

AIをケアの「主体」と呼ぶ前に、 誰の価値をどう確かめるのか

Hwangら (2025) の13人・12時間の共創モデルを読む

論文とETIC #5

AIをケアの 「主体」と呼ぶ前に

誰の価値を、
どう確かめるのか。

13人・12時間の共創モデルを読む

共創は、検証の代わりではない。



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

批判的論評 Version 1.0 2026年8月20日

DOI: 10.5281/zenodo.22004022

未査読 専門職・研究方法確認済み CC BY 4.0

要旨

Hwangらは、年齢20～88歳の13人からなるWorking Groupで、2.5か月に3回、Zoom 5時間と対面7時間の共創を行った。3つのペルソナと、スマートホーム、支援ロボット、インテリジェント車椅子を議論のプロープに使い、尊厳・所属・固有性の3属性、6価値、5設計方針からなる予備的モデルを提示した。

ただし、介入、比較群、アウトカム、効果量、統計解析はなく、論文自身も報告するデータはないとする。募集、役割別人数、分析過程、異論、費用、継続可能性も十分追えない。本稿は、共創された価値を検証済みモデルと扱わず、AIを関係へ作用する技術的主体と記述することと、説明・承認・救済の責任主体にすることを分ける。ETIC 2.0への更新案として、価値責任表、少数意見の保存、5段階検証を提案する。

論文とETIC #5 Version 1.0 / 2026年8月20日公開

AIが転倒を検知し、家族へ知らせ、生活の変化を予測するようになれば、それは単なる道具ではなく、ケアの関係を变える存在になります。

しかし、「関係を変える」とこと、「責任を負える」ことは同じではありません。AIが本人、家族、専門職の間で情報や選択肢を動かしても、誰の価値を優先したのか、誤ったとき誰が止めるのか、本人が異議を申し立てられるのかは、人と組織が決めなければなりません。

Hwangらは、旧来のAI倫理がプライバシーや偏りなど個人と社会の問題に寄りがちだったとして、ケアを複数の人と関係が変化する「動的ケア生態系」として捉えました。13人の学際的ワーキンググループが3回、合計12時間の共創を行い、3つの人間的属性、6つの価値、5つの設計方針をまとめています[1]。

これはETIC 2.0の重要な先行研究です。同時に、効果を検証した研究ではありません。論文自身が「報告するデータはない」とし、モデルの妥当性、使いやすさ、見落としの削減、生活への効果を測っていません。共創された指針を、検証済みの設計モデルとして扱わないことが今回の中心です。

30秒で分かる研究の要点

項目	内容
論文	Hwangらの人間中心AI AgeTechに関する予備的指針モデル
共創者	ワーキンググループ13人。全員が共著者
構成	年齢20-88歳、女性11人・男性2人。高齢者、認知症当事者、ケアパートナー、研究者、医療者、起業家等の経験を含む
期間	2.5か月。3回の生成的セッション
時間	Zoom 5時間、対面7時間、合計12時間
方法	ビジュアル素材、3つのペルソナ、スマートホーム・支援ロボット・知能化車椅子の仮想場面を使う共創
成果物	3属性、6価値、5設計方針
比較・効果	介入群、比較群、臨床・生活アウトカム、効果量はない
読み方	関係と価値を可視化する設計仮説。妥当性や有効性を実証した尺度ではない

English Summary

Hwang and colleagues convened a 13-member working group, aged 20 to 88, for three co-creation sessions over 2.5 months: five hours online and seven hours in person. Using three personas and three AgeTech probes, the group produced a preliminary model with three humanistic attributes, six values, and five design orientations. The paper reports no intervention, comparator, outcome measure, effect size, or statistical test, and explicitly states that there are no data to report. Recruitment, role counts, analytic trace, dissent, costs, and sustainability remain insufficiently reported. This commentary treats co-creation as a source of testable questions, not as validation, and separates AI as a relational technical actor from accountable human and organizational actors.

