

令和 7 年度 厚生労働科学研究費
(医薬品・医療機器等レギュラトリーサイエンス政策研究事業)

AI を用いた医療情報の医薬品安全への活用に向けた諸要件の調査研究
総括研究報告書

研究代表者 荒川 憲昭
国立医薬品食品衛生研究所 医薬安全科学部 室長

研究要旨

近年、ファーマコビジランス (PV) 活動では、人工知能 (AI) 技術を活用した業務支援への期待が高まっている。本研究では、市販後安全対策および個別症例安全性報告 (ICSR) を中心とする PV 業務において、AI を適切に利活用するための課題、留意点および運用上の要件を整理し、医薬品安全管理の向上に資する知見を提供することを目的としている。

最終年度である令和 7 年度は、令和 5 年度からの調査・検討結果を統合し、日本の PV 実務に即した AI 活用の基本的考え方を取りまとめた。具体的には、CIOMS や欧米規制当局等の国際動向、国内の製薬企業・IT 企業・医療機関等を対象とした調査、文献調査、AI 活用シミュレーション、医療情報ソースおよび因果関係評価、個人情報保護・ELSI に関する検討結果を総合的に整理した。さらに、国内の医療機関、製薬企業、規制当局における PV プロセスについて、どのような情報がどの主体から収集され、どのように評価・報告・分析・措置等に活用されるのかを体系的に整理し、PV 業務全体における情報の流れを可視化した。その上で、各プロセスにおいて想定される AI 活用場面ごとに、課題、留意点および対応策案を検討した。

本検討から、PV 領域における AI は、人の最終判断を代替するものではなく、情報抽出、要約、分類、候補提示、論点整理等を通じて、専門職の意思決定を支援するツールとして位置づけることが重要であると考えられた。AI 導入にあたっては、利用目的と適用範囲の明確化、ユースケースごとのリスク層別化、Human-in-the-Loop を前提とした人による確認・修正・承認、品質保証・検証、継続的な性能監視、個人情報保護、記録管理および責任分界を含むガバナンス体制の整備が重要な要件になることを整理した。

A. 序論

近年、医薬品安全性監視活動 (ファーマコビジランス : PV) の分野では、個別症例安全性報告 (Individual Case Safety Report : ICSR) の増加、医療情報ソースの

多様化、症例評価の複雑化を背景として、人工知能 (AI) 技術を活用した業務支援への期待が高まっている。PV 業務では、有害事象情報の収集、症例情報の整理、重複判定、標準用語へのコーディング、ナラティブ作成、既知・未知判断、重篤性評価、

因果関係評価、シグナル検出・評価など多くの工程が存在し、これらの一部に AI を活用することで、業務効率化、品質の均一化、安全性情報の早期把握に資する可能性がある

一方で、PV で扱う情報は患者安全や規制判断に直結するため、AI の判断誤り、重要情報の見落とし、出力根拠の不明確さ、学習データの偏り、個人情報保護、責任の所在など、一般的な業務効率化 AI とは異なる課題を有する。したがって、PV 領域における AI 活用では、AI を人の最終判断の代替ではなく、人の意思決定を支援するツールとして位置づけ、リスクベースアプローチおよび Human-in-the-Loop を前提とした運用設計が必要である。

本研究では、令和 5 年度から令和 7 年度にわたり、日本の PV 分野における AI 活用の推進と医薬品安全管理の向上に資する知見を提供することを目的として、国内外の規制動向、文献情報、製薬企業・IT 企業・医療機関等における AI 活用ニーズ、AI 活用シミュレーション、医療情報ソースの要件、個人情報保護および倫理的・法的・社会的課題について調査・検討を行ってきた。最終年度である令和 7 年度は、これまでの調査・検討結果を踏まえ、国際動向の更新、国内実務における課題の整理、AI 活用場面ごとの留意事項の抽出、品質保証・検証方法・使用環境・法令上の留意点の統合的整理を行い、市販後安全対策および個別症例安全性報告 (ICSR) を中心とする PV 業務において、AI 利活用の推進と医薬品安全管理の向上に資する基礎知見を取りまとめた。

B. 研究方法

B-1. 国内外の医薬品安全対策における AI 活用の検討状況の調査

B-1-1. 欧米規制当局等

PV 領域における AI 活用の国際的な動向を把握するため、CIOMS Working Group XIV、欧州医薬品庁 (EMA)、米国食品医薬品局 (FDA) 等が公表している AI 関連文書、ガイダンス、ディスカッションペーパー、規制機関における AI 導入状況に関する公開情報を更新調査した。また、日本における AI 関連法令、AI 事業者ガイドライン、規制業務への生成 AI 導入動向を整理し、日本の PV 実務への AI 活用に関する課題と留意点を検討した (担当：佐井君江)。

B-1-2. アカデミア等

前年度に引き続き、PubMed を用いてファーマコビジランスおよび医薬品安全性業務に関連する AI 活用論文を抽出し、タイトルおよび要旨を対象としたテキストマイニング、各文書に含まれる単語の重要度を定量化するため、Term Frequency-Inverse Document Frequency (TF-IDF) による特徴量生成、少量教師あり手法および独自辞書に基づく分類を行った。得られた文献を、安全性業務の各工程に対応する分類に整理し、AI の具体的なユースケースを抽出した (担当：今任拓也)。

B-1-3. 国内医療現場等

医療機関における副作用情報管理の実態、医薬品安全情報報告書作成時の負担、AI 活用への期待および不安点を把握するため、病院薬剤部門を主たる対象としたアンケート調査を実施した。Web フォームおよび FAX による回答を集計し、副作用情報の収集・管理、PMDA への報告、製薬企業への情報提供、報告書作成に要する時間、作成時の負担、AI 活用意向、AI 活用に関する懸念等を整理した (担当：瀬戸僚馬、協力：若林進)。

