

Keeping Algorithmic Medication Advice Open to Challenge in Routine Pharmacy Care

Zsolt Kovács^{1*}, Katalin Szabo¹, Gabor Nagy², Eszter Toth³

¹Department of Open-to-Challenge Medication Advice, Faculty of Pharmacy, Semmelweis University, Budapest, Hungary. ²Department of Algorithmic Accountability in Pharmacy, Faculty of Pharmacy, University of Debrecen, Debrecen, Hungary. ³Department of Clinical Challenge and AI, Faculty of Pharmacy, University of Pécs, Pécs, Hungary.

Abstract

Algorithmic systems increasingly generate medication alerts, risk estimates, prioritizations, and treatment suggestions that can influence pharmacists, prescribers, patients, and organizations. Their safety cannot be secured solely through predictive performance, the presence of a professional reviewer, or the provision of a technical explanation. When advice is difficult to question, interrupt, reassess, or correct, it may acquire practical authority beyond its evidential strength. This article develops an original contestability framework for algorithmic medication advice in routine pharmacy care. Contestability is defined as the practical capacity of an affected or responsible person to obtain decision-relevant information, raise a challenge, activate a proportionate protective response, secure accountable review, and achieve correction or justified confirmation. The proposed framework connects six functions: registration of the advice object; access to reasons, evidence, alternatives, uncertainty, and limitations; human deliberation; challenge and escalation; correction or supersession; and organizational learning. Accountability, accessibility, equity, traceability, response time, and auditability operate as cross-cutting conditions. The framework distinguishes model output from medication decisions, explanation from justification, nominal oversight from meaningful review, and patient-level correction from system-level learning. Evaluation should assess whether contestability can be located, understood, initiated, answered, and closed under realistic workload and consequence conditions. Validation requires prospective pharmacy-workflow studies, patient and professional usability research, safety evaluation, equity analysis, external assessment, and longitudinal monitoring. Contestability is presented as a necessary governance capability rather than proof of algorithmic validity, clinical benefit, regulatory acceptability, or deployment readiness.

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

INTRODUCTION

Clinical decision-support systems can improve access to relevant information, but their effects depend on data quality, workflow integration, alert design, user response, and implementation conditions. Medication-alert research also shows that technical optimization, alert acceptance, error interception, external validation, and routine deployment are distinct questions rather than interchangeable indicators of value [1-3].

In pharmacy practice, an algorithmic output may appear as a dose suggestion, interaction warning, dispensing alert, medication-review priority, adherence-risk estimate, or recommendation to suppress another alert. Each output enters a professional and organizational setting in which time pressure, incomplete records, hierarchy, interface design, and perceived machine authority may shape whether it is verified.

Algorithmic medication advice becomes practically authoritative when accepting it is easier than understanding, questioning, or delaying it. Authority in this sense does not necessarily arise from a formal rule. It may emerge from workflow defaults, inaccessible evidence, unexplained

confidence, limited override routes, or uncertainty about who is responsible for review. A pharmacist may technically remain free to disagree while lacking sufficient time, information, permission, or escalation support to do so meaningfully.

Machine-learning systems can also create unintended consequences that extend beyond incorrect prediction,

Address for correspondence: Zsolt Kovács, Department of Open-to-Challenge Medication Advice, Faculty of Pharmacy, Semmelweis University, Budapest, Hungary.
zsolt.kovacs@semmelweis.hu

Received: 28 April 2026; **Accepted:** 02 September 2026

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: Kovács Z, Szabo K, Nagy G, Toth E. Keeping Algorithmic Medication Advice Open to Challenge in Routine Pharmacy Care. Arch Pharm Pract. 2026;17(3):75-83.
<https://doi.org/10.51847/5ojs70Ftj2>

including altered professional reasoning, redistribution of responsibility, and reliance on outputs whose limitations are poorly visible [4]. Nevertheless, challengeability should not be equated with generalized distrust. Requiring every output to undergo exhaustive reconsideration could increase delay, workload, alert fatigue, or inappropriate override. The relevant question is therefore not whether algorithmic advice should always be rejected or independently reconstructed, but whether the sociotechnical system preserves a proportionate and usable route through which consequential advice can be examined.