Keywords: AgeTech, artificial intelligence, dynamic care ecosystem, co-creation, older adults, AI governance, ETIC 2.0

本稿の位置づけ

本稿は公開済み査読論文をUnknown Labが批判的に読み解いた未査読の論評記事である。著者の結論、共創過程から確認できる観察結果、因果効果として言えないこと、Unknown Labの論評、ETIC 2.0への支持・挑戦・更新案を分けて記載している。個別の診断、治療、リハビリテーション、AI製品の導入を指示するものではない。

補強記事06から、何を一步深めるのか

直前の補強記事「MOHO・PEO・ICFがあるのに、ETIC 2.0は何を足すのか」は、ETIC 2.0が既存理論を置き換える新理論ではなく、作業・環境・参加の理解をAIライフサイクルの統治へつなぐ設計・点検レンズだと位置づけました[6]。

その記事でHwangらの論文は、「AIを含む動的ケア生態系」の直接の先行概念として引用されています。つまり、ケア生態系、時間変化、AIが人間関係へ作用するという発想はETICだけのものではありません。

今回は、その土台になった論文自体を詳しく読みます。問うのは、考え方が魅力的かどうかだけではありません。

- 誰が、どのように参加したのか。
- 価値や設計方針は、どの手順で導かれたのか。
- 何が観察され、何がまだ仮説なのか。
- 「AIも主体になる」という表現が、責任の所在を曖昧にしないか。
- ETIC 2.0は、この先行モデルをどう検証可能な設計手順へ変えるべきか。

この順番で読むと、ETIC 2.0への支持と、ETIC自身が受けるべき批判を同時に見られます。

この論文は、先行するAI倫理へ何を足したのか

著者らは、AI AgeTechの倫理議論が、プライバシー、同意、自律、差別、アルゴリズムの偏りなど、個人と社会への影響を主に扱ってきたと整理します。これらは不可欠ですが、ケアは一对一の利用場面だけで完結しません。

たとえば、スマートホームが高齢者の活動、気分、健康行動を家族へ要約するとします。本人には見守りと連絡の利点がある一方、家族は新しい通知対応を担い、専門職は誤検知の判断を求められます。安全を高める設定が、本人のプライバシーや自由を縮めることもあります。

Hwangらの新しさは、AIを「高齢者が使う製品」とだけ見ず、次の組み合わせとして扱ったことです。

- 高齢者、家族、友人、専門職などのケア主体。
- 誰が誰をどの頻度で、どのように支えるかという関係。
- 同居、別居、遠距離などの生活配置。
- 役割、責任、日課、予定、技術利用を含むケア配置。
- 時間と状況による価値、能力、関係、技術への信頼の変化。

さらに、AIが予測し、選び、半自律的に反応することで、人間同士の関係を媒介し、ときに関係上の一主体のように振る舞う可能性を示しました。個人の権利だけでなく、価値の衝突と関係の変化を設計対象にした点が、先行議論との差です。

誰が共創したのか

ワーキンググループは13人で、全員が論文の共著者です。論文の図2には、参加者が自己申告した経験と背景がまとめられています[1]。

領域	原文で示された範囲
経験・役割	高齢者、認知症当事者、家族ケアパートナー、研究者、学生、医療者、起業家
分野	医学、看護、作業療法、リハビリテーション科学、ビジネス、コミュニケーション、哲学
年齢	20-88歳
性別	女性11人、男性2人
文化的背景	東アジア、東南アジア、南アジア、南アフリカ、北米・欧州・ユーラシア系白人等
生活・ケア配置	独居、配偶者との同居、核家族、多世代家族、地域居住、民間退職者住宅、公的長期ケア施設
地理的文脈	カナダ、米国、香港、フランス、トルコ

この構成は、単一職種だけで作った専門家パネルより広い視点を持ちます。当事者やケアパートナーが、評価対象ではなく共著者として参加したことも重要です。

ただし、各役割に何人いたのか、誰がどの属性を重ね持っていたのかは示されていません。たとえば、13人中何人が高齢者で、何人が認知症当事者または家族ケアパートナーだったかは確定できません。年齢の中央値、社会経済状況、障害、デジタル経験の分布も報告されていません。

「多様な13人」という説明は、参加者の幅を示しますが、想定される利用者全体を代表することを意味しません。

3回・12時間の共創で何をしたのか

1人の研究者が、2.5か月にわたり3回の生成的セッションを進行しました。Zoomで5時間、対面で7時間、合計12時間です。参加者はビジュアル素材を見ながら、言葉にしにくい経験、感情、将来の設計可能性を話し合いました[1]。

