

Competence before Confidence in Pharmacist Use of Generative Artificial Intelligence

Mikko Lahtinen^{1*}, Elina Salo¹, Juhani Virtanen²

¹Department of Pharmacy Education and Generative AI, Faculty of Pharmacy, University of Helsinki, Helsinki, Finland. ²Department of AI Competence and Professional Development, Faculty of Pharmacy, University of Eastern Finland, Kuopio, Finland.

Abstract

Generative artificial intelligence is entering pharmacy through medication-information searches, evidence synthesis, documentation, patient communication, clinical reasoning support, safety surveillance, education, and administrative work. Its fluent output can encourage confidence before users have demonstrated the ability to define an appropriate task, verify evidence, recognize errors, communicate uncertainty, and retain professional accountability. This article proposes an original Pharmacist Generative-AI Competency Framework that treats competence as task-specific, observable, context-dependent, and continuously requalifiable. The framework connects authorized pharmacy tasks with nine domains: task framing and authorization; model and data literacy; evidence retrieval and provenance; verification and error recognition; medication safety and uncertainty; human-AI and patient communication; workflow accountability; equity and representation; and monitoring and requalification. Four proposed levels extend from supervised constrained use to critical governance, with advancement controlled by case-based assessment gates rather than self-reported confidence or course completion. The framework distinguishes model performance, pharmacist task performance, pharmacist-AI team performance, workflow consequence, and patient-relevant outcome. It also places individual competence inside organizational conditions such as evidence access, privacy protection, supervision, auditability, incident response, and change control. Validation would require reliable scoring, external assessment across practice settings, medication-safety testing, human-factors evaluation, equity analysis, and longitudinal evidence after model or workflow change. The framework is non-empirical and does not establish clinical benefit, authorization, regulatory acceptance, or deployment readiness. Its contribution is a testable architecture for defining what pharmacists should be able to demonstrate before confidence in generative-AI use is treated as professionally meaningful.

Keywords: Digital pharmacy, Artificial intelligence, Pharmacy practice, Medication safety, Human-AI collaboration, Clinical decision support

INTRODUCTION

Generative artificial intelligence can produce plausible, rapid, and conversational responses that feel usable before their evidential adequacy has been established. Confidence may describe a pharmacist's subjective comfort, willingness to engage, or expectation of success; competence concerns whether the pharmacist can perform a defined task safely and accountably under specified conditions. These constructs may correlate, but they are not interchangeable. Research on resistance to medical AI also shows that acceptance is shaped by perceptions of personal uniqueness and the nature of machine involvement, indicating that willingness or reluctance may reflect psychological response rather than demonstrated capability [1].

The opposite problem is equally important. Familiarity and trust can increase reliance even when an AI-supported recommendation is wrong. Experimental evidence shows that clinicians may be susceptible to incorrect decision-aid advice, making calibrated reliance and the ability to override output central to competent use [2]. For pharmacy, this distinction is consequential because an apparently minor failure in

medication information, dose context, contraindication recognition, or risk communication may be propagated into a professional decision. Competence must therefore include justified non-use, escalation, and independent checking, not merely effective prompting.

Explanations generated by or attached to AI systems do not resolve this problem automatically. Explainable-AI methods

Address for correspondence: Mikko Lahtinen, Department of Pharmacy Education and Generative AI, Faculty of Pharmacy, University of Helsinki, Helsinki, Finland.
mikko.lahtinen@helsinki.fi

Received: 20 May 2025; **Accepted:** 12 September 2025

This is an open-access article distributed under the terms of the Creative Commons Attribution-Non Commercial-Share Alike 4.0 License, which allows others to remix, tweak, and build upon the work non commercially, as long as the author is credited and the new creations are licensed under the identical terms.

How to cite this article: Lahtinen M, Salo E, Virtanen J. Competence before Confidence in Pharmacist Use of Generative Artificial Intelligence. Arch Pharm Pract. 2025;16(3):96-93.
<https://doi.org/xxxxxxxxxxxxxxxxxx>

can create an appearance of transparency without establishing that the underlying output is reliable, clinically useful, or causally informative [3]. A pharmacist may understand why a system presents a statement yet still lack evidence that the statement is correct, current, applicable to the patient, or appropriate for the intended decision. The proposed framework consequently treats explanation as one possible input to appraisal rather than proof of trustworthiness.