B-2. AI 活用シミュレーション

PV 実務における AI 活用可能性および検証上の課題を把握するため、個別症例評価に関連する AI 活用シミュレーションを実施した。副作用の既知・未知判断支援モデルでは、医薬品副作用データベース (JADER) および添付文書情報を用い、令和 6 年度に作成した部分一致検索処理によるモデル v1.0 を基盤として、令和 7 年度は生成 AI を用いた複合語分解・抽出機能を追加したモデル v2.0 を構築した。GPT による抽出結果について、True Positive、False Positive、False Negative、再現率、適合率、F 値に基づき評価した。

また、報告初期段階から追加情報の蓄積に伴って内容が更新される段階的症例報告を対象として、生成 AI を用いた副作用因果関係評価支援の可能性を探索的に検討した。評価結果の変動、根拠説明、項目別評価と総合評価の整合性、プロンプトや使用モードによる差異を確認し、生成 AI を症例評価支援に用いる際の品質保証および検証上の留意点を整理した (担当：岡本里香、協力：太田実紀)。

B-3. 医療情報ソースおよび因果関係評価に関する検討

医薬品安全対策に必要な医療情報ソースの要件を整理するため、致死性脳出血症例を対象とした因果関係評価研究を実施した。単施設・後ろ向き観察研究として、死亡に至る経過で頭蓋内出血が確認され、出血発症前に投薬があった症例を抽出した。多職種 5 名が、同一の電子カルテ情報を用いて、ACAD-FCH、Naranjo algorithm、WHO-UMC scale の 3 手法により独立に評価し、評価者間一致度、「不明」回答の割合、ツール間の判定差、情報不足になりやすい項目を比較した。

また、PV 業務全体を対象としたプロセスマッピングを行い、「有害事象発生」「副

作用疑い症例の報告」「行政による収集・調査・分析」「行政による措置」の各プロセスにおいて、医療機関、薬局、患者、製薬企業、PMDA 等の主体間でどのような情報が受け渡されるかを整理した。これにより、AI をどの工程に適用するかを明確化し、プロセスごとの情報特性、AI 活用可能性、留意事項を検討した (担当：太田実紀)。

B-5. 個人情報保護および倫理的・法的・社会的課題の検討

PV 分野における AI 導入をめぐる個人情報・個人データ保護および倫理的・法的・社会的課題を整理するため、日本における医療 DX、医療データ利活用、個人情報保護法制、次世代医療基盤法、欧州ヘルスデータスペース (European Health Data Space : EHDS) 等の制度動向を調査した。あわせて、リアルワールドデータ (RWD) 活用、生成 AI、検索拡張生成 (RAG) 等を PV 業務に用いる場合に想定される、入力データの範囲、学習利用の有無、データ保持、ログ管理、委託先管理、アクセス制御、越境移転、説明責任、監査可能性等の論点を整理した。

さらに、学習と推論の分離、モデルの残存性、判断過程の不可視性、外部参照情報の管理、責任主体の分散といった AI 利用に伴う構造的課題を検討した。これらを踏まえ、PV 領域で AI を利用する際の基本原則として、利用目的の明確化、人による最終判断、入力情報の最小化、非学習・非保持、閉域運用、AI 出力の確認・修正記録、参照情報および判断根拠の保存を整理し、実務導入の出発点とな

る最小安全構成について検討した（担当：鈴木晶子）。

C. 結果および考察

C-1. 国内外のPV領域におけるAI活用の検討状況

C-1-1. 欧米規制当局等

CIOMS Working Group XIVでは、PVにおけるAI活用に関するレポートが最終化され、AIの責任ある導入に向けて、リスクベースアプローチ、人間による監視、妥当性・堅牢性、透明性、データプライバシー、公平性、ガバナンスと責任が重要な原則として示された。これらの原則では、AIは人間の能力を補完・拡張するツールであり、人間の判断を代替するものではないことが強調されていた。

EMAでは、EU AI Actの成立を踏まえ、医薬品ライフサイクル全般へのAI活用に関するリフレクションペーパーやAIワークプランに基づき、医薬品規制ネットワークにおけるAI利用、大規模言語モデル（LLM）利用ガイダンス、PV向けAIツールの開発が進められていた。FDAにおいても、医薬品・生物製剤開発におけるAI/ML活用、規制意思決定を支援するAI利用、ICSR評価プロセス支援ツールなどが検討されていた。これらの国際動向から、PV領域におけるAI活用では、リスク評価、人的監督、性能検証、継続的モニタリング、文書化、透明性、説明責任を組み合わせた運用が重視されていると考えられた。

C-1-2. アカデミア等

アカデミアにおける文献調査では、2026年4月時点でAIを用いたPV関連論文が135件抽出され、データ利用、シグナル検

出、概念・レビューに関する研究が多かった。一方、規制対応や妥当性確認に関する研究は少なく、AIを実際のPV業務に導入する際の品質保証、責任分界、監査可能性、実運用時のリスク管理に関する検討は十分とはいえなかった。AIの活用例としては、非構造化データの構造化、シグナルの早期検出、症例特性・重篤性評価の効率化、ICSR業務支援等が想定されたが、いずれも人による確認と判断を前提とした活用が必要であると考えられた。

C-1-3. 国内医療現場

病院薬剤部門を対象としたアンケート調査では、WebフォームおよびFAX回答の計124件を集計した。副作用情報の収集・管理、PMDAへの報告、製薬企業への情報提供はいずれも薬剤部門が中心であり、医療現場における副作用情報管理において薬剤部門が重要な役割を担っていることが確認された。

医薬品安全情報報告書1枚あたりの作成時間は「1時間以上2時間未満」が最も多く、報告書作成の負担として、作成時間を確保できないこと、作成者だけで判断できない項目が多いこと、カルテから転記する情報が多いことが上位に挙げられた。このことから、医療現場における負担は、報告様式そのものよりも、診療情報の探索、転記、要約、判断材料の整理に集中していると考えられた。

