

Can AI or Python help us catalogue books we can't read?

Adam Swann  0000-0003-1041-1308

Metadata Migration Assistant Librarian, Digital Library Team, University of Glasgow Library

Sally Bell  0000-0001-8647-305X

Head of Collection Development, Collection Development Team, University of Glasgow Library

Received: 14 August 2026 | Published: 16 September 2026

ABSTRACT

The potential for using tools like AI and Python to increase the efficiency of library cataloguing is being increasingly recognised. The University of Glasgow Library has a substantial backlog of uncatalogued non-English material, but few staff with the specialist language skills required to tackle it. Recent advances in computerised image analysis and transliteration have made it possible to begin cataloguing this backlog without requiring specialist language skills.

This article discusses a project which started by investigating the viability of using AI tools to help us copy catalogue Russian books. The project found that ChatGPT and Copilot can carry out OCR and transliteration of Russian text quickly and accurately. AI tools, however, come with serious environmental, legal, and ethical drawbacks.

We therefore started developing our own tool written in Python which has similar OCR and transliteration functionality to AI tools but without the associated environmental or ethical costs. The initial results from our Python tool are promising but more development is required.

Our project concluded that both AI and Python can be successfully used to help cataloguers without specialist language skills transliterate non-English texts. To maintain metadata quality, however, cataloguers with language skills must remain an integral part of our workflows.

KEYWORDS transliteration; non-English cataloguing; AI; Python, OCR

CONTACT Adam Swann  adam.swann@glasgow.ac.uk  University of Glasgow Library

The Centre for Russian, Central and East European Studies is a long-established and internationally renowned research partnership based at the University of Glasgow. To support the work of this unit the University of Glasgow Library (UoGL) has a substantial volume of stock in Russian, but no dedicated cataloguer for this material. Consequently, a backlog of uncatalogued Russian material has built up which is divided into two equal parts.

The first is approximately 1000 Russian books which require retrospective cataloguing, as they are on the shelves and recorded in a card catalogue but not yet added to our electronic catalogue. The second part is a donation of approximately 1000 Russian books which have not been appraised or catalogued. Some records from the card catalogue have been added to our electronic catalogue, but the Russian-language cataloguer working on this has subsequently moved to another team in the library.

This article discusses a pilot project which explored the viability of using AI and Python to help us address our Russian cataloguing backlog. It begins by situating the project in the broader context of AI use in the sector, then outlines the workflow and discusses the results of our project and the viability of using AI to help us catalogue non-English books. The project's initial scope was limited to using AI and so was originally intended to conclude there. However, the project was conducted against the backdrop of a broader debate about the tensions between what AI promises and what it delivers, and at what cost.

By the end of the AI phase of the project we, like many librarians, had become increasingly uncomfortable about AI use. We therefore started a second phase of the project which recreated the same workflow but using the Python programming language instead of AI. This article documents not just our project, but our evolving attitudes to AI use which are a microcosm of the shifting debates in our sector.

Recent discussions of AI use in academic libraries describe it as not only "rapidly changing the library landscape" ([Togia et al., 2026](#), p. 1), but even "the next frontier in library cataloguing" ([Ogungbenro et al., 2025](#), p. 105). It has particular potential for addressing cataloguing backlogs, since it offers "a level of productivity and efficiency in metadata creation that is unthinkable for humans" ([Dover and Grzegorski, 2025](#), p. 605).

However, this must be balanced against the reality that for all its apparent sophistication, AI is currently a somewhat blunt instrument for metadata work. It has phenomenal data-processing capability but lacks the nuanced understanding of context necessary for accurate cataloguing ([Oyighan et al., 2024](#), p. 24); simply put, AI alone "cannot consistently create quality metadata" ([Sussmeier and Henry, 2025](#), p. 245).

We find a similar pattern in recent experiments which used AI to catalogue non-English books. When asked to OCR images of book title pages to create RDA MARC records, ChatGPT did so effectively for major languages like Arabic, French, and Spanish, but struggled with minority languages like Tamil and Sinhala ([Gamage and Wanigasooriya, 2024](#), p. 241, 227). Similarly, it successfully converted Chinese book titles to Latin characters but was less accurate with authors "due to the intricacies of Chinese names" ([Jiang et al., 2024](#), p. 48).

Our pilot project's initial aim was to investigate whether AI tools can help staff without specialist language skills to copy catalogue non-English books. The scope of the pilot was limited to Russian Cyrillic texts, but there is potential to include both Latin non-English languages and other non-Latin scripts in the future. We aimed to be able to use AI tools to identify the correct catalogue record to use for a non-English book with at least 80% accuracy.