Clinical AI can also produce unintended consequences that arise from interaction among technology, workflow, incentives, professional roles, and local context rather than from model performance alone [4]. This makes competence partly relational and organizational. A pharmacist who verifies an output accurately may still work within a system that obscures provenance, prevents documentation, encourages automation bias, or assigns responsibility ambiguously. This article therefore proposes a non-empirical framework in which confidence becomes professionally meaningful only after task-bounded competence has been demonstrated and supported by suitable governance. The framework does not claim that competence eliminates error or guarantees benefit; it specifies constructs, levels, assessment gates, and boundary conditions that can be tested.

Generative-AI Tasks Encountered in Pharmacy

A pharmacy task is defined here by its purpose, information inputs, output type, medication-decision proximity, potential harm, authorized user, and required verification. Generative-AI use is therefore not a single activity. It may involve drafting or summarizing information, translating technical language, retrieving or organizing evidence, proposing explanations, producing documentation, generating educational cases, or supporting preliminary clinical reasoning. Comparative clinical-pharmacy testing shows that model performance varies across questions and should not be inferred from linguistic fluency [5]. Task definition must precede competency judgment because a pharmacist may perform adequately in one bounded use while remaining unprepared for another.

Medication-information work illustrates this variability. A generated answer may be factually plausible yet incomplete,

insufficiently contextualized, or inappropriate for the specific question. Evaluations of medication-related responses support the need to check accuracy, completeness, clinical relevance, and patient-specific applicability rather than treating a coherent answer as a finished professional product [6]. In the proposed taxonomy, medication information and evidence synthesis therefore require explicit provenance, source authentication, recency checking, and separation of retrieved evidence from generated interpretation.

Drug-information use also exposes fabricated or mismatched citations, unsupported claims, omissions, and inconsistent risk communication. Exploratory pharmacy evidence identifies these as practical hazards rather than merely stylistic defects [7]. A pharmacist using generative AI for documentation or patient communication must consequently verify both content and communicative consequence. A grammatically polished discharge explanation, counselling note, or interprofessional message can still be unsafe if it suppresses uncertainty, changes the evidential meaning, or implies professional endorsement that has not been earned.

Medication-safety applications include detection, prediction, triage, prioritization, and communication across the medication-use process, but technical evaluation remains more common than clinical evaluation [8]. Generative systems may assist with identifying potential interactions, organizing adverse-event narratives, or drafting safety communications, yet these uses differ in consequence and evidential demand. The article therefore proposes six task families: medication information and evidence synthesis; documentation and communication; clinical reasoning support; medication-safety detection and triage; education and simulation; and administrative workflow support. These categories organize assessment but do not establish that a particular system is safe or effective.

Table 1 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for definitions, components, relationships, and intended uses for the article Competence before Confidence in Pharmacist Use of Generative Artificial Intelligence.

Table 1. Definitions, components, relationships, and intended uses for the article Competence before Confidence in Pharmacist Use of Generative Artificial Intelligence.

Component or scientific dimension	Problem addressed	Inputs or determinants	Proposed mechanism or relationship	Expected contribution	Evidence required	Failure or uncertainty risk	Boundary statement
Medication information and evidence synthesis	Fluent but incomplete or unsupported answers	Question scope, patient context, source access, evidence recency	Proposed: generated synthesis must pass provenance and applicability checks before use	Faster organization of evidence without transferring evidential authority	Source authenticity, recency, claim-to-source fit, and expert review [5-7]	Fabricated citation, omitted contraindication, or outdated evidence	Output is not an evidence source or medication decision
Documentation and communication	Polished language may distort meaning	Intended recipient, health literacy, clinical record,	Proposed: content verification precedes adaptation of tone or format	Drafting support with retained pharmacist accountability	Fidelity to verified facts, disclosure rules, and recipient comprehension [7]	Hallucinated detail, misleading certainty, or	Linguistic quality does not establish clinical correctness

