

Dividing Responsibility When a Pharmacist and an Algorithm Co-Produce the Medication Decision

Wei Zhang^{1*}, Chen Hui², Michael Tan¹

¹Department of Shared Decision-Making in Pharmacy, Faculty of Pharmacy, University of Melbourne, Melbourne, Australia. ²Department of Pharmacist-Algorithm Collaboration, Faculty of Pharmacy, National University of Singapore, Singapore.

Abstract

Medication decisions supported by artificial intelligence are often described as if responsibility remains wholly with the pharmacist because a human formally approves the final action. This view overlooks how data selection, model design, interface presentation, organizational configuration, professional judgment, and follow-up jointly produce the decision. This article develops a proposed Responsibility-Allocation Model for co-produced medication decisions. The model treats each decision as a responsibility-bearing episode extending across data, recommendation, approval, and follow-up. At every stage, four dimensions are examined separately: contribution, meaning the input supplied to the decision; control, meaning the feasible capacity to inspect or alter the process; authority, meaning legitimate power to approve, reject, or escalate; and liability, meaning potential legal or institutional answerability. The model maps these dimensions across the pharmacist, algorithm developer or vendor, and deploying organization. It also introduces proposed rules for non-equivalence among responsibility dimensions, effective professional control, authority–control alignment, documented disagreement, and reassessment after material change. Evaluation requires evidence of technical validity, joint human–algorithm task performance, workflow effects, medication safety, equity, traceability, organizational response, and longitudinal stability. The model does not determine negligence, allocate legal liability, establish clinical benefit, or authorize deployment. Its original contribution is a structured vocabulary and testable governance architecture for examining responsibility without reducing it either to algorithmic influence or to the pharmacist's final click, enabling later empirical testing across pharmacy tasks and organizational contexts.

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

INTRODUCTION

A medication decision may appear to culminate in one professional act, yet the conditions shaping that act are distributed. Clinical machine-learning systems can introduce unintended consequences after entering practice because their effects depend not only on model construction but also on interpretation, workflow, and institutional use [1]. The relevant unit is therefore not the isolated algorithmic output or the pharmacist's final response. It is the whole decision-production process through which data are converted into a recommendation, presented within a work system, accepted or rejected, implemented, and later reviewed.

Clinical decision support illustrates this distribution. Its benefits are conditional on integration with clinical work, usable delivery, risk management, and organizational support [2]. An alert or prediction may narrow attention, rank prescriptions, suppress alternatives, frame uncertainty, or change the order in which cases are reviewed. These actions can materially influence the decision even when the system lacks formal authority. Conversely, a pharmacist may possess

Address for correspondence: Wei Zhang, Department of Shared Decision-Making in Pharmacy, Faculty of Pharmacy, University of Melbourne, Melbourne, Australia.
wei.zhang@unimelb.edu.au

Received: 03 May 2025; **Accepted:** 21 August 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: Zhang W, Hui C, Tan M. Dividing Responsibility When a Pharmacist and an Algorithm Co-Produce the Medication Decision. Arch Pharm Pract. 2025;16(3):46-53.
<https://doi.org/10.51847/bgArgCNMwK>

formal approval authority while having limited time, incomplete information, or no practicable means of interrogating the model.

Pharmacy-specific evidence makes the problem concrete. A hybrid system can use machine-learning risk scores to rule prescriptions into or out of pharmacist medication review, thereby distributing the work between algorithmic triage and professional assessment [3]. The resulting decision is not produced by either component independently: the model shapes which prescriptions receive attention, while the pharmacist interprets the selected cases and determines action. This does not show that responsibility should be divided equally, but it does show that the final professional act cannot describe the full causal and organizational pathway.

Artificial intelligence has also been used to optimize medication alerts, although the evidence remains heterogeneous across tasks, settings, outcomes, and evaluation designs [4]. Improved alert targeting would not by itself establish safer medication decisions, appropriate professional reliance, or adequate governance. The central problem is therefore one of distributed decision production: several actors and technical components contribute differently, exercise different degrees of control, hold different forms of authority, and may face different kinds of answerability. This article proposes a model for separating those dimensions and tracing them across the medication-decision lifecycle. It is a conceptual architecture requiring empirical, organizational, and jurisdiction-specific validation.

