

From Informed Consent to Informed Dependence: Rethinking Autonomy in AI-Mediated Clinical Care

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

¹Department of AI Ethics and Clinical Autonomy, Faculty of Medicine, University of Helsinki, Helsinki, Finland.

²Department of Informed Dependence and Patient-Provider Relationships, Faculty of Medicine, University of Eastern Finland, Kuopio, Finland.

*E-mail ✉ mikko.lahtinen@helsinki.fi

Abstract

Artificial intelligence is becoming embedded in diagnosis, documentation, communication, risk prediction, and clinical decision support. These uses complicate autonomy because patients may depend not only on a clinician's recommendation but also on models they cannot inspect, organizations that select and monitor those models, data infrastructures that may change, and clinicians whose own knowledge and communication are increasingly mediated by AI. Conventional informed consent remains indispensable, but its episodic emphasis on authorizing a recognizable intervention may incompletely address these continuing dependencies. This article develops informed dependence as a proposed bioethical concept for identifying when autonomy requires attention to the conditions under which patients must rely on AI-mediated clinical arrangements. It distinguishes epistemic, technical, organizational, and relational dependence and argues that materially consequential dependence may generate differentiated duties of disclosure, oversight, continuity, and feasible exit or recourse. The framework does not require patients to understand model mechanics or demand separate consent for every computational process. Instead, it directs attention toward information and safeguards that affect whether reliance is reasonably intelligible, contestable, monitored, and revisable. This account reframes consent as extending beyond a single authorization moment while preserving proportionality and clinical practicability. Its principal limitations are conceptual rather than empirical: the proposed categories and duties have not been validated as a clinical consent model, their relevance will vary across technologies and jurisdictions, and meaningful alternatives may sometimes be unavailable. Informed dependence is therefore offered as a normative lens for analyzing autonomy under sustained AI mediation, not as a universal legal rule or implementation-ready consent standard.

Keywords: Artificial intelligence, Autonomy, Informed consent, Dependence, Clinical ethics, Patient-centred care

Introduction

AI changes what patients depend upon

Artificial intelligence in health care is no longer confined to experimental image classifiers or isolated prediction tools. Contemporary systems can participate in documentation, diagnosis, triage, communication, prognostication, decision support, and organizational

workflow, with their consequences shaped by how they are embedded in clinical practice rather than by technical performance alone [1]. This expansion matters for bioethics because it changes the architecture through which patients encounter professional judgment. A recommendation may still be communicated by a clinician, yet the reasons available to that clinician, the data underlying the recommendation, the institutional processes governing the system, and even the language through which the recommendation is conveyed may increasingly depend on AI-mediated processes.

Patient responses already suggest that this architecture of reliance matters. Trust in clinical AI is associated with trust in the health system using it, while clinician involvement, governance cues, system performance,

Access this article online

<https://smerpub.com/>

Received: 18 January 2026; Accepted: 19 April 2026

Copyright CC BY-NC-SA 4.0

How to cite this article: Lahtinen M, Salo E, Virtanen J. From Informed Consent to Informed Dependence: Rethinking Autonomy in AI-Mediated Clinical Care. Asian J Ethics Health Med. 2026;6(1):114-23. <https://doi.org/10.51847/9yLVG44HBY>

representative data, and clinical or national context can influence acceptance and attitudes [2–4]. These findings should not be treated as proving a particular consent doctrine. They do, however, challenge an account in which the only ethically relevant question is whether a patient knows that an algorithm exists. What patients may need to understand is not simply *that* AI is present, but what role it occupies, what remains under human judgment, what uncertainties attach to its use, and what happens when the system or its surrounding organization changes.

This point is consistent with broader accounts of clinical AI as a sociotechnical enterprise whose safety and usefulness depend on data quality, workflow integration, evaluation, and human participation [5]. The bioethical implication developed here is more specific. AI can alter what patients must depend upon in order to exercise autonomy. The relevant dependencies may be epistemic, because reasons become difficult to access; technical, because reliability depends on models and data infrastructures; organizational, because institutions determine deployment and monitoring; and relational, because clinicians increasingly interpret or communicate through computational mediation.

The argument is not that dependence is itself incompatible with autonomy. Medicine has always involved reliance on professional expertise, institutions, laboratories, and technologies. The problem arises when materially important dependencies are obscured, unreviewable, unstable, or practically unavoidable while consent continues to be treated as though the patient were authorizing a discrete and sufficiently understood intervention. The task, therefore, is not to replace informed consent but to examine when autonomy requires a richer account of what ongoing reliance entails.