This article proposes practical contestability as an organizing principle for routine pharmacy care. Practical contestability means that a patient, pharmacist, prescriber, or other authorized participant can identify the advice, access information needed to appraise it, initiate a challenge without unreasonable burden, obtain an accountable response, and secure confirmation, correction, or supersession. The framework is conceptual. It does not establish that contestability improves outcomes, determines which disputed recommendation is correct, or substitutes for model validation, professional competence, cybersecurity, or medication-safety governance.

Defining Contestability in Medication Decisions

Contestability is broader than transparency. Transparency may reveal that an algorithm was used, whereas contestability concerns whether an affected or responsible person can question the resulting advice and receive an accountable response. Ethical analyses of machine learning in medicine support scrutiny, accountability, and protection of persons affected by consequential computational decisions [5]. In the proposed framework, contestability therefore combines an epistemic function—access to information needed for appraisal—with a procedural function—a usable route to review and response.

Contestability is also distinct from placing a nominal human “in the loop.” Meaningful oversight requires a defined division of labour: who may challenge, who must respond, whose judgment can pause or supersede the advice, and which responsibilities remain with the organization that selected, configured, or deployed the system [6]. A pharmacist who can view an output but cannot access its basis, record disagreement, or obtain timely review does not exercise complete contestability.

Explanation is one component rather than the whole mechanism. Explainability requirements vary by stakeholder, decision, consequence, and purpose; information useful to a developer may be unusable for a patient or pharmacist [7]. Contestability may therefore require reasons, provenance, alternatives, uncertainty, known limitations, and role information in different forms. Patient-centred work on contestable medical AI similarly identifies data use, bias, performance, and division of labour as important dimensions of an informed challenge [8].

This article proposes four connected capabilities. Epistemic access concerns reasons, evidence, alternatives, assumptions, applicability, and uncertainty. Procedural access concerns discoverable challenge, review, and appeal routes. Protective capacity concerns the ability to pause, defer, or safely override disputed advice when consequences warrant it. Corrective capacity concerns confirmation, amendment, supersession, notification, and learning. These categories organize inquiry; they are not validated legal rights, technical standards, or universal workflow requirements.

Table 1 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for definitions, components, relationships, and intended uses for the article Keeping Algorithmic Medication Advice Open to Challenge in Routine Pharmacy Care.

Table 1. Definitions, components, relationships, and intended uses for the article Keeping Algorithmic Medication Advice Open to Challenge in Routine Pharmacy Care.

Component or scientific dimension	Problem addressed	Inputs or determinants	Proposed mechanism or relationship	Expected contribution	Evidence required	Failure or uncertainty risk	Boundary statement
Algorithmic medication-advice object	Advice may be transient, untraceable, or confused with an order.	Output, intended action, source, version, user, time, applicability.	Proposed: register advice as an identifiable and versioned object; workflow dependence is evidence-informed [1].	Enables later review and correction.	Workflow and traceability studies.	Excess documentation or false precision.	Registration does not validate the advice.
Practical contestability	Formal oversight may exist without usable challenge.	Information access, authority, time, response route.	Proposed: combine epistemic, procedural, protective, and corrective capacities; accountability is evidence-informed [5].	Converts abstract oversight into observable functions.	Construct-validity and usability studies.	Symbolic compliance.	It is not equivalent to correctness or legal appeal.
Meaningful human review	Human presence may conceal unclear responsibility.	Role, competence, independence, authority, workload.	Proposed: review must permit informed action; division-of-labour	Clarifies responsibility boundaries.	Simulation and field observation.	Rubber-stamping or responsibility diffusion.	Human involvement alone does not ensure control.