When cataloguing books written in non-Latin scripts, it is standard practice to transliterate them into Latin script, e.g. 'Государственные финансы царской России в эпоху империализма' becomes 'Gosudarstvennye finansy tsarskoi Rossii v epokhu imperializma'. This allows readers of the secondary script "to comprehend the orthography and the approximate pronunciation of the original text" ([Tour et al., 2025](#), p. 1) and enhances discoverability by providing both original script and Latinised catalogue entry points. When transliterating non-Latin scripts like Russian Cyrillic into Latin characters, we use the Library of Congress Romanisation standard¹.

We began by selecting 50 random Russian books from the card catalogue and photographing their title pages, ready for the AI tool to transliterate this Russian Cyrillic text into Library of Congress-compliant Latin characters. We would then use the transliterated author, title, and publication details to identify the correct record on Worldcat to download to copy catalogue the book.

We decided to use ChatGPT for the project because it supports both OCR from images and LC-compliant transliteration. Our initial workflow started by uploading a title page image to ChatGPT and asking for an LC-compliant transliteration of its bibliographic details. We would then format and collate these into a spreadsheet and search for the appropriate record on Worldcat.

However, the repeated formatting and collation required soon became laborious. Two important and intertwined principles for "obtaining precise and useful responses from AI models" are specifying the desired output format and reviewing the output to refine the input prompt ([Geroimenko, 2025](#), p. 24, 31). We therefore provided ChatGPT with a bibliographic details template to use in its responses, and after some iteration and refinement we were satisfied with the format of the output.

We then streamlined the start of the process by asking ChatGPT to respond to an input image of a Cyrillic title page by immediately supplying the transliterated version in our preferred format, with an English translation for reference. This simplified workflow greatly increased the speed of the transliteration process.

After we had collated all 50 transliterations, a colleague fluent in Russian reviewed them alongside the original images and confirmed their accuracy. We then felt confident proceeding with the project and using the transliterations to search for records on Worldcat.

¹ See the Library of Congress's Russian Romanisation Table ([Library of Congress, 2012a](#)) and its source document ([Library of Congress, 2012b](#)).

We searched for all 50 books on Worldcat using the transliterated bibliographic details and the results were excellent:

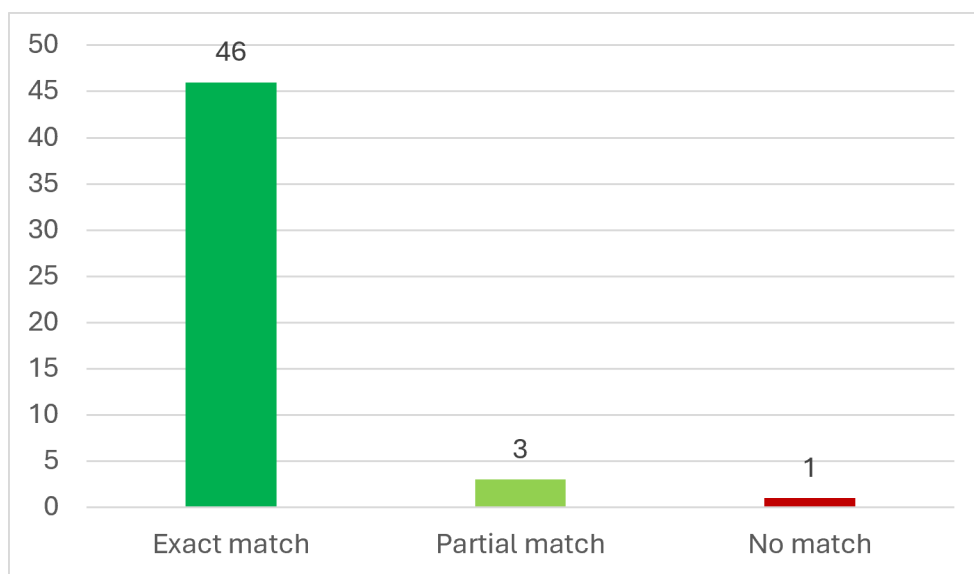


Figure 1: Retrospective cataloguing results

We were able to identify 46 exact matches on Worldcat (a 92% success rate), and only one book was not identified because it had no Worldcat record. The remaining three records were partial matches. They weren't exact matches as ChatGPT had transliterated them due to minor transliteration differences or errors:

Our ChatGPT Transliteration	Transliteration on Worldcat Record
ekonomike	ekonomiko
razvitogo	razvitoi
- [en-dash]	- [hyphen]

While these are minor differences a human cataloguer can still identify the correct catalogue record, so with some tweaking of search queries the final success rate was 98%.