Distinguishing Contribution, Control, Authority, and Liability

Responsibility is often used as an undifferentiated term, obscuring whether the speaker means causal involvement, practical control, professional duty, moral accountability, organizational ownership, or legal liability. Ethical analysis of medical machine learning emphasizes transparency, trust, data protection, and identifiable accountable actors [5]. For co-produced medication decisions, these concerns require a vocabulary that does not convert every form of participation into the same kind of responsibility.

Contribution is the informational, computational, clinical, or organizational input supplied to the decision. A dataset contributes through recorded variables and omissions; a

model contributes through prediction or ranking; an interface contributes through timing and presentation; a pharmacist contributes through interpretation and contextual judgment; and an organization contributes through configuration, staffing, and escalation arrangements. Contribution may be substantial without conferring legitimate authority or feasible control.

Control is the practical capacity to inspect, alter, stop, verify, configure, override, escalate, or monitor the process. Authority is different: it is the legitimate power, arising from professional scope or organizational rules, to approve, reject, implement, restrict, or escalate an action. Liability is different again. Legal analysis indicates that potential liability for use of medical artificial intelligence depends on factors such as professional conduct, product behavior, applicable standards, causation, and jurisdiction rather than on algorithmic involvement alone [6]. The present model therefore treats liability as a separate, context-dependent inquiry rather than the automatic endpoint of contribution.

Meaningful human control cannot be inferred merely because a pharmacist remains “in the loop.” Ethical analysis of algorithmic decision-making distinguishes genuine professional agency from nominal oversight [7]. A pharmacist who sees an output but lacks relevant data, understandable uncertainty, sufficient review time, or an effective path to override may possess formal authority without effective control. Conversely, a vendor or deploying organization may control model configuration, interface defaults, and update timing without holding authority to approve an individual medication order.

Safety accountability should consequently be distributed but not allowed to become diffuse. Healthcare-AI accountability requires identifiable responsibilities across development, deployment, clinical use, and oversight [8]. The proposed model uses contribution, control, authority, and liability as separate lenses. It does not assume that they carry equal weight, that they always belong to different actors, or that conceptual allocation determines legal responsibility.

Table 1 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for definitions, components, relationships, and intended uses for the article Dividing Responsibility When a Pharmacist and an Algorithm Co-Produce the Medication Decision.

Table 1. Definitions, components, relationships, and intended uses for the article Dividing Responsibility When a Pharmacist and an Algorithm Co-Produce the Medication Decision.							
Component or scientific dimension	Problem addressed	Inputs or determinants	Proposed mechanism or relationship	Expected contribution	Evidence required	Failure or uncertainty risk	Boundary statement
Contribution	Final-action analysis hides upstream influence.	Data, model output, interface, clinical context, workflow.	Proposed: trace each material input to the decision episode.	Makes distributed production visible.	Provenance, logs, and workflow observation.	Causal influence may remain uncertain.	Contribution does not establish authority or liability [5].

Control	Formal oversight may be practically ineffective.	Access, time, competence, transparency, override and escalation capacity.	Proposed: assess feasible rather than nominal intervention capacity.	Identifies who could alter or interrupt risk.	Usability, access, workload, and override testing.	Apparent control may be unusable under actual conditions.	Control does not automatically establish liability [6].
Authority	Interface position may be mistaken for legitimate decision power.	Professional scope, policy, credentialing, and organizational rules.	Proposed: identify who may approve, reject, restrict, or escalate.	Clarifies legitimate decision rights.	Role, scope, and policy analysis.	Authority may be ambiguous or shared.	Authority without feasible control may create a responsibility gap [7].
Liability	Causal contribution is conflated with legal answerability.	Law, contract, standard of care, product conduct, and causation.	Maintain liability as a separate jurisdiction-specific assessment.	Prevents unsupported legal conclusions.	Case-specific legal analysis.	Rules and standards vary among jurisdictions.	The model cannot decide negligence or liability [6].
Co-produced medication decision	Human final approval conceals joint production.	Algorithmic triage, professional judgment, and workflow conditions.	Proposed: evaluate the combined decision pathway.	Shifts analysis from the model alone to the joint task.	Comparative and process-based studies.	Individual contributions may be difficult to isolate.	Co-production does not imply equal responsibility [8].
Responsibility-Bearing Decision Episode	Analysis of one moment omits upstream and downstream duties.	Data, recommendation, approval, implementation, and follow-up.	Proposed: use a longitudinal unit connecting stages and actors.	Supports responsibility reconstruction.	Prospective process tracing and incident review.	Episode boundaries may vary by task and setting.	This is an analytic construct, not a legal category.
Effective override capacity	An override button may be mistaken for meaningful control.	Information, time, competence, usable interface, and escalation route.	Proposed: treat oversight as substantive only when dissent is practicable.	Tests whether professional control is genuine.	Simulation and real-world human-factors evidence.	Override may be technically possible but operationally unsafe.	The condition requires empirical validation.

