

ETIC 2.0には 何が足りないのか

原理を並べる段階から、測定・比較・外部検証へ進むための研究ロードマップ



Daiki Morimoto / Unknown Lab / ORCID 0009-0009-6444-750X

対象限定レビュー / Version 1.0 / 公開日 2026年8月31日

DOI: 10.5281/zenodo.22168116 / 未査読 / CC BY 4.0

原理を並べる段階から、測定・比較・外部検証へ進むための研究ロードマップ

- Unknown Lab Research Report 08
- Version: 1.0
- 公開日: 2026年8月31日
- 著者: Daiki Morimoto / Unknown Lab
- ORCID: 0009-0009-6444-750X
- ライセンス: CC BY 4.0
- DOI: 10.5281/zenodo.22168116
- 種別: 対象限定レビュー。PRISMA準拠の系統的レビューではない。

30秒で読む

ETIC 2.0は、AIを含むケアを設計するとき、本人の作業参加、家族・専門職・制度との関係、技術・データの統治を一緒に考えるための公開された設計仮説である。

しかし現時点では、検証済みの尺度でも、臨床ガイドラインでも、効果が確立した介入でもない。

Research Report 01〜07は、自動化バイアス、安全と参加の衝突、家族・専門職への負担移転、デジタル格差、作業療法AI研究の偏り、既存モデルとの重なり、AI相談の限界を明らかにした。残った課題は、原理をさらに増やすことではない。

必要なのは、概念を固定し、反証可能な予測を置き、本人・家族・専門職・開発者と評価項目を確かめ、小規模試行、既存手法との比較、外部検証、導入後の変更管理までつなぐことである。

本稿は、その工程を7つの検証ゲートとして提案する。ETIC 2.0が本当に見落としを減らし、生活をよくするかは、この先の研究で確かめる。

「ETICに沿って設計しました」の次に聞くこと

ある開発会議で、チームがETIC 2.0の項目を一つずつ確認したとする。

本人中心。安全。公平性。家族への影響。説明責任。人が止める仕組み。すべてに印が付いた。

それでも、まだ分からないことがある。

- 本人と開発者は、同じ言葉を同じ意味で理解したのか。
- 見落とされていた問題は、本当に見つかったのか。
- 見つけた問題は、仕様変更につながったのか。
- 既存のモデルやチェックリストだけでは見つからない問題だったのか。
- 会議、研修、記録の負担に見合う価値があったのか。
- 本人の作業参加、安全、公平性、家族・専門職負担は改善したのか。
- 別の施設、別のチーム、別の集団でも同じ結果になるのか。
- AIやデータ、画面、運用が変わった後も、以前の評価を引き継げるのか。

項目に印を付けられることは、実装の始まりである。妥当性や効果の証明ではない。

ETIC 2.0が次へ進むには、「大切な原理が書かれている」から、「誰が使っても検証でき、支持されない結果も受け入れられる」へ移らなければならない。

方法

2026年8月30日までに公開されたETIC 2.0 Version 2.1とResearch Report 01〜07を統合し、次の分野の方法論・公的資料を確認した。

- 複雑介入の開発・評価
- 評価項目の内容妥当性
- 患者・市民参画と共同設計の報告
- 医療AIの初期臨床評価とライフサイクル評価
- 実装アウトカム
- AIリスク管理と変更管理

一次資料、査読済み方法論、国際合意、公的ガイダンスを優先した。ETIC 2.0を直接評価した外部研究は見つっていないため、外部資料はETICの妥当性根拠ではなく、検証方法を設計する参照として用いた。