	or conceal uncertainty	communication purpose				responsibility diffusion	
Clinical reasoning support	Suggestions may invite automation bias	Case completeness, differential possibilities, user expertise, risk level	Proposed: AI suggestions are hypotheses to be challenged, not recommendations to be accepted	Structured consideration of alternatives	Independent reasoning, contradiction search, and escalation behavior [2, 5]	Premature closure, anchoring, or missed context	Model output does not possess professional authority
Medication- safety detection and triage	Technical signals may not translate into safe action	Medication list, timing, dose, comorbidity, workflow, alert burden	Proposed: signals require pharmacist adjudication and prioritization	Earlier attention to possible harm	Sensitivity to material hazards, false-positive burden, and workflow evaluation [8]	Missed harm, alert fatigue, or inappropriate prioritization	Detection is not causation, diagnosis, or validated benefit
Education and simulation	Correct answers may be confused with workplace competence	Learning objective, learner level, case realism, reference standard	Proposed: generative cases should test verification, uncertainty, and escalation	Scalable practice opportunities	Content validity, scoring reliability, and faculty review	Artificial cases, answer leakage, or overconfidence	Educational performance does not authorize clinical use
Administrative workflow support	Low-risk appearance may obscure privacy or propagation risks	Data sensitivity, authorization, downstream systems, audit trail	Proposed: task authorization and minimum-necessary data govern use	Reduced clerical burden where locally justified	Privacy review, process accuracy, and incident monitoring [4]	Data leakage, copied error, or undocumented system substitution	Administrative convenience does not override governance

The task taxonomy is intentionally functional rather than product-specific. Model names, interfaces, and vendor claims may change, whereas the professional questions remain stable: What task is being performed, what evidence is required, what can go wrong, who verifies the output, and who remains accountable? Competence should be assessed at this level of specificity.

Proposed Competency-Domain Framework

Clinician AI-competency literature indicates that responsible use requires more than technical familiarity. Relevant capabilities include foundational knowledge, critical appraisal, workflow integration, communication, and implementation awareness [9]. For pharmacists, these capabilities must be connected to medication evidence, patient-specific risk, professional scope, and the consequences of incorrect or unjustifiably certain output. The framework therefore defines competence as the demonstrated ability to complete an authorized generative-AI-supported pharmacy task, under specified conditions, while preserving evidence standards, medication safety, accountability, equity, and appropriate uncertainty.

Clinical competency models also provide precedent for separating foundational knowledge, appraisal, decision use, technical interaction, communication, and awareness of unintended consequences [10]. The proposed framework adapts this logic but does not assume that domains discussed in another profession are automatically sufficient for pharmacy. It adds explicit attention to evidence provenance, medication-safety escalation, task authorization, and continuing requalification because generative output can change with model version, prompt context, retrieval source, and workflow integration.

International digital-health competency work supports specifying observable outcomes and expectations appropriate to role and stage [11]. Accordingly, each proposed domain is expressed through actions that can be observed: defining the task and prohibited uses; identifying model and data limits; locating and authenticating evidence; detecting and correcting errors; communicating uncertainty; documenting accountability; recognizing inequitable performance; and responding to model or workflow change. Competence is thus not possession of knowledge alone but performance under conditions that expose foreseeable failure.

Reviews of AI education programs show heterogeneous content and limited evaluation, reinforcing the need to connect educational exposure to observable assessment [12]. Existing competency and education frameworks converge on multidomain, observable, role-specific competence rather than product familiarity [9, 11, 12]. On that basis, this article proposes nine interdependent domains: task framing and authorization; model and data literacy; evidence retrieval and provenance; verification and error recognition; medication safety and uncertainty; human-AI and patient communication; workflow accountability; equity and representation; and monitoring and requalification. No domain is assigned a universal numerical weight. Their relative importance is conditional on task, patient risk, setting, and authorized scope.