Proposed Responsibility-Allocation Model

Human and artificial intelligence can contribute complementary capabilities, but complementarity does not settle who is responsible for the resulting decision [9]. The proposed Responsibility-Allocation Model begins with a Responsibility-Bearing Decision Episode, defined as a bounded medication-decision process extending from data formation to recommendation, pharmacist approval or rejection, implementation, and follow-up. The episode is examined through contribution, control, authority, and liability rather than compressed into a single responsibility score.

Human–AI task performance may improve or worsen according to AI quality, interface design, and the user's response to advice [10–12]. Accordingly, the model rejects two symmetrical simplifications. The first is algorithmic determinism, in which model influence is treated as if it displaced professional judgment and absorbed responsibility. The second is human-finalism, in which the pharmacist's final approval is treated as if it erased upstream design,

configuration, and workflow contributions. Both positions overlook how influence and control may be distributed unequally.

The model contains three actor lanes: the algorithm developer or vendor, the pharmacist, and the deploying organization or oversight body. These lanes intersect four process stages. At the data stage, responsibility concerns provenance, target definition, representativeness, and local data quality. At recommendation, it concerns intended use, model behavior, uncertainty, presentation, and abstention. At approval, it concerns contextual assessment, disagreement, override, escalation, and final professional action. At follow-up, it concerns outcome capture, incident review, drift detection, and change control.

Figure 1 presents the original conceptual synthesis for responsibility-allocation model for co-produced medication decisions, showing how the article's principal components, evidence relationships, uncertainties, and decision boundaries are connected.

