



FG-AI4H会合報告

慶應義塾大学 大学院 政策メディア研究科 特任教授 **川森 雅仁** かわもり まさひと



1. はじめに

ITUとWHOが共同で開催しているFocus Group on AI for Health (FG-AI4H) は、コロナ禍で、再びオンライン開催となり、2021年9月28日から30日にM会合が開催された。

2. 今回の主要な結果

今回の会合での特筆すべきことは、WHOがAIに関するガイドラインを公表したことだろう。FGはWHOとアドホックグループを通じて議論に参加し、またそのためのワークショップを開催した。その結果を含めて今回寄与文書として公表された。この文書はFGの出力文書の倫理規定に反映することになるだろう。

またコロナ禍に対応するためのガイドラインの作成も特筆される。

2.1 AI4Hの倫理に関するガイドライン

AIの健康領域における使用についての倫理的検討はWHOにとっての重要課題である。その意味で、WHOが6月に発表した「健康のための人工知能の倫理とガバナンス (“Ethics and Governance of Artificial Intelligence for Health”)」を寄与文書として提案したことの意義は大きい。この文書は〈<https://www.who.int/publications/i/item/9789240029200>〉から入手可能になっている。

WHOのウェブサイトにはテドロスWHO長官からのビデオメッセージが載せられており、WHOがこのガイドラインを重要視していることがうかがえる。

このガイドラインのすべてあるいはその抄訳的なものが、



■ 図1. AIについてのWHOテドロス長官からのビデオメッセージ

AI技術を開発及び使用する際の倫理的要求条件に関する「DEL-01 AI4Hの倫理に関する考慮事項 (AI4H ethics considerations)」としてFG-AI4Hの出力文書に含まれることになるだろう。

このガイドラインは、160ページに及ぶ大部のものでWHOがこの分野をいかに重要視しているかを示すといえよう。



■ 図2. WHOのAI倫理ガイドラインの表紙

ガイドラインは全体で9つのセクションに分かれている。以下が各セクションの概説である。

セクション1では、このトピックへのWHOの関与の根拠と、レポートの調査結果、分析及び推奨事項の対象読者について説明している。

セクション2と3は、その方法とアプリケーションを通じて健康のためのAIの定義を与えている。

セクション2: AI。このセクションは、あまり技術的にならないレベルでの定義が与えられている。AI技術の一部として、機械学習を含めている。また、生物医学または健康のビッグデータを構成するデータのソースを含む「ビッグデータ」も定義している。このガイドラインでは、AIは以下のように定義されている。

「“Artificial intelligence” generally refers to the performance by computer programs of tasks that are commonly associated with intelligent beings (「人工知能」とは、一般に、インテリジェントな存在に一般的に関連付けられているタスクのコンピュータープログラムによるパフォーマンスを指す。)」



この定義では、一般のコンピュータープログラムとの区別がつかないが、幅広くAIを対象とするためかもしれない。執筆者にはAI専門家もいるので、この定義で問題がなかったのだと思われる。

セクション3：健康のためのAIの応用。このセクションは、医療、健康研究、医薬品開発、健康システムの管理と計画、公衆衛生監視など、健康のためのAI技術の包括的な分類と例を示す。これらは低中所得国で使用されるアプリケーションを含む。

セクション4：健康のためにAIを使用するための重要な倫理的原則。このセクションは、健康のためのAIの使用に適用される、あるいは適用される可能性のある法律、ポリシー及び原則をまとめている。これらの中には、AIに適用される人権義務、データ保護法の枠組みの役割、さらには、その他の健康データに関する法規やポリシーが含まれている。また、このセクションでは、健康のためにAIを使用するための倫理原則を推奨している。いくつかの枠組及び倫理規範の基となる生命倫理、法律、公共政策、規制などの役割について説明している。

セクション5：健康のためにAIを使用するための重要な倫理的原則。このセクションでは、専門家グループが健康のためのAIの開発と使用を導くものとして特定した倫理原則について説明している。

セクション6：健康のためにAIを使用することへの倫理的課題。このセクションは、これらの指針となる倫理原則を適用できる専門家グループによって特定及び議論された以下のような倫理的課題を示す。

- ・ AIの使用の必要性
- ・ AIとデジタルディバイド
- ・ データの収集と使用
- ・ AIによる意思決定に対する説明責任（アカウンタビリティ）を含む責任
- ・ 自律的な意思決定
- ・ AIに関連するバイアスと差別
- ・ 安全性とサイバーセキュリティに対するAIのリスク
- ・ ヘルスケアにおける労働と雇用に対するAIの影響
- ・ ヘルスケア向けAIの商業化における課題
- ・ AIと気候変動