本稿は系統的レビューではない。検索結果の全件スクリーニング、risk-of-bias評価、メタ分析は行っていない。

7本のResearch Reportで、何が残ったのか

これまでの7本は、ETIC 2.0を支持する証拠だけを集めたものではない。むしろ、ETIC 2.0が何を証明しなければならないかを厳しくした。

Research Report	明らかになったこと	次に検証すること
01 自動化バイアス	人はAI提案を信じすぎることがある	ETICレビューで見落としが減り、確認・書き・停止が実際に使われるか
02 転倒と参加	転倒が減ることと、出かけられることは同じではない	安全、参加、本人の意味を同時に測り、価値衝突を扱えるか
03 在宅AIリハ	便益と負担は本人だけでなく家族、専門職、制度へ分布する	負担移転と責任の偏りを検出できるか
04 デジタル包摂	端末、接続、認知、言語、費用、支援が利用可能性を左右する	非AI代替、到達、離脱、集団別の結果を評価できるか
05 作業療法AI	研究は一部の用途と評価法に偏る	近接成果と健康・生活成果を分け、長期実装まで追えるか
06 既存モデルとの比較	MOHO、PEO、ICFにも技術・環境を扱う力がある	ETICが何を追加し、どんな追加負担を生むか
07 AI相談	相談準備と専門職代替、実装と効果は別である	実装可能性、安全、対話、健康・生活成果を段階的に評価できるか

この整理は、各報告の公開版に基づく[16][17][18][19][20][21][22]。

Research Report 06が重要な境界を示した。既存モデルで扱えることを、新しい名前に置き換えただけでは新規性にならない[21]。

ETIC 2.0の追加価値を主張するなら、MOHO、PEO、ICF、作業的公正、WHOやNISTの枠組み、通常的设计レビューと比べ、重要な見落としを再現可能に減らす必要がある。さらに、その便益が研修、会議、記録、監査の追加負担を上回るかも測らなければならない。

現時点のETIC 2.0は、何であり、何ではないか

現時点で最も正確な呼び方は、公開された設計仮説または検証中の作業枠組みである。

ETIC 2.0 Version 2.1は、本人の作業参加、ケア生態系、技術・データ統治を同じ設計レビューへ置く問いを公開している[15]。Care Compassでは、相談準備、緊急時の人への離脱、紙の代替、データ最小化といった要求の一部を実装へ移している[22]。

これは、概念を実装可能な形へ変換できるという一例である。

しかし、次のことはまだ示していない。

- ETICの概念を異なる人が同じように理解できる。
- 評価項目が重要な論点を網羅している。
- ETICを使うと、既存手法より見落としが減る。
- ETICによる設計変更が、安全、参加、公平性、負担を改善する。
- 著者と無関係なチームや異なる文脈でも結果が再現する。
- AIやデータが変わった後も同じ評価が有効である。

したがって、ETIC 2.0を検証済み尺度、臨床ガイドライン、効果が確立した介入と呼ぶことはできない。

なぜ「原理を増やす」だけでは足りないのか

WHOは医療AIについて、自律、安全、透明性、責任、公平性、応答性・持続可能性の原則を示している[9]。NIST AI RMFは、Govern、Map、Measure、Manageを循環させ、役割、文脈、測定、監視を結ぶ[8]。

国際原則はすでに豊富である。

ETIC 2.0の課題は、もう一つ原理を足すことより、原理が設計判断をどう変え、その変更が誰の何を改善し、誰へどんな負担や害を生んだかを検証できるようにすることにある。

SkivingtonらのMRC枠組みは、複雑介入を開発・同定、実行可能性、評価、実装の反復過程として扱う。文脈、プログラム理論、関係者関与、主要な不確実性、改善、経済性を中核に置く[1]。

ETIC 2.0にも、同じ厳しさが必要である。

「この原理は大切だ」から、「この仕組みが、この条件で、この行動を変え、この成果につながるはずだ」へ進む。結果が支持しなければ、定義、項目、対象、使い方を修正する。

原理から検証済みモデルへ進む7つのゲート

本稿は、Research Report 01～07と外部方法論を統合し、7つの検証ゲートを提案する。これは国際標準ではなく、Unknown Labの研究ロードマップ案である。

Gate 1 対象と概念を固定する

最初に決めるのは、ETIC 2.0が何に使われ、誰が使い、何には使わないかである。

「本人中心」「作業参加」「ケア生態系」「技術」「統治」といった言葉は、立場によって意味が変わる。抽象語のままでは、評価者が違う判断をしても、その違いが価値観なのか定義の曖昧さなのか分からない。