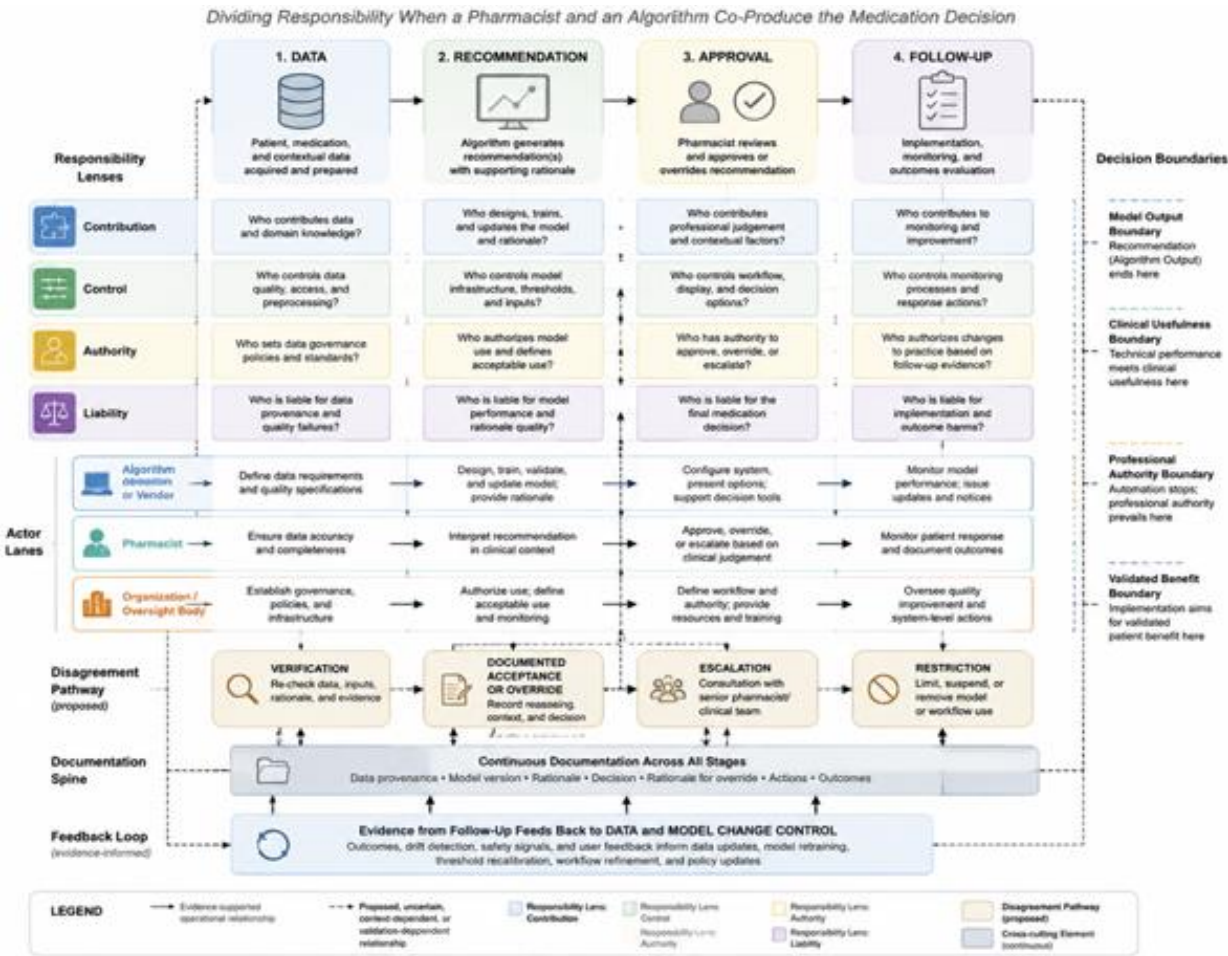


Figure 1. Responsibility-Allocation Model for Co-Produced Medication Decisions

Four proposed rules govern allocation. First, the non-equivalence rule prohibits inferring one responsibility dimension from another. Second, the effective-control condition requires assessment of information access, competence, review time, uncertainty communication, and practicable dissent. Third, the authority–control alignment rule identifies a possible gap when an actor possesses authority without feasible control, or control without accountable authority. Fourth, the dynamic-reallocation rule requires reassessment after material changes in data, model,

interface, workflow, staffing, task, or population. These rules generate testable propositions rather than validated causal laws.

Table 2 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for evaluation criteria, evidence requirements, and failure modes for the article Dividing Responsibility When a Pharmacist and an Algorithm Co-Produce the Medication Decision.

Table 2. Evaluation criteria, evidence requirements, and failure modes for the article Dividing Responsibility When a Pharmacist and an Algorithm Co-Produce the Medication Decision.

Architecture layer or process stage	Core function	Information flow	Interaction with other components	Evaluation criterion	Potential failure mode	Human or organizational responsibility	Readiness boundary
Data	Form the evidentiary substrate.	Records and labels enter the model.	Shapes every downstream output.	Provenance, relevance, and representativeness.	Missing, biased, or obsolete inputs.	Organization and model owner verify sources and local data connections.	The proposed allocation cannot compensate for invalid data.
Recommendation	Generate and present advice.	Model output reaches the pharmacist.	Interface affects influence and interpretation.	Intended-use fit, uncertainty, and usability.	Overconfident, misleading, or poorly timed advice [10].	Vendor documents behavior; organization	Standalone model performance is insufficient.

						configures presentation.	
Approval	Contextualize, accept, reject, or escalate.	Recommendation is combined with patient and medication context.	Depends on effective control and legitimate authority.	Independent review and defensible rationale.	Automation-biased acceptance [12].	Pharmacist assesses the case; organization enables verification.	A final click does not establish meaningful control.
Joint task	Produce the actionable medication decision.	Human and algorithmic inputs converge.	Joint performance may differ from either component alone.	Comparative safety and usefulness of the combined task.	Assistance degrades expert performance [11].	Evaluation is shared across professional and system owners.	Complementarity must be demonstrated rather than presumed [9].
Disagreement pathway	Manage conflicting assessments.	Difference triggers verification or escalation.	Connects approval to governance.	Traceable rationale and timely resolution.	Silent acceptance, unjustified override, or delay.	Pharmacist documents reasoning; organization supplies escalation.	No universal escalation threshold is proposed.
Follow-up	Detect consequences and changing conditions.	Outcomes and incidents return to oversight.	Feeds monitoring and change control.	Outcome capture, drift review, and incident response.	Harm or performance degradation remains invisible.	Organization owns surveillance; vendor supports technical review.	Monitoring requirements require empirical validation.
Reallocation gate	Reassess duties after material change.	Change information returns to all actor lanes.	May alter contribution, control, authority, and potential liability.	Timely responsibility remapping and revalidation.	Obsolete or unowned responsibility assumptions.	Governance body defines and applies change control.	Materiality must be specified locally.

