

**INTEGRATION OF ARTIFICIAL INTELLIGENCE INTO HEALTH
TECHNOLOGY ASSESSMENT: OPTIMIZING CLINICAL ANALYSIS AND
SYNTHESIS OF EVIDENCE ON EFFECTIVENESS AND SAFETY**

Keywords: Health Technology Assessment; Artificial Intelligence; Clinical Efficiency;
Data Mining; Machine Learning; Evidence-Based Medicine

ABSTRACT

Clinical assessment is one of the most labor-intensive components of Health Technology Assessment, as it requires systematic literature retrieval, critical appraisal, and synthesis of scientific evidence. Recent advances in artificial intelligence have created new opportunities to optimize these processes; however, comparative analyses of artificial intelligence-based tools supporting the clinical component of Health Technology Assessment remain limited.

The aim of the study was to systematize current evidence on the application of artificial intelligence in the clinical component of Health Technology Assessment and to identify its advantages, limitations, and prospects for evidence synthesis. The study was based on a structured review of scientific publications using descriptive, comparative, and content analysis methods.

The findings demonstrated that contemporary artificial intelligence-based tools support different stages of evidence synthesis, including literature retrieval, publication screening, structured data extraction, risk-of-bias assessment, and evidence synthesis. According to their functional purpose, the identified tools were grouped into two principal categories: specialized platforms developed to automate specific stages of clinical assessment and large language models capable of supporting a broad range of analytical tasks. Comparative analysis revealed substantial differences in their functional capabilities, degree of automation, and potential applicability to Health Technology Assessment. None of the identified artificial intelligence-based tools currently supports the complete clinical assessment workflow, indicating that they should be regarded as complementary rather than interchangeable technologies.

For Ukraine, the integration of artificial intelligence-based tools into national and hospital-based Health Technology Assessment offers promising opportunities to improve the efficiency of evidence synthesis, optimize the use of expert resources, and reduce the time required to prepare the clinical component of Health Technology Assessment reports. The implementation of artificial intelligence-based tools should be accompanied by mandatory expert validation of the obtained results. Further integration of artificial intelligence into Health Technology Assessment practice requires the development of methodological guidance together with clear principles governing the application and validation of artificial intelligence-based tools.

I. A. КОСТЮК (<https://orcid.org/0000-0002-3689-3379>), канд. фарм. наук, доцент
Національний медичний університет імені О. О. Богомольця, м. Київ, Україна

ІНТЕГРАЦІЯ ШТУЧНОГО ІНТЕЛЕКТУ В ОЦІНКУ МЕДИЧНИХ ТЕХНОЛОГІЙ: ОПТИМІЗАЦІЯ КЛІНІЧНОГО АНАЛІЗУ ТА СИНТЕЗУ ДОКАЗІВ ЕФЕКТИВНОСТІ Й БЕЗПЕКИ

Ключові слова: оцінка медичних технологій, штучний інтелект, клінічна ефективність, інтелектуальний аналіз даних, доказова медицина/

АНОТАЦІЯ

Клінічна оцінка є одним із найбільш трудомістких етапів оцінки медичних технологій, оскільки потребує систематичного пошуку, критичного аналізу та синтезу наукових доказів. Сучасний розвиток технологій штучного інтелекту створює нові можливості для оптимізації цих процесів, однак порівняльний аналіз інструментів штучного інтелекту, що застосовуються для підтримки клінічної частини в оцінці медичних технологій, залишається обмеженим.

Метою дослідження було систематизувати сучасні дані щодо застосування штучного інтелекту в клінічній частині оцінки медичних технологій визначити його переваги, обмеження та перспективи використання під час синтезу доказів. Дослідження ґрунтувалося на структурованому огляді наукових публікацій із застосуванням описового, порівняльного та контент-аналізу.

Встановлено, що сучасні інструменти штучного інтелекту підтримують різні етапи синтезу доказів, зокрема пошук літератури, скринінг публікацій, вилучення структурованих даних, оцінювання ризику систематичної похибки та синтез доказів. За функціональним призначенням їх згруповано у дві основні категорії: спеціалізовані платформи, розроблені для автоматизації окремих етапів клінічної оцінки, та великі мовні моделі, здатні підтримувати широкий спектр аналітичних завдань. Порівняльний аналіз засвідчив суттєві відмінності між ними за функціональними можливостями, рівнем автоматизації та потенціалом практичного застосування в оцінці медичних технологій. Водночас жоден із проаналізованих інструментів штучного інтелекту наразі не забезпечує повний цикл клінічної оцінки, що свідчить про доцільність їх комплексного, а не взаємозамінного використання.