Q1-Journal-Publication-Ready Figure Prompt

Create a publication-grade conceptual figure titled “Pharmacist Generative-AI Competency Framework” for the article “Competence before Confidence in Pharmacist Use of Generative Artificial Intelligence.” Build the composition around one unmistakable central visual anchor representing the article’s core proposed architecture, framework, model,

theory, review synthesis, or decision pathway. Organize the surrounding modules according to these manuscript sections: Introduction: Confidence Is Not Evidence of Competence; Generative-AI Tasks Encountered in Pharmacy; Proposed Competency-Domain Framework; Progressive Levels from Supervised Use to Critical Governance; Assessment through Cases, Verification, and Error Recognition; Organizational Support and Continuing Requalification; Educational and Professional Implications; Limitations. Show directional relationships with solid arrows only where the selected evidence supports the connection and dashed arrows where relationships are proposed, uncertain, context-dependent, or require validation. Include explicit boundary markers separating model output from medication decision, technical performance from clinical usefulness, association from causation, automation from professional authority, and implementation from validated benefit where relevant. Use abstract scientific icons and labelled modules only. Do not use corporate logos, journal logos, software interfaces, copied scientific figures, copyrighted imagery, stock photography, approval stamps, or certification marks. Do not invent numerical values, effect sizes, performance scores, clinical outcomes, regulatory thresholds, or readiness scores. Use a clean white background, precise geometric alignment, professional sans-serif typography, thin uniform borders,

strong visual hierarchy, a restrained colorblind-safe palette, and sufficient label size for journal-column reduction and 300 dpi export.

Exact Placement

Section 5, “Proposed Competency-Domain Framework,” immediately after the preceding substantive in-text callout. The framework’s central boundary separates model output from the pharmacist’s medication decision. Generative output may organize information or propose language, but it does not inherit professional authority. Verification connects evidence provenance to medication-safety judgment; communication connects technical appraisal to patient and team use; workflow accountability assigns final responsibility; equity cross-cuts task framing, validation, and monitoring; and requalification closes the lifecycle when material change occurs. These relationships are proposed and require empirical testing.

Table 2 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for evaluation criteria, evidence requirements, and failure modes for the article Competence before Confidence in Pharmacist Use of Generative Artificial Intelligence.

Table 2. Evaluation criteria, evidence requirements, and failure modes for the article Competence before Confidence in Pharmacist Use of Generative Artificial Intelligence.

Architecture layer or process stage	Core function	Information flow	Interaction with other components	Evaluation criterion	Potential failure mode	Human or organizational responsibility	Readiness boundary
Task authorization	Define purpose, user, risk, allowed data, and prohibited uses	Practice need → bounded task statement	Constrains all later domains	Task and decision proximity are explicit	Unbounded or unauthorized use	Organization defines scope; pharmacist confirms fit [4, 9]	No use outside authorized scope
Model and data appraisal	Identify capability, version, limits, and data dependence	System information → user appraisal	Informs task selection and uncertainty	User can state relevant limitations	Product familiarity mistaken for competence	Organization provides documentation; pharmacist appraises it [9, 10]	Appraisal does not validate performance
Evidence provenance	Authenticate claims and medication sources	Output claim → source → patient/context fit	Feeds verification and safety judgment	Claims are traceable and applicable	Fabricated, stale, or mismatched evidence	Pharmacist performs or confirms independent checking [5-7]	Generated text is not primary evidence
Verification and error recognition	Detect factual, reasoning, omission, and citation errors	Output → reference standard → correction, rejection, or escalation	Controls downstream communication and decision use	Material errors are recognized and prevented from propagation	Fluent unsafe output accepted	Pharmacist retains final responsibility [2, 5, 7]	Passing one case does not establish general competence
Medication safety and uncertainty	Identify harm, express limits, and escalate	Patient context plus evidence → risk judgment	Can stop progression or require fallback	Safe non-use and escalation occur when needed	False certainty, missed contraindication, or unsafe prioritization	Pharmacist and organization maintain escalation pathways [6, 8]	No clinical-readiness inference without contextual evaluation
Communication and accountability	Disclose role, limitations, and final decision ownership	Verified content → patient, record, or team	Converts appraisal into professional action	Meaning, uncertainty, and accountability remain clear	Responsibility diffusion or misleading certainty	Named pharmacist owns final use; organization supports audit [1, 9, 11]	Explanation does not transfer authority

