

One Person Lab

MAS 智能体白皮书

把医学研究组织成一条可推进、可审阅、可交付的研究线

MAS 沿着问题、证据、分析、稿件与独立审阅组成的连续研究线，持续形成可复查、可修订、可交付的医学研究成果。

Med Auto Science / Research Line / Evidence Chain / Manuscript Delivery

2026-07-13

Public whitepaper

目录

1 定位摘要 2

2 医学研究需要一条连续的研究线 2

3 MAS 的设计思想 2

4 设计原则 4

4.1 研究线优先 4

4.2 问题先行 4

4.3 证据成链 4

4.4 阶段推进 4

4.5 审稿式质量 4

4.6 人机协作 4

4.7 交付可复查 4

5 能力模块 4

6 用户如何开始 5

6.1 暂停课题如何继续 5

7 一个脱敏的微型案例 5

8 有成果就推进，有质量债就明确关闭承诺 6

9 用户会看到什么 6

10 为什么效果更可靠 6

11 人机协作边界 7

12 与 One Person Lab 的关系 7

13 本文边界 7

14 延伸阅读 8

15 结语 8

发布日期：2026-07-13

最近修订：2026-08-10

适用对象：希望理解 Med Auto Science 如何帮助医学研究团队推进真实课题，以及为什么 MAS 的设计适合长期科研任务的医生、PI、研究者、合作者、早期采用者和技术决策者。

核心判断：MAS 沿着问题、证据、分析、稿件与独立审阅组成的连续研究线，持续形成可复查、可修订、可交付的医学研究成果。

本文解释 MAS 如何看待医学研究、为什么这样设计，以及这些设计如何让用户获得更连贯、更可信、更容易接手的科研体验。

1 定位摘要

Med Auto Science (MAS) 是 One Person Lab 系列中的医学研究 Foundry Agent。它把医学研究从一次 AI 对话提升为一条持续推进的研究线：问题、数据、证据、分析、图表、稿件、审阅和交付始终围绕同一个课题组织。

MAS 的设计目标是让研究团队拥有更强的认知与执行能力：AI 承担整理、比较、分析、创作、审阅和修订，研究者负责临床解释、研究边界、主张接受、伦理合规、期刊策略与最终投稿决定。

这套设计有三个用户可感知的结果：

- **研究持续向前。**每一阶段都留下下一步能接住的成果，对话和任务状态围绕成果提供连续线索。
- **结论有来路。**论文主张能回到数据、分析、文献、图表与审阅意见，证据不足会被明确暴露。
- **质量问题可修。**独立审阅把问题转成下一轮补证、收紧表达或人工决策，并清楚区分“已经生成”与“已经可靠”。

2 医学研究需要一条连续的研究线

医学研究的核心难点，是让一个课题从数据和想法持续走向可以被审阅、修订和交付的论文。

研究团队常常遇到相似的断点：有队列时需要判断哪个问题值得做；有初步结果时需要收敛论文主线；图表、分析记录、验证结果和草稿需要跨文件夹与对话组织；几周后，团队还要确认当前稿件和有效证据，并把审阅意见、补充分析与 claim 收紧组织成清楚的下一步。

MAS 把这些问题看成同一件事：医学研究需要一条连续的研究线。

研究线支持迭代回环。团队可以在新证据出现时回到问题设计，在结果不稳定时回到分析，在主张过强时回到写作，在独立审阅发现缺口时回到补证。变化始终留在同一课题上下文中，研究者随时可以回答：

- 当前研究问题是什么，为什么值得继续？
- 数据和文献实际支持到哪里？
- 当前稿件和图表服务哪条主张？
- 哪些问题已解决，哪些仍是质量债？
- 下一步由 AI 推进，还是需要研究者决定？

3 MAS 的设计思想

MAS 把医学研究组织成三个相互约束的层次。

研究线保持连续。系统围绕同一个课题组织选题、分析、写作、审阅和交付。每一阶段都要产生可继续使用的结果、明确的缺口或需要人工判断的问题。

证据链保持诚实。关键 claim 应能回到数据来源、变量与队列定义、分析结果、文献依据、图表和审阅记录。支持不足时，MAS 继续补证、收紧表达、调整路线，或把决定交给研究者。

