*, 17 April. Available at: <https://www.atla.com/blog/cyborgs-and-centaurs-prophets-and-priests-anywhere-left-for-curators/> [Accessed: 28 May 2025]
- Saccucci, C. and Potter, A. (2024 a) Exploring Computational Description: LC Labs Planning Framework in Action, *OCLC RLP Webinar*, 12 March. Available at: <https://www.oclc.org/research/events/2024/ai-planning-framework-in-action.html> [Accessed: 15 May 2025]

- Saccucci, C. and Potter, A. (2024 b) Exploring Machine Learning: A Cataloging Experiment at the Library of Congress, *PCC Operations Committee Meeting*, 2 May. Available at: <https://www.loc.gov/aba/pcc/documents/OpCo-2024/Potter-Saccucci-Machine-Learning.pdf> [Accessed: 24 May 2025]
- Shojaee, P., Mirzadeh, I., Alizadeh, K., Horton, M., Benigo, S. and Farajtabar, M. (2025) *The Illusion of Thinking: Understanding the Strengths and Limitations of Reasoning Models via the Lens of Problem Complexity*. Available at: <https://ml-site.cdn-apple.com/papers/the-illusion-of-thinking.pdf> [Accessed: 10 June 2025]
- Strubell, E., Ganesh, A., and McCallum, A. (2020) Energy and policy considerations for modern deep learning research. *Proceedings of the AAAI conference on artificial intelligence*, 34(9), pp. 13693-13696. Available at: <https://doi.org/10.1609/aaai.v34i09.7123>
- Suominen, O., Inkinen, J., Virolainen, T., Fürneisen, M., Kinoshita, B. P., Veldhoen, S., Sjöberg, M., Zumstein, P., Neatherway, R. and Lehtinen, M. (2023) *Annif* (v1.0.0). Zenodo. <https://doi.org/10.5281/zenodo.8262313>
- Taniguchi, S. (2024) Creating and Evaluating MARC 21 Bibliographic Records Using ChatGPT. *Cataloging & Classification Quarterly*, 62(5), pp. 527-546. Available at: <https://doi.org/10.1080/01639374.2024.2394513>
- Uhlmann, S. and Grote, C. (2021) Automatic subject indexing with Annif at the German National Library (DNB). *Semantic Web in Libraries*, Online, 29 November – 3 December. Available at: <https://swib.org/swib21/slides/03-02-uhlmann.pdf> [Accessed: 24 May 2025]
- UN Environment Programme (2024) *AI has an environmental problem. Here's what the world can do about that*. Available at: <https://www.unep.org/news-and-stories/story/ai-has-environmental-problem-heres-what-world-can-do-about> [Accessed: 24 May 2025]
- United States Copyright Office (2025) *Copyright and artificial intelligence Part 3: Generative AI Training*. Pre-publication version. Washington: United States Copyright Office. Available at: <https://www.copyright.gov/ai/Copyright-and-Artificial-Intelligence-Part-3-Generative-AI-Training-Report-Pre-Publication-Version.pdf> [Accessed: 1 June 2025]
- Wang, P., Zhang, L.-Y., Tzachor, A. and Chen, W.-Q. (2024) E-waste challenges of generative artificial Intelligence. *Nature Computational Science*, 4, pp. 818-823. Available at: <https://doi.org/10.1038/s43588-024-00712-6>
- Weinryb-Grohsgal, L., Potter, A. and Saccucci, C. (2024) Could Artificial Intelligence Help Catalog Thousands of Digital Library Books? An Interview with Abigail Potter and Caroline Saccucci. *The Signal*, 19 November. Available at: <https://blogs.loc.gov/thesignal/2024/11/could-artificial-intelligence-help-catalog-thousands-of-digital-library-books-an-interview-with-abigail-potter-and-caroline-saccucci/> [Accessed: 24 May 2025]
- Wu, J., Gan, W., Chen, Z., Wan, S. and Yu, P. S. (2023) Multimodal Large Language Models: A Survey. *2023 IEEE International Conference on Big Data*, Sorrento, 15 – 18 December. Available at: <https://doi.org/10.1109/BigData59044.2023.10386743>
- York, E. and Hanegbi, D. (2024) *Metadata Enrichment using AI First Glance at Research and Findings*. Available at: https://knowledge.exlibrisgroup.com/@api/deki/files/166843/AI_enrichment_-_March_27.pdf?revision=1 [Accessed: 24 May 2025]
- Zewe, A. (2025) Explained: Generative AI's environmental impact. *MIT News*, 17 January. Available at: <https://news.mit.edu/2025/explained-generative-ai-environmental-impact-0117> [Accessed: 24 May 2025]

Zhang, M. (2024) *Data Center Water Usage: A Comprehensive Guide*. Available at: <https://dgtlinfra.com/data-center-water-usage/> [Accessed 24 May 2025]