

Pharmacy Practice in the Age of Agentic Systems That Can Plan, Act, and Escalate

Maria Gonzales^{1*}, Jose Santos¹, Anna Cruz², Carlos Reyes³

¹Department of Agentic AI and Pharmacy Practice, Faculty of Pharmacy, University of the Philippines Manila, Manila, Philippines. ²Department of Autonomous Systems in Pharmacy, Faculty of Pharmacy, University of Santo Tomas, Manila, Philippines. ³Department of Escalation and Human Oversight, Faculty of Pharmacy, Ateneo de Manila University, Quezon City, Philippines.

Abstract

Artificial intelligence in pharmacy has largely been framed as a source of predictions, alerts, classifications, or generated recommendations. Agentic systems create a different governance problem because they may maintain task state, formulate multistep plans, invoke tools, initiate workflow actions, monitor resulting conditions, and decide whether to revise, stop, or escalate. This article develops a non-empirical governance architecture for situating such capabilities within medication-use work. The proposed Bounded Agentic Pharmacy-Work Theory conceptualizes agentic performance as a relationship among a defined pharmacy-work mandate, authorized context, planning and action functions, permission constraints, consequence monitoring, escalation pathways, and accountable human and organizational control. Four interacting layers are proposed: the pharmacy-work mandate, agentic execution, permission and authority, and human and organizational control. These layers are connected through mandate specification, an action envelope, an authority contract, an escalation contract, and a trace-and-consequence ledger. A plan-act-monitor-escalate cycle is introduced to distinguish generation of a plausible plan from permission to cause a medication-related action. Evaluation must therefore extend beyond model accuracy to plan coherence, tool use, permission adherence, action correctness, reversibility, monitoring sensitivity, escalation performance, human factors, medication-safety consequences, equity, robustness, and local workflow effects. The architecture does not establish that agentic systems improve pharmacy practice or that any proposed permission level is clinically, legally, or regulatorily acceptable. Its original contribution is a testable framework for determining how agentic capability may be bounded, evaluated, supervised, restricted, or withdrawn before it is treated as legitimate pharmacy-work authority.

Keywords: Digital pharmacy, Artificial intelligence, Pharmacy practice, Medication safety, Human-AI collaboration, Clinical decision support

INTRODUCTION

Clinical decision-support systems have conventionally supplied alerts, predictions, reminders, classifications, or recommendations within a workflow in which a clinician remains responsible for interpreting the output and determining what occurs next [1]. Artificial intelligence can extend these functions, but its clinical value has generally been understood to depend on a productive relationship between computational capability and human expertise rather than on independent machine performance [2]. This framing assumes that the principal boundary lies between an output and the professional decision made from it.

Agentic systems alter that boundary. A predictive model estimates a state or outcome; an agent may organize a sequence of operations directed toward a goal. Machine-learning performance alone does not specify how an output should be converted into a clinical action, who may authorize that action, or how its consequences should be observed [3]. Responsible translation consequently requires attention to the

Address for correspondence: Maria Gonzales, Department of Agentic AI and Pharmacy Practice, Faculty of Pharmacy, University of the Philippines Manila, Manila, Philippines.
maria.gonzales@upm.edu.ph

Received: 24 October 2025; **Accepted:** 05 February 2026

This is an open-access article distributed under the terms of the Creative Commons Attribution-Non Commercial-Share Alike 4.0 License, which allows others to remix, tweak, and build upon the work non commercially, as long as the author is credited and the new creations are licensed under the identical terms.

How to cite this article: Gonzales M, Santos J, Cruz A, Reyes C. Pharmacy Practice in the Age of Agentic Systems That Can Plan, Act, and Escalate. Arch Pharm Pract. 2026;17(1):39-47.
<https://doi.org/10.51847/BDjsmWFfzU>

intended problem, data provenance, evaluation environment, workflow integration, monitoring, and downstream harm throughout the system lifecycle [4]. These requirements become more consequential when an AI system can select tools, alter records, prepare orders, communicate with other systems, or continue operating after its initial output.

The central challenge is therefore not simply whether an agent can produce a clinically plausible plan. It is whether its planning and action functions are situated within a pharmacy-work arrangement that defines what the system is permitted to know, propose, prepare, execute, monitor, and escalate. Technical performance frequently fails to produce clinical benefit when usability, local validation, workflow fit, and implementation conditions are inadequate [5]. For an agentic system, such failures may propagate across several connected steps rather than remain confined to a single prediction.