AI活用意向では、処方歴や検査結果の転記、経過記録の要約に対する肯定的回答が高く、情報抽出・転記・要約・報告書ドラフト作成支援は、医療現場におけるAI活用の初期的な対象として現実的であると考えられた。一方、被疑薬の特定については肯定的回答がやや低く、薬学的・臨床的判断を伴う工程については、専門職による確認を前提とした慎重な運用が必要であることが示唆された。

AI 活用に関する懸念としては、個人情報・セキュリティ・情報管理、正確性・誤記・信頼性、確認負担・責任所在・運用人材が多く挙げられた。したがって、医療現場に AI を導入する際には、AI の作業範囲、出力確認手順、誤出力時の修正手順、記録方法、責任分界をあらかじめ定めることが重要である。

C-2. AI 活用シミュレーション

副作用の既知・未知判断支援モデルでは、令和 6 年度に作成した部分一致検索処理モデル v1.0 に対し、令和 7 年度は生成 AI による複合語分解・抽出機能を追加したモデル v2.0 を構築した。従来の部分一致検索では、添付文書中の複合語表現、表記ゆれ、上位語・下位語・同義語を含む意味的包含関係を十分に扱えず、本来検出すべき関連事象の一部が検出されない場合があった。生成 AI を用いた複合語分解・抽出機能を追加することにより、こうした検出漏れを一定程度補完できる可能性が示された。

一方で、薬剤間で適合率や F 値に差が認められ、誤検出および検出漏れも確認された。このことから、PV 領域で AI を用いる際には、単一の正答率や一致率のみではなく、使用目的に応じて、再現率、適合率、F 値、False Negative、False Positive を区別して評価する必要があると考えられた。特に False Negative は重要な副作用情報の見落としにつながる可能性があるため、重点的な検証が必要である。一方、False Positive についても、人手による抽出漏れの補完につながる可能性があり、単純な誤検出として扱うのではなく、事例ごとの要因分析が必要と考えられた。

段階的症例報告に基づく因果関係評価の検討では、症例情報の追加に伴い AI の評価が変化し得ることが確認された。また、

使用するモードやプロンプトにより、根拠説明の詳細さ、評価の保守性、代替原因の扱い、項目別点数と総合評価の整合性に差が生じる可能性が示された。したがって、生成 AI を因果関係評価支援に用いる場合には、AI の出力結果をそのまま採用するのではなく、項目別評価、総合評価、判断根拠、参照情報の整合性を人が確認する必要がある。

致死性脳出血症例を対象とした因果関係評価研究では、多職種 5 名が同一の診療情報を用いて、3 種類の因果関係評価ツールにより評価を行った。その結果、評価者間一致度には限界があり、特に情報不足や職種間の判断視点の違いが再現性低下の要因となることが示された。Naranjo algorithm では、dechallenge/rechallenge など死亡例では確認困難な項目で「不明」と回答されることが多く、死亡例における一般的評価アルゴリズムの適用限界が示唆された。一方、致死性事象に特化した ACAD-FCH では、Naranjo algorithm と比較して「不明」回答が少なく、死亡例における評価支援に有用である可能性が示された。

これらの結果は、AI を因果関係評価に活用する場合にも、必要な情報項目をあらかじめ明確化し、情報不足になりやすい項目を把握したうえで、専門職による確認と評価根拠の記録を行う必要があることを示している。AI は、症例情報の整理、評価に必要な情報の抽出、不足情報の提示、判断根拠の可視化には有用である可能性があるが、最終的な因果関係判断は、人による専門的評価を前提とする必要がある。また、PV プロセスマッピングでは、有害事象の発生から行政措置に至る一連の流れを可視化し、各主体間で受け渡される情報を整理した。これにより、PV 業務における AI 活用を検討する際には、単に AI 技

術の性能を評価するだけでなく、どのプロセスで、どの情報処理に、どの程度 AI を用いるのかを明確化することが重要であると考えられた。

D. 国内実装に向けた課題と留意事項の整理

令和 5 年度および令和 6 年度に実施した製薬企業および IT 企業に対するアンケート調査・ヒアリング調査の結果から、国内では市販後 PV 業務、とりわけ個別症例安全性報告 (ICSR) に関連する業務において、AI 活用へのニーズが高いことが示唆された。これを踏まえ、本検討では、有害事象情報の収集、症例データの整理・処理、評価、報告等の各プロセスを主な対象範囲とした。また、本整理は、ICSR を中心とする PV 業務への AI 活用を検討する製薬企業、医療機関、規制当局、PV 向け AI システムの開発・提供に関わる IT 企業、ならびに共同研究等に関わる学術研究者等を主たる想定読者とした。

D-1. AI 導入における基本的な留意事項

本検討では、PV 用 AI を導入する際に共通して考慮すべき基本的な考え方を整理するとともに、想定されるユースケースごとに、利用目的、想定されるリスク、人による確認の必要性、検証方法、情報管理上の留意点を検討した。なお、本検討は、特定の AI 技術や導入形態の是非を判断するものではなく、今後、日本の PV 分野における AI 活用に関するガイドランス案を検討する際の基礎資料として、実務上の課題および留意点を整理す

ることを目的とした。

CIOMS、EMA、FDA 等における国際的な議論では、PV 領域における AI 活用において、リスクベースアプローチ、人による監督、透明性、妥当性・堅牢性、データプライバシー、ガバナンスおよび説明責任等が重要な原則として示されている。一方で、日本においては、PV 実務に特化した AI 活用の運用方法、検証方法、記録管理、責任分界等について、なお具体化すべき点が残されていると考えられた。

そこで本研究では、PV 領域に AI を導入する際の基本的な留意事項として、利用目的と適用範囲の明確化、ユースケースに応じたリスク層別化、責任分界、最小安全構成、Human-in-the-Loop を前提とした人による確認、品質保証・検証、継続的な性能監視、および個人情報保護・情報管理等に関する論点を整理した。