Equity and representation	Detect unequal applicability or access	Population and context information → subgroup scrutiny	Cross-cuts task, assessment, and monitoring	Relevant language, access, and representation risks are examined	Unrecognized disparity or inaccessible output	Shared professional and organizational duty [11]	Absence of observed disparity does not prove equity
Monitoring and requalification	Reassess after material change	Incidents, updates, or drift signals → review	Returns use to supervision or restricted scope	Change triggers documented reassessment	Stale authorization or unrecognized drift	Organization operates change control; pharmacist reports failures [4, 12]	Prior competence is version- and context-bounded

All architecture layers, interactions, and readiness boundaries in **Table 2** are proposed; the citations identify evidence informing their construction. The table specifies evaluation logic rather than a validated scoring instrument. Transition thresholds, weighting, assessor preparation, and acceptable evidence must be developed and tested for defined pharmacy tasks and settings.

Progressive Levels from Supervised Use to Critical Governance

Health-workforce AI education should be differentiated by professional role, responsibility, and stage of development rather than delivered as a uniform body of content [13]. The framework therefore proposes levels of entrusted activity rather than a single label of “AI competent.” A pharmacist may be entrusted at a higher level for administrative drafting but remain at supervised use for patient-specific medication reasoning. Progression is task-specific, reversible, and conditional on local authorization.

Entrustment offers a useful basis for deciding when supervision can be reduced. An AI-related entrustment framework in health professions education emphasizes that increasing autonomy should follow demonstrated performance and risk-sensitive judgment rather than confidence or exposure alone [14]. In the proposed Level 1, the pharmacist undertakes supervised, constrained use with pre-approved tools, non-sensitive inputs, and direct review. At Level 2, the pharmacist contributes verified output to a defined task but must document sources, corrections, uncertainty, and escalation.

Human–AI collaboration research shows that team performance can vary with user expertise and interaction, meaning that neither AI-alone nor clinician-alone performance predicts the combined system reliably [15]. Proposed Level 3 therefore concerns independent critical use only within an authorized scope. It requires the pharmacist to challenge output, identify situations in which the model degrades judgment, communicate limitations, and use fallback processes. Independence refers to absence of concurrent supervision for the specified task, not freedom from audit, governance, or professional standards.

Early clinical evaluation guidance emphasizes usability, human factors, workflow effects, and safety in the intended setting rather than accuracy alone [16]. Proposed Level 4 extends competence to critical governance: evaluating task

suitability, supervising others, defining escalation pathways, reviewing incidents, participating in change control, and determining when requalification is necessary. It does not confer regulator-like authority or imply that the pharmacist can certify a system as clinically beneficial.

Safe progression requires demonstrated task performance, risk-sensitive entrustment, and staged evaluation of human–AI use in context [14–16]. Movement between levels should therefore be controlled by assessment gates covering task framing, evidence verification, seeded-error recognition, uncertainty, safe non-use, escalation, communication, accountability, and equity. Failure at a material gate should restrict scope, return the user to supervision, or suspend the use case until remediation and reassessment. The proposed levels are an organizing hypothesis; their content validity, reliability, feasibility, and relationship to medication-safety outcomes remain to be established.

Assessment through Cases, Verification, and Error Recognition

Competence should be demonstrated through pharmacy cases that make the AI intervention, pharmacist actions, and downstream use visible. Reporting standards for AI interventions emphasize describing system inputs, outputs, human interaction, and error handling rather than reducing evaluation to one aggregate score [17]. A competency assessment should therefore identify the task, model version, information available to the pharmacist, expected reference sources, authorized scope, and consequences of accepting, correcting, rejecting, or escalating an output.

Assessment conditions should be specified before performance is judged. Protocol guidance for AI interventions supports prospective definition of inputs, outputs, fallback procedures, and human involvement [18]. Proposed cases should include routine requests, incomplete information, conflicting evidence, fabricated citations, clinically important omissions, inappropriate certainty, privacy risks, and situations requiring safe non-use. Observable actions should include source authentication, patient-context reconstruction, contradiction checking, correction, communication of uncertainty, documentation, and escalation.