We then applied the same workflow to 50 random books from the uncatalogued donation which had been identified by an academic as rare or unusual. As before, we uploaded photographs of their title pages into ChatGPT then used the transliterated details to search for records on Worldcat.

The results were not as successful as with the retrospective cataloguing:

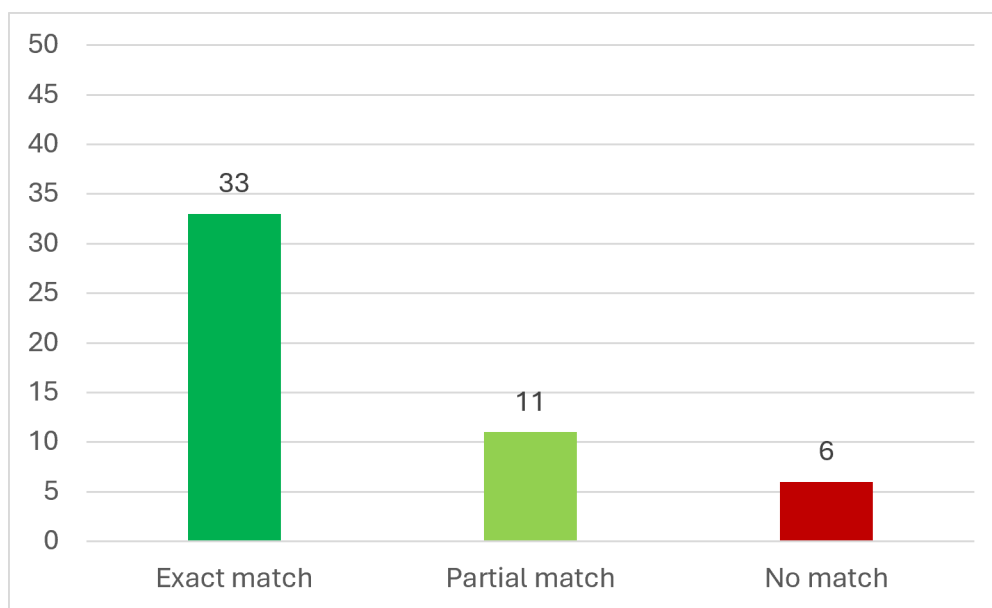


Figure 2: Donation cataloguing results

We had 33 exact matches (a 66% success rate), and 6 books had no Worldcat records. The remaining 11 books were partial matches; 5 titles had been catalogued on Worldcat using contractions (e.g. 'univers.' instead of 'universitov'), 3 titles had minor differences or errors in transliteration, and 3 titles were different editions to our own copies. The lower match rates are possibly attributable to these being rarer books, meaning there will be fewer holdings elsewhere for us to match against.

As before, with human input we could identify the correct catalogue record to use, so the final success rate for the books from the uncatalogued donation was 88%. Between retrospective cataloguing and donation cataloguing, therefore, we had an average final success rate of 93%.

A key part of the donation appraisal process is scarcity checking items against local and national holdings to assess rarity and research value. Using our ChatGPT transliterations, we scarcity checked this donation against both our own holdings and other Scottish and UK libraries on Library Hub just as easily as if the books were in English:

Author	Title	Publication details	GUL copies	Other copies	Scottish copies	UK copies
Alekseev, B. K. & M. N. Perfil'ev	Printsipy i tendentsii razvitiia predstavitel'nogo sostava mestnykh	Leningrad: Lenizdat, 1976	0	0	0	3
Vinogradov, A. N.	Verkhovnyi Sovet RSFSR	Moskva: Izdatel'stvo Iuridicheskoi literatury, 1967	0	0	0	4
Kurashvili, B. P.	Ocherk teorii gosudarstvennogo upravleniia	Moskva: Nauka, 1987	1	0	0	4
Kryshchanskaya, O. V.	Inzheneri : stanovlenie i razvitie professional'noi gruppy	Moskva: Nauka, 1989	1	0	0	4
	Pervichnaia partiinaia organizatsiia : opyt, formy i metody raboty	Moskva: Izdatel'stvo politicheskoi literatury, 1979	1	0	0	4
Iudin, I. N.	Sotsial'naya baza rosta KPSS	Moskva: Izdatel'stvo politicheskoi literatury, 1973	1	0	0	5
Tkachenko, N. G.	Izbratel'nyi uchastok	Moskva: Iuridicheskaiia literatura, 1975	0	0	0	2
Grigor'ev, V.	Poriadok provedeniia vyborov v Verkhovnye Sovety soiuzykh, a	Moskva: Izdatel'stvo Izvestiia, 1975	0	0	0	1
Smirnov, V. P.	Referendum v pechati	Moskva: Mysl, 1978	1	0	0	0