			concerns are evidence-derived [6].				
Reason and evidence access	An output may be visible while its basis is inaccessible.	Explanation, provenance, population, assumptions, recency.	Proposed: provide task-appropriate information; stakeholder dependence is evidence-supported [7].	Supports appraisal rather than passive receipt.	Comprehension and decision-process studies.	Plausible but misleading explanations.	Explanation does not prove validity.
Alternatives and uncertainty	Advice may appear singular or more certain than warranted.	Competing options, no-action option, confidence, calibration, limitations.	Proposed: connect alternatives and uncertainty to review choices.	Makes disagreement actionable.	Prospective human-factors testing.	Confusion, overload, or false reassurance.	More information is not necessarily better information.
Challenge and protective response	Disputed advice may become irreversible before review.	Consequence, urgency, challenge reason, escalation route.	Proposed: permit clarification, pause, reasoned override, or second review.	Preserves opportunity for correction.	Safety and timing studies.	Delay, inappropriate override, or challenge fatigue.	Challenge does not presume algorithmic error.
Correction and supersession	Edits may erase the previous advice or fail to reach affected users.	Review outcome, rationale, responsible actor, notification scope.	Proposed: preserve prior version while recording confirmation, correction, or supersession.	Creates traceable closure.	Audit and notification studies.	Incomplete propagation or privacy burden.	A corrected case does not establish a general model defect.
Organizational learning	Patient-level challenges may not influence recurring system problems.	Categorized challenge, recurrence, model, data, policy, interface, workflow.	Proposed: connect closed challenges to periodic system review.	Supports detection of recurring failure patterns.	Longitudinal implementation studies.	Under-reporting or punitive use of challenge data.	Learning effects require empirical demonstration.

Proposed Contestability Framework

The proposed framework treats contestability as a property of the entire medication-advice pathway rather than a feature attached only to the predictive model. Early clinical evaluation guidance and human–AI studies indicate that the system, interface, professional user, workflow, and resulting decision should be examined as a coupled intervention [9–11]. Accordingly, the framework contains six linked layers.

First, the advice object is registered with its source, version, intended action, target user, timestamp, and applicability conditions. Second, a contestability-information package supplies decision-relevant reasons, evidence provenance, alternatives, uncertainty, and limitations. Third, human deliberation integrates current medicines, laboratory information, symptoms, preferences, contraindications, and professional judgment. Fourth, challenge and protective response allow clarification, temporary pause, reasoned override, independent review, or escalation proportionate to urgency and consequence. Fifth, correction and supersession preserve the previous advice while recording the review

outcome and notifying affected actors. Sixth, organizational learning examines whether the event reflects a patient-specific exception, missing data, interface problem, policy conflict, model failure, inequity, or recurrent workflow defect.

The pathway is embedded in a work system. Human-factors evidence indicates that safety emerges through interactions among people, tasks, technologies, organizations, environments, and time rather than through a technical component in isolation [12]. Contestability must therefore remain usable during routine workload, shift changes, cross-setting communication, and time-critical decisions. Accountability, accessibility, equity, traceability, and response time operate across all six layers.

Figure 1 presents the original conceptual synthesis for contestability framework for algorithmic medication advice, showing how the article’s principal components, evidence relationships, uncertainties, and decision boundaries are connected.