Responsibility across Data, Recommendation, Approval, and Follow-Up

At the data stage, responsibility begins before any recommendation appears. Electronic health record models can inherit missingness, measurement error, misclassification, and representation bias [13]. Data responsibility therefore concerns who selected the variables, defined labels and targets, established inclusion rules, connected local records, and monitored data quality. The pharmacist may identify implausible inputs at the point of use but cannot reasonably be treated as controlling upstream data-generation practices that are inaccessible or technically unreviewable.

Target selection is itself responsibility-bearing. An algorithm can produce inequitable results when a proxy outcome poorly represents the clinical need that the system is intended to identify [14]. In medication decision support, analogous risks may arise when predicted utilization, historical intervention, or recorded cost substitutes for medication need, preventable harm, or therapeutic benefit. The proposed model therefore assigns upstream contribution and control to those who define targets and construct datasets, while leaving legal attribution open to case-specific analysis.

At the recommendation stage, the model owner and deploying organization should make intended use, uncertainty, limitations, version status, and abstention behavior available in a form the pharmacist can use. Communicating uncertainty is necessary because a point estimate can conceal whether a case lies outside familiar data or whether alternative actions remain plausible [15]. Uncertainty display does not itself make a recommendation

safe. Its value depends on calibration, task-specific interpretation, workflow timing, and a defined response when confidence is insufficient.

At approval, the pharmacist contributes clinical context, medication knowledge, patient-specific judgment, and professional authority. The proposed model does not weaken that role; it specifies the conditions under which it can be exercised meaningfully. Approval responsibility is strongest when the pharmacist can inspect relevant data, understand the recommendation's intended use, independently assess consequences, disagree without penalty, and escalate unresolved uncertainty. Where these conditions are absent, the organization retains responsibility for the work design that constrained professional control.

Follow-up extends the responsibility-bearing episode beyond implementation. Delivering clinical impact from artificial intelligence requires external validation, workflow integration, usable outputs, and monitoring beyond model development [16]. Under the proposed allocation, the deploying organization should establish outcome capture, incident review, drift surveillance, and feedback routes to pharmacists and model owners. Developers or vendors should support version traceability, communication of known limitations, and technical investigation of changed behavior. Pharmacists should document clinically meaningful disagreement, unexpected recommendations, and observed consequences through a practicable reporting process.

Responsibility across the four stages is neither equal nor static. A pharmacist may dominate authority at approval while an organization controls configuration and escalation;

a vendor may dominate technical contribution while lacking authority over an individual medication order; and follow-up information may require all three actors to revise earlier assumptions. The proposed allocation must therefore be reconsidered whenever material changes alter what an actor contributes, can control, is authorized to decide, or may be answerable for. It remains a conceptual governance tool rather than a determination of professional fault, clinical appropriateness, or legal liability.

Disagreement, Override, and Failure Scenarios

Disagreement is not inherently evidence that either the pharmacist or the algorithm has failed. It may reveal missing context, an unsuitable model boundary, a defensible professional exception, or an unsafe departure from a valid recommendation. The proposed model therefore distinguishes six scenarios: justified override, insufficiently supported override, justified acceptance, automation-biased acceptance, unresolved disagreement, and inability to verify. Automation bias becomes more likely when verification is cognitively difficult or operationally burdensome [17]. Responsibility analysis must consequently examine the conditions under which agreement or disagreement occurred, not only the direction of the final action.