This article proposes Bounded Agentic Pharmacy-Work Theory (BAPWT) as an original governance architecture. The theory treats agentic pharmacy as a relationship among a bounded medication-use mandate, context-sensitive planning, permissioned action, consequence monitoring, escalation, and accountable human control. It distinguishes technical capability from professional authority, human presence from meaningful oversight, and implementation from validated benefit. The proposed architecture is intended to organize future empirical evaluation; it does not establish that agentic systems are safe, effective, or appropriate for autonomous pharmacy practice.

What Makes an AI System Agentic in Pharmacy

An AI system becomes meaningfully agentic when it moves beyond generating a discrete output and participates in pursuing a goal over time. Healthcare-agent architectures commonly include goal interpretation, memory or maintained state, planning, reflection, tool interaction, and action [6]. However, terminology remains inconsistent, and the healthcare evidence base does not yet support a stable boundary between systems described as AI agents and those described as agentic AI [7]. The term should therefore be

assigned according to observable functions rather than promotional labels.

In pharmacy, agenticity should be assessed at the level of the **task trajectory**. A system that answers a medication question is not necessarily agentic. A system that gathers authorized clinical context, formulates a monitoring plan, requests additional information, revises the plan when new results arrive, and routes an unresolved concern to a pharmacist exhibits a sequence of goal-directed operations. Evaluations of clinical agents accordingly need to examine longitudinal interactions and intermediate decisions rather than isolated response accuracy [8].

Action capability is another defining dimension, but technical action must not be confused with legitimate authority. Sandboxed medical agents have demonstrated the ability to navigate records, initiate investigations, interpret results, and construct treatment plans across multistep scenarios [9]. Such demonstrations show that agentic functions can be operationalized, but they do not establish safe performance in live pharmacy environments, where actions may affect medication supply, administration, monitoring, documentation, communication, or treatment continuity.

For this article, an agentic pharmacy system is proposed as a computational system that pursues a bounded medication-use objective through authorized context assembly, multistep planning, tool selection, action initiation or preparation, state observation, plan revision, and escalation. Evaluation must consequently examine reasoning trajectories, tool use, action sequences, human assessment, safety, and clinical consequences in addition to final-output quality [10]. The definition remains functional and governance-oriented: it does not confer authority merely because the system can act. **Table 1** organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for definitions, components, relationships, and intended uses for the article Pharmacy Practice in the Age of Agentic Systems That Can Plan, Act, and Escalate.

Table 1. Definitions, components, relationships, and intended uses for the article Pharmacy Practice in the Age of Agentic Systems That Can Plan, Act, and Escalate.

Component or scientific dimension	Problem addressed	Inputs or determinants	Proposed mechanism or relationship	Expected contribution	Evidence required	Failure or uncertainty risk	Boundary statement
Predictive tool	Estimating a clinical state or outcome	Structured or unstructured patient data	Produces a score, class, forecast, or recommendation [1, 3]	Informs professional assessment	External and local predictive validation	Miscalibration, dataset shift, inappropriate interpretation	Prediction does not authorize action.
Clinical decision support	Delivering relevant information within workflow	Patient data, rules, models, knowledge sources	Presents alerts or recommendations to a user [1]	May improve information availability	Workflow and clinical-outcome evaluation	Alert burden, interruption, non-use, overreliance	Availability of advice does not establish benefit.