最初に、ケア生態系とAI AgeTechの共通理解をつくりました。その後、異なる生活機能レベルを持つ3つのペルソナを小グループで具体化しました。扱った技術は、スマートホーム、支援ロボット、知能化車椅子です。

仮想場面では、技術の利点だけでなく、価値の衝突を検討しています。

- 安全のための活動監視は、本人のプライバシーと選択を狭めないか。
- 家族の不安を減らす通知が、家族の確認負担を増やさないか。
- 知能化車椅子が移動の自由を広げても、職員の責任や他者の私的空間への侵入をどう扱うか。
- AIは本人、家族、施設の誰に忠実であるべきか。
- 体調や関係が変わったとき、自動化の程度を変えられるか。

この方法は、既に起きた出来事の頻度を測るのではなく、将来起こり得る問題を具体的な場面で考えるための方法です。したがって、ここで出た事例は実証された効果や事故ではなく、設計上の仮説と論点です。

共創から生まれた3属性・6価値・5方針

著者らは、人間的なAI AgeTechの指針を三層で整理しました。

層	構成要素	要点
人間的属性	尊厳、所属、固有性	何を守る設計か
価値	保護、選択、共感、協働、意味、適時性	関係の中で何が衝突・補強するか

層	構成要素	要点
設計方針	今誰に何が重要か理解する、変化する生態系を前提にする、AIを含む共感的協働を育てる、利用しながら意味を共創する、固有性に合わせ包摂的に設計する	開発と評価で何を行うか

6つの価値は独立したチェック項目ではありません。保護を強めると選択が狭まり、家族の安心が本人の監視感につながるように、一方の価値を最大化すると別の価値へ費用が生じます。

「適時性」は特に重要です。本人の体調、家族の余力、技術への慣れ、生活の優先順位は変わります。導入時に同意を得ただけで終わらず、設定、自動化、役割、データ共有を見直す必要があります。

一方、この三層図は評価尺度ではありません。各価値の得点法、合格基準、再評価時期、停止条件は定義されていません。価値の一覧があることと、現場で一貫して判断できることを分けて考える必要があります。

どのように「検証」したのか

結論から言うと、この論文はモデルの効果を検証していません。

介入群も比較群もなく、モデルを使ったチームと使わないチームを比べていません。設計上の見落とし、本人の理解、意思決定の質、導入時間、事故、生活の変化、費用、公平性を測るアウトカムもありません。統計検定、効果量、信頼区間もありません。

また、著者らはワーキンググループでアイデアを生成、統合、共同執筆したため、「報告するデータはなく、データ公開と事前登録は該当しない」と記しています[1]。

したがって、「13人を対象に12時間の介入を行い、モデルの効果を示した研究」と読むのは誤りです。正確には、13人の共著者が12時間の生成的セッションを行い、既存理論と経験を統合して予備的指針を提案した論文です。

数字が示すものと、示さないもの

この論文の数値は、臨床効果ではなく制作過程と成果物の規模を示します。

数値	直接示すこと	示さないこと
13人	ワーキンググループの人数	高齢者やケアパートナー全体の代表性
20-88歳	参加者の年齢範囲	年齢群別の意見差、中央値
女性11・男性2	性別構成	性別多様性が十分であること
3回・2.5か月	共創の回数と期間	長期利用で価値がどう変化するか
5時間＋7時間	Zoomと対面のセッション時間	準備、分析、執筆、人件費の総量
3ペルソナ	検討した仮想の生活場面	実在する利用者での妥当性
3属性・6価値・5方針	提案されたモデルの構成	内容妥当性、信頼性、有効性

この区別は、数値が少ないから論文の価値が低いという意味ではありません。概念モデルの初期段階では、問題の捉え方と言語をつくること自体に価値があります。問題は、初期モデルを実証済みの根拠として引用するときに起きます。

著者の結論、観察結果、因果効果を分ける

著者の結論

著者らは、人間的なAI AgeTech設計が倫理的要請を満たすだけでなく、高齢者とケアパートナーにとって重要なことを支え、包摂的な設計過程を通じて技術リテラシーと協働を高め、社会的・経済的資源を最適化し得ると論じます。その上で、モデルは予備的であり、多様なケア生態系を含む今後の研究で妥当性確認と発展が必要だと明記しています[1]。