Для України інтеграція інструментів штучного інтелекту у державну та госпітальну оцінку медичних технологій відкриває перспективи підвищення ефективності синтезу доказів, оптимізації використання експертних ресурсів і скорочення тривалості підготовки клінічної частини звітів з оцінки медичних технологій. Впровадження інструментів штучного інтелекту має здійснюватися із обов'язковою експертною верифікацією результатів. Подальша інтеграція технологій штучного інтелекту в практику оцінки медичних технологій потребує розроблення методичних рекомендацій і чітких принципів застосування та їх валідації.

Introduction

The modern healthcare system operates under conditions of limited financial and human resources, evolving population needs, and the rapid implementation of new medical technologies. In this context, Health Technology Assessment (HTA) serves as an effective tool for evidence-based management decision-making, grounded in clinical evidence and economic feasibility for both patients and healthcare facilities [1, 2]. The HTA procedure is complex, multi-stage, and resource-intensive, as it involves the analysis of vast volumes of scientific, clinical, and statistical data [3, 4]. The significant workload on experts, the duration of the process, and the potential risk of subjective assessment underscore the need for innovative solutions to optimize individual stages and improve evaluation accuracy [5, 6].

Recent advances in artificial intelligence (AI) have created new opportunities to optimize evidence synthesis and support HTA. Methodological approaches to AI-assisted evidence synthesis and systematic reviews have been explored by Marshall et al. [7-9], Chai et al.

[10], O'Mara-Eves et al. [11], and Guo et al. [12]. Organizational, methodological, and regulatory aspects of AI implementation in HTA were investigated by Zemplényi et al. [13], who identified key barriers and proposed recommendations for integrating AI-generated evidence into HTA process. At the national level, growing interest in the application of AI in healthcare is also reflected in recent Ukrainian studies, which emphasize its increasing role in clinical practice, biomedical research, and healthcare education [14-18].

Despite these advances, existing studies mainly focus on individual AI methods or organizational aspects of their implementation. Comprehensive comparative analyses of contemporary AI tools for the clinical assessment of HTA, including their strengths, limitations, applicability to the Ukrainian healthcare system, methodological transparency, reproducibility, potential algorithmic bias, and the need for expert oversight, remain limited.

Accordingly, the scientific and practical value of this study lies in the systematic analysis of current AI applications in the clinical assessment of HTA, identification of their advantages and limitations, and evaluation of the prospects for integrating these technologies into HTA methodology and evidence-informed decision-making in healthcare and pharmacy in Ukraine.

Aim. To systematize current evidence on the application of artificial intelligence in the clinical assessment of HTA and to identify its advantages, limitations, and prospects for implementation in evidence synthesis.

Materials and methods

The study was based on a structured literature review of scientific publications addressing the application of AI to support clinical assessment in HTA. Descriptive, comparative, and semantic content analyses were applied to summarize current evidence, identify the main areas of AI application, and comparatively evaluate contemporary AI-based solutions.

A literature search was conducted in the PubMed and Cochrane Library databases covering the period from 2015 to 2025. The search strategy combined keywords related to artificial intelligence (“artificial intelligence”, “machine learning”, “large language models”, “natural language processing”) with terms related to HTA (“health technology assessment”, “HTA”, “evidence synthesis”, “systematic review”). The search identified 156 publications, including 145 records from PubMed and 11 from the Cochrane Library. The reference lists of the selected publications were additionally screened to identify other relevant studies.

Publications describing or evaluating specific AI platforms or tools for literature searching, evidence identification, literature screening, data extraction, risk-of-bias assessment, or evidence synthesis applicable to clinical assessment in HTA were eligible for inclusion. Publications focused exclusively on diagnostic AI, medical imaging, drug discovery, or other clinical applications unrelated to evidence synthesis, as well as studies lacking a description or evaluation of specific AI-based tools, were excluded.

Following title, abstract, and full-text screening, nine publications meeting the predefined eligibility criteria were included in the final analysis. Based on these studies, a comparative evaluation of contemporary AI platforms and tools was performed according to their functional capabilities, advantages, limitations, and potential applicability to clinical assessment in HTA.

Results and discussion

Clinical analysis is a fundamental stage of HTA, the quality of which depends on the comprehensive identification of relevant studies within evidence-based medicine databases. The Cochrane Collaboration recommends conducting full duplicate screening of publications (by two reviewers independently at all stages), which requires significant time and financial resources, effectively doubling the workload on experts [6, 19]. Furthermore, during the abstract screening stage, up to 90% of the materials retrieved after the initial search are excluded. The subjectivity and variability of decisions between reviewers often necessitate additional time for reconciliation, creating difficulties given the limited resources available for assessments [5, 20].

The literature review identified several AI-based tools designed to support different stages of evidence synthesis within HTA. Their functionality ranges from literature screening and structured data extraction to risk-of-bias assessment and evidence synthesis.

One of the most widely described tools is RobotReviewer, an open-source platform developed to automate and standardize the analysis of biomedical literature. Its architecture is based on a three-stage workflow comprising preprocessing, natural language processing, and report generation. This structure enables the sequential transformation of unstructured full-text publications into standardized analytical outputs suitable for integration into HTA reports.