Verification must also be traceable. Updated reporting guidance for AI-enabled prediction models emphasizes defined data, outcomes, validation methods, and intended-use

conditions [19]. Applied here, competence evidence should identify the reference standard, assessor qualifications, scoring logic, inter-rater reliability, model and prompt conditions, and limits of applicability. Passing should support only the assessed task and supervision level. Content validity, reliability, external validity, equity, and association with safer pharmacy work require separate empirical study.

Organizational Support and Continuing Requalification

Competence is not produced by the pharmacist alone. Implementation depends on the characteristics of the intervention, the individuals involved, the inner and outer setting, and the implementation process [20]. Organizations must therefore define authorized tasks and tools, provide access to appropriate evidence, protect confidential information, establish escalation routes, preserve audit trails, and prevent productivity expectations from overriding verification.

Responsible health-AI use also requires lifecycle monitoring, stakeholder participation, accountability, and response to changing performance [21]. A competence judgment should consequently be version-specific and time-bounded. Material changes in the model, retrieval source, interface, workflow, evidence base, policy, patient population, or task should trigger review. Incidents, near misses, repeated correction patterns, and new inequities may justify restricted scope, renewed supervision, or reassessment.

Clinical impact cannot be inferred from technical accuracy because generalizability, workflow integration, behavior, regulation, and post-deployment performance may alter outcomes [22]. Sustained competence therefore depends on implementation context, lifecycle monitoring, and evidence of clinical integration [20-22]. Training cannot compensate for an unsafe configuration, unavailable evidence, ambiguous responsibility, or absent change control. Continuing requalification is proposed as a shared professional and organizational process, not as a permanent credential.

Educational and Professional Implications

Pharmacy education on AI remains heterogeneous, supporting a shift from exposure-based teaching toward demonstrable, pharmacy-specific competence [23]. Pre-licensure curricula could introduce task framing, evidence provenance, error recognition, uncertainty, communication, and accountability before learners undertake patient-proximal uses. Experiential assessment should use locally relevant medication cases rather than rely solely on examinations of AI terminology or prompt construction.

Implementation in education also requires capable faculty, ethical framing, suitable assessment, and technical infrastructure [24]. Faculty development should include calibration of reference standards, recognition of model-specific failure, assessment of safe non-use, and management

of learner dependence on generated answers. Institutions should disclose authorized educational uses and distinguish assistance with learning from substitution for independent professional reasoning.

Professional formation should emphasize risk, responsibility, and reflection rather than product-specific fluency [25]. Continuing professional development could apply the same framework to new roles, tools, and practice settings, while employers could use it to define supervision and requalification needs. These applications remain proposed. The framework is not an accreditation requirement, licensing standard, or universal curriculum, and educational adoption should be piloted for feasibility, reliability, equity, and unintended effects.

Limitations

The framework is a non-empirical synthesis, and its domains, levels, gates, and requalification triggers have not been validated. Proxy choices can conceal inequity even when technical performance appears strong [26]. Comparisons between clinicians and AI frequently contain design, reporting, and claim limitations, restricting transfer from non-pharmacy tasks [27]. Missingness, data-generation processes, and underrepresentation can further limit fairness and transportability [28]. The framework may require modification across countries, pharmacy roles, languages, technologies, and risk environments. It does not establish causal benefit, patient safety, regulatory acceptability, or readiness for deployment.

CONCLUSION

Competence before confidence reframes pharmacist use of generative artificial intelligence as an observable, task-specific, and continuously governed professional capability. The proposed framework connects authorized pharmacy tasks with evidence provenance, verification, medication-safety judgment, communication, accountability, equity, and requalification. Its progressive levels require demonstrated performance and calibrated reliance before supervision is reduced, while its assessment architecture makes error recognition, uncertainty, escalation, and safe non-use central rather than peripheral.