D-1-1. リスク層別化と責任分界

PV 用 AI を実務に導入するには、まず、AI をどの業務工程に、どの範囲で用いるのかを明確にする必要があると考えられた。例えば、単なる転記支援、情報抽出、要約、候補提示のように、人による確認を前提とした前処理的な業務と、既知・未知判断、重篤性評価、因果関係評価、シグナル評価の優先度付けのように、患者安全や規制判断により近い業務では、求められる検証水準や人による確認の程度が異なると考えられる。このため、PV 用 AI のリスクを一律に評価するのではなく、ユースケースごとに、患者安全への影響、規制判断

への影響、業務上の影響、出力誤りが生じた場合の影響範囲、人による修正可能性を踏まえて、低リスク、中リスク、高リスクに層別化することが望ましい。特に、AI の出力がそのまま報告期限、重篤性判断、既知・未知判断、因果関係評価、行政措置の判断等に影響する場合には、より慎重な検証と人による確認が求められると考えられた。

また、AI システムを導入した場合であっても、AI 自体が PV 上の責任主体となるわけではない。そのため、AI の出力を誰が確認し、誰が最終判断を行い、誤出力があった場合にどのように修正し、どの記録を残すのかを、あらかじめ明確にしておく必要があると考えられる。特に、製薬企業、医療機関、規制当局、IT ベンダー、研究機関など複数の主体が関与する場合には、システム開発者、運用者、利用者、最終判断者の役割分担を文書化しておくことが望ましいと考えられた。

D-1-2. 最小安全構成

個人情報保護および倫理的・法的・社会的課題（Ethical, Legal and Social Issues, ELSI）の観点からは、PV 領域で AI を利用する際の基本的な安全構成として、「閉域運用」「非学習・非保持」「人による判断記録」を中核とする考え方が有用な出発点となる。すなわち、医療情報や副作用症例情報を扱う場合には、外部への不要な情報流出を避けるため、可能な限り閉じた環境で運用し、入力情報が AI モデルの追加学習や外部目的に利用されない設定を確認し、人がどのように AI 出力を確認・修正・採否判断したかを記録すること

が望ましいと考えられる。

また、RAG 等の外部参照型 AI を利用する場合には、AI がどの情報源を参照し、どの時点の情報に基づいて出力を生成したのかを確認できることが重要となる。PV 領域では、添付文書、RMP、症例報告、文献、海外規制情報など、参照情報の版数や時点が判断に影響する可能性があるため、参照文書、検索条件、参照日時、出力内容、利用者による確認結果を記録することが望ましいと考えられた。

このような最小安全構成は、すべての AI 利用形態に対する最終的な要件を意味するものではないが、PV 領域において AI 導入を段階的に進める際の基準点として活用できる可能性がある。特に、初期導入段階では、AI を自動判断ツールとしてではなく、情報整理・候補提示・要約支援の範囲に限定し、人による確認と記録を前提とすることで、相対的にリスクを抑えた導入が可能になると考えられた。

D-1-3. 医療情報ソースの要件

AI を PV 業務に活用する際には、AI の性能のみならず、入力される医療情報ソースの質、構造、網羅性、標準化の程度が重要となる。病院薬剤部門を対象とした調査では、副作用情報の収集・管理、PMDA への報告、製薬企業への情報提供の多くを薬剤部門が担っており、医薬品安全情報報告書の作成には、医師記録、検査結果、投薬状況、薬剤師記録、看護記録など複数の診療情報源が参照されていた。

このことから、副作用報告支援に AI を用いる場合、単一の記録を要約するだけでは十分でない可能性があり、複数の情報源

から必要情報を抽出し、時系列や因果関係評価に必要な情報として整理する機能が求められると考えられた。一方、看護記録や自由記載などの非構造化情報は、副作用の端緒情報を含む可能性がある一方で、記録の粒度や表現が一定でないこと、患者状態や副作用の有無が明確に記載されていないことも多い。そのため、AI により候補情報を抽出した後、専門職が確認する運用が望ましいと考えられた。

また、因果関係評価に関する検討では、同一の診療情報を用いた場合でも、職種や評価視点により判断が揺らぐ可能性が示されている。特に、死亡例では dechallenge/rechallenge 等の情報が得られにくく、一般的な評価アルゴリズムでは「不明」となる項目が多いことが示唆された。したがって、AI を因果関係評価支援に用いる場合には、投与時期、発現時期、検査値、画像所見、併存疾患、代替原因、薬理学的妥当性、既知情報など、評価に必要な情報項目をあらかじめ整理し、不足情報を AI が提示できるようにすることが有用である。

D-1-4. AI の品質保証、信頼性、検証方法

AI システムの品質保証においては、単に出力精度を評価するだけではなく、使用目的に応じた評価指標を設定する必要があると考えられた。既知・未知判断支援モデルの検討では、従来の部分一致検索では検出困難であった複合語表現や意味的包含関係について、生成 AI を用いた複合語分解・抽出機能により一定程度補完できる可能性が示された。一方で、薬剤や添付文書の記載様式により、適合率、再現率、F

値に差が生じ、誤検出や検出漏れも確認された（岡本の研究分担報告書）。

この結果から、PV 領域における AI 評価では、全体の正答率だけでなく、False Negative と False Positive を区別して評価することが重要と考えられた。特に、False Negative は重要な副作用情報の見落としにつながる可能性があるため、安全性上の観点から重点的に検証すべき項目と考えられる。一方、False Positive は業務負荷を増加させる可能性があるが、人手による抽出漏れを補完する場合もあり得るため、単純に誤検出として扱うのではなく、出力内容の分類と要因分析が必要になる考えらえる。