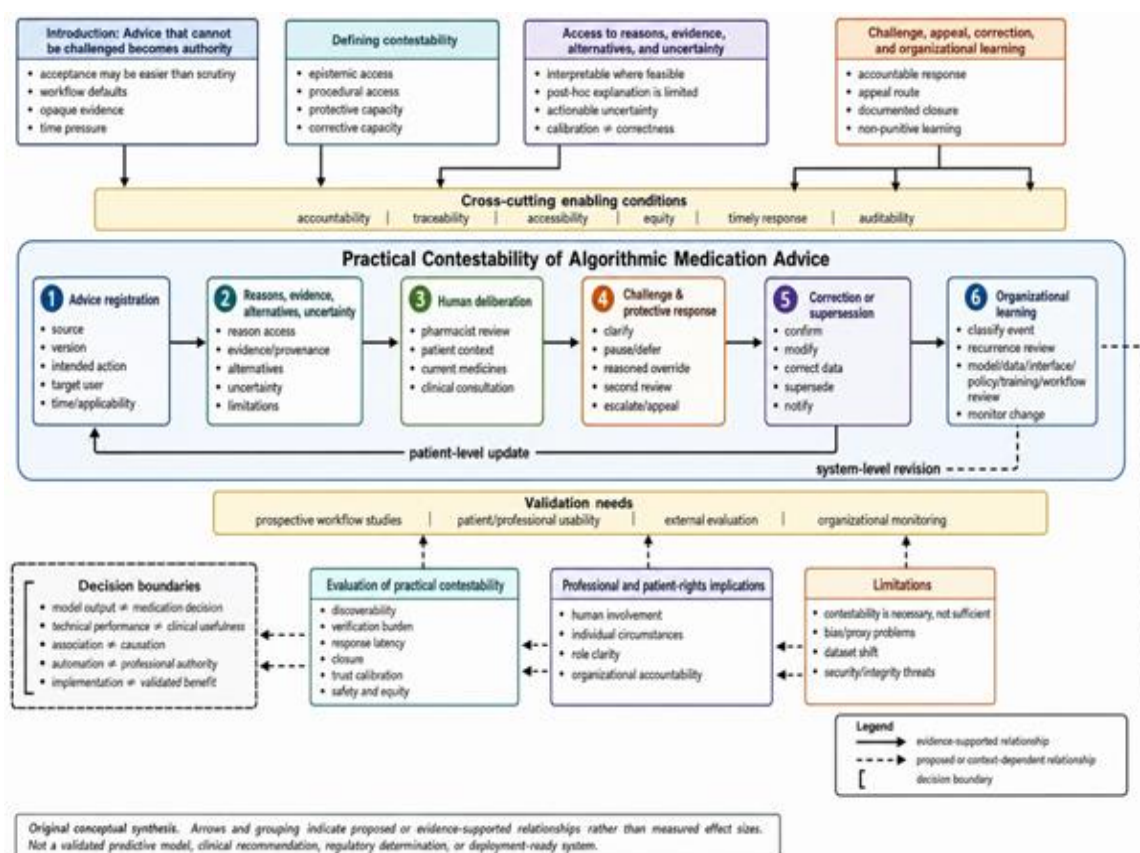


Figure 1. Contestability Framework for Algorithmic Medication Advice

The framework does not prescribe the same response for every output. Low-consequence advice may require rapid clarification and documentation, whereas advice capable of immediately changing a high-risk medication may require independent review or a protected pause. Proportionality is therefore a proposed governance rule, not a validated threshold. Similarly, challenge should not be treated as evidence that the algorithm is wrong. It may identify missing

context, stale data, a patient preference, uncertainty, policy conflict, or a justified exception.

Table 2 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for evaluation criteria, evidence requirements, and failure modes for the article Keeping Algorithmic Medication Advice Open to Challenge in Routine Pharmacy Care.

Table 2. Evaluation criteria, evidence requirements, and failure modes for the article Keeping Algorithmic Medication Advice Open to Challenge in Routine Pharmacy Care.

Architecture layer or process stage	Core function	Information flow	Interaction with other components	Evaluation criterion	Potential failure mode	Human or organizational responsibility	Readiness boundary
Advice registration	Identify the advice and intended medication action.	Source and version to user and audit record.	Enables provenance, review, and supersession.	Completeness and retrievability.	Missing version, stale advice, or ambiguous status.	Vendor, deploying organization, local owner.	Unregistered advice should not be treated as fully contestable.
Contestability-information package	Support informed appraisal.	Reasons, evidence, alternatives, uncertainty, limitations.	Feeds deliberation and escalation.	Comprehension and task relevance.	Information overload or misleading explanation.	System designer and clinical governance team.	Availability alone does not establish usability.
Human deliberation	Integrate advice with patient-specific evidence.	Algorithmic information plus current clinical and preference data.	May confirm or activate challenge.	Appropriate verification and calibrated reliance.	Rubber-stamping, omission, or overcorrection.	Pharmacist and relevant clinical team.	Professional presence does not prove meaningful review.