レポートの最後のセクションでは、適切なガバナンスフレームワークを含む、健康のためのAIの倫理的使用を促進するための法的、規制及び非法的措置を特定している。また勧告が述べられている。

セクション7：健康のためにAIを使用するための倫理的アプローチの構築。このセクションは、様々な利害関係者が、倫理的規範及び法的義務を予測または満たすための倫理的慣行、プログラム及び措置をどのように導入できるかを検討する。これらは以下を含む。

- ・ AI技術の倫理的で透明性のある設計
- ・ 公衆の関与と役割のためのメカニズム及び医療提供者と患者との信頼性を実証するためのメカニズム
- ・ インパクト調査；ヘルスケアのためのAIの倫理的使用に関する研究課題

セクション8：健康のためのAIの責任体制。健康でのAIの使用が増えるにつれて責任体制がどのように進化するかについての議論である。これには、医療に責任を割り当てる方法が含まれる。

プロバイダー、テクノロジープロバイダー及びAIテクノロジーを選択する医療システム、または病院及び責任のルールが開業医のAIの使用法にどのように影響するか。このセクションでは、機械学習アルゴリズムが製品であるかどうか、AIテクノロジーによって被害を受けた個人を補償する方法、規制当局の役割及び中低所得国特有の課題についても検討している。

セクション9：健康のためのAIのガバナンスのための枠組みの要素。このセクションは、健康のためのAIのガバナンスフレームワークの要素を示している。「健康のガバナンス」とは、国民皆保険（Universal Health Coverage）を助ける国の健康政策目標を達成するために、政府及び国際保健機関を含む他の意思決定者による運営及び規則制定のための一連の機能を指す。このセクションでは、開発中または既に成熟しているいくつかのガバナンスフレームワークを分析している。議論されるフレームワークは、データのガバナンス、制御と利益の共有、民間セクターのガバナンス、公共セクターのガバナンス、規制上の考慮事項、政策監視と立法の役割及びAIのグローバルガバナンスである。

ここで「健康のためのAI」とされている用語はAI for Healthそのものであり、このガイドラインがFG-AI4Hの活動を意識していることは明らかである。

2.2 健康のためのAIの規制について

DEL-02 AI4H規制のベストプラクティス（AI4H regulatory best practices）も今回大きくアップデートされた。

この文書も倫理に関する出力文書同様、WHOが中心になって作成しているAI4Hのものである。



現在、以下の6つの項目が議論されている。

1. ドキュメントと透明性
2. リスク管理とライフサイクル
3. データ品質
4. 分析的及び臨床的検証
5. エンゲージメントとコラボレーション
6. プライバシーとデータ保護

ここでは、かなり明確化され、推奨されている3、4、5、6の要求条件を以下に報告する。

2.2.1 データ品質に関する推奨事項

データ品質に関してはビッグデータでよく引用される「10のV」といわれるデータの属性に従って要求条件を考えている。WHOが挙げている「10のV」とはVelocity（速度）；Vocabulary（語彙）；Validity（妥当性）；Veracity（信憑性）；Volume（ボリューム）；Vagueness（曖昧性）；Variability（ばらつき）；Venue（場所）；Variety（バラエティ）；Value（価値）という10の性質である。これは必ずしも一般的な10の性質ではないがWHOがこれらを重要視しているということを示していると思われる。



■図3. WHOが挙げているビッグデータの10Vs

以下の4つの要求条件が挙げられている。

1. 開発者は、利用可能なデータが、意図した目標を達成できるシステムの開発をサポートするのに十分な品質であるかどうかを判断すること。
2. トレーニングデータ、アルゴリズム、その他のシステム設計上の問題によって、バイアスやエラーといった問題が増幅されることがないように、開発者はAIソリューションのプレリリース実証実験を厳密に行うことを検討すること。
3. 丁寧な設計や迅速なトラブルシューティングは、データ品質の問題の早期特定に役立つ。これにより、結果として生じる可能性のある危害を防止または改善できる可能性がある。

4. ヘルスケアデータとそれに関連するリスクで発生するデータ品質の問題を軽減するために、利害関係者は、高品質のデータソースの共有を促進するデータエコシステムの作成に引き続き取り組むこと。