审稿式质量推动改进。医学意义、统计可靠性、引用真实性、图文一致性与 claim restraint 在研究过程中持续进入独立审阅与修订。

这三层共同形成一条清晰的因果链：问题先决定研究方向，证据决定能说什么，独立审阅决定还需要修什么，研究者决定最终接受什么。

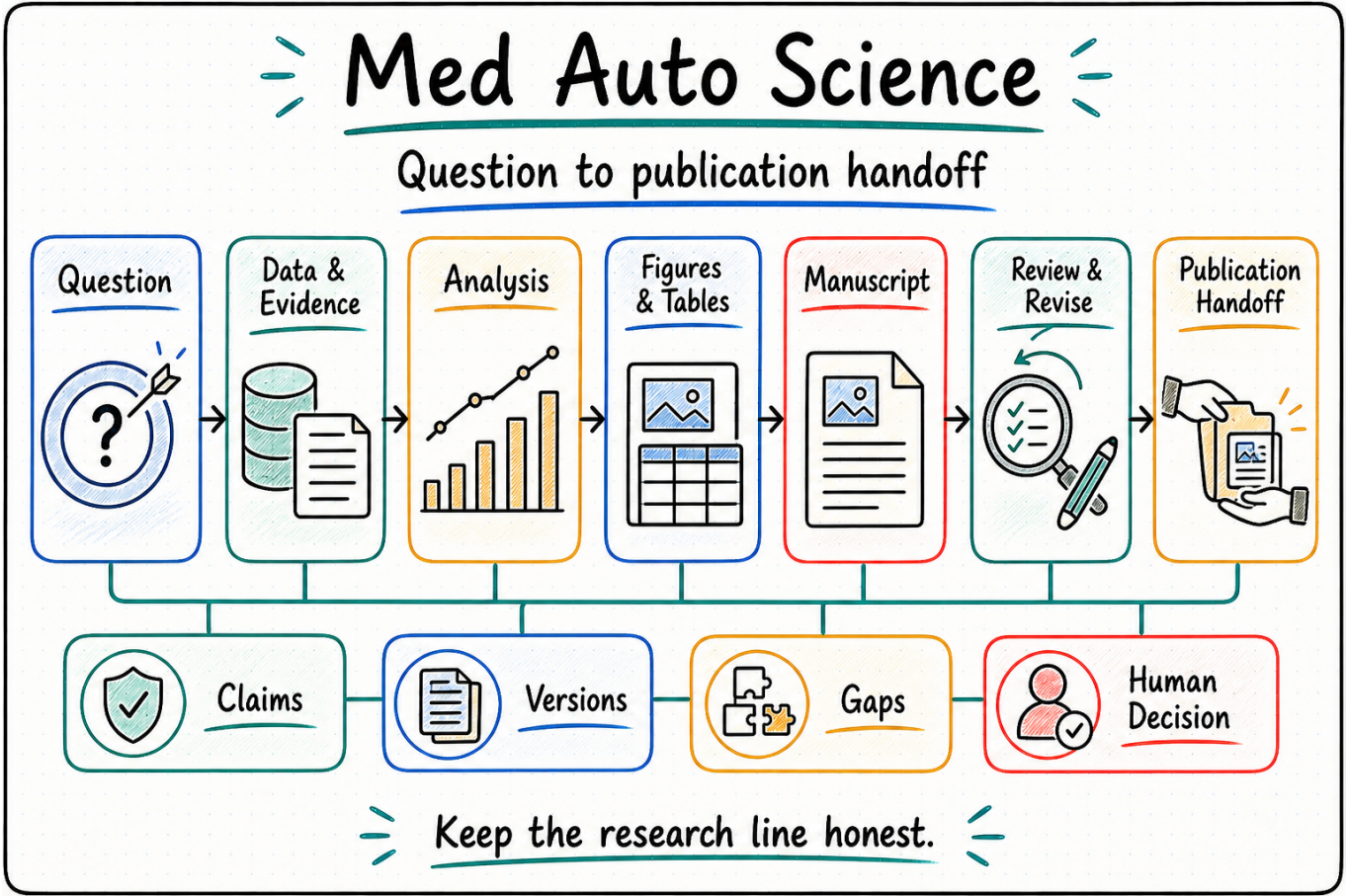


图 1: Med Auto Science 从研究问题到发表交接的连续研究线。

4 设计原则

4.1 研究线优先

MAS 围绕课题组织对话、任务和文件。用户离开几天后回来，仍然能够看到当前问题、已有证据、稿件版本、卡点和下一步。

4.2 问题先行

医学研究先判断问题是否值得做，再决定分析路线。MAS 围绕临床意义、数据支撑、文献空白、endpoint 可行性和发表价值组织思考，让分析从一开始就服务于清楚的医学问题。