Challenge initiation	Make disagreement visible and actionable.	Challenge reason to accountable reviewer.	Activates protective response.	Discoverability, burden, and appropriate use.	Hidden route, retaliation concern, or challenge fatigue.	Organization must provide a non-punitive route.	Higher challenge volume is not automatically better.
Protective response	Prevent irreversible action while review occurs.	Urgency and consequence determine pause, override, or escalation.	Connects challenge to review.	Timeliness and medication-safety impact.	Harmful delay or inappropriate override.	Authorized clinician and escalation owner.	Thresholds require setting-specific validation.
Correction or supersession	Close the disputed advice transparently.	Outcome, rationale, prior version, and notification.	Updates patient-level decision and audit trail.	Closure and propagation completeness.	Silent alteration or incomplete notification.	Clinical owner and information-system owner.	A closed case does not prove system-level correction.
Organizational learning	Detect recurrent system problems.	Categorized events to governance review.	May revise data, model, interface, policy, or training.	Recurrence detection and corrective follow-through.	Under-reporting, blame, or unused feedback.	Multidisciplinary governance body.	Learning benefit must be demonstrated longitudinally.
Cross-cutting boundary control	Preserve equity, integrity, and accountability.	Stratified use, audit, shift, and security information.	Constrains every framework layer.	Differential access, robustness, and traceability.	Unequal usability, dataset shift, or manipulation.	Organization, developer, security, and quality teams.	Contestability cannot substitute for validation, equity audit, or cybersecurity.

Access to Reasons, Evidence, Alternatives, and Uncertainty

A challenge cannot be informed when the user sees only a recommendation and a confidence-like visual cue. The proposed contestability-information package should communicate what action is suggested, why it was generated, which evidence or data influenced it, which population and setting support its use, what alternatives remain available, and what uncertainties or limitations matter for the current decision. The package should be proportional: a pharmacist confronting an urgent dosing alert needs a concise, actionable account, whereas retrospective governance review may require detailed provenance and model documentation.

For high-consequence decisions, inherently interpretable models should be preferred when they can perform the relevant task adequately, because direct inspection may offer stronger grounds for scrutiny than a separate explanation layered onto an opaque model [13]. This preference is conditional rather than absolute. Interpretability does not guarantee correct data, appropriate targets, calibration, fairness, or workflow benefit.

Post-hoc explanations require additional caution. They may describe patterns associated with an output without establishing that the explanation is faithful, causal, clinically meaningful, or sufficient to justify the recommendation for an individual patient [14]. A visually persuasive explanation may therefore increase authority rather than contestability. The framework proposes that explanations be labelled by purpose—for example, technical debugging, professional review, patient communication, or governance audit—and that their known limitations remain visible.

Evidence access should also include provenance. A reviewer should be able to determine whether advice arose from a

guideline rule, statistical model, learned association, local policy, medication database, or combination of sources. Relevant provenance includes version, date, reference population, data recency, exclusions, and known applicability limits. This information does not require disclosure of every technical detail at the point of care; it requires a traceable route to the level of detail needed by the responsible reviewer.

Alternatives are necessary because a single recommendation can conceal the decision structure. Depending on the context, the information package may identify reasonable medication alternatives, additional monitoring, consultation, deferral, or no change. Alternatives should not be fabricated merely to create an appearance of choice. They should be clinically plausible, distinguish evidence from preference-sensitive judgment, and identify when the system was not designed to compare options.

Uncertainty should be connected to action. Communication of uncertainty in medical machine learning is particularly relevant when a second opinion, additional evidence, or human review may be warranted [15]. The framework therefore proposes that uncertainty displays indicate what is uncertain, why it matters, and which response is available. A probability without an action pathway may confuse rather than support challenge.

Calibration must also be separated from discrimination. A model may rank patients effectively while producing probabilities that are poorly aligned with observed risk in the relevant setting [16]. Even well-calibrated population estimates do not establish individual correctness, causal benefit, or transportability. Access to uncertainty should therefore be accompanied by information about calibration setting, subgroup performance where available, missingness, and changes in population or workflow. These requirements

remain proposed until tested in routine pharmacy environments.

Challenge, Appeal, Correction, and Organizational Learning

Practical contestability begins when concern can be converted into an identifiable action. A challenge may arise because advice conflicts with current medication history, laboratory values, symptoms, patient preferences, professional knowledge, or another source of evidence. It may also arise because the advice is incomprehensible, unsupported, internally inconsistent, or generated from incomplete information. The challenge mechanism should permit the user to state the reason without requiring complete proof of algorithmic error.

Availability alone is insufficient. Verification complexity can promote automation bias when checking an output requires disproportionate time, navigation, memory, or data retrieval [17]. The proposed framework therefore treats challenge burden as a safety-relevant property. A challenge route should be visible at the point of use, accessible within the available workflow, and capable of preserving relevant context. It should also permit later challenge when immediate action was necessary and retrospective review becomes the safer option.