Why conventional informed consent is insufficient

Informed consent ordinarily asks whether a patient possesses decision-making capacity, receives material information, understands what is proposed, acts voluntarily, and authorizes an intervention. Those commitments remain essential. The difficulty is that AI can distribute ethically relevant information and responsibility across systems that extend beyond the immediate clinician-patient exchange. Recent bioethical analysis already distinguishes being notified that AI is used from receiving an explanation of its role and from giving consent to that use [6]. Collapsing these functions

into a single requirement to “tell the patient about AI” therefore obscures several different autonomy interests. Trust provides a further complication. Patients and clinicians may reasonably care about model validation, training data, provenance, performance, and the relationship between algorithmic output and professional judgment [7]. A patient need not understand model architecture to have an autonomy interest in knowing, for example, that a recommendation depends substantially on an externally developed model, that the clinician cannot independently verify a critical inference, or that uncertainty remains clinically important. Conversely, a technically elaborate explanation may contribute little if it leaves these practical conditions of reliance obscure. Consent requirements must therefore remain proportionate. Ethical frameworks for clinical AI appropriately resist the assumption that every algorithmic contribution requires an identical notification or consent process [8]. An administrative optimization tool, a diagnostic model materially affecting treatment, and a conversational system performing portions of consent cannot plausibly create the same informational burden. The insufficiency identified here is thus conditional rather than universal.

The deeper problem is temporal. Conventional consent is commonly organized around an identifiable decision: a procedure, medication, diagnostic intervention, or research participation. AI-mediated care may instead create continuing dependencies that persist after the initial authorization. Models may be updated; data distributions may change; clinicians may become increasingly reliant on automated recommendations; documentation systems may begin capturing encounters continuously; institutional oversight may weaken or change. A patient can therefore have been adequately informed at one point while later receiving care under materially altered conditions. An autonomy framework directed only toward the initial decision risks overlooking this transformation.

The appropriate response is not perpetual reconsent. It is to distinguish the moment of authorization from the continuing conditions under which reliance remains ethically acceptable.

The concept of informed dependence

This article proposes informed dependence as the condition in which a patient can reasonably understand and respond to the materially significant forms of reliance through which AI-mediated care is delivered.

The concept begins from a familiar insight: autonomy is not exercised in informational isolation. Digital environments can shape the possibilities through which individuals understand, choose, and act, making technological conditions relevant to health-related autonomy [9]. In clinical AI, this insight becomes especially important because patients may depend upon systems whose epistemic contribution is difficult for either party to reconstruct.

Patient-centred decision-making can be strained when an AI-supported recommendation changes which reasons are available and how those reasons can be explained or interrogated [10]. Informed dependence therefore asks a different question from conventional disclosure: what does this patient have to rely on for this decision to remain meaningfully theirs? The answer may include the clinician's continuing judgment, the reliability of an AI system, institutional monitoring, data integrity, availability of human review, or a pathway for contesting an output.

The concept is deliberately narrower than a demand for complete transparency. Patients need not learn how an algorithm mathematically generates every prediction.

Nor does informed dependence imply that dependence should be eliminated. Sophisticated medicine inevitably requires epistemic and technical reliance. The normative concern is whether a materially important dependency has become so hidden, unstable, noncontestable, or substitutive of professional judgment that the patient cannot adequately understand the conditions under which care is being recommended.

Informed dependence is also longitudinal. Ethical responsibility surrounding machine learning does not end at laboratory validation but extends through clinical translation and deployment [11]. The same temporal insight should inform autonomy. A patient who authorizes AI-mediated care under one set of assumptions may later encounter a meaningfully different system, institutional arrangement, or degree of clinician reliance.

The contrast developed in this section is summarized in **Table 1**. The table does not replace informed consent with a competing doctrine; it identifies where the proposed concept adds an additional question about continuing reliance.