AI agent	Pursuing a goal through multiple state, memory, tools, operations	Goal, maintained state, memory, tools, feedback	Coordinates planning, tool use, observation, and revision [6-8]	Supports sequential task completion	Trajectory-level and tool-level evaluation	Error propagation across steps	Agenticness is functional, not a product label. This is a proposed definition, not a recognized regulatory category.
Agentic pharmacy system	Organizing medication-use work	Authorized clinical context, mandate, available tools	Proposed: connects planning and action to permission, monitoring, escalation, and accountability	Provides an analyzable unit for pharmacy governance	Multilevel technical, clinical, workflow, and governance evidence	Capability may exceed authorized or validated use	
Pharmacy-work mandate	Preventing open-ended goal pursuit	Patient-specific purpose, task scope, time horizon, exclusions	Proposed: constrains what the agent may optimize and when the task ends	Clarifies intended work and decision boundaries	Evidence that the mandate is understandable, feasible, and clinically appropriate	Ambiguous goals, conflicting objectives, scope expansion	The agent may not redefine its own mandate.
Planning	Converting a mandate into sequenced operations	Context, constraints, available tools, uncertainty	Generates and revises a task plan [6, 9]	Makes intermediate reasoning and dependencies reviewable	Plan coherence, completeness, constraint adherence	Plausible but clinically inappropriate plans	A coherent plan is not necessarily an authorized plan.
Tool-mediated action	Connecting generated plans to workflow effects	Interfaces, permissions, authentication, action parameters	Invokes a tool or prepares an external action [8, 9]	Enables task execution rather than advice alone	Tool-selection, parameter, access, and execution testing	Wrong tool, wrong target, unauthorized action	Technical access does not create professional authority.
Monitoring and escalation	Detecting mismatch or inability to proceed safely	Expected state, observed state, uncertainty, exception rules	Proposed: compares consequences with expectations and routes unresolved states to qualified humans	Supports revision, pause, rollback, or handback	Detection accuracy, timeliness, recipient response, workload effects	Missed deterioration, excessive escalation, responsibility gaps	Escalation is protective only when a capable recipient can respond.

Table note: Evidence-derived descriptions are cited. Entries identified as proposed are original conceptual constructs requiring empirical and setting-specific validation.

Proposed Agentic Pharmacy-Work Theory

BAPWT begins from the premise that agentic pharmacy is not an isolated property of a model. It is a property of a sociotechnical arrangement in which people, technologies, tasks, organizational structures, environments, and patient journeys interact [11]. This perspective is important because pharmacists may accept, interrogate, reinterpret, or reject opaque AI outputs according to local work practices and their ability to compare the system's reasoning with other evidence [12]. The theory therefore treats professional judgment and organizational conditions as constitutive components rather than external safeguards added after technical development.

Four interacting layers are proposed. The pharmacy-work mandate layer specifies the patient-related purpose, task, temporal horizon, exclusions, intended recipient, and closure condition. The agentic execution layer includes context assembly, planning, tool selection, action preparation or execution, observation, and revision. The permission and authority layer limits access and action according to consequence, reversibility, professional scope, and required authorization. The human and organizational control layer assigns supervision, escalation response, incident review, change control, restriction, and retirement responsibilities.

Human–AI performance varies with task, user, interface, and error configuration, so complementarity cannot be presumed from the stand-alone performance of either participant [13].

Five proposed mechanisms connect the layers. Mandate specification converts a broad organizational objective into a bounded task. The action envelope states which information sources, tools, action types, recipients, and time periods are permitted. The authority contract specifies what the agent may observe, recommend, prepare, or execute and which actions require active human authorization. The escalation contract defines triggers, recipients, required context, response expectations, and contingency routes. The trace-and-consequence ledger records the mandate, plan, tool calls, permissions, actions, observations, authorizations, deviations, and closure state. Together, these mechanisms prevent the system's ability to produce or transmit an action from being treated as evidence that it should be allowed to do so.

Figure 1 presents the original conceptual synthesis for agentic pharmacy-work architecture, showing how the article's principal components, evidence relationships, uncertainties, and decision boundaries are connected.