生成 AI を用いる場合には、出力の非決定性、プロンプト依存性、モデル更新による結果変動が課題となり得る。そのため、モデル名、バージョン、プロンプト、入力データ、辞書、添付文書の版数、前処理・後処理ロジックを記録し、変更があった場合には再検証を行う仕組みが望ましいと考えられる。また、運用後も、性能低下、データ分布の変化、誤出力の傾向、利用者による修正内容を継続的に確認することが望ましい。

C-1-5. 個人情報保護および倫理的・法的・社会的課題（ELSI）

PV 分野への AI 導入をめぐる個人情報・個人データ保護の観点から、医療 DX、医療データ利活用、RWD、欧州 EHDS 等の制度動向を整理した。近年、全国医療情報プラットフォーム、公的データベース、仮名化情報の利活用等を通じて、医療・介護データの二次利用は制度的に進展しつつあり、

PV 分野においても RWD を活用した市販後安全対策の高度化が期待されている。一方で、医療データの利活用と AI の導入は、単なる技術的効率化にとどまらず、判断主体、責任構造、データの取扱い、公共性の位置づけなど、PV 実務の前提に関わる課題を生じさせる可能性があると考えられた。

従来の PV 業務は、症例情報の収集、評価、因果関係判断、安全対策の検討等において、人間の専門的判断を中核として構成されてきた。これに対し、AI を導入する場合には、学習と推論の分離、モデルの継続的利用・残存性、判断過程の不可視性、責任主体の分散といった構造的課題が生じ得る。特に生成 AI や RAG (Retrieval-Augmented Generation: 検索拡張生成) のような外部参照型 AI では、判断に用いられる知識がモデル内部にあるのか、外部文書やデータベースにあるのかが分かれ、知識の所在や参照範囲を明確にしなければ、後から判断根拠を確認することが困難となる可能性がある。

また、PV 領域で扱う情報には、患者の健康情報、症例経過、医療機関情報、企業が保有する未公開情報等が含まれるため、AI 利用に伴う個人情報・機微情報の取扱いには特に慎重な検討が必要である。入力情報が AI モデルの追加学習に利用されるか、サービス提供者側に保持されるか、ログとして保存されるか、外部環境へ送信されるかによって、情報管理上のリスクは大きく異なると考えられる。このため、AI を PV 業務に用いる際には、利用目的、入力データの範囲、学習利用の有無、保持期間、アクセス権限、ログ管理、委託先管理、利

用規約変更時の対応等を、あらかじめ運用設計として整理しておくことが望ましい。

さらに、生成 AI は欠損情報や曖昧な情報を文脈上もっともらしく補完し、自然な文章として説明し得るという特性を有する。このため、実際の根拠に基づく記述と、AI が推測または補完した記述が混在した場合、利用者が AI 出力を過信するリスクがある。人は自らが期待する、または望ましいと感じる回答を受け入れやすい傾向があることも踏まえると、PV 領域で生成 AI を用いる場合には、欠損情報を補完して記述させるのではなく、「不明」「要確認」「追加照会候補」として明示する設計が重要になると考えられた。

以上のことから、PV 領域で AI を活用する際には、AI の利便性のみならず、患者安全、個人情報保護、説明責任、監査可能性を確保することが不可欠である。特に、目的を明確にし、人間責任、監査可能性を基本原則とし、AI がどの情報を参照し、どのような処理を行い、人がどのように確認・修正・判断したのかを記録する仕組みが必要と考えられる。その実務的な出発点として、閉域運用、非学習・非保持、人による判断記録を中核とする「最小安全構成」は、PV 領域における AI 導入を段階的に検討する上で有用な視点になり得ると考えられた。

D-2. 想定ユースケースにおける課題や留意点、対策案の整理

本研究班では、国内の医療機関、製薬企業、規制当局における PV プロセスについて、どのような情報がどの主体から収集され、どのように評価・報告・分析・措置等に活用されるのかを体系的に整理

し、PV 業務全体における情報の流れ可視化するために、PV プロセスマップを作成してきた（太田分担研究報告書）。これを踏まえ、各プロセスにおいて想定される AI 活用場面を抽出し、課題、留意点および対応策案を検討した。具体的には、医療機関、製薬企業、行政の各主体において想定される 9 つのユースケースを設定し、使用目的、入力データ、想定される AI、情報処理内容、リスクレベル、および CIOMS 7 原則に対応する課題と対応策を整理した（表 1）。以下では、そのうち主要なユースケースを中心に概説する。

(1) 副作用報告のデータベース登録支援

副作用報告のデータベース登録では、AI により、症例報告書、医療機関からの情報、製薬企業が保有する症例情報等から、患者基本情報、投与薬、発現事象、検査値、経過情報、転帰等を抽出し、入力案を提示することが想定される。また、MedDRA コーディング候補の提示、重複症例の候補抽出、欠損項目の指摘、時系列情報の整理にも活用できる可能性がある。

このユースケースでは、AI は最終登録内容を自動確定するのではなく、入力案や確認すべき項目を提示する支援ツールとして位置づけることが適切と考えられる。主なリスクとしては、必要情報の見落とし、誤った項目への転記、MedDRA コーディング候補の誤提示、重複症例の見落とし、緊急性や重篤性に関する誤分類が考えられる。対応策としては、AI が抽出した情報について根拠箇所を表示すること、入力前に担当者が確認すること、修正履歴を残すこと、一定期間ごとに抽出精度や誤分類の傾向を評価することが望ましいと考え

られる。

(2) 副作用疑い症例の初期評価支援

副作用疑い症例の初期評価では、AI により、症例経過、投与時期、発現時期、検査値、既往歴、併用薬、既知情報、添付文書記載、文献情報等を整理し、評価担当者が確認すべき論点を提示することが想定される。いわゆる Evidence Pack や論点整理の作成支援として用いることで、評価者が必要な情報を見落としにくくなる可能性がある。