During the preprocessing stage, the algorithm performs structural decomposition of the document, identifying its key components and conducting tokenization. Subsequently, natural language processing techniques are applied to automatically identify study design, extract information according to the PICO (Population, Intervention, Comparator, Outcome) framework, and assess the risk of bias. In addition, RobotReviewer applies Principal Component Analysis to support data visualization and clustering, facilitating the identification of clinically relevant evidence. At the final stage, the extracted information is integrated into standardized HTML templates with export options including HTML, DOCX, and JSON formats. Overall, RobotReviewer is primarily intended to facilitate structured evidence extraction and risk-of-bias assessment, thereby supporting preparation of the clinical section of HTA reports [7, 21, 22].

While RobotReviewer primarily supports the critical appraisal of clinical evidence, other AI-based tools have been developed to address different stages of evidence identification and management. One of the most representative examples is TrialStreamer, which functions as a dynamic, continuously updated database specializing in collecting and structuring reports on randomized controlled trials. This system was designed to ensure continuous monitoring of biomedical publications for the rapid identification of new clinical evidence. The application of AI algorithms minimizes the time gap between the appearance of a publication and its subsequent analysis, supporting the use of the most up-to-date information for clinical analysis and HTA.

The technological advantage of TrialStreamer lies in the application of BioBERT, a language model pretrained on PubMed publications. Its architecture consists of two consecutive stages: automated aggregation and filtering of publications from evidence-based medicine databases, followed by intelligent screening and structured information

extraction. The platform automatically extracts PICO elements together with additional metadata, including sample size, treatment effects, and study conclusions, which are subsequently stored in the TrialStreamer Database.

A comprehensive performance evaluation of this system confirms its high operational efficiency and reliability in handling big data. It has been established that the intelligent classifier for randomized controlled trials demonstrates a significant ability to filter out information noise, successfully rejecting 94-97% of irrelevant content. At the same time, the accuracy of identifying target studies exceeds 50%, which is a high result for fully automated systems. This significantly reduces the labor intensity of the scientific publication screening stage while maintaining the high quality of evidence synthesis [8, 9].

A different approach to supporting systematic reviews is implemented in Research Screener, which focuses on optimizing the primary screening of scientific publications. The functional architecture of the system implements a cyclical active learning process based on the interaction between the algorithm and the researcher. Unlike static filters, this platform provides for continuous updates to the internal relevance model based on the user's expert decisions, allowing the system to dynamically adapt to highly specialized research topics.

The platform's operation begins with the stage of uploading the general information array generated by the initial search strategy and providing at least one representative standard – a publication that perfectly meets the inclusion criteria. The researcher is presented with an array of the 50 highest-ranking publications for manual verification. Each decision made by the researcher regarding the inclusion or exclusion of an article is instantly processed by the system, leading to an immediate recalculation of weight coefficients and a re-ranking of the remaining publications.

This approach in the Research Screener system allows for the accumulation of the most significant sources at the top of the list, which significantly reduces the volume of manual labor. Practical application of this structure demonstrates that a researcher can identify up to 95% of relevant studies by processing only a small fraction of the total search results [10, 11, 23].

Alongside the development of specialized AI-based tools, recent research has increasingly focused on the application of large language models, particularly ChatGPT, to evidence synthesis. Owing to their versatility, large language models can support a broad range of analytical tasks, including publication screening, scientific text analysis, evidence summarization, and preparation of narrative sections of HTA reports.

The validation of this tool for the preparation of the clinical section of the report was conducted by a group of Canadian researchers. They developed a database of titles and abstracts that had previously undergone independent verification by two experts according to strict inclusion criteria. Using the results of this double screening as a “gold standard” allowed for an objective assessment of the network's ability to recognize clinically relevant evidence.

The technical implementation of the process involved using a software script for direct interaction with the OpenAI API (GPT-3.5/GPT-4). The algorithm's operation was based on the principle of stream data processing: the network sequentially analyzed titles and

abstracts according to a formalized screening protocol. This approach allowed for the standardization of analytical operations and minimized the variability of decisions resulting from the human factor in manual screening.

According to the process architecture, each publication was evaluated by the language model within a single instructional prompt. The model classified the article as “relevant” for the clinical section of the HTA report based solely on the annotated data, fully simulating the real-world working conditions of an expert during the pre-screening stage.

Across six independent clinical datasets, ChatGPT demonstrated an overall accuracy of 0.91, sensitivity of 0.76, and specificity of 0.91 for excluding irrelevant publications. These findings suggest considerable potential for accelerating evidence synthesis while reducing technical errors associated with large-scale literature screening [12]. Unlike specialized platforms developed for individual stages of evidence synthesis, ChatGPT demonstrates the potential to support multiple analytical tasks throughout the clinical assessment process. Nevertheless, its outputs require mandatory expert validation before incorporation into HTA reports.