Figure 3: Transliterated scarcity checks

While ChatGPT OCR transliteration works well, its batch image processing requires a paid subscription. As part of our library's existing Microsoft subscription we already have full access to Microsoft Copilot, an AI tool based on ChatGPT. We had tested an earlier GPT-4-based version of Copilot early in the project and the results were not very

accurate. In August 2025 a new version of Copilot based on GPT-5 was released, so we decided to rerun the entire project using Copilot. The results closely matched those from ChatGPT:

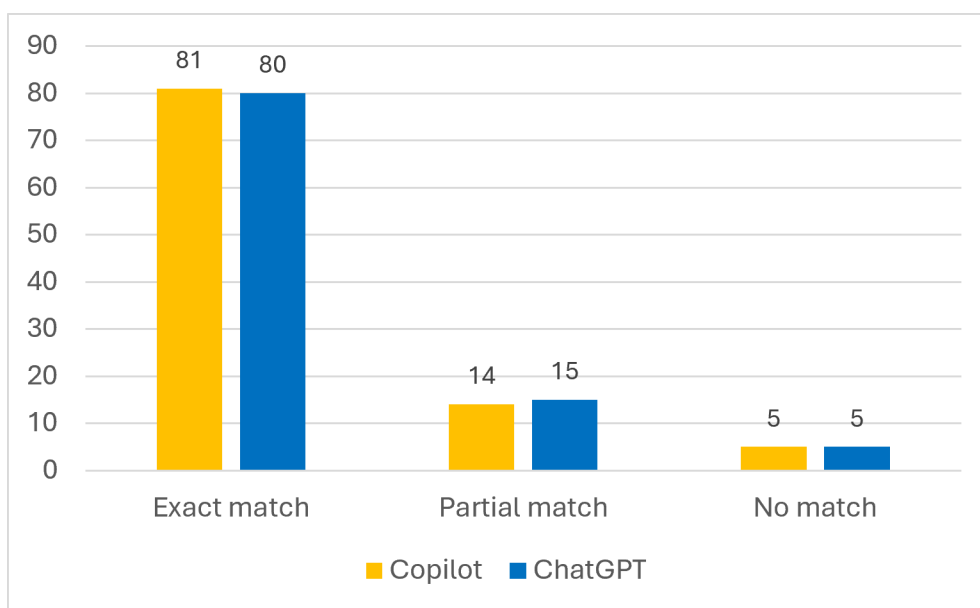


Figure 4: Copilot & ChatGPT overall results

From our study, we can therefore conclude that using AI transliterations to help catalogue non-English material works very well. AI can OCR non-English text with high accuracy, and AI transliteration is both effective and LC-compliant with around a 90% success rate in this pilot. With the help of AI tools, cataloguers can now catalogue books written in languages they can't read.

We can do this, but should we? AI is a powerful tool but it comes with serious environmental, ethical, and legal risks which we cannot ignore. For AI, analysing training data and responding to queries is an extremely intensive process, leading to "expanding demand for computing power, larger carbon footprints, shifts in patterns of electricity demand, and an accelerated depletion of natural resources" ([Bashir et al., 2024](#), p. 2). According to a recent United Nations report on AI, these environmental impacts "are growing significantly" ([United Nations Independent International Scientific Panel on Artificial Intelligence, 2026](#), p. 34). With England having just recorded its hottest June on record ([Carrington and Kassam, 2026](#)), we urgently need to start reducing, not increasing, our energy usage and carbon footprints.