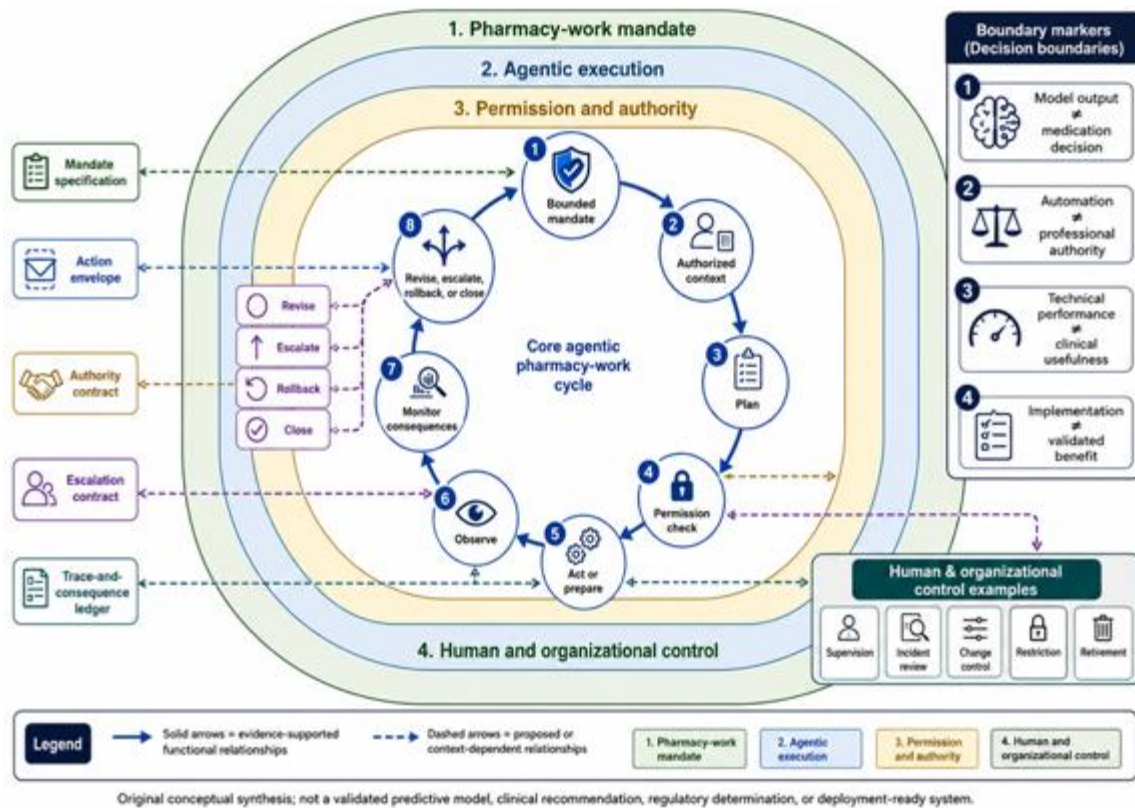


Figure 1. Agentic Pharmacy-Work Architecture

A person's nominal presence within the workflow does not by itself establish meaningful control. Clinicians can remain susceptible to erroneous AI advice even when they retain formal responsibility for the decision [14]. Under BAPWT, human control is therefore meaningful only when the responsible person has access to the relevant evidence and trace, sufficient competence and time to assess the action, practical authority to intervene, and a viable escalation or rollback mechanism.

The architecture also anticipates emergent effects. Machine-learning systems can introduce new error pathways, altered responsibilities, deskilling, and workflow consequences that were not represented in model-development metrics [15].

BAPWT consequently treats performance as a joint property of the agent, permission structure, pharmacist, organization, task, and environment. It is an organizing theory for empirical evaluation, not evidence that any particular configuration is safe or clinically useful.

Table 2 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for evaluation criteria, evidence requirements, and failure modes for the article Pharmacy Practice in the Age of Agentic Systems That Can Plan, Act, and Escalate.

Table 2. Evaluation criteria, evidence requirements, and failure modes for the article Pharmacy Practice in the Age of Agentic Systems That Can Plan, Act, and Escalate.

Architecture layer or process stage	Core function	Information flow	Interaction with other components	Evaluation criterion	Potential failure mode	Human or organizational responsibility	Readiness boundary
Work mandate	Defines purpose, scope, exclusions, and closure	Organization and pharmacist to agent	Constrains every subsequent step	Mandate clarity and task appropriateness	Ambiguous or conflicting objective	Approve, revise, or reject the proposed use case	No operation without a bounded mandate.
Context assembly	Collects authorized patient and workflow information	Approved sources to maintained task state	Supplies planning and monitoring	Completeness, provenance, timeliness, access compliance	Missing, stale, excessive, or unauthorized data	Define accessible sources and review exceptions	Data access does not imply action permission.