The framework also establishes boundaries that confidence alone cannot cross. Model output is not a medication decision; technical performance is not clinical usefulness; automation does not transfer professional authority; and implementation does not establish benefit. Individual capability remains conditional on organizational evidence access, privacy protection, workflow design, incident response, and change control.

The framework offers a structured basis for curriculum design, workplace assessment, governance, and empirical research. Its value depends on testing content validity, scoring reliability, feasibility, equity, external applicability,

human-AI team performance, and relationships with medication-safety and workflow outcomes. Until such evidence exists, it should be interpreted as a proposed scholarly architecture rather than a validated standard, clinical recommendation, or deployment-ready system.

ACKNOWLEDGMENTS: None

CONFLICT OF INTEREST: None

FINANCIAL SUPPORT: None

ETHICS STATEMENT: None

REFERENCES

- Longoni C, Bonezzi A, Morewedge CK. Resistance to medical artificial intelligence. *J Consum Res*. 2019;46(4):629-50. doi:10.1093/jcr/ucz013
- Gaube S, Suresh H, Raue M, Merritt A, Berkowitz SJ, Lerner E, et al. Do as AI say: Susceptibility in deployment of clinical decision-aids. *NPJ Digit Med*. 2021;4(1):31. doi:10.1038/s41746-021-00385-9
- Ghassemi M, Oakden-Rayner L, Beam AL. The false hope of current approaches to explainable artificial intelligence in health care. *Lancet Digit Health*. 2021;3(11). doi:10.1016/S2589-7500(21)00208-9
- Cabitz F, Rasoini R, Gensini GF. Unintended consequences of machine learning in medicine. *JAMA*. 2017;318(6):517-8. doi:10.1001/jama.2017.7797
- Huang X, Estau D, Liu X, Yu Y, Qin J, Li Z, et al. Evaluating the performance of ChatGPT in clinical pharmacy: A comparative study of ChatGPT and clinical pharmacists. *Br J Clin Pharmacol*. 2024;90(1):232-8. doi:10.1111/bcp.15896
- Grossman S, Nathan JP. Appropriateness of ChatGPT as a resource for medication-related questions. *Br J Clin Pharmacol*. 2024;90(10):2691-5. doi:10.1111/bcp.16212
- Morath B, Chiriac U, Jaszowski E, Deiß C, Nürnberg H, Hörth K, et al. Performance and risks of ChatGPT used in drug information: An exploratory real-world analysis. *Eur J Hosp Pharm*. 2024;31(6):491-7. doi:10.1136/ejpharm-2023-003750
- Syrowatka A, Song W, Amato MG, Foer D, Edrees H, Co Z, et al. Key use cases for artificial intelligence to reduce the frequency of adverse drug events: A scoping review. *Lancet Digit Health*. 2022;4(2). doi:10.1016/S2589-7500(21)00229-6
- Garvey KV, Thomas Craig KJ, Russell R, Novak LL, Moore D, Miller BM. Considering clinician competencies for the implementation of artificial intelligence-based tools in health care: Findings from a scoping review. *JMIR Med Inform*. 2022;10(11). doi:10.2196/37478
- Liaw W, Kueper JK, Lin S, Bazemore A, Kakadiaris I. Competencies for the use of artificial intelligence in primary care. *Ann Fam Med*. 2022;20(6):559-63. doi:10.1370/afm.2887
- Car J, Ong QC, Erlikh Fox T, Leightley D, Kemp SJ, Švab I, et al. The Digital Health Competencies in Medical Education Framework: An international consensus statement based on a Delphi study. *JAMA Netw Open*. 2025;8(1). doi:10.1001/jamanetworkopen.2024.53131
- Charow R, Jeyakumar T, Younus S, Dolatabadi E, Salhia M, Al-Mouaswas D, et al. Artificial intelligence education programs for health care professionals: Scoping review. *JMIR Med Educ*. 2021;7(4). doi:10.2196/31043