4.3 证据成链

数据、文献、分析、图表和审阅记录共同支持论文主张。弱结果、反向结果、不稳定分析与失败路径也被保留，让团队吸收每次尝试的完整证据并减少重复试错。

4.4 阶段推进

长任务被拆成可接力的研究阶段。选题形成候选问题与取舍理由，分析形成结果与证据缺口，写作形成稿件主线，审阅形成修订意见，交付形成可检查材料。

4.5 审稿式质量

执行与审阅是两个独立的智能体任务。撰写或分析工作的 executor 与 reviewer/auditor 分别使用独立调用、独立任务记录和独立回执。这样的外部视角能捕捉创作者的盲区，并让质量门由独立审阅确认。

4.6 人机协作

AI 主动承担高负荷工作，研究者保留高价值判断。系统把临床解释、伦理边界、投稿授权等需要人类决定的问题明确提出。

MAS 以论文作者的立场起草读者可见稿件。科学证据不足时，内容回到证据、分析或研究设计环节补齐，不能伪装成等待作者填空。只有伦理批准号、中心名称、作者贡献或资助信息这类只有作者能确认的客观事实，才以最小的局部 [AUTHOR INPUT: ...] 标记保留理想稿件结构。候选稿进入正式独立审阅后，reviewer 只消费已冻结的同一版快照，避免边改边审让结论失去共同基准。

4.7 交付可复查

正式交付是一组可以复查的材料。稿件、图表、表格、证据、审阅意见、质量债和交付状态一起出现，团队可以理解结果来自哪里、经过什么检查、还缺什么。

5 能力模块

用户直接看到一条围绕论文主张展开的研究旅程，七个概念能力模块在背后协同。它们解释 MAS 会做哪些事，不是七个孤立工具或七个安装选项。

真实课题通过六个公开 Stage action（研究阶段动作）进入：方向与路线选择、基线与证据准备、有边界的分析推进、稿件写作、审阅与质量门、定稿与发表交接。六个阶段动作组成一条可前进、可回退的连续研究路线；七个能力模块则在相应阶段提供专业能力，其中证据、图表和审阅会跨越多个阶段：

概念能力模块	主要承载阶段	用户得到的结果
研究问题设计	方向与路线选择	有取舍理由的候选问题与清楚的医学主线
证据构建	基线与证据准备，并贯穿后续阶段	数据、文献、变量定义、失败路径与缺口被组织在一起
分析推进	有边界的分析推进	主结果、补充分析、稳健性与限制围绕同一问题展开
稿件叙事	稿件写作	摘要、方法、结果与讨论形成克制一致的论证
图表表达	分析推进、稿件写作与质量门	图表、caption、正文数字与 claim 相互对应
审阅修订	审阅与质量门，并按需回到前序阶段	独立 reviewer 把医学、统计、引用和一致性问题转成修订任务
交付准备	定稿与发表交接	稿件、图表、证据、审阅记录与未关闭问题形成可复查材料

MAS 将 mas-scholar-skills 作为必需能力依赖。医学写作、独立审阅、统计审查、文献研究、图表与表格设计、投稿准备和数据治理等专业能力由这一依赖提供；普通用户只需选择 MAS，OPL 即确保该依赖 identity 存在且所需能力可调用。依赖缺失或不可调用时只阻断 MAS，并由托管安装/修复恢复；它不能被降级为可选依赖，也不影响无关 Package。MAS 继续持有医学研究与质量判断。

6 用户如何开始

用户可以从疾病方向、已有数据、初步结果或明确的论文目标开始。例如：

- 我有一个结直肠癌队列，帮我判断哪些问题值得写成论文。
- 我已经有初步结果，帮我整理成一条医学论文主线，并告诉我还缺什么证据。
- 围绕这个预测模型结果，检查校准、临床效用、亚组和图表表达。
- 把当前稿件、图表和审阅意见组织成下一轮修订计划。

MAS 先理解研究对象与目标，再沿研究线推进：

疾病方向 / 数据资产 / 论文目标

- > 方向与路线选择
- > 基线与证据准备
- > 有边界的分析推进
- > 稿件写作
- > 审阅与质量门
- > 定稿与发表交接

这六个公开阶段不是必须机械走完的菜单。MAS 从当前最早未解决的研究环节接手；需要重新选择 endpoint、补充验证、降级 claim 或替换图表时，可以带着已有上下文、证据和决策理由回到前序阶段，再继续推进。