Table 1. From authorizing an intervention to understanding the conditions of reliance

Analytical dimension	Conventional consent focus	Additional question introduced by informed dependence	Status in this article
Object of ethical attention	Proposed intervention or decision	Which human, technical, and organizational arrangements materially condition the decision	Proposed synthesis informed by distinctions among notice, explanation, and consent
Information required	Material benefits, risks, alternatives, and procedure	Material facts about AI role, clinician reliance, uncertainty, monitoring, and recourse	Evidence-informed extension; not a requirement for full technical disclosure
Temporal horizon	Authorization at a recognizable decision point	Whether materially relevant dependencies remain stable or change during care	Proposed longitudinal extension supported by lifecycle ethics
Understanding	Sufficient comprehension of the choice being authorized	Sufficient comprehension of what must be relied upon for that choice to remain meaningful	Proposed synthesis drawing on digitally mediated autonomy
Agency	Ability to accept or refuse the proposed intervention	Ability, where feasible, to question, seek human review, reconsider, or use an alternative pathway	Proposed; does not establish a universal right to AI-free care
Proportionality	Disclosure varies with materiality and risk	Dependence-related safeguards should likewise intensify only when AI mediation becomes materially consequential	Consistent with context-sensitive consent analysis

Note: "Informed dependence" and the comparison presented here are proposed conceptual contributions. The table does not establish legal entitlements, empirical thresholds, or a validated consent instrument.

Epistemic, technical, organizational, and relational dependence

Informed dependence becomes analytically useful only if different forms of reliance can be distinguished without

pretending that they operate independently. Four dimensions emerge from the preceding argument.

Epistemic dependence concerns reliance on claims, classifications, or recommendations whose grounds

cannot readily be reconstructed. In high-stakes settings, opacity can create an epistemic cost for clinicians themselves when they remain responsible for decisions they are less able to know or justify [12]. Patients can consequently become dependent not only on the AI but also on a clinician's ability to judge when an output deserves trust. This does not make opaque AI intrinsically unethical. It makes the limits of warranted reliance ethically relevant.

Technical dependence concerns the operational conditions that allow an AI system to perform as expected: data inputs, software, model behavior, interfaces, and the relationship between model output and the task for which it is being used. Explainability cannot simply dissolve this dependence. Current approaches to explainable AI may produce interpretations that are unstable, misleading, or insufficiently connected to clinical reasoning [13]. The autonomy-relevant question is therefore not whether every internal computation can be exposed, but whether the system's consequential limitations and uncertainties can be communicated in a form that supports reasonable reliance.

Organizational dependence arises because institutions decide which systems are purchased, validated, integrated, monitored, updated, and withdrawn. Governance models for health-care AI accordingly locate responsibility across organizational processes rather than

within the algorithm alone [14]. A patient may thus depend on practices that are largely invisible at the bedside: local validation, incident response, procurement standards, escalation pathways, and the capacity to detect degraded performance.

Relational dependence concerns what happens when clinicians mediate these technical and organizational arrangements. Patients commonly rely on clinicians not merely to transmit information but to interpret uncertainty, contextualize recommendations, assume professional responsibility, and remain responsive to disagreement. AI changes this relationship when clinicians cannot independently interrogate an output, when automated content shapes communication, or when institutional workflow makes deviation unusually difficult. Relational dependence is therefore not reducible to whether a “human is in the loop.” It asks whether human involvement remains substantively capable of judgment, explanation, and response.

These dimensions form a stack rather than four sealed categories: technical instability can produce epistemic uncertainty; organizational practices determine whether instability is detected; and clinicians mediate its significance for patients. **Figure 1** represents this manuscript-derived dependency structure without implying causal effect sizes or a validated hierarchy.