Planning	Produces a sequenced approach	Context and constraints to proposed steps	Precedes tool selection and action	Clinical coherence, constraint adherence, alternative recognition	Hallucinated dependency, omitted check, premature closure	Review high-consequence or uncertain plans	Plausibility alone cannot justify execution.
Permission gate	Compares a proposed step with the action envelope	Proposed action and authority contract to allow, block, or route	Intervenes before consequential action	Correct authorization and absence of self-expansion	Unauthorized tool use or permission bypass	Define permissions and investigate violations	Any unauthorized consequential action blocks readiness.
Action or preparation	Produces an external workflow effect or reviewable draft	Authorized plan to tool or human reviewer	Changes or prepares a workflow state	Intended-versus-actual action, reversibility, target correctness	Wrong patient, drug, dose, recipient, or timing	Authorize, supervise, reverse, or restrict	High-consequence actions require active qualified-human authorization.
Monitoring	Compares expected and observed states	Post-action signals to task state	Supports revision, closure, or escalation	Detection of mismatch, delay, deterioration, or incomplete closure	False reassurance or unrecognized downstream harm	Maintain monitoring capacity and define response ownership	Output monitoring alone is insufficient.
Escalation and handback	Transfers unresolved work and responsibility	Agent trace and exception state to qualified recipient	Connects failure detection with human control	Trigger appropriateness, information completeness, recipient response	Escalation overload, wrong recipient, unaccepted handback	Ensure staffed and competent escalation destinations	Escalation is not protective without response capacity.
Trace and governance review	Preserves reconstructable evidence	All plans, actions, permissions, observations, and interventions to audit record	Supports incident review and change control	Trace completeness, interpretability, version linkage	Unreconstructable action pathway or unclear responsibility	Audit, investigate, restrict, requalify, or retire	Traceability supports but does not prove safety.

Table note: The architecture stages and readiness boundaries are original proposed constructs. The evaluation logic is informed by evidence on clinical AI translation, agent functions, sociotechnical work, human–AI interaction, and unintended consequences [3–15]. No universal quantitative threshold or deployment score is proposed.

Planning, Acting, Monitoring, and Escalation Cycles

The operational core of BAPWT is a proposed plan–act–monitor–escalate cycle. The cycle begins when a bounded mandate is translated into a plan containing required information, intended steps, dependencies, constraints, and closure conditions. Before each consequential step, the system compares the proposed action with its current permission and authority contract. This design is necessary because human review can be weakened by automation bias and by the practical complexity of verifying an automated recommendation [16]. Passive oversight may also reduce situation awareness and leave the responsible professional poorly prepared to intervene when control must be resumed [17].

Planning is therefore distinct from permission. A plan may be clinically coherent yet outside the agent’s authorized task, dependent on unavailable information, or inappropriate for the patient’s current state. Acting is similarly divided into preparing a reviewable action and executing an external action. After action, the system must observe whether the expected state occurred, whether new information changes the plan, and whether the task has reached a legitimate closure condition. Clinical decision-support effects vary across interventions and workflow arrangements, demonstrating that consequences must be measured within the actual work system rather than inferred from algorithm availability [18].

Monitoring must extend beyond checking whether a tool call succeeded. It should compare expected and observed clinical, informational, and workflow states. A technically successful transmission can still reach the wrong recipient, omit required context, produce a delayed response, or cause an inappropriate downstream action. Monitoring may therefore result in continuation, revision, pause, rollback, escalation, or closure. Local validation is essential because a widely implemented model can behave materially differently when independently evaluated in another clinical environment [19]. Under the proposed architecture, authority is conditional on the system version, population, data sources, workflow, integration, and action envelope actually evaluated.

Escalation is not a generic warning function. It is a structured transfer of unresolved work and relevant responsibility. The agent should identify why it cannot continue, provide the current task state and evidence trace, name actions already taken, and specify the response required. The receiving professional or team must have the competence, time, information, and authority to respond. Otherwise, escalation may redistribute rather than control risk.

Agent benchmarks can examine planning, execution, verification, and task completion across multistep clinical scenarios, but their findings remain sensitive to task construction, agent architecture, available tools, and evaluation conditions [20]. Benchmark success cannot

establish that the complete plan–act–monitor–escalate cycle will function safely within pharmacy practice. The proposed cycle specifies observable functions and failure points for prospective testing; it does not demonstrate that monitoring, rollback, or escalation will prevent medication-related harm.

Permission, Authority, and Human-Control Layers

Agentic capability must be separated from permission and authority. Governance should connect system design, validation, implementation, monitoring, accountability, and withdrawal across the technology lifecycle [21]. BAPWT therefore proposes an action envelope specifying the information sources, tools, recipients, action types, time limits, and conditions available to an agent.

Authority allocation must also account for bias, transparency, professional obligations, and effects on clinical relationships [22]. Responsibility, fairness, data use, and patient protection remain relevant throughout development and operation rather than only at initial approval [23]. Because action authority may affect professional responsibility and liability, a system's technical ability to execute an action cannot itself legitimize that action [24].