既存モデルと重なる部分、補完する部分、未証明の新規候補も分ける。境界事例を複数の関係者が読み、概ね同じ範囲へ分類できなければ、次へ進まず定義へ戻る。

Gate 2 プログラム理論と反証命題を置く

ETIC 2.0を使うと、誰のどの見落としが、どんな設計行動を介して減るのかを示す。

例えば、「通常の設計レビューより、本人の作業参加を損なう仕様を多く発見し、修正案へつなげる」という予測なら、比較対象、発見の定義、重要度、再現性、所要時間を事前に決められる。

結果が出なかったときに、「正しく使われなかった」と説明するだけでは反証にならない。どの結果なら仮説を弱め、項目を変え、対象を狭め、研究を止めるかも先に置く。

Gate 3 評価項目と共同設計の妥当性を確かめる

COSMINの内容妥当性方法論は、構成概念、対象集団、利用文脈を明確にし、項目の関連性、網羅性、理解可能性を評価する[2]。COSMINだけでETIC 2.0全体の理論妥当性は判断できないが、設計レビュー項目を作る段階には使える。

本人、家族、専門職、開発者、管理者が、項目を理解できるか。重要な害や価値衝突が抜けていないか。回答負担は許容できるか。ここを確かめる。

共同設計も、参加人数や会議回数だけでは評価しない。

GRIPP2は、患者・市民参画の目的、方法、結果、影響、振り返りを報告する[3]。Carmanらは、参加を相談から関与、パートナーシップ・共同リーダーシップまでの連続として整理した[4]。Lloydらが93報をまとめた系統的レビューでは、アウトカムを評価した報告は39報（43.3%）で、方法と報告も一貫していなかった[5]。

必要なのは、誰が議題を決め、どの意見で仕様が変わり、どの提案が不採用になり、誰に拒否・停止権があったかを残すことである。GRIPP2は透明性を助けるが、権力共有の質を自動的に保証しない。

Gate 4 小規模に使えるかを試す

次に、評価者間一致、再検査、所要時間、研修負担、欠測を調べる。

重要なリスクを検出できても、評価者ごとに結論が変わり、数時間の訓練と長い会議が毎回必要なら、日常利用は難しい。反対に、使いやすくても、一般的な倫理項目をなぞるだけで具体的な修正につながらなければ、目的を果たしていない。

医療AIの初期評価を扱うDECIDE-AIは、実利用に近い段階で、性能だけでなく人間要因、安全、ワークフロー、変化、不平等、一般化可能性を報告する[6]。FUTURE-AIも、117人・50か国の合意を通じ、公平性、普遍性、追跡可能性、使いやすさ、堅牢性、説明可能性をライフサイクル全体で扱う30の推奨を示した[7]。

これらはETIC 2.0の直接証拠ではない。小規模試行でも、AI単体の精度だけでなく、人と組み合わせた仕事、負担、危害、格差を測る必要があることを示す。

Gate 5 既存手法より何を足すかを比較する

ETIC 2.0の新規性と有用性は、比較なしには判断できない。

同じ設計事例を、ETIC 2.0、既存の作業療法モデル・AIガバナンスチェックリスト、通常の設計レビューで評価する。各群が見つけた重要問題、重複、具体的修正案、所要時間、研修、記録負担を比べる。

ETICだけが見つけた論点が第三者にも重要と評価され、別の評価者でも再現し、設計変更につながるなら、追加価値の信号になる。

反対に、新しい発見が再現せず、会議と記録だけが增えるなら、追加価値の主張を弱める。それはETIC 2.0を否定するためではない。名前や構造を守るより、使える部分を残すためである。

Gate 6 外部チームと異なる文脈で再現する

著者が作り、著者が教え、著者が評価する研究だけでは、期待や解釈の影響を分離できない。

著者と無関係なチーム、複数施設、異なる対象集団で再現する。平均値だけでなく、年齢、障害、言語、デジタル利用、地域、経済条件などの分布を見る。

全体平均が改善しても、支援を受けにくい人の参加が減り、家族負担が増え、安全上の問題が集中するなら前進しない。Research Report 03と04が示した負担移転とデジタル包摂は、外部検証の中心課題である[18][19]。