Taken together, the reviewed AI-based tools support different stages of clinical assessment and perform complementary rather than interchangeable functions throughout the evidence synthesis process. Their potential application across the principal stages of clinical assessment in HTA is summarized in Figure 1.

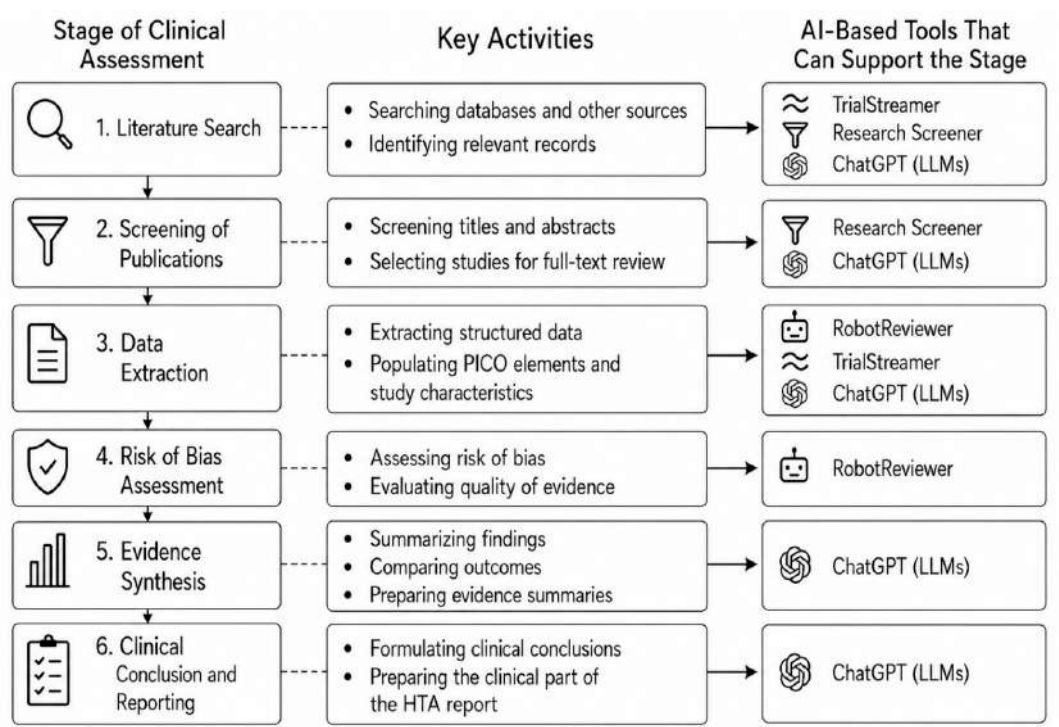


Figure 1. Application of AI-based tools across the stages of clinical assessment in HTA

The functional diversity of these AI tools prompted their further classification according to their principal role in clinical assessment.

Analysis of the selected publications enabled the identification of two principal categories of AI-based tools. The first comprises specialized platforms developed to automate individual stages of clinical assessment, including RobotReviewer, TrialStreamer, and Research Screener. Their functionality is primarily focused on literature searching and screening, structured data extraction, risk-of-bias assessment, and organization of scientific evidence. The second category consists of large language models, primarily ChatGPT, which, owing to their versatility, can support a broad spectrum of analytical activities, including scientific text analysis, evidence summarization, preparation of narrative reports, and evidence synthesis.

To provide a structured comparison of the identified AI-based tools, their principal functions, potential contribution to clinical assessment within HTA, and major limitations are summarized in Table 1.

Table 1

Comparative analysis of contemporary AI-based tools supporting clinical assessment in HTA

AI-based tool	Primary purpose	Functions	Potential contribution to clinical assessment in HTA	Limitations / considerations
Robot-Reviewer	Automation of biomedical literature analysis	Extraction of study characteristics, identification of PICO elements, automated risk-of-bias assessment	May reduce experts' workload during randomized controlled trial assessment and standardize preparation of the clinical section of HTA reports	Requires expert validation; primarily developed for randomized controlled trials and risk-of-bias assessment
Trial-Streamer	Automated identification and continuous monitoring of clinical trials	Automated identification, classification, and structuring of randomized controlled trial reports; extraction of key study characteristics	May facilitate timely identification of new clinical evidence for HTA	Primarily focused on randomized controlled trials; does not replace critical appraisal by experts
Research Screener	Semi-automated title and abstract screening	Machine learning-based ranking of publications according to relevance and support for literature screening	May substantially reduce manual workload during systematic reviews	Performance depends on model training quality and expert input
ChatGPT (LLMs)	Support for scientific text analysis and literature screening	Classification of publications, evidence summarization, support for evidence synthesis	Has the potential to accelerate evidence synthesis and preparation of analytical reports	Requires independent expert validation; susceptible to inaccuracies and response variability

The comparison presented in Table 1 revealed several important trends regarding the current application of artificial intelligence within HTA. Contemporary AI-based tools primarily automate the most labor-intensive stages of clinical assessment, including literature searching, screening, and evidence processing. However, none of the identified tools supports the complete clinical assessment workflow within HTA, indicating that their application should be considered complementary and selected according to the specific objectives of each assessment.