論文から直接確認できる観察結果

13人の共著者が3回の共創を行い、仮想場面と既存文献を統合して、3属性、6価値、5設計方針を作りました。安全と選択、本人と家族、現在の能力と将来の変化など、価値が衝突する設計質問を具体化しました。

因果効果として言えないこと

- このモデルを使うとAI AgeTechが高齢者に採用されやすくなる。
- 本人と家族の関係や生活の質が改善する。
- プライバシー侵害、事故、偏り、責任の見落としが減る。
- 技術リテラシーや協働が高まる。
- 社会的・経済的資源が最適化される。
- ETIC 2.0の有効性が支持された。

これらは今後検証すべき仮説です。論文の結論にある強い期待表現と、実際に測定された範囲を分けて読む必要があります。

Unknown Labの批判的論評

1. 共創者の選び方と構成が十分に追えない

全員が共著者であることは、当事者を単なる情報提供者にしない点で強みです。しかし、募集経路、選定基準、参加を断った人、謝礼、各回の出席者数、小グループの構成は報告されていません。

参加者と研究者の役割が重なる共創では、誰が議題を決め、誰の経験がモデルへ残り、誰の異論が残らなかったかが重要です。患者・市民参画の報告指針GRIPP2は、参画の目的、方法、影響、限界を追える形で報告することを求めています[3]。Hwang論文は経験の幅を図示していますが、参画が各要素へどう影響したかは十分追えません。

2. 価値を導いた分析の足跡が薄い

論文は、ビジュアル素材、ペルソナ、仮想場面を使ったことを説明します。一方、発言を記録・逐語化したか、どのように整理したか、誰がコード化・統合したか、意見の不一致をどう残したかは記載されていません。

質的研究の報告指針COREQは、研究者の立場、対象者選定、データ収集、分析、引用による裏づけなどを確認します[4]。この論文は質的面接研究を名乗らず、データも報告しないため、COREQへ機械的に当てはめるべきではありません。それでも、6価値がどの経験からどう抽出されたかを第三者が監査できないという問題は残ります。

少数意見を合意へ丸めると、最も保護が必要な人の異論が消えることがあります。共創では「全員が合意した価値」だけでなく、解けなかった対立と少数意見を成果物に残す必要があります。

3. 「AIも主体」という比喩は有用だが、責任主体とは分けるべき

AIを関係上の主体として見ると、AIが人の行動、役割、信頼、会話へ与える影響を見落としにくくなります。単なる機器として扱うより、設計上は有用です。

しかし、AIが行動するように見えることと、道徳的・法的責任を負うことは別です。AIは家族へ通知できますが、通知基準を決め、誤りを監視し、停止し、損害へ応答する責任は、人と組織に残ります。

原論文もAIの責任、説明責任、忠実性を明確にすべきだと述べます。それでも「AIの責任」と書くと、開発者、導入組織、専門職、家族の責任がAIへ移ったように見える危険があります。ETIC 2.0では、AIを関係へ作用する技術的主体として記述しても、説明・承認・救済を担う責任主体には数えない方がよいでしょう。

4. 費用、人員、継続可能性はほぼ測られていない

論文から分かる人員情報は、主に13人の共著者と1人の進行役です。セッションは12時間ですが、準備、ビジュアル制作、分析、執筆、調整の時間は示されていません。研究資金はSchwartz Reisman InstituteのFaculty Fellowshipですが、金額と用途は記載されていません。利益相反はないと報告されています[1]。

したがって、この共創を別の地域や組織で再現する費用、必要職種、参加者負担、継続的な再評価体制は判断できません。「社会的・経済的資源の最適化」は測定結果ではなく将来仮説です。

ETIC 2.0への支持・挑戦・更新案

区分	この論文から受け取ること	ETIC 2.0で明確にすること
支持	本人だけでなく家族、専門職、組織、技術の関係を見る	ケア生態系マップに役割、関係、支援頻度、生活配置を残す
支持	安全、選択、意味、関係は衝突し、時間で変わる	導入時だけでなく変更時・生活変化時に再評価する
支持	当事者とケアパートナーを共著者として含める	本人参加を情報提供で終わらせず、決定権と異論記録を設計する
挑戦	予備的モデルで効果検証がない	ETICも「有効なモデル」ではなく未検証の設計仮説と明記する
挑戦	AIを主体と呼ぶと責任が曖昧になり得る	技術的作用と人・組織の説明責任を別列で記録する
挑戦	価値の抽出過程と少数意見が追にくい	主張と根拠、対立、未解決点、変更理由の監査証跡を残す
更新	価値の問いを統治手順へ変換する必要	責任者、データ、判断権、検証、救済、停止、代替を結びつける