Gate 7 導入後の変化を管理する

AIは固定された介入ではない。

モデル、学習データ、prompt、閾値、画面、確認者、対象集団が変われば、以前の評価をそのまま引き継がない可能性がある。

NIST AI RMFは、統治、文脈把握、測定、管理を循環させ、ライフサイクルの試験・評価・検証・妥当性確認と継続監視を扱う[8]。FDAのAI医療機器向けPCCP最終ガイダンスは、予定された変更、開発・検証・実装方法、影響評価を結ぶ[12]。日本でもPMDAのIDATEN関連制度が、規制対象医療機器の変更計画を扱う[13]。

これらの規制資料は、すべてのAIサービスへ同じ法的義務を課すものではない。ただし、変更を記録し、影響を予測し、再検証し、必要なら差し戻すという考え方は、ETIC 2.0自身とETICを使う製品の双方に必要である。経済産業省のAI事業者ガイドライン第1.2版も、日本の横断的な統治参照点になる[14]。

研究ロードマップ

7つのゲートは、次の研究段階へ整理できる。

段階	主目的	主な方法	次へ進む条件
0 公開仕様	定義、適用外、理論、版を固定	概念分析、文献照合、公開プロトコル	第三者が概念と反証命題を追える
1 内容妥当性	項目の関連性・網羅性・理解可能性	多主体の質的調査、認知インタビュー、内容妥当性評価	重大な欠落と理解差が許容範囲へ減る
2 実行可能性	使えるか、負担は許容可能か	複数評価者の事例レビュー、時間・一致・欠測・有害事象	リスク検出と改善案が再現し、負担が許容可能
3 比較予備試験	追加価値の信号を探る	ETIC対既存手法対通常設計	見落とし減少と具体的改善が追加負担を上回る可能性
4 外部検証	文脈を越えて再現するか	独立チーム、多施設、層別分析	効果と害が一部集団へ偏らず再現する
5 実装・監視	日常運用と変更に対応するか	実装アウトカム、RE-AIM、監視、版管理	維持可能で、変更後も安全性と公平性を再確認できる

Proctorらは、受容性、採用、適切性、実行可能性、忠実度、費用、浸透、持続可能性を実装アウトカムとして区別した[10]。RE-AIMは、到達、効果、採用、実装、維持を分け、適応と費用も追う[11]。

「ETICレビューを実施できた」「研修を終えた」「項目を埋められた」は重要な実装成果である。しかし、本人の作業参加、安全、公平性、家族・専門職負担が改善したこととは別である。

近接成果、実装成果、健康・生活成果を分ける。これが、効果を過大評価しないための最低条件になる。

何がETIC 2.0を支持し、何が修正を求めるのか

ETIC 2.0の追加価値を支持するには、少なくとも次の結果が必要である。

- 既存手法より、重要な見落としを再現可能に減らす。
- 指摘が具体的な設計変更へつながる。
- 本人の作業参加、安全、公平性、負担の少なくとも一部が改善する。
- 追加時間、研修、記録、会議の負担が測られ、便益との釣り合いを説明できる。
- 著者と無関係な外部チームで結果が再現する。

反対に、次の結果が出れば、定義、項目、対象、使い方を修正し、場合によっては前進を止める。

- 評価者が同じ項目を大きく異なる意味で使う。
- 既存手法と比べ、新しい重要発見が再現しない。
- 会議や記録負担が増えるが、設計変更や成果につながらない。
- 一部集団の参加、アクセス、安全、負担が悪化する。
- 著者チームでは良いが、外部チームで再現しない。
- 版変更後に、以前の安全性・公平性を再確認できない。

支持されない結果を出せない理論は、検証できない。

ETIC 2.0が守るべきものは、現在の名前や図ではない。本人の生活にとって重要な問いを、設計と研究へ残すことである。

Care Compassを「証明」にしない

Care Compassは、ETIC 2.0から生まれた公開実装の一つである。相談準備、緊急時の人への離脱、AIを使わない紙経路、データ最小化といった要求を、画面とワークフローへ変えている。

