



臨床医として外来に立つ日々のあいだに、生成 AI (ChatGPT など) の効用と危うさを繰り返し体験している。結論を先に言えば、AI は結論を決める装置ではなく、段取りを速く正確に整える装置と考えている。

患者のあいだでも AI を駆使しているようだ。先日の外来のことである。強迫症の患者が「車で人を轢いたのでは」と不安になり、AI に尋ねて「その状況なら轢いていない」と返され、安堵して受診した。一次的には役立ったが、保証を求める行動が症状を強化することは臨床の常識である。そこで著者は、AI に「保証要求には応じず、曝露反応妨害 (exposure and response prevention : ERP) を想起させる」というカスタム指示を設定し、患者とは一回確認ルール (安全な場所での外観確認を一度のみ→戻らない・検索しない・AI に聞かない→不安はタイマーでやり過ごす) を合意した。その後の再診では、走行後の U ターン回数と AI への質問回数がともに減少し、患者は「“安心” より“練習” を優先できた」と述べた。AI の有用性は、安心の供給源ではなく、曝露を支える段取り係に位置づけたことだった。

教育でも「使える」場面が続いた。客観的臨床能力試験 (Objective Structured Clinical Examination : OSCE) の問診シナリオを「疾患・難易度・所要時間」でパラメータ化して一括生成し、ループリック案と模範解答を併置した。その後の実習では、学生の事前練習の質が揃い、当日の指導は臨床的な“詰め”に集中できた。看護学校の講義で時間が余るときには AI で仮想の看護師国家試験問題を作成し、それを解説することに充てた。研究では、抄録からの PICO 抽出、統計解析計画書の雛形化、倫理申請の説明文整序、参考文献の体裁点検を AI に任せ、著者は仮説の吟味と解釈に時間を割いた。査読応答では、AI が示した反論パターンの雛形を叩き台に、表現のニュアンスを素早く整えることができた。

一方、危険性は確かに存在する。AI は誤情報・バイアス拡大、個人情報漏えい、依存による判断力低下、責任所在の曖昧化、透明性不足と説明可能性の欠如、セキュリティ侵害や著作権・人権侵害の懸念を孕み、社会的信頼の毀損など深刻な影響を招きうる。著者は最近のケースレビューで、AI が架空の文献をもっともらしく提示した誤りを確認した。プライバシー、バイアス、依存、責任の空洞化、プロンプト注入——いずれも現実の脅威である。強迫症では「AI ショッピング (複数の AI を渡り歩いて保証を集める)」が起りやすく、症状を強化する。ゆえに著者は、用途制限 (学習・計画に特化、保証の再確認には不使用)、人による最終判断、生成物への出典・更新日の付記、利用箇所の記録・開示を院内とするという基本原則を徹底している。実害が疑われる事案 (強い衝撃、第三者の指摘、車体損傷) では、AI の関与を停止し、警察・保険会社などの標準手順へ直結するべきであろう。

総じて、生成 AI の最大の有用性は、要約・骨子作成・表記統一など“文書の下ごしらえ”を迅速かつ均質にし、患者・チームとの合意形成を助ける点にある。一方の危険は、設計で最小化できる。例えば外来では、患者メモを「睡眠・服薬・出来事・困りごと」で要約して提示するだけでも、診察の立ち上がりが滑らかになる。病棟では、退院サマリーの骨子とチェックリストを先に生成し、不足の指摘を受けて追記するだけで、質のばらつきが減る。精神科ではなじまないかもしれないが、一般診療科のサマリーは AI を活用している。用途は学習・計画に限定し、保証の再確認には使わない。生成物には出典・更新日を付け、最終判断は医師が担い、個人情報は最小限にする。この“AI が段取り、人が判断と責任”の分業を守るかぎり、臨床の匂いは損なわれないどころか、臨床を助ける強力な仲間となりうる。

古郡規雄