The reviewed studies consistently demonstrate that current AI-based tools are designed to automate routine analytical tasks, whereas the critical interpretation of clinical evidence, formulation of conclusions, and development of recommendations remain the responsibility of HTA experts. Accordingly, all reviewed publications emphasize the necessity of expert

validation of AI-generated outputs before their incorporation into HTA reports and healthcare decision-making.

Another important finding is the transition from highly specialized AI platforms toward versatile large language models. Whereas earlier systems were developed to automate individual stages of systematic reviews, large language models increasingly support multiple interconnected activities ranging from literature screening and scientific text analysis to evidence synthesis and preparation of analytical reports. However, this broader functionality is accompanied by new methodological challenges related to reproducibility, algorithmic transparency, and verification of generated outputs.

AI-based tools have the potential to substantially optimize the most labor-intensive stages of HTA. However, their effective implementation also depends on an appropriate legal and regulatory framework. Currently, the methodological documents governing HTA in Ukraine do not include provisions regulating the use of AI-based tools in HTA report preparation. Although the Ministry of Digital Transformation of Ukraine has developed the AI Roadmap and the White Paper on AI Regulation, these documents establish only a general regulatory framework and require further adaptation to the healthcare and pharmaceutical sectors.

For Ukraine, the application of AI-based tools within HTA appears particularly promising given the need to optimize expert, time, and financial resources, especially in the context of the ongoing development of national and hospital-based HTA. Automation of literature searching, publication screening, structured data extraction, and preparation of preliminary analytical materials could substantially reduce report preparation time, decrease expert workload, and improve the reproducibility of individual assessment procedures. These advantages may be particularly valuable for healthcare institutions implementing hospital-based HTA, where multidisciplinary expert capacity remains limited.

Nevertheless, the integration of AI into HTA practice should be accompanied by the development of national methodological guidance defining appropriate areas of application, requirements for mandatory expert validation, and principles governing the use of AI throughout the clinical assessment process. Particular attention should also be paid to algorithmic transparency, reproducibility, data protection, and the clear allocation of responsibilities between AI systems and human experts. The reviewed evidence suggests that AI should currently be regarded as a decision-support technology rather than an autonomous decision-making system within HTA. Its gradual integration into the Ukrainian HTA framework may become an important step toward strengthening evidence-informed decision-making and improving the efficiency and methodological consistency of both national and hospital-based HTA.

Conclusion

1. This study systematized the current evidence on the application of artificial intelligence in the clinical component of HTA. The findings demonstrated that contemporary AI-based tools support different stages of evidence synthesis, including literature retrieval, publication screening, structured data extraction, risk-of-bias assessment, and evidence synthesis, thereby improving the efficiency of the most labor-intensive stages of clinical assessment.

2. The comparative analysis demonstrated that specialized AI-based tools (RobotReviewer, TrialStreamer, and Research Screener) and large language models, particularly ChatGPT, differ substantially in their functional capabilities and should be

considered complementary rather than interchangeable. None of the identified AI-based tools currently supports the entire clinical assessment workflow within HTA, highlighting the continuing importance of expert involvement throughout the HTA process.

3. For Ukraine, the integration of AI-based tools into national and hospital-based HTA represents a promising direction for improving the efficiency of evidence synthesis and optimizing the use of expert resources. Future integration of AI into Ukrainian HTA will require the development of methodological guidance for AI-assisted HTA together with clear principles governing the application and validation of AI-based tools.