ここから言えるのは、ETIC的な問いを実装仕様へ変換できるということまでである。

Care Compassが相談行動、健康、生活、作業参加、安全、公平性を改善するかは、まだ分らない。第三者評価もない。自作事例がうまく動くことを、ETIC 2.0全体の内容妥当性、優越性、臨床効果、外部再現性へ広げてはならない。

今後Care Compassを研究に使うなら、製品の評価とETIC 2.0の評価を分ける必要がある。

誰が、次の研究に何を持ち込むか

本人・家族に求めるのは、AI理論の理解ではない。

何が大切か。どんな失敗が生活を壊すか。何を説明してほしいか。どこで訂正、拒否、停止したいか。AIを使わない経路は必要か。こうした判断を、実際の仕様と評価項目へ反映する役割である。

専門職は、適用範囲、確認・上書き、緊急離脱、作業参加、生活上の害、家族・専門職負担を評価する。開発者は、データ来歴、モデル版、設定、性能分布、変更履歴を追跡可能にする。管理者・行政は、調達、責任、異議申立て、監査、費用、停止・撤退を扱う。

研究者は、都合のよい成功指標だけを選ばず、比較、主要評価項目、層別分析、停止条件を事前に固定する。

そして、ETIC 2.0の著者は、自分で作った理論を自分だけで採点しない。統計、倫理、作業療法・臨床、当事者・家族、外部研究者の独立した確認を受ける。

まだ言えないこと

- ETIC 2.0は検証済みのモデルである。
- ETIC 2.0はMOHO、PEO、ICF、既存のAI倫理・リスク管理枠組みより優れている。
- ETIC 2.0を使えば、自動化バイアス、負担移転、デジタル格差、危害を減らせる。
- 共同設計を行えば、本人の意思が反映され、権力差が解消する。
- 評価表を埋めれば、本人中心、安全、公平なAIになる。
- Care CompassはETIC 2.0の効果と安全性を証明している。
- 一度検証したAIの結果を、モデル、データ、設定、対象が変わった後も引き継げる。