一方で、初期評価は、その後の追加調査、報告要否、報告期限、優先度判断に影響し得るため、AI による誤要約や論点の偏りには注意が必要である。AI が提示した論点が過度にもっともらしく見える場合、評価者が出力を過信する可能性も否定できない。したがって、AI が提示した内容はあくまで確認補助とし、評価者が原資料に戻って確認できるよう、参照元、引用箇所、出力生成時点を記録することが望ましいと考えられる。

(3) 文献・外部情報からの安全性情報抽出

文献、学会抄録、海外規制当局情報、WHO 等のデータベース、企業が入手する外部安全性情報から、有害事象情報やシグナル候補を抽出する用途では、AI による自然言語処理や要約機能が有用となる可能性がある。特に、膨大な文献や非構造化情報の中から、安全性評価に関連し得る記載を抽出し、担当者に提示する用途では、業務負担の軽減が期待される。

このユースケースでは、見落としが発生

した場合に安全性情報の把握が遅れる可能性があるため、検索条件、対象情報源、抽出ルール、除外基準を明確化することが重要である。また、AI が要約した情報のみで判断するのではなく、原文確認を前提とする運用が望ましい。特に、文献情報では、症例報告、レビュー、動物実験、in vitro 試験、規制措置情報など、情報の性質が異なるため、AI により分類された情報の意味づけを人が確認する必要があると考えられる。

(4) 既知・未知判断支援

既知・未知判断支援では、報告された有害事象名と添付文書、RMP、インタビューフォーム、海外添付文書、既存の安全性データベース等の記載を照合し、既知情報に該当する可能性のある記載を AI が提示することが想定される。令和 7 年度のシミュレーションでは、部分一致検索のみでは扱いにくい複合語表現や意味的包含関係について、生成 AI による補完が有用となる可能性が示唆された。

一方で、既知・未知判断は報告期限や規制対応に影響し得るため、AI の出力を最終判断として扱うことは慎重であるべきと考えられる。特に、添付文書の記載様式、用語の粒度、上位語・下位語・同義語、複合語表現、医学的に類似する概念の扱いにより、AI の判断が揺らぐ可能性がある。対応策としては、AI により提示された候補と根拠文を人が確認すること、判断に用いた文書の版数を記録すること、False Negative を重点的に検証すること、薬剤ごとの性能差を確認することが望ましいと考えられる。

(5) 因果関係評価支援

因果関係評価支援では、AI により、投与開始日、発現日、投与中止後の経過、再投与情報、代替原因、併用薬、既往歴、検査値、画像所見、文献情報等を整理し、因果関係評価に必要な情報を提示することが想定される。また、評価ツールに基づく項目別情報の整理や、不足情報の提示にも活用できる可能性がある。

ただし、因果関係評価は、単純な情報抽出ではなく、医学的・薬学的判断を含む評価であるため、AI による自動判定には慎重な検討が必要である。致死性の脳出血症例を対象とした検討では、同一情報を用いても評価者間一致度には限界があり、情報不足や職種間の判断視点の違いが再現性に影響する可能性が示されている。したがって、AI を因果関係評価に用いる場合には、最終判定ではなく、評価に必要な情報の整理、根拠候補の提示、不足情報の明確化、評価項目ごとの下書き作成にとどめることが現実的であると考えられる。

また、生成 AI を用いた場合、項目別評価と総合評価の整合性、点数計算の正確性、根拠説明の一貫性を確認する必要がある。AI が提示した因果関係評価案については、評価者が原資料と照合し、採用・修正・却下の理由を記録することが望ましいと考えられる。

(6) シグナル検出・シグナル評価の優先度付け

シグナル検出およびシグナル評価では、AI により、JADER 等の副作用報告データベース、RWD、文献情報、海外規制情報

等を統合的に整理し、統計的シグナル候補、症例群の特徴、重篤性、発現頻度、既知情報との差異、海外措置状況等を踏まえて、評価対象の優先度を提示することが想定される。文献調査では、PV 領域における AI 研究として、データ利用やシグナル検出に関する研究が比較的多く、非構造化データの構造化やシグナルの早期検出への応用が検討されていることが示されている。

一方で、シグナル評価は、安全対策措置の検討に近い工程であり、AI の誤分類や優先度付けの偏りが、評価の遅れや過剰対応につながる可能性がある。したがって、AI による優先度付けは、判断の自動化ではなく、評価対象の整理、レビュー順の提案、関連情報の要約支援として用いることが望ましいと考えられる。また、シグナル候補の抽出条件、統計指標、データソース、モデル更新履歴、評価者による最終判断を記録し、後から検証可能な形にしておく必要がある。

D-3. ユースケース検討から抽出された共通論点

以上のユースケース検討を通じて、PV 用 AI の導入において共通して考慮すべき論点を整理した（表 2）。

第一に、AI の位置づけとして、AI は人の最終判断を代替するものではなく、情報抽出、構造化、候補提示、要約、論点整理、根拠提示等を支援するツールとして用いることが望ましいと考えられた。特に、因果関係評価、重篤度判断、既知・未知判断、報告要否、リスク最小化策の必要性、ベネフィット・リスク評価など、医学的・規制的判断を伴う工程で

は、AI 出力を最終判断として扱わず、人が評価・確認・承認する運用が必要と考えられた。

第二に、AI のリスクは技術そのものではなく、使用される業務工程、入力データの性質、出力が後続判断に与える影響、人による修正可能性等に応じて評価する必要があると考えられた。したがって、リスク層別化はユースケースごとに行い、患者安全や規制判断に近い工程では、より慎重な検証と人による確認を設定することが重要と考えられた。

第三に、Human-in-the-Loop を形式的に設けるだけでは十分ではなく、リスクに応じて医師、薬剤師、安全性担当者、調査専門員等が AI 出力を確認・修正・承認する工程を設ける必要があると考えられた。また、誰が、どの出力を、どの根拠に基づいて確認・修正・承認したのかを記録し、後から検証可能な形にしておくことが望ましいと考えられた。