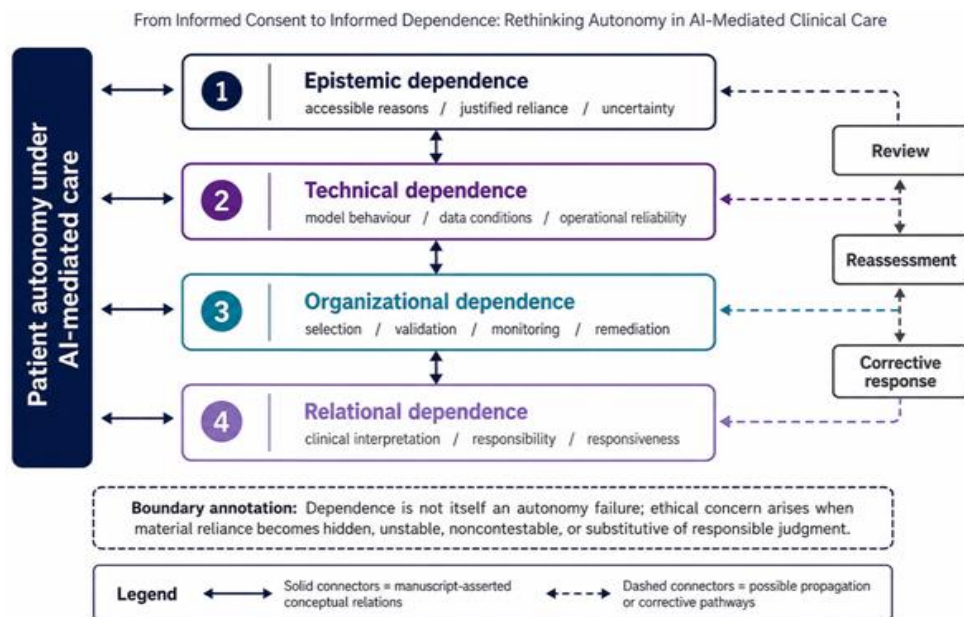


Figure 1. The dependency stack of AI-mediated clinical autonomy

Duties of disclosure, oversight, continuity, and exit

If informed dependence identifies ethically significant reliance, it must also clarify what follows normatively.

Four duties are proposed: disclosure, oversight, continuity, and exit. They are not asserted as universally established legal rights. They are autonomy-oriented obligations whose force should vary with the materiality of the dependence, the stakes of the clinical decision, and the feasibility of alternative arrangements.

Disclosure concerns the aspects of AI mediation a reasonable patient would need in order to understand the conditions of care. Public-attitude evidence indicates substantial interest in being notified when AI is used in health care [15]. Yet disclosure should not become an inventory of every computational process. Its focus should be functional: what role the AI plays, whether it materially influences diagnosis, treatment, communication, or monitoring, how clinician judgment relates to its output, and which important uncertainties or alternatives bear on the patient's decision.

Oversight requires more than nominal human presence. A clinician who simply relays an output without meaningful authority or capacity to question it provides little protection against dependence becoming domination. Patient-centred work on contestable AI supports the importance of mechanisms through which algorithmically mediated decisions can be questioned or reviewed [16]. The proposed duty therefore concerns access to substantively capable human review, especially where AI meaningfully structures high-stakes decisions. Continuity responds to the temporal character of AI systems. A consent process can accurately describe a

system at one moment while its performance environment later changes. Empirical work on deployed clinical models demonstrates that harmful data shifts can arise and require detection and remediation [17]. Such evidence does not establish that every model update requires renewed consent. It does support a distinction between initial authorization and continuing institutional responsibility for conditions on which that authorization relied. Material changes in system role, performance, data use, or clinician reliance may justify renewed explanation or reassessment.

Exit is the most constrained duty. In some settings an equivalent non-AI pathway may be realistic; in others, AI may be inseparable from routine infrastructure or may offer clinically important advantages that cannot be duplicated. Informed dependence therefore cannot promise a universal right to AI-free treatment. The defensible requirement is narrower: where feasible, patients should retain meaningful avenues to question an AI-mediated decision, request human reconsideration, or choose a clinically reasonable alternative. Where no meaningful alternative exists, that constraint may itself become material information.

Together, these duties convert informed dependence from a demand for more information into a demand for preserved agency under reliance. **Table 2** distinguishes what each duty protects, when it becomes ethically salient, and what the proposed account does not require.

Table 2. Autonomy-preserving duties under material AI dependence

Proposed duty	Dependence that makes it salient	Minimum autonomy-protecting requirement	Boundary of the proposed duty
Disclosure	AI materially changes the informational or decisional conditions of care	Explain the AI's consequential role, clinician involvement, important uncertainty, and relevant alternatives; notification preferences support the salience of disclosure	No requirement to enumerate every background algorithm or disclose source code
Oversight	Patient or clinician must rely on an output that may materially affect care	Preserve access to a person or process capable of substantive review, challenge, and justified departure; contestability supports this concern	Nominal "human-in-the-loop" status is insufficient, but continuous duplication of every AI task is not required
Continuity	Reliability, use conditions, data environment, or degree of reliance can change over time	Monitor material changes and reopen explanation or deliberation when the basis of reasonable reliance substantially changes	Model updates do not automatically trigger formal consent
Exit or recourse	AI mediation becomes difficult to avoid or contest	Provide a reasonable alternative, human reconsideration, or escalation route where clinically and operationally feasible	No universal entitlement to AI-free care; absence of a feasible alternative may itself require disclosure

Note: The four-duty structure is a proposed normative synthesis. Evidence supports the underlying concerns—notification, contestability, and post-deployment change—but does not validate this table as a legal standard, clinical protocol, or universal implementation rule.