Five analytical permission levels are proposed: P0 observe, P1 recommend, P2 prepare a reversible action for review, P3 execute a predefined reversible action, and P4 prepare a consequential action requiring active qualified-human authorization. These levels are not legal classifications. An agent must not expand its own permission.

Meaningful human control requires more than a nominal reviewer. The responsible person must receive sufficient evidence and trace information, possess the competence and time to assess the action, and retain practical authority to pause, reject, reverse, or escalate it. Human deliberation must also preserve contextual values and patient preferences that cannot be reduced safely to a single encoded objective [25].

Evaluation and Governance Requirements

Evaluation must treat the agent as a clinical-work intervention rather than a response generator. Reports should

specify the intervention, inputs, outputs, human interaction, errors, and operating conditions [26]. Prospective protocols should additionally predefine data handling, failure management, human–agent interaction, and analysis procedures [27].

Early clinical evaluation should examine technical performance together with safety, workflow integration, human factors, and system errors before authority is expanded [28]. BAPWT therefore proposes transition gates covering local model validity, plan coherence, correct tool use, permission adherence, action correctness, reversibility, monitoring, escalation, workload, equity, security, and clinical consequences. No gate has a universal quantitative threshold.

Implementation also requires organizational ownership, training, transparency, equity assessment, incident response, and continuing monitoring [29]. A responsible governance system must connect lifecycle accountability, early clinical evaluation, and postdeployment monitoring rather than treat them as separate activities [21, 28, 29]. Authorization remains conditional on the specified task, population, workflow, integration, system version, and permission envelope.

Implications for Pharmacy Roles and Work Design

Current pharmacy applications of artificial intelligence are concentrated substantially in workflow and productivity functions, while direct evidence of patient-outcome improvement remains limited [30]. Agentic systems may therefore change the organization of pharmacy work before they justify independent clinical authority. Potential role shifts include mandate definition, exception management, authorization, investigation, patient communication, incident review, and system governance.

Figure 2 presents the original conceptual synthesis for plan–act–monitor–escalate cycle with permission and human-control layers, showing how the article's principal components, evidence relationships, uncertainties, and decision boundaries are connected.