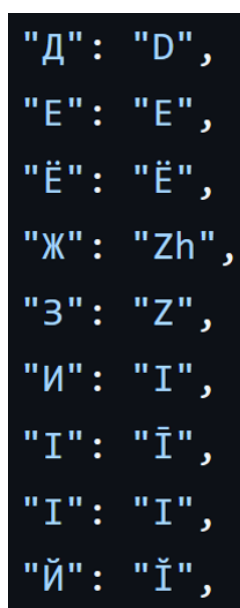
There are also ethical and potential legal issues surrounding AI use. The training data used for AI models is based on a "vast amount of publicly available data, some of which may include copyrighted content or unpublished proprietary information" ([Dover and Grzegorski, 2025](#), p. 607). ChatGPT's developer OpenAI claim their ingestion of copyrighted material falls under fair use ([The New York Times Company v. Microsoft Corporation', 2023](#), p. 4), but Hyung Doo Nam has argued that while fair use "was originally designed as an exceptional system for little users, ... AI developers and Big Tech companies have abused this doctrine in the name of 'public interest'"

([Nam, 2026](#), p. 606). As librarians who observe both the letter and spirit of copyright law, it is problematic if we use tools which arguably do not.

There are also practical implications to integrating AI into our workflows. AI models are opaque since "vendors do not always disclose the sources of their training data and how their models work, making their output difficult to understand, explain, and replicate" ([Dover and Grzegorski, 2025](#), p. 607). The difficulty of replicating results is compounded by the observed phenomenon of AI tools "mimic[ing] human-like fatigue when performing too many similar or tedious tasks in succession" ([Sun, 2025](#), p. 276), which makes them unreliable for metadata tasks which demand consistency. Integrating programs provided by external companies into our core workflows also exposes us to both pragmatic and financial risks, since the service offered by these companies could at any time become prohibitively expensive or even withdrawn entirely ([Michalak and Ellixson, 2024](#), p. 317).

So while it has been recently asserted that "there is an urgent need for the adoption of AI by libraries" ([Harisanty et al., 2023](#), p. 520), for this workflow we disagree. Instead, we propose using Python. Python is a programming language designed for data analysis and used extensively in data science. At the University of Glasgow Library, we have developed and started using our own Python tools to carry out automated scarcity checking and large-scale metadata analysis. We therefore decided to investigate whether Python could be used for OCR and transliteration of non-English texts.

Well-established open-source Python libraries are already available for both OCR and LC-compliant transliteration of Cyrillic text; our own project uses Pytesseract ([Lee, 2023](#)) for OCR and Domovyk ([Goodale, 2023](#)) for transliteration. One of the key benefits of using open-source software is that in contrast to AI's black box, we can inspect the code and see exactly what is going on:



"Д":	"D",
"Е":	"E",
"Ё":	"Ё",
"Ж":	"Zh",
"З":	"Z",
"И":	"I",
"І":	"Ī",
"І":	"I",
"Й":	"Ĭ",

Figure 5: Excerpt of Domovyk's transliteration table

If the LC transliteration tables were to be updated, we could amend our code immediately to reflect this rather than having to wait months or potentially years for an AI tool to be updated.

Python has already been successfully used to transliterate languages as diverse as Arabic, Manipuri, Middle Persian, and even Hittite cuneiform ([Farsi et al., 2025](#); [Tour et al., 2025](#); [Moirangthem and Nongmeikapam, 2026](#)). Monika Kilić has used Python to transliterate Russian and Serbian Cyrillic texts. Kilić used text rather than images as input, but she nevertheless found that "the computer's ability to parse and transliterate hundreds of titles in a few seconds offers a rate of efficiency that is unmatched" ([Kilić, 2025](#), p. 90). She used Copilot in her project to generate and revise the Python code and found that several stages of iteration were necessary to achieve accurate results ([Kilić, 2025](#), p. 82). Several other projects have also used Python for both OCR and transliteration ([Shrividya et al., 2025](#); [Roy, 2025](#)).

We have now written an experimental Python program which follows the same workflow we previously used with ChatGPT. It takes an image of a Russian title page, uses OCR to convert it to Cyrillic text, then transliterates this into LC-compliant Latin text which can then be used for copy cataloguing or scarcity checking.

This has several benefits compared to using AI:

- The program requires no subscription, account, or even internet access.
- It comes at a much lower environmental cost than AI, using negligible computing power while still processing an image in a few seconds.
- Being based on open-source code which was made freely available for use and reuse, it sidesteps the ethical and legal concerns of using AI tools.
- All the code is inspectable and modifiable, so we can adapt the Python code to our evolving requirements.
- Perhaps most valuably, developing our own in-house Python solution is an opportunity to upskill our own staff rather than outsourcing expertise to external companies.