結論

ETIC 2.0には何が足りないのか。

足りないのは、もう一つの美しい原理ではない。

定義、反証可能な予測、内容妥当性、共同設計の影響記録、評価者間一致、既存手法との比較、独立外部検証、変更後の再評価を、一つの研究工程としてつなぐことである。

Research Report 01〜07は、ETIC 2.0が答えるべき問いを増やした。本稿は、その問いを7つの検証ゲートへ並べた。

ここから先は、理論を説明する段階ではない。

本人、家族、専門職、開発者、管理者、外部研究者とともに測り、比較し、支持されない結果も公開する段階である。

ETIC 2.0の新規性と価値は、主張の強さではなく、第三者が同じ問いを使い、同じ問題を見つけ、生活に意味のある改善へつなげられるかで判断される。

参考文献

- [1] Skivington K, et al. A new framework for developing and evaluating complex interventions: update of Medical Research Council guidance. *BMJ*. 2021;374:n2061. doi:10.1136/bmj.n2061
- [2] Terwee CB, et al. COSMIN methodology for assessing the content validity of PROMs. *Quality of Life Research*. 2018;27:1159-1170. doi:10.1007/s11136-018-1829-0
- [3] Staniszewska S, et al. GRIPP2 reporting checklists: tools to improve reporting of patient and public involvement in research. *BMJ*. 2017;358:j3453. doi:10.1136/bmj.j3453
- [4] Carman KL, et al. Patient and family engagement: a framework for understanding the elements and developing interventions and policies. *Health Affairs*. 2013;32(2):223-231. doi:10.1377/hlthaff.2012.1133
- [5] Lloyd N, Kenny A, Hyett N. Evaluating health service outcomes of public involvement in health service design in high-income countries: a systematic review. *BMC Health Services Research*. 2021;21:364. doi:10.1186/s12913-021-06319-1
- [6] Vasey B, et al. Reporting guideline for the early-stage clinical evaluation of decision support systems driven by artificial intelligence: DECIDE-AI. *Nature Medicine*. 2022;28:924-933. doi:10.1038/s41591-022-01772-9
- [7] Lekadir K, 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
- [8] National Institute of Standards and Technology. Artificial Intelligence Risk Management Framework (AI RMF 1.0). 2023. doi:10.6028/NIST.AI.100-1
- [9] World Health Organization. Ethics and governance of artificial intelligence for health. WHO; 2021. ISBN 978-92-4-002920-0. Official publication: <https://www.who.int/publications/i/item/9789240029200>
- [10] Proctor E, et al. Outcomes for implementation research: conceptual distinctions, measurement challenges, and research agenda. *Administration and Policy in Mental Health*. 2011;38:65-76. doi:10.1007/s10488-010-0319-7
- [11] Glasgow RE, Estabrooks PE. Pragmatic applications of RE-AIM for health care initiatives in community and clinical settings. *Preventing Chronic Disease*. 2018;15:E02. doi:10.5888/pcd15.170271
- [12] U.S. Food and Drug Administration. Marketing Submission Recommendations for a Predetermined Change Control Plan for Artificial Intelligence-Enabled Device Software Functions. Final Guidance. 2025. FDA-2022-D-2628. Official guidance: <https://www.fda.gov/regulatory-information/search-fda-guidance-documents/marketing-submission-recommendations-predetermined-change-control-plan-artificial-intelligence>
- [13] 独立行政法人医薬品医療機器総合機構. 医療機器の変更計画確認手続制度 (IDATEN) 関連情報. 2026年8月30日確認. 公式ページ: <https://www.pmda.go.jp/review-services/drug-reviews/about-reviews/devices/0039.html>
- [14] 経済産業省. AI事業者ガイドライン第1.2版. 2026. 公式ページ: https://www.meti.go.jp/shingikai/mono_info_service/ai_shakai_jisso/20260331_report.html
- [15] Morimoto D. ETIC 2.0 Version 2.1. Unknown Lab; 2026. doi:10.5281/zenodo.21848456

- [16] Morimoto D. AIの提案を人はなぜ信じすぎるのか. Unknown Lab Research Report 01. 2026. doi:10.5281/zenodo.21327355
- [17] Morimoto D. 転ばない生活と、出かけられる生活は同じではない. Unknown Lab Research Report 02. 2026. doi:10.5281/zenodo.21500071
- [18] Morimoto D. 在宅AIリハは、本人だけを見てはいけない. Unknown Lab Research Report 03. 2026. doi:10.5281/zenodo.21679161
- [19] Morimoto D. 便利なケアからこぼれる人を、どう設計に戻すか. Unknown Lab Research Report 04. 2026. doi:10.5281/zenodo.21752940
- [20] Morimoto D. 作業療法でAIはどこまで使われているのか. Unknown Lab Research Report 05. 2026. doi:10.5281/zenodo.21855906
- [21] Morimoto D. MOHO・PEO・ICFがあるのに、ETIC 2.0は何を足すのか. Unknown Lab Research Report 06. 2026. doi:10.5281/zenodo.21962791
- [22] Morimoto D. AI相談は、専門職の代わりになれるのか. Unknown Lab Research Report 07. 2026. doi:10.5281/zenodo.22065352

利益相反・限界・更新

著者はETIC 2.0とCare Compassの開発者であり、自己評価上の利益相反がある。本稿は査読済み論文ではなく、独立した統計、倫理、作業療法・臨床、当事者・家族レビューをまだ受けていない。

複数分野の方法論を7つのゲートへ対応づけた部分はUnknown Labの統合推論であり、国際合意や規制要件ではない。COSMIN、GRIPP2、DECIDE-AI、FUTURE-AI等に沿うだけで、ETIC 2.0の妥当性、有効性、安全性は証明されない。

AI、規制、ガイドライン、ETIC 2.0の版は変化する。本稿は2026年8月30日までに確認した公開資料に基づく。