Redesigning consent for AI-mediated clinical care

If informed dependence is taken seriously, consent should be redesigned around material changes in reliance

rather than multiplied into repeated signatures. Ambient generative-AI documentation illustrates why: the ethically relevant feature is not simply that software is present, but that an encounter may be persistently captured, processed, and transformed into the clinical record [18]. A proportionate model would therefore begin with concise baseline disclosure and add further deliberation when AI assumes a role that could materially alter diagnosis, treatment, communication, privacy, or the patient's ability to question what is happening.

AI can also migrate into the consent process itself. Delegating informational or conversational components of procedural consent to a language model may improve accessibility, but it also redistributes responsibility for explanation and understanding [19]. In paediatric contexts, acceptable assistance must remain sensitive to capacity, developmental status, parental authority, and professional obligations [20]. Systems that simulate empathy introduce an additional issue: communicative fluency can shape the patient's experience of being heard without establishing a genuine reciprocal relationship [21].

The proposed redesign is therefore layered and trigger-based. Baseline disclosure should locate AI within the

clinical relationship: whether it advises, records, drafts, predicts, or communicates; how consequentially its output bears on care; and who remains answerable for the resulting decision. Additional explanation should match the patient's questions and clinical stakes rather than become a technically exhaustive script.

Renewed discussion becomes ethically salient when a dependency changes in a way that could reasonably matter to authorization: for example, when AI moves from clerical assistance to substantive recommendation, when clinician review becomes more attenuated, when data use expands, or when previously available human reconsideration disappears. These are proposed materiality triggers, not empirically validated thresholds. Human review should remain available where it can affect the decision. This is not perpetual consent. It is a commitment to keeping authorization connected to the conditions under which reliance remains reasonable, and to reopening deliberation when those conditions are substantially transformed.

Figure 2 translates this manuscript-derived redesign into a patient-centred interaction architecture.

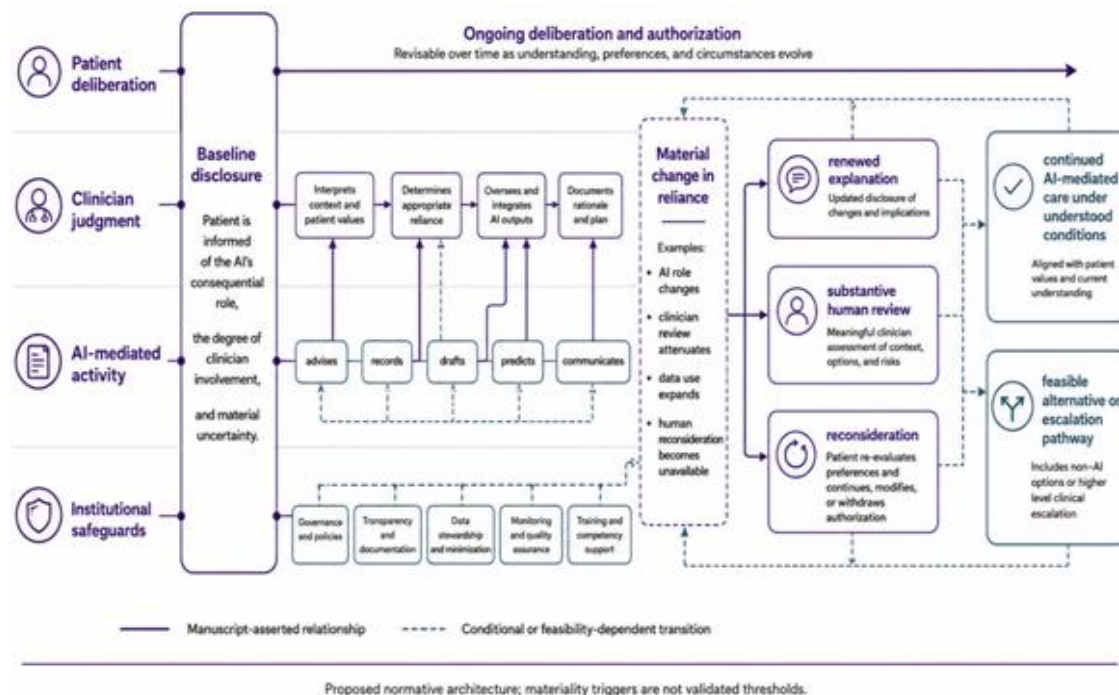


Figure 2. Consent as a revisable pathway under AI-mediated care

Application to clinical scenarios

The value of informed dependence becomes clearer when the proposed lens is applied to distinct uses rather than to

“AI” as a single category. In diagnostic decision support, patients may regard information about AI involvement as material, but preferences vary and cannot by themselves determine the ethical rule [22]. The central question is whether the model materially structures the diagnosis and whether the clinician can critically assess, explain, and if necessary depart from its output.