6.1 暂停课题如何继续

已暂停、已交付暂停或已停止的课题不会因为运行环境、执行器或 provider 又变成活跃就自动恢复。只有用户明确选择继续该课题，并经过 MAS 与 OPL 的权威恢复路径后，研究线才重新进入活跃状态。仅用于确认安装、运行环境或任务传输的 qualification-only 工作项，不构成开始或继续真实医学研究的授权。

7 一个脱敏的微型案例

下面是一个虚构的说明性场景，用于展示设计如何影响结果。

问题。某多中心术后队列希望研究一种常见围术期指标与短期不良结局的关系。团队最初想直接写成“该指标可以预测不良结局”。MAS 先把问题收紧为明确人群、暴露、时间窗和 endpoint，并明确观察性数据支持关联表达，因果或

临床预测承诺需要额外设计和验证。

数据。系统整理数据来源、纳排标准、变量定义、缺失模式与中心差异，让每个分析对象都能回到队列构成和处理规则。待确认变量进入显式决策清单。

分析。MAS 围绕主问题组织效应量、置信区间、混杂调整、敏感性分析和中心分层。主结果存在关联，但某个亚组不稳定，外部验证也尚未完成。

图表。系统生成一张主结果图和一张敏感性分析图，并检查图中数字、caption、正文与分析输出是否一致。图表同时显示效应大小与不确定性，对单个显著性结果保持克制。

Claim。executor 根据现有证据形成稿件草案，主张是“该指标与短期不良结局独立相关，并可能用于风险预测”。每句话都连接到对应结果、图表或文献依据。稿件中缺失的伦理批准号是作者才能确认的客观事实，因此只在方法部分保留一处局部 [AUTHOR INPUT: 请补充伦理批准号]；亚组结果不稳定属于科学证据缺口，必须回到分析环节，不能交给作者填空。

独立 reviewer。候选稿冻结后，一个与 executor 上下文分离的 reviewer 只审阅这一版快照，并发现两点：观察性设计只支持关联表达；缺少外部验证时，“用于风险预测”过强。reviewer 同时指出摘要、讨论和图注的表述强度不一致。

修订。MAS 将主张改为“该指标与短期不良结局存在独立关联，值得在外部队列进一步验证其风险分层价值”，同步修订摘要、讨论和图注，并把外部验证登记为明确质量债。团队获得一版更诚实、可审阅、与证据强度一致的稿件增量。

这个过程展示了 MAS 的核心设计：问题约束分析，证据约束语言，图表服务主张，独立审阅发现盲区，修订把质量问题重新接回研究线。

8 有成果就推进，有质量债就明确关闭承诺

可修复的质量问题进入下一轮，已有结果保留为可消费成果；投稿承诺则等待相应质量门关闭。MAS 把推进与就绪分开处理。

当一轮工作已经产生可消费的稿件、分析、图表或证据增量，但预设的审阅与修订预算耗尽时，研究线可以带着明确质量债继续到下一阶段。这一状态在运行语义中称为 `completed_with_quality_debt`。它仅表示“已有下一步可用的成果”；质量通过仍由相应质量门证明。

质量债会关闭质量通过、publication-ready、export-ready 和 submission-ready 等承诺，直到对应问题被补证、修订、独立复核或由有权者决定。当一轮工作缺少可消费成果，或遇到真实的权限、安全、身份、凭据、不可逆操作与人工决策边界时，研究线进入硬阻断。

这种设计同时保护持续推进与如实完成状态：局部瑕疵进入明确修复路径，用户看到的“完成”始终有对应证据。

9 用户会看到什么

打开一个课题后，用户首先看到研究工作本身，底层任务名和内部诊断按需展开：

- 当前问题与主线是否已经清楚；
- 最近新增了哪些分析、图表、稿件段落或审阅意见；
- 哪些主张已被证据支持，哪些需要收紧；
- 哪些质量债正在等待修订；
- 下一步由 AI 继续，还是需要研究者判断。

进度也用研究语言表达。比起“任务完成 60%”，更有用的是“候选问题已收敛为一条主线”“主结果已形成，敏感性分析仍待补充”“图表支持关联 claim，预测承诺仍需更多证据”“独立审阅发现两处待确认引用”。

用户专注于研究，底层运行与包依赖由系统管理。MAS 以声明式医学研究 Pack 表达研究阶段、知识、质量门、可用行动与环境需求；OPL Framework 根据这些声明生成或托管通用入口、工作台、状态、运行、恢复和回执。用户看到统一体验，医学判断仍由 MAS 与研究者持有。