第四に、透明性、説明可能性および監査証跡の確保が重要と考えられた。AI がどの入力情報、参照文書、判定ルール、プロンプト、モデル設定等に基づいて出力したのかを確認できるようにするとともに、出力履歴、修正履歴、承認記録を保存し、監査時に説明できる運用とする必要があると考えられた。

第五に、個人情報・機微情報を扱うユースケースでは、閉域またはローカル環境での運用、非学習・非保持、アクセス制御、入力情報の最小化、マスキング、ログ管理等を基本とする方向性が適切と考えられた。医療情報や ICSR には患者の健康情報や症例経過等が含まれるため、情報管理上のリスクを踏まえた運用設計が必要である。

第六に、品質保証・検証の観点からは、単なる正答率だけでなく、False Negative（偽陰性）、False Positive

(偽陽性)、症例属性別の性能、データソースの変更、添付文書改訂、モデル・プロンプト更新時の影響等を評価する必要があると考えられた。特に、False Negative は重要な安全性情報の見落としにつながる可能性があるため、重点的な検証が必要と考えられた。

第七に、ガバナンスおよび責任分界の明確化が重要と考えられた。AI 開発者、提供者、利用者、確認者、最終判断者、運用責任者の役割を明確にし、SOP、変更管理、監査対応、教育体制等を整備することにより、PV 実務における AI 利用の信頼性と説明責任を確保する必要があると考えられた。

E. まとめ

令和 7 年度は、これまでの調査・検討結果を踏まえ、PV 用 AI の導入に関する留意事項の方向性を整理するとともに、想定されるユースケースごとに、AI 活用の範囲、リスク、人による確認、品質保証、情報管理上の課題を検討した。

その結果、PV 領域における AI は、情報抽出、転記、要約、候補提示、論点整理、既知・未知判断支援、因果関係評価支援、シグナル評価の優先度付けなどに活用できる可能性があると考えられた。一方で、AI は人の最終判断を代替するものではなく、専門職や規制当局の判断を支援するツールとして位置づけることが適切と考えられる。特に、患者安全や規制判断に影響し得る工程では、Human-in-the-Loop を前提とし、AI 出力の根拠確認、人による修正、判断記録、監査可能性を確保することが重要である。

また、AI 導入にあたっては、ユースケースごとのリスク層別化、利用目的と適用

範囲の明確化、データ品質の確認、False Negative および False Positive を含む性能評価、継続的なモニタリング、モデル・プロンプト・参照情報の変更管理、個人情報保護、閉域運用、非学習・非保持、ログ管理等を組み合わせた実務的なガバナンスが求められると考えられた。

本年度の検討結果を含め、本研究班の成果は、「総合研究報告書」において最終的に取りまとめることとした。本研究により整理された知見は、PV 領域における AI 活用の基本的考え方、想定ユースケースごとの課題および留意点、ならびに実務導入時に求められる運用上の要件を示すものである。今後、日本の PV 分野における AI 活用ガイダンス案の作成や、医療機関、製薬企業、規制当局、IT 企業等における実務導入を検討する際の基礎資料として活用されることが期待される。

F. 研究発表

1. 論文発表

該当なし

2. 学会発表

- 1) 荒川 憲昭、他「ファーマコビジランス領域における AI 活用のガイダンス案作成に向けて」第 15 回レギュラトリサイエンス学会学術大会 (September 5, 2025)

表 1. PV 用 AI の想定ユースケースと導入時の主な留意事項案

No.	利用者・ 官署者	想定ユースケース ・使用プロセス	AI の主な役割	リスク 層目安	主な課題・留意点	主な対応策・統 制案
1	医療機関	カルテ情報から副 作用関連情報を収 集	診療録、看護記 録、投薬履歴、 検査値等から、 副作用報告に関 連し得る症状、 発現時期、使用 薬剤、経過、転 帰等を抽出・整 理する。	低～中	抽出結果が後続の 報告書作成に影響 する可能性がある。 自由記載カル テの文脈誤解釈、 診療科・記載者に よる表現差、抽出 漏れ・誤抽出が課 題となり得る。	AI の役割を「情 報抽出・整理」 に限定し、評価 判断や報告要否 判断には用いない。 抽出元の記 載箇所を提示 し、医師・薬剤 師が確認・修正 する。医療機関 内の閉域環境ま たはオンプレミ ス環境で運用 し、入力データ を必要最小限と する。
2	医療機関	抽出した副作用情 報を PMDA 副作 用報告様式へマッ ピング	抽出済み情報を PMDA 報告様 式や記載要領に 沿って各項目に 対応付け、記載 案を作成する。	中	規制当局へ提出さ れる公式文書の作 成に関与するた め、マッピング誤 り、記載漏れ、不 要な個人情報の残 存、様式改訂への 追従不足が問題と なり得る。	AI の使用範囲を 様式への割付、 文章整形、記載 漏れ・不整合の 検出補助に限定 する。AI 生成部 分を識別可能に し、作成者・確 認者・承認者の 役割を明確化す る。最終的な内 容確認と提出責 任は人が担う。
3	製薬企業	副作用報告のデー タベース登録	自由記述、 PDF、電子フォ ーム等の副作用 報告から、有害 事象名、発現時 期、薬剤情報、 重篤性等を抽 出・構造化し、 MedDRA 対応 を含む初期入力 案を作成する。	低～中	登録内容は後続の 集積解析、シグナ ル検出、規制対応 の基盤となる。誤 抽出、誤分類、文 脈の単純化、重篤 性ニュアンスの平 準化がデータ品質 に影響し得る。	AI は入力案提示 に限定し、最終 登録は人が確認 する。原文箇 所、プロンプ ト、処理履歴、 修正履歴を保存 する。個人識別 情報をマスキ ングし、閉域環境 で運用する。症 例属性別の性能 評価も検討す る。