While the development of our Python program for OCR and transliteration of Russian books is still in its early stages, the initial results have been promising. The program can transliterate Russian text very effectively, likely because this is a straightforward process of converting one character to another according to a crossmapping table.

We have, however, found the OCR part of the process more challenging. In our testing, Python OCR works very well with full pages of text but is less effective on pages with areas of blank space (such as title pages). However, if we crop out smaller sections of text to use as the input images the results are much more accurate:

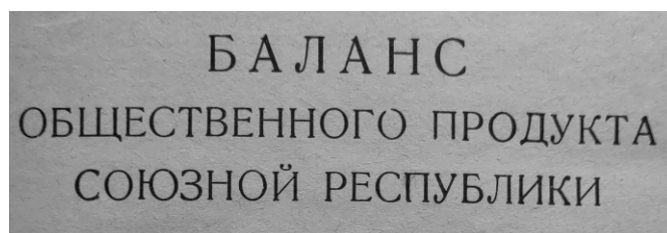


Figure 6: UoGL Python programme input

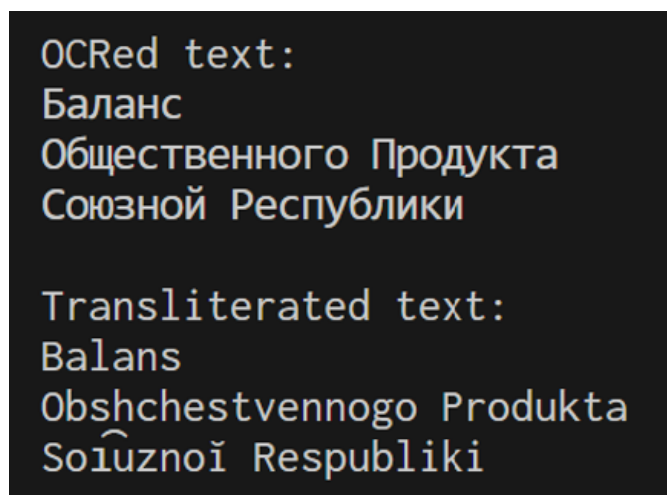


Figure 7: UoGL Python programme output

Python OCR is also much more sensitive to input image quality. While AI tools are generally quite forgiving of poorer image quality, effective Python OCR requires images which are in sharp focus and black and white or grayscale.

Our future plans for this Python OCR project are initially centred on further development and refinement of the process. We want to carry out more testing to determine the optimal image preprocessing pipeline for the input images, because other projects using the same Tesseract OCR engine as us have improved output quality by applying image filtering ([Shrividyia et al., 2025](#), p. 2). Our program could then automatically apply the optimal preprocessing pipeline to each input image.

We will carry out more extended testing to compare the accuracy of our Python OCR and transliteration tool against both AI tools and human cataloguers. Once we are satisfied with the accuracy of the output, we will investigate expanding the tool to support other languages and scripts. When our Python tool is feature-complete and working reliably, we will make it open-source and freely available online for anyone to use and modify.

In conclusion, AI tools such as ChatGPT and Copilot are capable of accurate OCR and transliteration of Russian texts. They can provide high-quality output even with low-quality input images and are very easy to use. However, they come with serious environmental, legal, and ethical drawbacks.

Python can transliterate accurately but its OCR quality can be more variable without image preprocessing. Our Python tool's great strength is that it is fully customisable, transparent, efficient, and does not have any associated environmental or ethical costs. It is, however, still in its early experimental stages and requires more development before it is ready to be integrated into daily work.

AI tools are currently the best technological option to help cataloguers without specialist language skills transliterate non-English texts. But with further development, we are confident that a Python tool will be a viable alternative which for us is preferable to using AI.

But whether we use Python or AI tools, we must remember that for all their sophistication they are just that: tools. The Program for Cooperative Cataloguing's recent argument that "AI cannot replace human expertise" ([PCC Task Group on AI and Machine Learning for Cataloging and Metadata, 2025](#), p.5) applies to Python just as much as AI. If we rely solely on these tools for transliteration and don't have cataloguers with language skills to check the output, errors will accumulate over time with a consequent gradual degradation of metadata quality.

These tools can be useful to help us create a basic catalogue record to make an uncatalogued item discoverable, but we cannot and should not assume their outputs are correct. Cataloguers with language skills must remain an integral part of our workflows, because "while the fantasy of being able to delegate our responsibilities to the touch of a button is potentially alluring, we must not sacrifice the quality and depth of our work in leaping to chase it" ([Engel et al., 2025](#), p. 273).