更新案1: 価値から統治までを一行で追える表をつくる

価値の一覧だけでは、誰が何をするかが決まりません。ETIC 2.0のレビュー記録に、次の最小列を追加します。

列	記録する問い
価値・作業目標	本人にとって今何が重要か
影響を受ける人	本人、家族、専門職、同居者、地域の誰か
AI・データの作用	何を収集、推定、通知、自動化するか
決定権	誰が設定、共有、上書き、拒否を決めるか
責任者	誰が正確性、安全、説明、事故対応を担うか
根拠	何をもって設計が妥当と判断したか
再評価契機	体調、関係、住環境、モデル、目的の何が変わったら見直すか
救済・代替	訂正、異議申立て、停止、非AI経路へどう戻るか

この表は、Hwangらの価値と関係の問いを、NIST AI RMFが求めるGovern、Map、Measure、Manage、人間監督、継続監視、停止へつなぐ役割を持ちます[5]。

更新案2: 少数意見を成果物として残す

共創の結論を一つにまとめる前に、次を記録します。

- 参加者の役割と意思決定権。
- 合意した点。
- 解けなかった価値の衝突。
- 少数意見と、その意見を採用しなかった理由。
- 本人または影響を受ける人が後から修正できる経路。

合意率の高さを目標にすると、権力の弱い参加者が沈黙しただけでも「合意」に見えます。ETIC 2.0は、合意よりも異論が安全に残ることを品質指標に含めるべきです。

更新案3: 予備的モデルを5段階で検証する

段階	検証すること	例となる指標
1. 内容妥当性	必要な論点が含まれるか	多様な当事者・専門職による不足項目、理解一致
2. 使用可能性	現場で理解し実行できるか	所要時間、迷い、記録漏れ、負担
3. 比較性能	既存の組合せより何を追加できるか	PEO+NIST等と比べた見落とし、重複、誤警報
4. 前向き実装	実際の設計判断が変わるか	設定変更、責任明確化、異議申立て、停止判断
5. 生活・公平性	本人の作業参加と負担へどう影響するか	意味ある活動、信頼、家族・職員負担、格差、害

ETIC 2.0は、少なくとも第1・第2段階を通るまで「検証済みモデル」と呼ばない。第3段階では、ETICを使わなくても同じ見落としを防げる可能性を公平に比較します。

現場でこのモデルを使うなら、最初に確認すること

Hwangらの6価値を、会議の理念確認で終わらせないための質問です。

確認領域	現場での問い
今の目標	本人は今、何を続けたいのか。家族・専門職の目標とどこが違うか
保護と選択	安全を高める代わりに、何の自由・秘密・尊厳を失うか
意味と関係	AI導入で誰との会話、役割、頼り方が変わるか
適時性	体調、慣れ、生活配置が変わったとき、設定を誰が見直すか
AIの作用	AIは何を推定し、誰へ通知し、何を自動で行うか
責任	誤検知、見逃し、漏えい、過剰介入へ誰が応答するか
異議と停止	本人や同居者が拒否、訂正、停止できるか
代替	AIを使わない、または支援付きの経路へ不利益なく戻れるか
評価	導入率ではなく、作業参加、関係、負担、公平性、害を測るか

AIが「賢いか」より先に、誰の生活をどの関係の中で支え、何が変わったらやめるかを決めます。

次に読む論文

次に読むべき論文は、スマートホームの倫理を導入前の価値確認で終わらせず、開発、評価、実装、継続利用の全期間で管理する必要を論じたWangらの論文です。

Wang RH, Tannou T, Bier N, Couture M, Aubry R. Proactive and Ongoing Analysis and Management of Ethical Concerns in the Development, Evaluation, and Implementation of Smart Homes for Older Adults With Frailty. JMIR Aging. 2023;6:e41322. doi:10.2196/41322 (https://doi.org/10.2196/41322)[2].