2.2.2 分析的及び臨床的検証のための推奨事項

以下の5つの要求条件が挙げられている。

1. ツールの使用目的の明確な文書を提供すること。データサイズ、設定と母集団、入力と出力のデータ、母集団の人口構成など、AIツールを支えるトレーニングデータセットの構成の詳細を透過的に文書化し、ユーザーに提供すること。WG-RCは、専門知識とベストプラクティスを共有するためのフォーラムを提供する。
2. トレーニングデータを超越するパフォーマンスを実証するには、独立したデータセットでの外部の分析的検証が必要である。このために、ツールの使用対象となる母集団とその設定を反映する必要があり、外部データセットとパフォーマンスの評価法について詳細なドキュメントを提供すること。
3. 臨床的検証には、リスクに基づいた段階的な要件のセットが推奨される。ランダム化臨床試験は、臨床成績を比較評価するための最善の判断基準であり、最もリスクの高いツールや、最高水準のエビデンスが必要な場合に適していると思われる。
4. 予測的バリデーション（Prospective Validation）は、現実世界での実用展開と実装トライアルであるが、他のリスククラスに関しては、指定されたエンドポイント上で有意な改善結果が認められた比較集団が適切である可能性もある。
5. 実装展開以降に有害事象が発生しないよう、より徹底的な監視期間をおくようにすること。本項目については、臨床評価WGがさらなる検討を行っている。

2.2.3 コラボレーションとエンゲージメントに関する推奨事項

ここでのエンゲージメントとは、健康領域でのAI使用に関わるステークホルダー（利害関係者）間での意思疎通や仕事に対するコミットメント、というようなことである。この領域では以下の要求条件が推奨されている。

AIのイノベーションと発展段階での主要なステークホルダー間で、適切な形でエンゲージメントとコラボレーションを促進するような、アクセシブルなプラットフォームを検討す

ること。

これらステークホルダーには、AIやMLの開発者、デバイスメーカー、医療従事者、政策立案者、監督機関が含まれるが、これらに限定されない。

エンゲージメント及びコラボレーションプラットフォームは、AI規制の監視プロセスを合理化する上で重要な役割を果たす一方で、ユーザーコミュニティ現場でのAI利用の仕方を変化させるような進歩を加速する可能性がある。

2.2.4 プライバシーとデータ保護に関する推奨事項

この重要な領域では、以下の3つの一般的要件条件が挙げられている。

1. AIソリューションの設計と展開の際には、プライバシーとデータ保護を考慮すること
2. 開発プロセスの早い段階で、開発者は適用されるデータ保護規制とプライバシー法を理解し、開発プロセスがそのような法的要件を満たしていることを確認すること
3. コンプライアンスプログラムは、リスクに対処すること。また潜在的な危害及び施行環境を考慮したプライバシーとサイバーセキュリティの実践方法と優先順位を開発すること

2.3 データ・アノテーション仕様

データのアンノテーションは機械学習にとっては重要な項目であり、また標準化が一番期待できる分野でもある。今回、データアノテーションの仕様文書DEL05.3データ・アノテーション仕様（Data annotation specification）がアップデートされたことは重要な意味を持つといえる。

本文書の目的は、

- 標準的な操作手順によってデータアノテーション品質管理を支援
- 一貫性のないデータアノテーションによって引き起こされるモデルのパフォーマンス問題を軽減
- 多様なデータフォーマットと複数のアノテーターによる大規模なデータセットによるプロジェクトを可能化
- 非専門家のアノテーターのトレーニングと教育を促進するとともに共通の理解を向上

図4は標準的なアノテーションプロセスの概略図である。

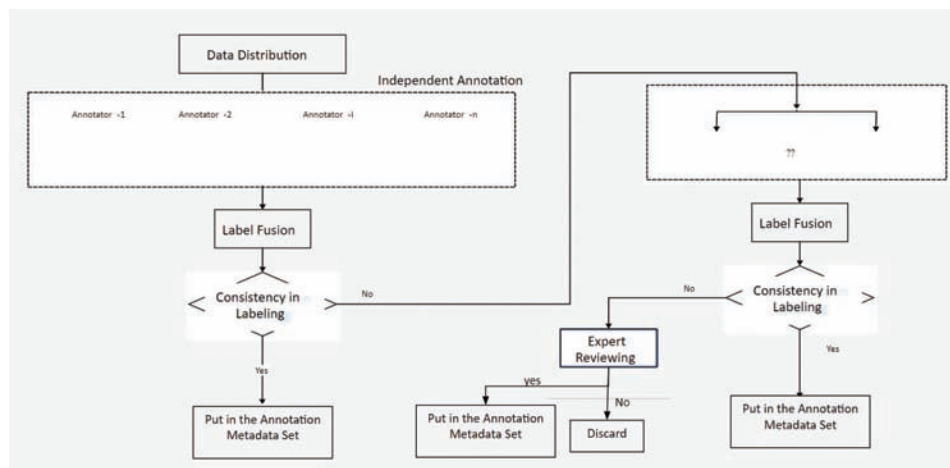
この出力文書は中国のエディターが中心に作成されている。これがどれほどの影響を与えるかは未知であるが、中国の医療データサイエンスにおける地位の向上につながる可能性はある。