References

1. World Health Organization. *Health technology assessment*. URL: https://www.who.int/health-topics/health-technology-assessment#tab=tab_1
2. O'Rourke, B., Oortwijn, W., & Schuller, T. (2020). The new definition of health technology assessment: A milestone in international collaboration. *International Journal of Technology Assessment in Health Care*, 36(3), 187–190. <https://doi.org/10.1017/S0266462320000215>
3. Institutionalizing health technology assessment mechanisms: A how-to guide. (2021). World Health Organization. URL: <https://www.who.int/publications/i/item/9789240020665>
4. Gylmark, M., Lampe, K., Ruof, J., Pöhlmann, J., Hebborn, A., & Kristensen, F. B. (2018). Is the EUnetHTA HTA Core Model® fit for purpose? Evaluation from an industry perspective. *International Journal of Technology Assessment in Health Care*, 34(5), 458–463. <https://doi.org/10.1017/S0266462318000594>
5. Polanin, J. R., Pigott, T. D., Espelage, D. L., & Grotzinger, J. K. (2019). Best practice guidelines for abstract screening in large-evidence systematic reviews and meta-analyses. *Research Synthesis Methods*, 10(3), 330–342. <https://doi.org/10.1002/jrsm.1354>
6. Cochrane Library Training Hub: User Guide. (2025). URL: <https://www.wiley.com/en-ie/solutions-partnerships/customer-success-hub/cochrane-library-training-hub>
7. Marshall, I. J., Kuiper, J., Banner, E., & Wallace, B. C. (2017). Automating biomedical evidence synthesis: RobotReviewer. In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics: System Demonstrations*, 7–12. <https://doi.org/10.18653/v1/P17-4002>
8. Marshall, I. J., Nye, B., Kuiper, J., Noel-Storr, A., Marshall, R., Maclean, R., Soboczenski, F., Nenkova, A., Thomas, J., Wallace, B. C. (2020). Trialstreamer: A living, automatically updated database of clinical trial reports. *Journal of the American Medical Informatics Association*, 27(12), 1903–1912. <https://doi.org/10.1093/jamia/ocaa163>
9. Nye, B. E., Nenkova, A., Marshall, I. J., & Wallace, B. C. (2020). Trialstreamer: Mapping and Browsing Medical Evidence in Real-Time. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics: System Demonstrations*, 63–69. <https://doi.org/10.18653/v1/2020.acl-demos.9>
10. Chai, K. E. K., Lines, R. L. J., Gucciardi, D. F., & Ng, L. (2021). Research Screener: A machine learning tool to semi-automate abstract screening for systematic reviews. *Systematic Reviews*, 10(1), Art. 93. <https://doi.org/10.1186/s13643-021-01635-3>
11. O'Mara-Eves, A., Thomas, J., McNaught, J., Miwa, M., & Ananiadou, S. (2015). Validity of using a semi-automated screening tool in a systematic review assessing non-specific effects of respiratory vaccines. *Systematic Reviews*, 4, Art. 11. <https://doi.org/10.1186/2046-4053-4-11>
12. Guo, E., Gupta, M., Deng, J., Park, Y. J., Paget, M., & Naugler, C. (2024). Automated paper screening for clinical reviews using large language models: Data analysis study. *Journal of Medical Internet Research*, 26, e48996. <https://doi.org/10.2196/48996>
13. Zemplényi, A., Tachkov, K., Balkanyi, L., Németh, B., Petykó, Z. I., Petrova, G., Czech, M., Dawoud, D., Goettsch, W., Gutierrez-Ibarluzea, I., Hren, R., Knies, S., Lorenzovici, L., Maravic, Z., Piniashko, O., Savova, A., Manova, M., Tesar, T., Zerovnik, S., & Kaló, Z. (2023). Recommendations to overcome barriers to the use of artificial intelligence-driven evidence in health technology assessment. *Frontiers in Public Health*, 11, 1088121. <https://doi.org/10.3389/fpubh.2023.1088121>
14. Shostakovych-Koretska, L. R., & Kopcha, V. S. (2025). Use of artificial intelligence in clinical medicine and scientific research. *Medical Education*, 1, 99–107. <https://doi.org/10.11603/m.2414-5998.2025.1.15382> [in Ukrainian].
15. Matviienko, M. M. (2024). Artificial intelligence technologies as a component of the digital competence of future doctors. *Medicine and Pharmacy: Educational Discourses*, 1, 23–29. <https://doi.org/10.32782/eddiscourses/2024-1-4> [in Ukrainian].
16. Kostiuk, I. A., & Kosyachenko, K. L. (2026). Hospital-based health technology assessment as a tool for managerial decision-making: Methodological analysis. *Medical Science of Ukraine*, 22(1), 145–152. <https://doi.org/10.32345/2664-4738.1.2026.17> [in Ukrainian].