- Gray K, Slavotinek J, Dimaguila GL, Choo D. Artificial intelligence education for the health workforce: Expert survey of approaches and needs. *JMIR Med Educ*. 2022;8(2). doi:10.2196/35223
- Gin BC, O'Sullivan PS, Hauer KE, Abdunour RE, Mackenzie M, ten Cate O, et al. Entrustment and EPAs for artificial intelligence (AI): A framework to safeguard the use of AI in health professions education. *Acad Med*. 2025;100(3):264-72. doi:10.1097/ACM.0000000000005930
- Tschandl P, Rinner C, Apalla Z, Argenziano G, Codella N, Halpern A, et al. Human-computer collaboration for skin cancer recognition. *Nat Med*. 2020;26(8):1229-34. doi:10.1038/s41591-020-0942-0
- Vasey B, Nagendran M, Campbell B, Clifton DA, Collins GS, Denaxas S, et al. Reporting guideline for the early-stage clinical evaluation of decision support systems driven by artificial intelligence: DECIDE-AI. *Nat Med*. 2022;28(5):924-33. doi:10.1038/s41591-022-01772-9
- Liu X, Cruz Rivera S, Moher D, Calvert MJ, Denniston AK, Ashrafian H, et al. Reporting guidelines for clinical trial reports for interventions involving artificial intelligence: The CONSORT-AI extension. *Nat Med*. 2020;26(9):1364-74. doi:10.1038/s41591-020-1034-x
- Cruz Rivera S, Liu X, Chan AW, Denniston AK, Calvert MJ, Ashrafian H, et al. Guidelines for clinical trial protocols for interventions involving artificial intelligence: The SPIRIT-AI extension. *Nat Med*. 2020;26(9):1351-63. doi:10.1038/s41591-020-1037-7
- Collins GS, Moons KGM, Dhiman P, Riley RD, Beam AL, Van Calster B, et al. TRIPOD+AI statement: Updated guidance for reporting clinical prediction models that use regression or machine learning methods. *BMJ*. 2024;385. doi:10.1136/bmj-2023-078378
- Damschroder LJ, Reardon CM, Widerquist MAO, Lowery J. The updated Consolidated Framework for Implementation Research based on user feedback. *Implement Sci*. 2022;17(1):75. doi:10.1186/s13012-022-01245-0
- Wiens J, Saria S, Sendak M, Ghassemi M, Liu VX, Doshi-Velez F, et al. Do no harm: A roadmap for responsible machine learning for health care. *Nat Med*. 2019;25(9):1337-40. doi:10.1038/s41591-019-0548-6
- Kelly CJ, Karthikesalingam A, Suleyman M, Corrado G, King D. Key challenges for delivering clinical impact with artificial intelligence. *BMC Med*. 2019;17(1):195. doi:10.1186/s12916-019-1426-2
- Abdel Aziz MH, Rowe C, Southwood R, Nogid A, Berman S, Gustafson K. A scoping review of artificial intelligence within pharmacy education. *Am J Pharm Educ*. 2024;88(1):100615. doi:10.1016/j.ajpe.2023.100615
- Chan KS, Zary N. Applications and challenges of implementing artificial intelligence in medical education: Integrative review. *JMIR Med Educ*. 2019;5(1). doi:10.2196/13930
- Sawyer K, Dave VS, Gonyeau MJ, Cain J. Teaching tomorrow's pharmacists in an AI world: Risk, responsibility, and reflection. *Am J Pharm Educ*. 2025;89(12):101905. doi:10.1016/j.ajpe.2025.101905
- Obermeyer Z, Powers B, Vogeli C, Mullainathan S. Dissecting racial bias in an algorithm used to manage the health of populations. *Science*. 2019;366(6464):447-53. doi:10.1126/science.aax2342
- Nagendran M, Chen Y, Lovejoy CA, Gordon AC, Komorowski M, Harvey H, et al. Artificial intelligence versus clinicians: Systematic review of design, reporting standards, and claims of deep learning studies. *BMJ*. 2020;368. doi:10.1136/bmj.m689
- Gianfrancesco MA, Tamang S, Yazdany J, Schmajuk G. Potential biases in machine learning algorithms using electronic health record data. *JAMA Intern Med*. 2018;178(11):1544-7. doi:10.1001/jamainternmed.2018.3763