10 为什么效果更可靠

MAS 的可靠性来自一组相互制约、持续发现和修正问题的设计判断：

- **问题先约束分析。**先明确临床意义、研究边界和数据能力，再投入分析。
- **证据持续约束主张。**数据、文献、分析和图表共同决定能说到哪里。
- **失败路径同样保留。**弱结果、反向结果和不稳定分析进入完整证据链。
- **执行与审阅相互独立。**reviewer 用不同上下文检查执行者的盲区，质量门由独立审阅关闭。
- **质量债与 readiness 分别表达。**可消费成果可以继续推进，投稿和发布承诺在质量债解决后开放。
- **人工判断位置明确。**临床解释、伦理合规、主张接受和最终投稿始终有清楚负责人。
- **交付材料可以复查。**文件同时附带依据、审阅和未关闭问题。

这让产品“好用”和“可信”来自同一套设计：系统主动承担复杂工作，同时把证据、不确定性和决策权放在正确的位置。

11 人机协作边界

MAS 适合承担高负荷、反复性、需要上下文组织的研究工作：整理材料、比较方向、形成候选问题、推进分析、组织证据、起草稿件、改写段落、设计图表、汇总审阅意见和准备交付材料。

研究者保留高价值判断：临床问题是否重要，数据使用是否合规，研究边界是否合理，主张是否接受，期刊策略是否合适，以及最终是否投稿。涉及患者安全、伦理、外部系统写入和不可逆操作时，系统必须等待相应授权。

MAS 的最小权威函数集中持有需要专业裁决的医学判断与责任：研究和数据来源真相、独立审阅与 publication quality、artifact 或 memory 的接受/拒绝、owner receipt、typed blocker、human gate 以及必要的医学原生处理。通用运行和界面统一由 OPL Framework 与工作台承载。

12 与 One Person Lab 的关系

Med Auto Science 是 One Person Lab 系列中的医学研究 Foundry Agent，也是一个标准 OPL Agent：

Declarative Medical Research Pack + OPL generated/hosted surfaces + minimal authority functions

这条边界让每一层只负责自己最擅长的事：

用户的医学研究目标

- > OPL App / hosted workbench：统一入口、进度、产物、卡点和下一步
- > OPL Framework：阶段运行、恢复、状态投影、回执与包生命周期
- > MAS declarative pack：研究阶段、知识、质量门与行动声明
- > MAS minimal authority：医学 truth、质量判断与交付授权
- > 研究者：临床解释、伦理合规、主张接受与最终投稿

mas-scholar-skills 作为 MAS 的 required dependency，由统一 Package 入口维护存在性与能力可调用性；普通组合不要求跨包版本锁、ABI 求解或原子 closure。用户从 MAS 进入，OPL 承接通用运行、生命周期和界面，ScholarSkills 提供医学研究方法能力，MAS 始终对医学质量和权威结果负责。当前优先使用 Codex 作为正式 executor 与 carrier 生态，但 Package identity、科研任务与 typed views 保持 OPL-owned、executor-neutral。

这种分工减少重复平台，让界面、运行状态、通用工具与医学结论各归其主。用户感知到一条已经组织好的连续科研工作线。

13 本文边界

本文聚焦 MAS 的产品思想、设计原则和用户价值。安装说明、运行状态证明、论文质量证明、投稿授权和发表承诺分别由对应权威来源提供。

具体课题的 paper progress、publication-ready 和 submission-ready，必须由对应 workspace、稿件、图表、证据记录、独立审阅记录、controller 记录、owner receipt、human decision 与最新运行输出共同证明。

MAS 可以组织和推进医学研究工作。临床解释、伦理和数据合规、主张接受、目标期刊选择、最终投稿与外部系统交互，继续由研究者、PI 或团队负责人决定。

14 延伸阅读

- [Med Auto Science 公开入口](#)：产品概览与源码入口。
- [项目概览](#)：MAS 当前角色、边界与目标。
- [架构概览](#)：MAS、OPL Framework 与 owner 边界。
- [AI-first Quality Boundary Policy](#)：医学论文质量判断边界。
- [Stage-Led Research Autonomy Policy](#)：研究阶段与自治推进原则。

15 结语

MAS 的目标，是让医学研究团队拥有一条持续推进的智能研究线。

它从问题出发，把数据、证据、分析、图表、稿件、独立审阅、修订和交付连接起来。它让 AI 承担重活，让研究判断保持有依据；让可消费成果持续前进，让质量债与投稿就绪保持清楚区分；让用户专注研究本身，由 OPL 和 MAS 在背后保持清楚的运行与权威边界。

这就是 Med Auto Science 的设计理念：主动推进，证据约束，独立审阅，边界清楚，交付可复查。