Medication-alert evidence shows that overrides vary in appropriateness and cannot automatically be classified as error [18]. A justified override may occur when the pharmacist identifies patient-specific information unavailable to the system. An unsafe override may occur when relevant advice is dismissed without adequate review. Conversely, algorithm acceptance may be justified by concordant evidence or may reflect excessive deference. Poorly targeted renal medication alerts illustrate how frequent low-value recommendations can normalize overriding while leaving a smaller number of consequential alerts vulnerable to inappropriate dismissal [19].

Work complexity and repeated alerts also affect responsiveness to decision support [20]. An organization cannot therefore attribute every failed verification to an individual pharmacist while ignoring staffing, interruption, interface, or escalation conditions. The proposed escalation pathway is activated when material data are missing, uncertainty cannot be resolved, disagreement concerns a high-consequence action, an override lacks defensible rationale, or repeated patterns suggest systematic failure. Escalation may involve a second professional review, temporary restriction of the recommendation, technical investigation, or organizational incident assessment. No universal quantitative trigger is proposed, and the classification does not determine negligence or causation.

Documentation and Accountability Evaluation

Responsibility allocation must be reconstructable. Evaluation guidance for clinical artificial intelligence emphasizes documentation of intended use, human–AI interaction,

workflow integration, errors, and the handling of system outputs [21–23]. For a co-produced medication decision, the proposed minimum accountability record includes relevant data provenance, model and version, recommendation, communicated uncertainty, pharmacist assessment, acceptance or override rationale, escalation, final action, and available follow-up information.

Model-side transparency should also describe data, methods, performance evaluation, and applicability [24]. However, documentation completeness is not equivalent to decision quality. A thoroughly documented decision may still be unsafe, while incomplete documentation may prevent reliable reconstruction even when the action was clinically appropriate.

Evaluation should therefore separate model performance from joint task performance. Required domains include technical validity, pharmacist–algorithm performance, workflow effects, medication safety, equity, traceability, escalation effectiveness, and longitudinal stability. Validation should examine whether the proposed constructs can be applied consistently, whether authority–control mismatch predicts recognizable failure patterns, and whether documentation improves incident reconstruction. These are research requirements rather than demonstrated properties of the model.

Implications for Professional and Organizational Governance

Implementation is a sociotechnical and longitudinal process requiring organizational adaptation, lifecycle governance, and human-centred work design [25–27]. Pharmacists should retain authority appropriate to professional scope, but organizations must provide the information, time, training, escalation routes, and monitoring infrastructure needed to exercise that authority. Developers or vendors remain responsible for accurate technical documentation, version traceability, known limitation disclosure, and communication of material changes.

Governance should allocate responsibilities across lifecycle actors rather than treating frontline oversight as the sole safeguard [28]. Procurement should establish intended use and accountability ownership; local validation should examine task and population fit; routine use should include incident review and monitoring; and material model or workflow changes should trigger renewed responsibility mapping. Restriction or retirement should be available when safe use cannot be supported.

Table 3 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for governance responsibilities, implementation conditions, and monitoring requirements for the article *Dividing Responsibility When a Pharmacist and an Algorithm Co-Produce the Medication Decision*.

Table 3. Governance responsibilities, implementation conditions, and monitoring requirements for the article Dividing Responsibility When a Pharmacist and an Algorithm Co-Produce the Medication Decision.

Governance domain	Responsible actor	Required authority	Documentation requirement	Monitoring indicator	Escalation trigger	Corrective action	Accountability boundary
Problem definition and procurement	Organization with pharmacy and technical representation	Approve intended use and contractual conditions	Intended task, excluded uses, ownership, and vendor duties	Misalignment between purchased function and clinical need	Unclear ownership or unsupported intended use	Delay, redefine, or reject implementation	Procurement does not establish clinical benefit [25].
Model and data configuration	Developer or vendor and deploying organization	Configure technical operation and local data connections	Data provenance, model version, target, and configuration	Missing, biased, or obsolete inputs	Material data or configuration change	Revalidate, restrict, or revert configuration	Pharmacists do not control inaccessible upstream design.
Professional use	Pharmacist	Accept, reject, contextualize, or escalate within scope	Recommendation, assessment, rationale, and final action	Unexplained reliance, repeated disagreement, or inability to verify	High-consequence unresolved uncertainty	Second review, escalation, or temporary non-use	Formal approval does not prove effective control.
Work-system support	Deploying organization	Set staffing, access, training, and escalation arrangements	Training, role assignments, downtime, and escalation pathway	Alert burden, verification delay, or inaccessible support	Oversight becomes operationally infeasible	Redesign workflow or restrict system use	Safety is produced by the work system, not the individual alone [27].
Postdeployment monitoring	Organization with vendor support	Investigate incidents and authorize operational changes	Outcomes, incidents, complaints, overrides, and version history	Changed performance, subgroup disparity, or recurrent failure	Drift, harmful pattern, or unannounced update	Investigate, recalibrate, suspend, or retire	Monitoring does not establish causal liability.
Change control	Governance body	Approve material model or workflow changes	Change description, impact assessment, and renewed allocation map	Altered contribution, control, or authority	Material technical or organizational change	Repeat validation and redefine responsibilities	Proposed: materiality thresholds require local specification [28].
Restriction or retirement	Organization and authorized oversight body	Restrict, suspend, or discontinue use	Rationale, affected decisions, communication, and remediation	Persistent unmitigated risk or unavailable monitoring	Safe operating conditions cannot be maintained	Restrict, replace, or retire the system	Continued availability does not establish readiness.