Challenge and override are not equivalent. A challenge requests examination; an override changes or prevents the advised action. Some disagreements can be resolved through clarification, retrieval of missing evidence, or confirmation that the advice is applicable. Others may require a temporary pause, an alternative medication decision, or escalation to another professional. The response should be proportionate to the likely consequence of acting, delaying, or overriding.

Appeal becomes relevant when the first response is unavailable, procedurally inadequate, or itself disputed. The proposed appeal pathway should identify who can provide independent or higher-authority review, how urgency is communicated, and when unresolved disagreement must be transferred to clinical or organizational governance. Lifecycle approaches to responsible health-care machine learning support explicit monitoring, failure response, and accountability rather than reliance on individual vigilance alone [18]. Local governance must nevertheless determine who holds authority in each setting.

Correction should distinguish among several outcomes. The original advice may be confirmed; modified because of patient-specific context; superseded by another recommendation; withdrawn because of defective data; or referred for system-level investigation. A corrected record should preserve the previous version, the reason for change, the responsible actor, and the users or processes requiring notification. Silent deletion would weaken traceability and

obscure whether downstream decisions relied on the earlier advice.

Organizational learning requires more than collecting challenge events. Real-world implementation experience indicates that user feedback becomes consequential only when operational teams can interpret it and change system behaviour or workflow [19]. The framework therefore proposes classifying closed challenges as patient-specific exceptions, missing or stale data, interface problems, policy conflicts, explanation failures, model-performance concerns, inequitable effects, or unresolved professional disagreements. These categories are provisional and require empirical revision.

Not every accepted challenge demonstrates model failure, and not every rejected challenge demonstrates correct advice. Learning requires examination of recurrence, consequence, affected populations, and the adequacy of the response. Challenge records must not become punitive surveillance of pharmacists, patients, or prescribers. Non-retaliation, confidentiality, and proportionate data use are necessary if the learning loop is to capture genuine concerns rather than only institutionally acceptable ones.

Evaluation of Practical Contestability

Practical contestability should be evaluated as a multidimensional capability rather than reduced to model performance, explanation availability, or the presence of an override button. Clinical impact, study quality, workflow fit, usability, and calibrated reliance must be examined together while remaining analytically distinct [20-22]. No single composite score is proposed because aggregation could conceal a serious failure in one domain.

Evaluation should first determine whether relevant users can locate and understand the challenge pathway. Candidate measures include discoverability, comprehension of the advice basis, time required to retrieve supporting evidence, ability to identify alternatives, and recognition of uncertainty or applicability limits. These measures should be assessed under realistic workload, interruption, staffing, and time-pressure conditions.

A second domain concerns initiation and response. Measures may include whether seeded or naturally occurring problems are recognized, whether appropriate challenges are initiated, response latency, availability of an accountable reviewer, and whether urgent cases receive a protective response. Higher challenge volume should not automatically be interpreted as better performance. It may indicate effective scrutiny, poor advice quality, excessive ambiguity, or a burdensome system that stimulates defensive disagreement.

A third domain concerns closure. Evaluation should examine whether the disputed advice is confirmed, corrected, or superseded; whether the rationale is recorded; whether

affected users and records are updated; and whether recurrent concerns reach organizational review. Closure should be distinguished from mere acknowledgement. A system that records objections but leaves them unresolved remains procedurally open but practically non-contestable.

Safety evaluation must include harms created by the contestability process itself. Possible failures include inappropriate override, treatment delay, challenge fatigue, duplication of work, diffusion of responsibility, adversarial use of appeals, or overreliance on a second reviewer. Prospective studies should therefore compare both failures to challenge and harms arising from poorly designed challenge mechanisms.

Evaluation should also examine calibrated reliance. The desired outcome is not maximal trust in the algorithm or maximal distrust. It is context-sensitive use in which professionals and patients can accept well-supported advice, question uncertain or inapplicable advice, and understand the limits of their own review. Measures should assess whether confidence changes appropriately when advice quality, explanation quality, evidence availability, or uncertainty changes.

Validation will require simulated medication scenarios, prospective pharmacy-workflow studies, patient and professional usability research, external testing across settings, and longitudinal examination of organizational response. Results should be stratified by language, health literacy, disability, professional role, workload, setting, and medication consequence. Proposed measures will require reliability, validity, and sensitivity testing before they can support comparison or readiness judgments.