4	製薬企業	副作用疑い症例の初期評価	投与開始・中止、有害事象発現、転帰、デチャレンジ・リチャレンジ、代替要因、不足情報を整理し、Evidence Packや論点整理を提示する。	高	因果関係評価や報告要否判断に影響し得る。AIの一次判定が後続判断を固定化する可能性や、automation biasが懸念される。評価者間でも判断が揺らぐ領域であるため、AI単独の結論付けは慎重であるべきと考えられる。	AIは結論ではなく、根拠箇所、反証候補、不足情報、時系列整理、代替要因候補を提示する役割に限定する。死亡例、未知重篤疑い、非典型例では二重レビューやエスカレーションを設定する。最終判断者、確認記録、監査証跡を明確化する。
5	製薬企業	PSUR、評価会議資料等のレビュー用報告書作成	過去報告書、安全性データ、症例要約、シグナル評価履歴等を参照し、定型構成に沿った草案、差分抽出、文体統一、用語標準化を支援する。	中	集積情報の偏った要約や解釈により、安全性評価や規制対応文書の内容が歪む可能性がある。AIが欠測や曖昧な記述を推測で補完するリスクもある。	AI利用は草案作成、章立て整形、差分抽出、文体統一に限定する。シグナル有無、ベネフィット・リスクバランス、追加リスク最小化策の要否等の結論は人が作成・確認する。RAGを用いる場合は参照資料を限定し、根拠を明示する。
6	製薬企業等	文献情報の定期モニタリング、外部情報抽出、非英語文献翻訳	学術論文、学会抄録、規制当局公開情報等から副作用関連情報を抽出・翻訳・要約し、レビュー候補を提示する。	低～中	偽陰性によるシグナル見落とし、ハルシネーション、翻訳時の医学的文脈欠落、外部ソース利用時の情報管理が課題となり得る。	AI出力を過信せず、リスクに応じて、Human-in-the-LoopまHuman-on-the-Loopを設定する。また、重要業績評価指標による継続的な性能モニタリング、信頼度スコアに基づく優先度付け、個人識別情報の匿名化、閉域環境での運用、ベンダーとの責任分担の明確化を行う。

7	行政	副作用報告受理後のデータ品質管理	欠損、齟齬、重複、コーディング、マスキング、時系列情報、優先度等を事前チェックし、レビュー作業の効率化を支援する。	低～中	必要情報の見落とし、優先度の誤分類、類似症例間での分類ばらつき、分類根拠の不透明性が課題となり得る。	AI 出力はレビュー前の補助情報とし、最終判断は人が行う。分類理由と根拠情報を提示し、確認・更新履歴を保存する。閉域環境で処理し、マスキング確認、内部監査、行政内の AI 利用管理規定を整備する。
8	行政	適正使用・既知/未知判断の一次スクリーニング	症例情報の自由記載から投与理由、用量、投与期間、禁忌、報告事象等を抽出し、添付文書や RMP と照合する。	中 高リスク症例では追加統制	添付文書の記載様式のばらつき、MedDRA との不一致、用語の表現ゆれにより、未知事象や不適正使用を見逃す可能性がある。	疑わしい症例は人による二次評価へ回すよう閾値を低く設定する。高リスク事象は自動的に人の評価フローへ入れる。照合対象、文書版数、判定ロジックのトレースログを残す。一定割合を人が再評価し、精度を継続的に監視する。
9	行政	規制当局による重篤度判断支援	ICSR 構造化項目、ナラティブ、退院サマリ、検査結果等から、死亡、障害、入院・入院延長、先天異常等の根拠を抽出し、重篤性基準へ対応付ける。	高	偽陰性により重篤症例を非重篤として扱うと、安全対策の遅れにつながり得る。一方、偽陽性が多い場合は人的負担が増加する。入院目的など文脈判断が必要な項目では誤判定リスクが残る。	偽陰性を最小化する閾値設計とし、重篤可能性があれば上位レビューに回す。死亡・障害等は二重レビューとする。AI は理由と根拠候補を提示するにとどめ、最終優先度・重篤性判断は人が確定する。欠測情報は推測せず、追加照会候補として提示する。

表 2. ユースケース検討から抽出された共通論点

観点	共通して整理された方向性
AI の位置づけ	AI は人の最終判断を代替するものではなく、情報抽出、構造化、候補提示、要約、論点整理、根拠提示を支援するツールとして位置づけることが望ましい。
リスク層別化	AI のリスクは技術そのものではなく、使用工程、入力データ、出力が後続判断に与える影響、人による修正可能性に応じて評価する必要がある。
Human-in-the-Loop	AI 出力を自動的に採用せず、リスクに応じて医師、薬剤師、安全性担当者、調査専門員等が確認・修正・承認する工程を設けることが望ましい。特に、因果関係評価、重篤度判断、既知・未知判断、報告要否等に影響し得る工程では、人による最終判断を前提とした運用が必要と考えられる。
透明性・説明可能性 ・ 監査証跡	AI がどの入力情報、参照文書、判定ルール、プロンプト、モデル設定等に基づいて出力したのかを確認できるようにし、さらに、誰が、どの出力を、どの根拠に基づき確認・修正・承認したかを記録することが重要と考えられる。後から検証可能にすることが望ましい。
個人情報・機微情報	医療情報や ICSR を扱う場合は、閉域またはローカル環境、非学習・非保持、アクセス制御、入力情報の最小化、マスキング、ログ管理を基本とする方向性が考えられる。
品質保証・検証	正答率だけでなく、偽陰性、偽陽性、症例属性別性能、データソース変更、添付文書改訂、モデル・プロンプト更新時の再検証が必要になる可能性がある。
ガバナンス・責任分界	AI 開発者、提供者、利用者、確認者、最終判断者、運用責任者の役割を明確化し、SOP、変更管理、監査対応、教育体制を整備することが望ましい。