References

- Bashir, N. et al. (2024) 'The Climate and Sustainability Implications of Generative AI', in *An MIT Exploration of Generative AI*. Available at: <https://doi.org/10.21428/e4baedd9.9070dfe7> [Accessed: 18 August 2026]
- Carrington, D. and Kassam, A. (2026) 'UK and Switzerland record hottest ever June day as health emergencies surge in Europe', *The Guardian*, 25 June. Available at: <https://web.archive.org/web/20260629171559/https://www.theguardian.com/environment/2026/jun/25/highest-june-minimum-temperature-record-broken-cardiff-savage-heatwave-continues> [Accessed: 18 August 2026]
- Dover, A. and Grzegorski, J. (2025) 'Artificial Intelligence Through the Lens of the Cataloguing Code of Ethics', *Cataloging & Classification Quarterly*, 63(6–7), pp. 600–620. Available at: <https://doi.org/10.1080/01639374.2025.2544137> [Accessed: 18 August 2026]
- Engel, J. Y. et al. (2025) 'Artificial Intelligence in Library Cataloging: A Review of Literature', *Journal of Library Metadata*, 25(4), pp. 261–276, Available at: <https://doi.org/10.1080/19386389.2025.2526913> [Accessed: 18 August 2026]

- Farsi, F. et al. (2025) 'ParsiPy: NLP Toolkit for Historical Persian Texts in Python', *Second Workshop on Ancient Language Processing*, Albuquerque, 3–4 May. Association for Computational Linguistics, pp. 238–250. Available at: <https://doi.org/10.18653/v1/2025.alp-1.17> [Accessed: 18 August 2026]
- Gamage, R. and Wanigasooriya, P. (2024) 'Using Generative AI for Bibliographic Description: A Study with ChatGPT 4', *Journal of the University Librarians Association of Sri Lanka*, 27(2), pp. 257–284. Available at: <https://doi.org/10.4038/jula.v27i2.8083> [Accessed: 18 August 2026]
- Geroimenko, V. (2025) *The Essential Guide to Prompt Engineering: Key Principles, Techniques, Challenges, and Security Risks*. Springer.
- Goodale, I. (2023) *Domovyk* (Version 1.0.1) [Computer program]. Available at: <https://github.com/ian-nai/domovyk> [Accessed: 18 August 2026]
- Harisanty, D. et al. (2023) 'Is Adopting Artificial Intelligence in Libraries Urgency or a Buzzword? A Systematic Literature Review', *Journal of Information Science*, 51(2), pp. 511–522. Available at: <https://doi.org/10.1177/01655515221141034> [Accessed: 18 August 2026]
- Jiang, Y. et al. (2024) 'Exploring the Potential Applications of Generative AI in East Asian Librarianship: Use Cases, Reflections, and Inspirations', *Journal of East Asian Libraries*, 179, pp. 43–58. Available at: <https://scholarsarchive.byu.edu/jeal/vol2024/iss179/5/> [Accessed: 18 August 2026]
- Kilić, M. (2025) 'Using AI Tools for Slavic Transliteration to Support Cataloging Workflows: Potential Use Cases and Limitations', *Slavic & East European Information Resources*, 26(1–2), pp. 80–91. Available at: <https://doi.org/10.1080/15228886.2025.2518369> [Accessed: 18 August 2026]
- Lee, M. A. (2023) *Pytesseract* (Version 0.3.13) [Computer program]. Available at: <https://github.com/madmaze/pytesseract> [Accessed: 18 August 2026]
- Library of Congress (2012a) *Russian Romanization Table*. Available at: <https://www.loc.gov/catdir/cpso/romanization/russian.pdf> [Accessed: 3 September 2026]
- Library of Congress (2012b) *Russian Romanization Table* [Source document]. Available at: <https://www.loc.gov/catdir/cpso/romanization/russian.doc> [Accessed: 3 September 2026]
- Michalak, R. and Ellixson, D. (2024) 'Buy Versus Build: Navigating Artificial Intelligence (AI) Tool Adoption in Academic Libraries', *Information Services and Use*, 44(4), pp. 316–326. Available at: <https://doi.org/10.1177/18758789241296755> [Accessed: 18 August 2026]
- Moirangthem, G. and Nongmeikapam, K. (2026) 'Optimizing Low-Resource Machine Transliteration: A Case of Script Transition from Manipuri in Bengali Script to Meetei-Mayek', *ACM Transactions on Asian and Low-Resource Language Information Processing*, 25(5). Available at: <https://doi.org/10.1145/3806198> [Accessed: 18 August 2026]
- Nam, H. D. (2026) 'The Paradox of Fair Use – A Giant on a Seesaw: Restoring the Balance in Copyright and Media-Specific Issues', *GRUR International*, 75(7), pp. 605–606. Available at: <https://doi.org/10.1093/grurint/ikag001> [Accessed: 18 August 2026]
- '*The New York Times Company v. Microsoft Corporation*' (2023) United States District Court for the Southern District of New York, 1:23-cv-11195. Court Listener. Available at: <https://www.courtlistener.com/docket/68117049/the-new-york-times-company-v-microsoft-corporation/#entry-1> [Accessed: 18 August 2026]
- Ogunbenro, O. D. et al. (2025) 'Revolutionizing Library Services: The Impact of Artificial Intelligence on Cataloguing and Access to Information in Nigeria Academic Libraries',