Professional and Patient-Rights Implications

Algorithmic medication advice enters relationships already structured by professional duties, patient autonomy, institutional policies, and unequal access to expertise. Contestability should strengthen these relationships without implying that every preference overrides evidence or that every disagreement constitutes a system failure.

Patients may resist medical AI when they believe that an algorithm does not recognize their individual circumstances [23]. In medication care, relevant circumstances may include previous adverse effects, treatment burden, cultural preferences, affordability, adherence barriers, pregnancy intentions, or priorities that are not represented in structured data. A contestable process should allow such information to be introduced and considered. Individual circumstances do not automatically invalidate population-based evidence, but advice that cannot receive relevant context remains epistemically incomplete.

Patients also differ in how much technical information and participation they desire. Patient research identifies apprehensions concerning privacy, loss of human

involvement, accountability, and the capacity of AI to understand personal circumstances [24]. The framework therefore proposes layered communication: concise explanations for immediate decisions, optional access to additional evidence, and a human contact for questions or escalation. Requiring every patient to interpret technical model documentation would transfer an unreasonable burden rather than create meaningful participation.

For pharmacists, contestability protects neither unrestricted discretion nor automatic immunity from responsibility. Pharmacists remain responsible for using professional competence, documenting relevant reasoning, and acting within their authority. At the same time, organizations remain responsible for procurement, data quality, configuration, access controls, monitoring, staffing, escalation routes, and corrective action. Ethical analysis of implementation supports maintaining accountability for these upstream choices rather than transferring it entirely to the professional who receives the output [25].

Professional rights are therefore connected to organizational obligations. A pharmacist should be able to question advice without retaliation, identify who owns the response, and obtain access to sufficient information. Organizations should not describe a system as human-supervised when professionals lack time, authority, or evidence to intervene. Conversely, professionals should not treat contestability as permission to reject advice without reason or to avoid documentation.

Patient rights and professional responsibilities may sometimes conflict. A patient may request rejection of advice that a pharmacist considers important for safety, or a pharmacist may challenge advice that a patient prefers. The framework does not determine the substantive outcome of such disputes. It proposes a process in which reasons, evidence, authority, and unresolved disagreement remain visible.

Limitations

Contestability is necessary but insufficient. Biased targets or proxy outcomes can produce inequitable recommendations even when individual users can challenge them [26]. Patient-level objections may reveal some inequities, but they cannot replace systematic examination of data, labels, objectives, and subgroup effects.

Dataset shift may also reduce reliability when patient populations, prescribing practices, formularies, documentation, or workflows differ from development conditions [27]. A usable challenge pathway cannot establish transportability without external and continuing evaluation.

Adversarial manipulation, corrupted inputs, and other integrity threats may evade ordinary professional review [28]. Independent security, validation, and monitoring controls therefore remain necessary.

The framework is conceptual and has not been prospectively tested. Its components may impose workload, delay, privacy burdens, or unintended responsibility shifts. Legal rights, professional authority, and feasible escalation routes will differ across jurisdictions and pharmacy settings. The proposed categories and evaluation criteria should consequently be treated as hypotheses for empirical refinement rather than universal standards.

CONCLUSION

Algorithmic medication advice remains professionally and ethically bounded only when it can be examined without unreasonable burden, interrupted when consequences warrant caution, and corrected through an accountable process. This article proposes practical contestability as a sociotechnical capability connecting advice registration, epistemic access, human deliberation, challenge, protective response, correction, supersession, and organizational learning.

The framework distinguishes access to an explanation from access to reasons sufficient for action; professional presence from meaningful authority; disagreement from demonstrated algorithmic error; and patient-level correction from system-level improvement. It also places accountability, accessibility, equity, response time, traceability, and auditability across the full medication-advice pathway rather than assigning them to a single model or user.

Contestability cannot make invalid advice valid, remove biased objectives, guarantee transportability, prevent manipulation, or establish clinical benefit. Poorly designed challenge mechanisms may themselves delay treatment, encourage inappropriate override, increase workload, or diffuse responsibility. Validation therefore requires realistic pharmacy-workflow research, patient and professional evaluation, safety assessment, equity analysis, external testing, and longitudinal study of whether challenges produce accountable closure and organizational learning.