AI-drafted portal messages create a different dependency. Qualitative evidence indicates that patients can care about disclosure and clinician review when messages are generated with AI [23]. Related work suggests that disclosure may also change how AI-drafted communication is evaluated [24]. These findings make the tension visible: transparency can carry experiential costs, yet preference for an undisclosed output does not settle whether patients have a legitimate interest in knowing who or what shaped clinical communication. The ethically salient issue is whether the patient can identify who reviewed, endorsed, and assumed responsibility for a message that may influence care.

Generative systems capable of emulating different clinician communication styles raise a further concern. Proof-of-concept work shows that such styles can be reproduced by large language models [25]. The ethical problem is therefore not limited to factual accuracy. A system may influence tone, reassurance, framing, or apparent empathy in ways that alter deliberation while leaving authorship and professional responsibility unclear. If a patient interprets simulated responsiveness as evidence that a clinician personally considered a concern, the relational condition of reliance may be altered even when the underlying medical content remains correct.

The framework yields different conclusions across these scenarios. Diagnostic AI may call primarily for explanation of role, uncertainty, and reviewability; AI-mediated communication may additionally require clarity about authorship and clinician responsibility; ambient systems may foreground persistent data capture and organizational stewardship. In each case, informed dependence asks whether the patient can identify the material conditions of reliance, whether responsible human judgment remains substantively available, whether relevant institutional safeguards persist, and whether changed conditions can be reconsidered. The concept therefore functions as a discriminator among AI uses rather than as a uniform demand for additional consent.

Objections concerning burden, comprehension, and innovation

The strongest objection is that informed dependence could turn consent into an exhausting inventory of algorithmic details. That would defeat its purpose. Black-box opacity does not automatically make valid consent impossible; patients may rationally authorize reliance on processes whose internal mechanics they do not understand [26]. The relevant standard should therefore be material understanding, not mechanistic transparency. A clinically appropriate disclosure may sometimes be very short if the AI role is limited, independently reviewed, stable, and unlikely to alter available choices. A second objection concerns comprehension. Clinically useful explanation depends on user needs and task context rather than on maximal technical detail [27]. The proposed framework consequently favors layered explanation: what the system does, how strongly it influences care, what uncertainty matters, who remains responsible, and what review or alternatives exist. Deeper technical information may be appropriate for some patients or decisions but should not become a universal prerequisite. Crucially, simplification must not become concealment: reducing cognitive burden is compatible with identifying the few dependencies that materially shape choice.

A third objection is innovation burden. The literature does not establish that every clinical AI system requires full mechanistic explainability, and arguments linking explainability, autonomy, accountability, and safety remain contested [28]. Informed dependence should therefore intensify safeguards only when reliance becomes materially consequential. Institutions should not be encouraged to surround low-consequence automation with ritualistic consent processes merely to display ethical caution. If applied indiscriminately, the framework would become bureaucratic rather than autonomy-promoting; if applied too narrowly, it could reproduce the opacity it is intended to expose. Proportionality is therefore internal to the proposal rather than an external exception to it.

Scope and limitations

This account is normative and conceptual. Its categories and proposed duties have not been validated as a clinical consent instrument, and their relevance will vary by technology, clinical stakes, institutional resources, and patient population. Evidence on AI inclusivity also remains uneven, limiting assumptions that the same

systems, disclosures, or alternatives serve all groups comparably [29]. Digital literacy, language, disability, socioeconomic constraints, and differential access to clinicians may change whether an ostensibly available review or exit pathway is meaningful in practice.

Stakeholder preferences differ across patients, clinicians, and other actors [30]. The framework therefore cannot assume a single desired amount of explanation or a uniform preference for human-only care. Jurisdictional differences likewise constrain claims about legal duties, and meaningful exit may be impossible where AI is embedded in infrastructure or where alternatives are clinically inferior or unavailable.