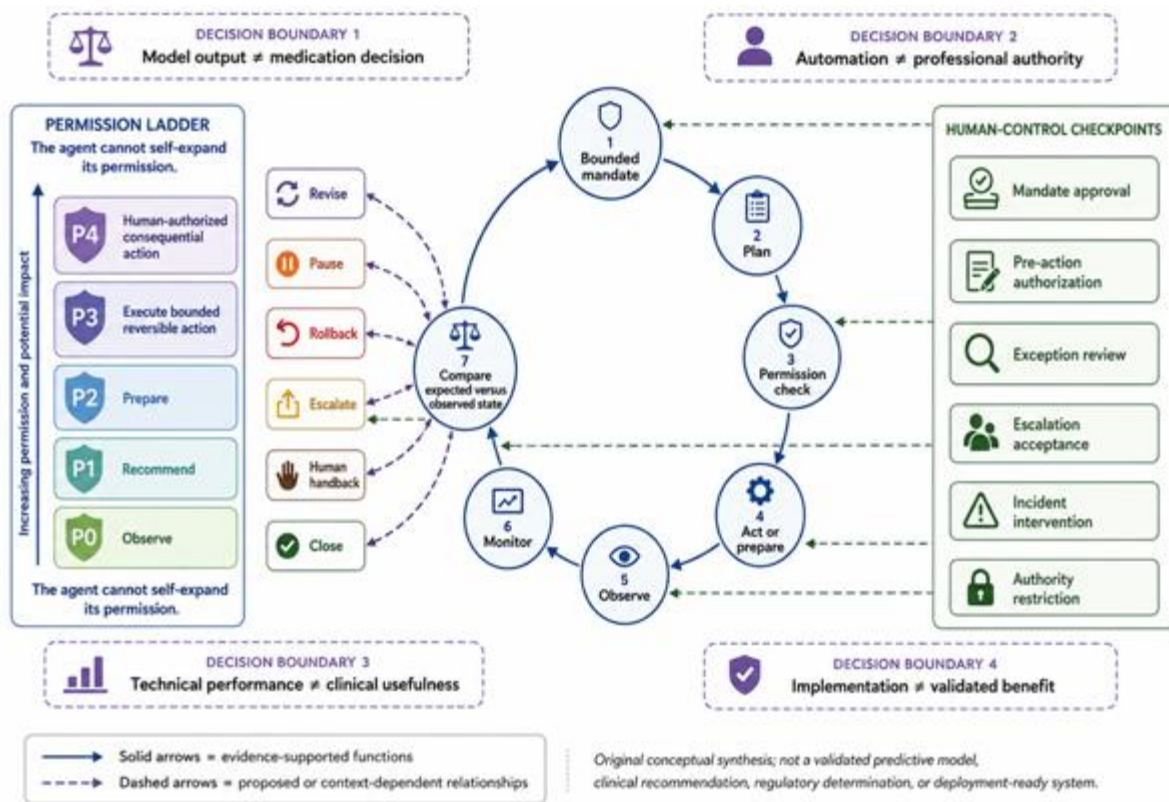


Figure 2. Plan-Act-Monitor-Escalate Cycle with Permission and Human-Control Layers

Medication-alert optimization illustrates how AI may alter prioritization and review burden, although many systems remain limited by small datasets or insufficient formal validation [31]. Medicine-related AI should consequently be evaluated for clinical utility, safety, equity, and implementation rather than as a technical product alone [32]. Clinical-pharmacy evidence covers decision support, adverse-event detection, personalized treatment, pharmacometrics, and operational functions, but its heterogeneity does not justify generalized agentic authority

[33]. Pharmacy role redesign should be evaluated as a work-system intervention rather than inferred from model capability alone [11, 30, 32].

Table 3 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for governance responsibilities, implementation conditions, and monitoring requirements for the article Pharmacy Practice in the Age of Agentic Systems That Can Plan, Act, and Escalate.

Table 3. Governance responsibilities, implementation conditions, and monitoring requirements for the article Pharmacy Practice in the Age of Agentic Systems That Can Plan, Act, and Escalate.

Governance domain	Responsible actor	Required authority	Documentation requirement	Monitoring indicator	Escalation trigger	Corrective action	Accountability boundary
Mandate authorization	Pharmacy governance body and clinical lead	Approve task purpose and exclusions	Proposed mandate specification	Scope deviations	Goal ambiguity or task expansion	Pause and redefine mandate	Agent cannot redefine purpose.
Permission management	Organization and accountable pharmacist	Assign and restrict P0–P4 permission	Proposed authority contract	Blocked or unauthorized actions	Permission bypass or self-expansion	Disable action access and investigate	Capability does not create authority [24].
Human readiness	Pharmacy service leadership	Confirm reviewer competence and capacity	Training and readiness record	Review delay, override pattern, workload	Inability to verify or intervene	Reduce authority or redesign workflow	Human presence alone is insufficient.
Early evaluation	Multidisciplinary evaluation team	Authorize staged testing	Protocol, error log, workflow assessment	Technical, safety, and human-factor findings	Unresolved high-	Return to simulation or shadow mode	Early evaluation does not establish

					consequence failure		deployment readiness [28].
Live monitoring	Operational owner	Maintain or restrict authorization	Proposed trace-and-consequence ledger	Action discrepancies, incomplete closure, delayed response	Drift, incident, or workflow change	Pause, rollback, requalify, or withdraw	Authorization is version- and context-specific [29].
Medication-safety oversight	Pharmacist and medication-safety team	Review medication-related consequences	Incident and near-miss record	Wrong target, timing, recommendation, or action	Actual or credible patient harm	Contain, disclose, investigate, and redesign	Agent trace does not replace professional review.
Equity oversight	Governance and patient-safety bodies	Require subgroup and proxy assessment	Equity impact record	Differential errors, access, or escalation	Persistent subgroup disadvantage	Restrict use and revise objectives or data	Aggregate accuracy cannot excuse inequity.
Retirement and handback	Organization and service owner	Withdraw system authority	Withdrawal and responsibility-transfer record	Unsupported version or unavailable oversight	Inability to monitor, maintain, or govern	Deactivate and return work to an approved pathway	Responsibility must remain identifiable after withdrawal.

Table note: The governance artifacts, permission levels, corrective actions, and accountability boundaries are proposed constructs. They require local legal, professional, technical, organizational, and empirical evaluation.