The principal contribution is not a claim that algorithmic medication advice should always be resisted. It is the proposition that advice should never acquire authority merely because the surrounding system makes acceptance easier than scrutiny. Keeping such advice open to proportionate, informed, and consequential challenge may provide a clearer basis for responsible human–AI collaboration, but its value and limits must be established empirically.

ACKNOWLEDGMENTS: None

CONFLICT OF INTEREST: None

FINANCIAL SUPPORT: None

ETHICS STATEMENT: None

REFERENCES

- 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
- Graafsma J, Murphy RM, van de Garde EMW, Karapinar-Çarkit 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
- Chen CY, Chen YL, Scholl J, Yang HC, Li YCJ. Ability of machine-learning based clinical decision support system to reduce alert fatigue, wrong-drug errors, and alert users about look alike, sound alike medication. *Comput Methods Programs Biomed.* 2024;243:107869. doi:10.1016/j.cmpb.2023.107869
- Cabitza F, Rasoini R, Gensini GF. Unintended consequences of machine learning in medicine. *JAMA.* 2017;318(6):517-8. doi:10.1001/jama.2017.7797
- Vayena E, Blasimme A, Cohen IG. Machine learning in medicine: Addressing ethical challenges. *PLoS Med.* 2018;15(11). doi:10.1371/journal.pmed.1002689
- 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
- Amann J, Blasimme A, Vayena E, Frey D, Madai VI, Precise4Q Consortium. Explainability for artificial intelligence in healthcare: A multidisciplinary perspective. *BMC Med Inform Decis Mak.* 2020;20(1):310. doi:10.1186/s12911-020-01332-6
- Ploug T, Holm S. The four dimensions of contestable AI diagnostics—A patient-centric approach to explainable AI. *Artif Intell Med.* 2020;107:101901. doi:10.1016/j.artmed.2020.101901
- 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
- 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
- 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
- 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
- Rudin C. Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nat Mach Intell.* 2019;1(5):206-15. doi:10.1038/s42256-019-0048-x
- 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
- Kompa B, Snoek J, Beam AL. Second opinion needed: Communicating uncertainty in medical machine learning. *NPJ Digit Med.* 2021;4(1):4. doi:10.1038/s41746-020-00367-3
- Van Calster B, McLernon DJ, van Smeden M, Wynants L, Steyerberg EW. Calibration: The Achilles heel of predictive analytics. *BMC Med.* 2019;17(1):230. doi:10.1186/s12916-019-1466-7
- 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
- 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
- Sendak MP, Ratliff W, Sarro D, Alderton E, Futoma J, Gao M, et al. Real-world integration of a sepsis deep learning technology into routine clinical care: Implementation study. *JMIR Med Inform.* 2020;8(7). doi:10.2196/15182
- 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
- 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

- Sutton RT, Pincock D, Baumgart DC, Sadowski DC, Fedorak RN, Kroeker KI. An overview of clinical decision support systems:

22. Asan O, Bayrak AE, Choudhury A. Artificial intelligence and human trust in healthcare: Focus on clinicians. *J Med Internet Res.* 2020;22(6). doi:10.2196/15154
23. Longoni C, Bonezzi A, Morewedge CK. Resistance to medical artificial intelligence. *J Consum Res.* 2019;46(4):629-50. doi:10.1093/jcr/ucz013
24. Richardson JP, Smith C, Curtis S, Watson S, Zhu X, Barry B, et al. Patient apprehensions about the use of artificial intelligence in healthcare. *NPJ Digit Med.* 2021;4(1):140. doi:10.1038/s41746-021-00509-1
25. Char DS, Shah NH, Magnus D. Implementing machine learning in health care—addressing ethical challenges. *N Engl J Med.* 2018;378(11):981-3. doi:10.1056/NEJMp1714229
26. 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
27. 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
28. Finlayson SG, Bowers JD, Ito J, Zittrain JL, Beam AL, Kohane IS. Adversarial attacks on medical machine learning. *Science.* 2019;363(6433):1287-9. doi:10.1126/science.aaw4399