Limitations

The model is conceptual and has not been empirically validated. Explainability does not guarantee causal understanding, trustworthy advice, or meaningful professional control [29]. Dataset shift may change model behavior and render a previously plausible responsibility allocation obsolete [30]. Much comparative clinical-AI evidence has methodological and reporting limitations, restricting transfer from diagnostic studies to pharmacy practice [31]. Ethical requirements concerning fairness, privacy, transparency, and accountability are also context-dependent and longitudinal [32].

The model does not determine legal liability, professional negligence, regulatory compliance, or appropriate treatment. Its actor categories may require expansion where patients, prescribers, caregivers, data providers, or multiple organizations materially shape the decision. Empirical studies must establish construct reliability, task-specific applicability, and appropriate transition gates.

CONCLUSION

Medication decisions supported by artificial intelligence are neither produced by algorithms alone nor adequately explained by assigning all responsibility to the pharmacist who performs the final approval. The proposed Responsibility-Allocation Model treats the decision as a longitudinal episode spanning data, recommendation, approval, and follow-up. It separates contribution, control, authority, and liability so that influence is not mistaken for legal answerability and formal authority is not mistaken for feasible control.

The model offers a structured basis for documenting disagreement, examining override conditions, identifying authority–control gaps, allocating organizational duties, and reassessing responsibility after technical or workflow change. Its scientific value depends on later evaluation of technical performance, joint human–algorithm decision quality, medication safety, equity, workflow effects, traceability, and longitudinal governance.

The framework remains an original conceptual proposal. It does not establish that responsibility can be divided

numerically, that any actor is legally liable, or that a particular system should be deployed. Its principal contribution is to make the distributed production of medication decisions visible and empirically examinable without weakening professional judgment or allowing organizational and technical responsibilities to disappear behind the pharmacist's final action.