17. Babenko, M. M., & Kosyachenko, K. L. (2025). Expert evaluation of priority areas for the implementation of medical and pharmaceutical technologies in Ukraine. *Zaporizhzhye Medical Journal*, 27(1), 31–38. <https://doi.org/10.14739/2310-1210.2025.1.315813> [in Ukrainian].
18. Csanádi, M., Inotai, A., Oleshchuk, O., Lebega, O., Balta, A., Piniashko, O., Kaló, Z., & Kaló, Z. (2019). Health Technology Assessment Implementation in Ukraine: Current Status and Future Perspectives. *International Journal of Technology Assessment in Health Care*, 35(5), 393–400. <https://doi.org/10.1017/S0266462319000679>
19. Stoll, C. R. T., Izadi, S., Fowler, S., Green, P., Suls, J., & Colditz, G. A. (2019). The value of a second reviewer for study selection in systematic reviews. *Research Synthesis Methods*, 10(4), 539–545. <https://doi.org/10.1002/jrsm.1369>
20. Belur, J., Tompson, L., Thornton, A., & Simon, M. (2021). Interrater reliability in systematic review methodology: Exploring variation in coder decision-making. *Sociological Methods & Research*, 50(2), 837–865. <https://doi.org/10.1177/0049124118799372>
21. Tian, Y., Yang, X., Doi, S. A., et al. (2024). Towards the automatic risk of bias assessment on randomized controlled trials: A comparison of RobotReviewer and humans. *Research Synthesis Methods*, 15(6), 1111–1119. <https://doi.org/10.1002/jrsm.1761>
22. Cheng, L., Katz-Rogozhnikov, D. A., Varshney, K. R., & Baldini, I. (2021). Automated meta-analysis: A causal learning perspective. In *Proceedings of the ACM Conference on Health, Inference, and Learning*, April 08–10, 2021, New York, NY, USA, 11 p. <https://doi.org/10.48550/arXiv.2104.0463>
23. Research Screener. <https://researchscreener.com/>
24. Macia, G., Liddell, A., & Doyle, V. (2024). Conversational AI with large language models to increase the uptake of clinical guidance. *Clinical eHealth*, 7, 147–152. <https://doi.org/10.1016/j.cej.2024.12.001>
25. Fleurence, R. L., Bian, J., Wang, X., Xu, H., Dawoud, D., Higashi, M., Chhatwal, J., et al. (2025). Generative artificial intelligence for health technology assessment: Opportunities, challenges, and policy considerations: An ISPOR Working Group Report. *Value in Health*, 28(2), 175–183. <https://doi.org/10.1016/j.jval.2024.10.3846>

Список використаної літератури

1. World Health Organization. Health technology assessment. URL: https://www.who.int/health-topics/health-technology-assessment#tab=tab_1 (дата звернення 05.04.2026)
2. O'Rourke B., Oortwijn W., Schuller T. The new definition of health technology assessment: A milestone in international collaboration. *International Journal of Technology Assessment in Health Care*. 2020. Vol. 36, No 3. P. 187–190. <https://doi.org/10.1017/S0266462320000215>
3. Institutionalizing health technology assessment mechanisms: a how-to guide. Geneva : World Health Organization, 2021. URL: <https://iris.who.int/server/api/core/bitstreams/1ca70520-c2a1-4ce6-b2ee-ec5e488e8da1/content> (дата звернення 05.04.2026)
4. Gyldmark M., Lampe K., Ruof J., Pöhlmann J., Hebborn A., Kristensen, F. B. Is the EUnetHTA HTA Core Model® fit for purpose? Evaluation from an industry perspective. *International Journal of Technology Assessment in Health Care*. 2018. Vol. 34, No 5. P. 458–463. <https://doi.org/10.1017/S0266462318000594>
5. Polanin J. R., Pigott T. D., Espelage D. L., Grotperter J. K. Best practice guidelines for abstract screening in large-evidence systematic reviews and meta-analyses. *Research Synthesis Methods*. 2019. Vol. 10, No 3. P. 330–342. <https://doi.org/10.1002/jrsm.1354>
6. Cochrane Library Training Hub: User Guide. URL: <https://www.wiley.com/en-ie/solutions-partnerships/customer-success-hub/cochrane-library-training-hub> (дата звернення 05.04.2026)
7. Marshall I. J., Kuiper J., Banner E., Wallace B. C. Automating biomedical evidence synthesis: RobotReviewer. In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics: System Demonstrations*. 2017. P. 7–12. <https://doi.org/10.18653/v1/P17-4002>
8. Marshall I. J., Nye B., Kuiper J., Noel-Storr A., Marshall R., Maclean R., Soboczenski F., Nenkova A., Thomas J., Wallace B. C. Trialstreamer: A living, automatically updated database of clinical trial reports. *Journal of the American Medical Informatics Association*, 2020. Vol. 27, No 12. P. 1903–1912. <https://doi.org/10.1093/jamia/ocaa163>
9. Nye B. E., Nenkova A., Marshall I. J., Wallace, B. C. *Trialstreamer*: Mapping and Browsing Medical Evidence in Real-Time. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics: System Demonstrations*, 2020. P. 63–69. <https://doi.org/10.18653/v1/2020.acl-demos.9>
10. Chai K. E. K., Lines R. L. J., Gucciardi D. F., Ng L. Research Screener: A machine learning tool to semi-automate abstract screening for systematic reviews. *Systematic Reviews*. 2021. Vol. 10, No 1. Art. 93. <https://doi.org/10.1186/s13643-021-01635-3>
11. O'Mara-Eves A., Thomas J., McNaught J., Miwa M., Ananiadou S. Validity of using a semi-automated screening tool in a systematic review assessing non-specific effects of respiratory vaccines. *Systematic Reviews*. 2015. Vol. 4. Art. 11. <https://doi.org/10.1186/2046-4053-4-11>