Hwangらが3属性、6価値、5設計方針を示したのに対し、Wangらはプライバシーと安全、個人的・関係的自律、同意と意思決定支援、社会的包摂、偏見・差別、公平なアクセスという6領域を、プロジェクト全期間で扱う枠組み、問い、報告資源、チーム教育、当事者教育へ接続します。

この論文も見解論文であり、効果検証ではありません。それでも、「何を大切にするか」から「誰が、いつ、何を管理するか」へ問いを進めるため、ETIC 2.0の次の読書として適切です。

結論

Hwangらの論文は、AI AgeTechを本人と機器の対一関係から解放し、家族、専門職、生活配置、価値の衝突、時間変化を含む動的ケア生態系として捉えました。13人の共著者が12時間の共創を行い、3属性、6価値、5設計方針を示したことは、AIを人間の関係と生活の中で設計するための重要な出発点です。

しかし、この論文はモデルの効果を検証していません。参加者選定と分析手順の報告は限られ、比較群、アウトカム、効果量、費用、継続可能性のデータはありません。「本人に重要なことを届ける」「資源を最適化する」は、まだ結論ではなく仮説です。

ETIC 2.0が受け取るべきなのは、ケア生態系という言葉だけではありません。

価値の衝突を残すこと。AIの技術的作用と、人・組織の責任を分けること。本人の目標から、データ、決定権、根拠、再評価、救済、停止、代替までを一行で追えること。そして、既存の枠組みと比べて本当に見落としを減らすかを検証することです。

共創は、検証の代わりではありません。共創でつくった問いを、比較可能で修正可能な設計手順へ変える。そこからETIC 2.0の次の段階が始まります。

記事の位置づけ

本稿は、公開済み査読論文をUnknown Labが批判的に読み解いた未査読の論評記事です。原論文の提案、共創過程から確認できること、因果効果として言えないこと、Unknown LabによるETIC 2.0への支持・挑戦・更新案を区別しています。個別の診断、治療、リハビリテーション、AI製品やスマートホームの導入を指示するものではありません。

執筆: Unknown Lab 編集部 (AI支援) / 専門職・研究方法確認済み

参考文献

- 1 Hwang AS, Tannou T, Nanthakumar J, et al. Co-creating Humanistic AI AgeTech to Support Dynamic Care Ecosystems: A Preliminary Guiding Model. *Gerontologist*. 2025;65(1):gnae093. doi:[10.1093/geront/gnae093](https://doi.org/10.1093/geront/gnae093)
- 2 Wang RH, Tannou T, Bier N, Couture M, Aubry R. Proactive and Ongoing Analysis and Management of Ethical Concerns in the Development, Evaluation, and Implementation of Smart Homes for Older Adults With Frailty. *JMIR Aging*. 2023;6:e41322. doi:[10.2196/41322](https://doi.org/10.2196/41322)
- 3 Staniszewska S, Brett J, Simera I, et al. GRIPP2 reporting checklists: tools to improve reporting of patient and public involvement in research. *BMJ Open*. 2017;7:e013318. doi:[10.1136/bmjopen-2017-013318](https://doi.org/10.1136/bmjopen-2017-013318)
- 4 Tong A, Sainsbury P, Craig J. Consolidated criteria for reporting qualitative research (COREQ): a 32-item checklist for interviews and focus groups. *Int J Qual Health Care*. 2007;19(6):349-357. doi:[10.1093/intqhc/mzm042](https://doi.org/10.1093/intqhc/mzm042)
- 5 Tabassi E. Artificial Intelligence Risk Management Framework (AI RMF 1.0). NIST AI 100-1. 2023. doi:[10.6028/NIST.AI.100-1](https://doi.org/10.6028/NIST.AI.100-1)
- 6 Morimoto D. MOHO・PEO・ICFがあるのに、ETIC 2.0は何を足すのか: 作業・環境・参加の蓄積と、AI・データ・統治をつなぐ設計レンズ. Unknown Lab Research Report 06. 2026. doi:[10.5281/zenodo.21962791](https://doi.org/10.5281/zenodo.21962791)

本稿は公開済み査読論文を批判的に読み解いた未査読の論評記事であり、ETIC 2.0の有効性を検証した研究ではない。個別の診断、治療、リハビリテーション、AI製品の導入を指示するものではない。

執筆: Unknown Lab編集部 (AI支援) / 専門職・研究方法確認済み

Copyright (C) 2026 Daiki Morimoto. Licensed under CC BY 4.0.