ACKNOWLEDGMENTS: None
CONFLICT OF INTEREST: None
FINANCIAL SUPPORT: None
ETHICS STATEMENT: None

REFERENCES

- 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
- Sutton RT, Pincock D, Baumgart DC, Sadowski DC, Fedorak RN, Kroeker KI. An overview of clinical decision support systems: Benefits, risks, and strategies for success. *npj Digit Med*. 2020;3:17. doi:10.1038/s41746-020-0221-y
- Levivien C, Cavagna P, Grah A, Buronfosse A, Courseau R, Bézé Y, et al. Assessment of a hybrid decision support system using machine learning with artificial intelligence to safely rule out prescriptions from medication review in daily practice. *Int J Clin Pharm*. 2022;44(2):459-65. doi:10.1007/s11096-021-01366-4
- Graafsma J, Murphy RM, van de Garde EMW, Karapinar-Carkit F, Derijks HJ, Hoge RHL, et al. The use of artificial intelligence to optimize medication alerts generated by clinical decision support systems: A scoping review. *J Am Med Inform Assoc*. 2024;31(6):1411-22. doi:10.1093/jamia/ocae076
- Vayena E, Blasimme A, Cohen IG. Machine learning in medicine: Addressing ethical challenges. *PLoS Med*. 2018;15(11):e1002689. doi:10.1371/journal.pmed.1002689
- Price WN 2nd, Gerke S, Cohen IG. Potential liability for physicians using artificial intelligence. *JAMA*. 2019;322(18):1765-6. doi:10.1001/jama.2019.15064
- Grote T, Berens P. On the ethics of algorithmic decision-making in healthcare. *J Med Ethics*. 2020;46(3):205-11. doi:10.1136/medethics-2019-105586
- Habli I, Lawton T, Porter Z. Artificial intelligence in health care: Accountability and safety. *Bull World Health Organ*. 2020;98(4):251-6. doi:10.2471/BLT.19.237487
- Topol EJ. High-performance medicine: The convergence of human and artificial intelligence. *Nat Med*. 2019;25(1):44-56. doi:10.1038/s41591-018-0300-7
- 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
- Kiani A, Uyumazturk B, Rajpurkar P, Wang A, Gao R, Jones E, et al. Impact of a deep learning assistant on the histopathologic classification of liver cancer. *npj Digit Med*. 2020;3:23. doi:10.1038/s41746-020-0232-8
- 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;4:31. doi:10.1038/s41746-021-00385-9
- 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
- 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
- Kompa B, Snoek J, Beam AL. Second opinion needed: Communicating uncertainty in medical machine learning. *npj Digit Med*. 2021;4:4. doi:10.1038/s41746-020-00367-3
- Kelly CJ, Karthikesalingam A, Suleyman M, Corrado G, King D. Key challenges for delivering clinical impact with artificial intelligence. *BMC Med*. 2019;17:195. doi:10.1186/s12916-019-1426-2
- Lyell D, Coiera E. Automation bias and verification complexity: A systematic review. *J Am Med Inform Assoc*. 2017;24(2):423-31. doi:10.1093/jamia/ocw105
- Nanji KC, Seger DL, Slight SP, Amato MG, Beeler PE, Her QL, et al. Medication-related clinical decision support alert overrides in inpatients. *J Am Med Inform Assoc*. 2018;25(5):476-81. doi:10.1093/jamia/ocx115
- Shah SN, Amato MG, Garlo KG, Seger DL, Bates DW. Renal medication-related clinical decision support alerts and overrides in the inpatient setting following implementation of a commercial electronic health record: Implications for designing more effective alerts. *J Am Med Inform Assoc*. 2021;28(6):1081-7. doi:10.1093/jamia/ocaa222
- Ancker JS, Edwards A, Nosal S, Hauser D, Mauer E, Kaushal R. Effects of workload, work complexity, and repeated alerts on alert fatigue in a clinical decision support system. *BMC Med Inform Decis Mak*. 2017;17(1):36. doi:10.1186/s12911-017-0430-8
- 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; SPIRIT-AI and CONSORT-AI Working Group. 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; SPIRIT-AI and CONSORT-AI Working Group. 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:e078378. doi:10.1136/bmj-2023-078378
- Greenhalgh T, Wherton J, Papoutsis C, Lynch J, Hughes G, A'Court C, et al. Beyond adoption: A new framework for theorizing and evaluating nonadoption, abandonment, and challenges to the scale-up, spread, and sustainability of health and care technologies. *J Med Internet Res*. 2017;19(11):e367. doi:10.2196/jmir.8775
- 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
- Carayon P, Wooldridge A, Hoonakker P, Hundt AS, Kelly MM. SEIPS 3.0: Human-centered design of the patient journey for patient safety. *Appl Ergon*. 2020;84:103033. doi:10.1016/j.apergo.2019.103033
- Reddy S, Allan S, Coghlan S, Cooper P. A governance model for the application of AI in health care. *J Am Med Inform Assoc*. 2020;27(3):491-7. doi:10.1093/jamia/ocx192
- 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):e745-e750. doi:10.1016/S2589-7500(21)00208-9
- Subbaswamy A, Saria S. From development to deployment: Dataset shift, causality, and shift-stable models in health AI. *Biostatistics*. 2020;21(2):345-52. doi:10.1093/biostatistics/kxz041
- 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:m689. doi:10.1136/bmj.m689
- Chen IY, Pierson E, Rose S, Joshi S, Ferryman K, Ghassemi M. Ethical machine learning in healthcare. *Annu Rev Biomed Data Sci*. 2021;4:123-44. doi:10.1146/annurev-biodatasci-092820-114757