12. Guo E., Gupta M., Deng J., Park Y. J., Paget M., Naugler C. Automated paper screening for clinical reviews using large language models: Data analysis study. *Journal of Medical Internet Research*. 2024. Vol. 26. Art. e48996. <https://doi.org/10.2196/48996>
13. Zemplényi A., Tachkov K., Balkanyi L., Németh, B., Petykó Z. I., Petrova G., Czech M., Dawoud D., Goettsch W., Gutierrez Ibarluzea I., Hren R., Knies S., Lorenzovici L., Maravic Z., Piniashko O., Savova A., Manova M., Tesar T., Zerovnik S., Kaló Z. Recommendations to overcome barriers to the use of artificial intelligence-driven evidence in health technology assessment. *Frontiers in Public Health*. 2023. Vol. 11. Art. 1088121. <https://doi.org/10.3389/fpubh.2023.1088121>
14. Шостакович-Корецька Л. Р., Копча В. С. Використання штучного інтелекту в клінічній медицині й наукових дослідженнях. *Медична освіта*. 2025. № 1. С. 99–107. <https://doi.org/10.11603/m.2414-5998.2025.1.15382>
15. Матвієнко М. М. Технології штучного інтелекту як складова цифрової компетентності майбутніх лікарів. *Медицина та фармація: освітні дискурси*. 2024. № 1. С. 23–29. <https://doi.org/10.32782/eddiscourses/2024-1-4>
16. Костюк І. А., Косяченко К. Л. Госпітальна оцінка медичних технологій як інструмент прийняття управлінських рішень: методологічний аналіз. *Медична наука України*. 2026. Т. 22 (1). С. 145–152. <https://doi.org/10.32345/2664-4738.1.2026.17>
17. Бабенко М. М., Косяченко К. Л. Експертне оцінювання пріоритетних напрямів впровадження медичних і фармацевтичних технологій в Україні. *Запорізький медичний журнал*. 2025. Т. 27 (1). С. 73–79. <https://doi.org/10.14739/2310-1210.2025.1.315813>
18. Csanádi M., Inotai A., Oleshchuk O., Lebeda O., Balta A., Piniashko, O., Németh B., Kaló Z. Health Technology Assessment Implementation in Ukraine: Current Status and Future Perspectives. *International Journal of Technology Assessment in Health Care*. 2019. Vol. 35, No 5. P. 393–400. <https://doi.org/10.1017/S0266462319000679>
19. Stoll C. R. T., Izadi S., Fowler S., Green P., Suls J., Colditz G. A. The value of a second reviewer for study selection in systematic reviews. *Research Synthesis Methods*. 2019. Vol. 10, No 4. P. 539–545. <https://doi.org/10.1002/jrsm.1369>
20. Belur J., Tompson L., Thornton A., Simon M. Interrater reliability in systematic review methodology: Exploring variation in coder decision-making. *Sociological Methods & Research*. 2021. Vol. 50, No 2. P. 837–865. <https://doi.org/10.1177/0049124118799372>
21. Tian Y., Yang X., Doi S. A., Marshall I. J., Wallace B. C., Beller E. M. Towards the automatic risk of bias assessment on randomized controlled trials: A comparison of RobotReviewer and humans. *Research Synthesis Methods*. 2024. Vol. 15, No 6. P. 1111–1119. <https://doi.org/10.1002/jrsm.1761>
22. Cheng L., Katz-Rogozhnikov D. A., Varshney K. R., Baldini I. Automated meta-analysis: A causal learning perspective. *Proceedings of the ACM Conference on Health, Inference, and Learning*. 2021. <https://doi.org/10.48550/arXiv.2104.0463>
23. Research Screener. *Research Screener*. URL: <https://researchscreener.com/> (дата звернення 05.04.2026)
24. Macia G., Liddell A., Doyle V. Conversational AI with large language models to increase the uptake of clinical guidance. *Clinical eHealth*. 2024. Vol. 7. P. 147–152. <https://doi.org/10.1016/j.ceh.2024.12.001>
25. Fleurence R. L., Bian J., Wang X., Xu H., Dawoud D., Higashi M., Chhatwal J., et al. ISPOR Working Group on Generative AI. Generative artificial intelligence for health technology assessment: opportunities, challenges, and policy considerations: an ISPOR Working Group Report. *Value in Health*. 2025. Vol. 28, No 2. P. 175–183. <https://doi.org/10.1016/j.jval.2024.10.3846>

Conflict of interest: The author declares that there are no potential or apparent conflicts of interest related to the manuscript. This article didn't receive financial support from a government, non-governmental or commercial organization.

Author Contributions:

Kostiuk I. A. – conceptualization, methodology, software, formal analysis, resources, writing, project administration.

Надійшла до редакції 27 травня 2026 р.

Прийнято до друку 31 липня 2026 р.

Опубліковано 28 серпня 2026 р.

Електронна адреса для листування з авторами: kostiuk.iryana@nmu.ua

(Костюк І. А.)