International guidance increasingly treats trustworthy AI as a lifecycle and sociotechnical governance problem [31]. That supports continued attention to monitoring and institutional responsibility, but it does not validate informed dependence. Future work should combine qualitative co-design, comparative experiments, longitudinal implementation research, equity-sensitive evaluation, and normative analysis to test whether the proposed distinctions improve comprehension, trust calibration, contestability, or meaningful choice without creating disproportionate burden.

Conclusion

AI-mediated care does not make autonomy obsolete, nor does dependence itself negate autonomy. The ethical problem is that patients may increasingly rely on epistemic, technical, organizational, and relational arrangements that are difficult to see, difficult to contest, and capable of changing after an initial consent encounter.

Informed dependence is proposed as a way to make those conditions ethically visible. It supplements rather than replaces informed consent by asking whether materially important reliance is sufficiently intelligible, monitored, reviewable, and revisable for patient choice to remain meaningful. Its associated duties of disclosure, oversight, continuity, and feasible exit are therefore proportional safeguards, not universal legal entitlements.

The framework remains conceptual and requires empirical, normative, and implementation testing across technologies, populations, and jurisdictions. Its value will depend on whether it helps distinguish ethically significant dependence from harmless background automation without converting consent into technical overload or assuming that every patient seeks the same

degree of human involvement. The central claim is modest but consequential: where AI changes what patients must rely upon, autonomy may require attention not only to what they authorize, but also to the conditions under which they continue to depend.

Acknowledgments: None

Conflict of Interest: None

Financial Support: None

Ethics Statement: None

References

1. Angus DC, Khera R, Lieu T, Liu V, Ahmad FS, Anderson B, et al. AI, health, and health care today and tomorrow: The JAMA Summit report on artificial intelligence. *JAMA*. 2025;334(18):1650-64. doi:10.1001/jama.2025.18490
2. Nong P, Platt J. Patients' trust in health systems to use artificial intelligence. *JAMA Netw Open*. 2025;8(2):e2460628. doi:10.1001/jamanetworkopen.2024.60628
3. Bracic A, Spector-Bagdady K, Towle S, Zhang R, James CA, Price WN. Factors for patient trust and acceptance of medical artificial intelligence. *JAMA Netw Open*. 2026;9(3):e260815. doi:10.1001/jamanetworkopen.2026.0815
4. Busch F, Hoffmann L, Xu L, Zhang LJ, Hu B, García-Juárez I, et al. Multinational attitudes toward AI in health care and diagnostics among hospital patients. *JAMA Netw Open*. 2025;8(6):e2514452. doi:10.1001/jamanetworkopen.2025.14452
5. Rajpurkar P, Chen E, Banerjee O, Topol EJ. AI in health and medicine. *Nat Med*. 2022;28(1):31-8. doi:10.1038/s41591-021-01614-0
6. Hurley ME, Lang BH, Kostick-Quenet KM, Smith JN, Blumenthal-Barby J. Patient consent and the right to notice and explanation of AI systems used in health care. *Am J Bioeth*. 2025;25(3):102-14. doi:10.1080/15265161.2024.2399828
7. Kostick-Quenet K, Lang BH, Smith J, Hurley M, Blumenthal-Barby J. Trust criteria for artificial intelligence in health: Normative and epistemic considerations. *J Med Ethics*. 2024;50(8):544-51. doi:10.1136/jme-2023-109338