- Journal of Library Metadata*, 25(2), pp. 99–118. Available at: <https://doi.org/10.1080/19386389.2025.2475418> [Accessed: 18 August 2026]
- Oyighan, D. et al. (2024) 'The Role of AI in Transforming Metadata Management: Insights on Challenges, Opportunities, and Emerging Trends', *Asian Journal of Information Science and Technology*, 14(2), pp. 20–26. Available at: <https://doi.org/10.70112/ajist-2024.14.2.4277> [Accessed: 18 August 2026]
- PCC Task Group on AI and Machine Learning for Cataloging and Metadata (2025) *PCC Task Group on AI and Machine Learning in Cataloging and Metadata: Final Report*. Available at: <https://web.archive.org/web/20260708094646/https://www.loc.gov/aba/pcc/taskgroup/AI-and-Machine-Learning-TG-final-report.pdf> [Accessed: 18 August 2026]
- Roy, S. (2025) 'Global Scriptualization of Sanskrit: AI-Assisted OCR, Transliteration, and Translation Across 50+ Languages and Scripts' [Preprint]. Available at: <https://doi.org/10.13140/RG.2.2.35119.60326> [Accessed: 18 August 2026]
- Shrividya, M. L. et al. (2025) 'Multilingual Digitization of Ancient Manuscripts: OCR-Based Preservation and Translation', 5th International Conference on Emerging Research in Electronics, Computer Science and Technology (ICERECT), Mandya, 12–13 September. IEEE. Available at: <https://doi.org/10.1109/ICERECT65215.2025.11377349> [Accessed: 18 August 2026]
- Sun, L. (2025) 'Enhancing Cataloging with Generative AI: Converting Wade-Giles to Pinyin', *Cataloging & Classification Quarterly*, 63(4), pp. 267–283. Available at: <https://doi.org/10.1080/01639374.2025.2508956> [Accessed: 18 August 2026]
- Sussmeier, S. and Henry, J. A. (2025) 'Mind the Gap!: How Do We Ethically Bridge the Divide Between the Cataloging/Metadata Community and the World of AI?', *Journal of Library Metadata*, 25(4), pp. 241–259. Available at: <https://doi.org/10.1080/19386389.2025.2525720> [Accessed: 18 August 2026]
- Togia, A. et al. (2026) 'Artificial Intelligence in Greek Libraries: Perceptions, Challenges, and Readiness for Adoption', *Journal of Librarianship and Information Science*. Available at: <https://doi.org/10.1177/09610006251410240> [Accessed: 18 August 2026]
- Tour, E. et al. (2025) 'Potnia: A Python Library for the Conversion of Transliterated Ancient Texts to Unicode', *Journal of Open Source Software*, 10(108). Available at: <https://doi.org/10.21105/joss.07725> [Accessed: 18 August 2026]
- United Nations Independent International Scientific Panel on Artificial Intelligence (2026) *Preliminary Report of the Independent International Scientific Panel on AI: Evidence-Based Assessment of Opportunities, Risks and Impacts of Artificial Intelligence*. United Nations. Available at: [https://web.archive.org/web/20260710091141/https://www.un.org/independent-international-scientific-panel-ai/sites/default/files/2026-07/en_Preliminary%20Report .pdf](https://web.archive.org/web/20260710091141/https://www.un.org/independent-international-scientific-panel-ai/sites/default/files/2026-07/en_Preliminary%20Report.pdf) [Accessed: 18 August 2026]