2.4 COVID-19対策のガイドライン

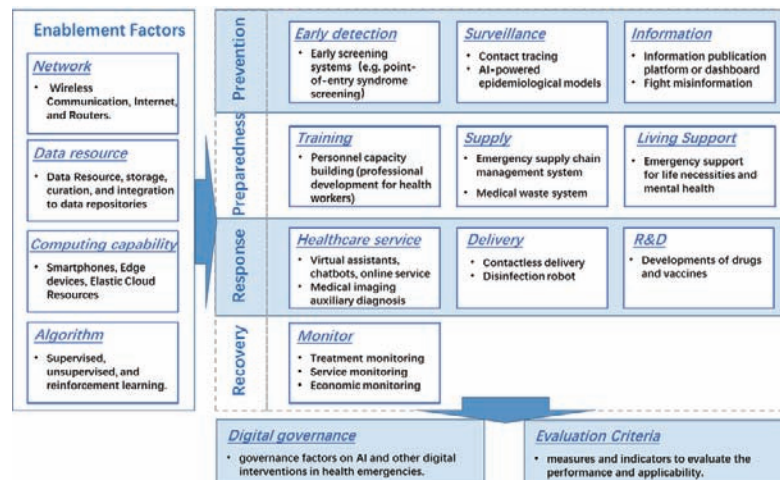
FG-AI4Hは、新型コロナウイルスの全世界的な蔓延による緊急事態のためのAI及びその他のデジタル技術を使用したベストプラクティスを集めたレポートを作成し、今回そのアップデートが成された。この文書は、「Ad-hoc group on digital technologies for COVID health emergency」というアドホックグループが作成している。現在、このグループは各国のユースケースを収集するためにアンケート調査を行っているということである。

文書構造は、AI applications、Enablement factors、Digital governance、Value assessmentの4つの章を中心に構成されるとされている。しかし、現在の章立てはAI Applicationsのみを含んでいる。

全体の枠組みは、OECDのCOVID-19に対応するAIアプリケーションの使用の相関図を参考にしているとのこと



■図4. データアノテーションの標準的プロセス



■ 図5. COVID-19に対応するAIソリューションの相関図 (OECDを参考に作成)

ある。しかし、図5と本文書の章立てとの関係などについては特に明確にされていない。

また、付録には、予防 (Prevention)、準備 (Preparedness)、対応 (Response)、回復 (Recovery) からなる緊急事態管理の各ステージにおいてどのようなAIアプリケーションが利用可能か、という例を表にしている。これも現場での参考になると思われる。

■ 表. 緊急事態ステージに対応したAIアプリケーション使用の例

Stage	Application
1. Prevention	1.1 Early detection
	1.2 Surveillance
	1.3 Information
2. Preparedness	2.1 Training
	2.2 Supply
	2.3 Living support
3. Response	3.1 Health service
	3.2 Delivery
	3.3 R&D
4. Recovery	4.1 Monitor

Annex Aに「Literature review on stage definition of emergency management」という緊急事態のステージ定義等に関する文献調査が挙げられており、参考として役立つと思われる。しかしながら、全体としては、文書の完成度が低く、そのままガイドラインとしては使用できない状態である。全体の枠組みやプロセスの議論が中心で、具体的な内容はあまり書けていない。ユースケースを集める、という趣旨であったが、具体的なものはまだ挙がっていない。

い。コロナ対策のために緊急性を持って作成されたにもかかわらず、アンケート調査という手間のかかる作業から始めており、長期化する可能性があるのが残念である。

3. 今後の予定と課題

FG-AI4H会合には、毎回50以上の寄与文書が寄せられ、プレゼンテーションを中心に進められている。そこでは、各WGやTGの発表があるが、細かい内容についての議論は時間の関係上できない。また、TG (例えば内視鏡や眼科、脳神経科など) もある団体や大学が中心になって進めているもので、その医科業界のコンセンサスを反映しているわけではない。一方で、WHOが関わっている倫理規定や品質に関するガイドラインはWHOのガイドラインともなる可能性が高く、実際、今回の倫理に関するWHOからのガイドライン発表のようにWHO側が先行しているものもある。そういう意味で、WHO側から見ると、倫理規定についての出力文書が確保できたことは、一つのステージの完成を意味すると思われる。

一方、より技術的なAIに対してのガイドラインについては、コンセンサスを得ることは難しく、乱立するAI関係の団体のOne of themとなる可能性もある。

望ましいことは、WHOのガイドラインを実装されるAI技術にどう反映するか、という観点からの技術的議論が進められると、出力文書も単なる参考文献にならないと思われる。

また、ITUから考えると標準化とどうつながるか、ということはいまだ明確になっていない課題である。

FG-AI4Hの次の会合は、2022年2月15日から17日にN会合が開催されることになっている。