8. Rose SL, Shapiro D. An ethically supported framework for determining patient notification and informed consent practices when using artificial intelligence in health care. *Chest*. 2024;166(3):572-8. doi:10.1016/j.chest.2024.04.014
9. Laacke S, Mueller R, Schomerus G, Salloch S. Artificial intelligence, social media and depression. A new concept of health-related digital autonomy. *Am J Bioeth.* 2021;21(7):4-20. doi:10.1080/15265161.2020.1863515
10. Bjerring JCK, Busch J. Artificial intelligence and patient-centered decision-making. *Philos Technol.* 2021;34(2):349-71. doi:10.1007/s13347-019-00391-6
11. McCradden MD, Anderson JA, Stephenson EA, Drysdale E, Erdman L, Goldenberg A, et al. A research ethics framework for the clinical translation of healthcare machine learning. *Am J Bioeth.* 2022;22(5):8-22. doi:10.1080/15265161.2021.2013977
12. Schmidt E, Putora PM, Fijten R. The epistemic cost of opacity: How the use of artificial intelligence undermines the knowledge of medical doctors in high-stakes contexts. *Philos Technol.* 2025;38(1):5. doi:10.1007/s13347-024-00834-9
13. 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-e50. doi:10.1016/S2589-7500(21)00208-9
14. 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/ocz192
15. Platt J, Nong P, Carmona G, Kardia S. Public attitudes toward notification of use of artificial intelligence in health care. *JAMA Netw Open.* 2024;7(12):e2450102. doi:10.1001/jamanetworkopen.2024.50102
16. 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
17. Subasri V, Krishnan A, Kore A, Dhalla A, Pandya D, Wang B, et al. Detecting and remediating harmful data shifts for the responsible deployment of clinical AI models. *JAMA Netw Open.* 2025;8(6):e2513685. doi:10.1001/jamanetworkopen.2025.13685
18. Lawrence K, Kuram VS, Levine DL, Sharif S, Polet C, Malhotra K, et al. Informed consent for ambient documentation using generative AI in ambulatory care. *JAMA Netw Open.* 2025;8(7):e2522400. doi:10.1001/jamanetworkopen.2025.22400
19. Allen JW, Earp BD, Koplin J, Wilkinson D. Consent-GPT: Is it ethical to delegate procedural consent to conversational AI? *J Med Ethics.* 2024;50(2):77-83. doi:10.1136/jme-2023-109347
20. Allen JW, Earp BD, Wilkinson D. AI-assisted consent in paediatric medicine: Ethical implications of using large language models to support decision-making. *J Med Ethics.* 2026;52(6):391-9. doi:10.1136/jme-2024-110624
21. Rudra P, Balke WT, Kacprowski T, Ursin F, Salloch S. Large language models for surgical informed consent: An ethical perspective on simulated empathy. *J Med Ethics.* 2026;52(2):85-90. doi:10.1136/jme-2024-110652
22. Park HJ. Patient perspectives on informed consent for medical AI: A web-based experiment. *Digit Health.* 2024;10:20552076241247938. doi:10.1177/20552076241247938
23. Owens K, Jayaram A, Chowdhury A, Pollak KI, Klotman S, Griffen Z, et al. Patient perspectives on AI-drafted electronic portal messages. *JAMA Netw Open.* 2026;9(7):e2622463. doi:10.1001/jamanetworkopen.2026.22463
24. Cavalier JS, Goldstein BA, Ravitsky V, Bélisle-Pipon JC, Bedoya A, Maddocks J, et al. Ethics in patient preferences for artificial intelligence-drafted responses to electronic messages. *JAMA Netw Open.* 2025;8(3):e250449. doi:10.1001/jamanetworkopen.2025.0449
25. Zohny H, Allen JW, Wilkinson D, Savulescu J. Which AI doctor would you like to see? Emulating healthcare provider-patient communication models with GPT-4: Proof-of-concept and ethical exploration. *J Med Ethics.* 2026;52(e1):e36-e43. doi:10.1136/jme-2024-110256
26. Director S. Does black box AI in medicine compromise informed consent? *Philos Technol.* 2025;38(2):62. doi:10.1007/s13347-025-00860-1
27. Bienefeld N, Boss JM, Lüthy R, Brodbeck D, Azzati J, Blaser M, et al. Solving the explainable AI conundrum by bridging clinicians' needs and developers' goals. *NPJ Digit Med.* 2023;6(1):94. doi:10.1038/s41746-023-00837-4

28. Blackman J, Veerapen R. On the practical, ethical, and legal necessity of clinical artificial intelligence explainability: An examination of key arguments. *BMC Med Inform Decis Mak.* 2025;25(1):111. doi:10.1186/s12911-025-02891-2
29. Marko JGO, Neagu CD, Anand PB. Examining inclusivity: The use of AI and diverse populations in health and social care: A systematic review. *BMC Med Inform Decis Mak.* 2025;25(1):57. doi:10.1186/s12911-025-02884-1
30. Vo V, Chen G, Aquino YSJ, Carter SM, Do QN, Woode ME. Multi-stakeholder preferences for the use of artificial intelligence in healthcare: A systematic review and thematic analysis. *Soc Sci Med.* 2023;338:116357. doi:10.1016/j.socscimed.2023.116357
31. Lekadir K, Frangi AF, Porras AR, Glocker B, Cintas C, Langlotz CP, et al. FUTURE-AI: International consensus guideline for trustworthy and deployable artificial intelligence in healthcare. *BMJ.* 2025;388:e081554. doi:10.1136/bmj-2024-081554