Limitations

BAPWT is a conceptual architecture assembled from emerging agent research, clinical AI evaluation, human-factors evidence, governance scholarship, and pharmacy literature. An agent may accurately optimize an inequitable proxy, so objective design and subgroup consequences require explicit examination [34]. Action-capable systems may also be vulnerable to manipulated inputs, compromised tools, or adversarial behavior [35].

Explanations can support review but do not establish causal correctness, validity, safety, or appropriate reliance [36]. Moreover, performance and governance findings may not transfer across institutions, populations, workflows, integrations, or authority arrangements [37]. The proposed architecture therefore requires prospective, local, multidisciplinary validation and cannot establish universal permission levels, staffing requirements, legal responsibility, or clinical benefit.

CONCLUSION

Agentic systems change the pharmacy-governance problem because they may organize and initiate multistep work rather than merely produce predictions or recommendations. BAPWT proposes that such capability should be governed through a bounded work mandate, an explicit action envelope, permission and authority controls, consequence monitoring, escalation pathways, and accountable human and organizational oversight.

The architecture makes technical capability necessary but insufficient. Pharmacy value must be assessed through action correctness, medication-safety consequences, workflow effects, meaningful human control, equity, security, and local implementation conditions. Authority should remain conditional, reversible, and linked to reconstructable responsibility.

The proposed theory, permission levels, transition gates, and governance mechanisms are intended to structure empirical inquiry and professional debate. They do not demonstrate that agentic systems improve pharmacy practice, replace pharmacists, or warrant autonomous medication decisions. The central decision boundary is that an agent may plan or act only within evidence-supported, explicitly authorized, locally governed conditions, with a viable path for human intervention, restriction, and withdrawal.

ACKNOWLEDGMENTS: None

CONFLICT OF INTEREST: None

FINANCIAL SUPPORT: None

ETHICS STATEMENT: None

REFERENCES

- Sutton RT, Pincock D, Baumgart DC, Sadowski DC, Fedorak RN, Kroeker KI. An overview of clinical decision support systems: Benefits, risks, and strategies for success. *NPJ Digit Med.* 2020;3:17. doi:10.1038/s41746-020-0221-y
- Topol EJ. High-performance medicine: The convergence of human and artificial intelligence. *Nat Med.* 2019;25(1):44-56. doi:10.1038/s41591-018-0300-7
- Rajkomar A, Dean J, Kohane I. Machine learning in medicine. *N Engl J Med.* 2019;380(14):1347-58. doi:10.1056/NEJMr1814259
- Wiens J, Saria S, Sendak M, Ghassemi M, Liu VX, Doshi-Velez F, et al. Do no harm: A roadmap for responsible machine learning for health care. *Nat Med.* 2019;25(9):1337-40. doi:10.1038/s41591-019-0548-6
- Kelly CJ, Karthikesalingam A, Suleyman M, Corrado G, King D. Key challenges for delivering clinical impact with artificial intelligence. *BMC Med.* 2019;17(1):195. doi:10.1186/s12916-019-1426-2
- Liu F, Niu Y, Zhang Q, Wang K, Dong Z, Wong IN, et al. A foundational architecture for AI agents in healthcare. *Cell Rep Med.* 2025;6(10):102374. doi:10.1016/j.xcrm.2025.102374
- Collaco BG, Haider SA, Prabha S, Gomez-Cabello CA, Genovese A, Wood NG, et al. The role of agentic artificial intelligence in healthcare: A scoping review. *NPJ Digit Med.* 2026;9(1):345. doi:10.1038/s41746-026-02517-5
- Mehandru N, Miao BY, Almaraz ER, Sushil M, Butte AJ, Alaa A. Evaluating large language models as agents in the clinic. *NPJ Digit Med.* 2024;7(1):84. doi:10.1038/s41746-024-01083-y
- Ferber D, Hilgers L, Höper C, Kinny-Köster B, Eckardt JN, Egger-Heidrich K, et al. Towards autonomous medical artificial intelligence

- agents. *Nature*. 2026;655(8125):1282-91. doi:10.1038/s41586-026-10675-5
10. Chen X, Xiang J, Lu S, Liu Y, He M, Shi D. Evaluating large language models and agents in healthcare: Key challenges in clinical applications. *Intell Med*. 2025;5(2):151-63. doi:10.1016/j.imed.2025.03.002
11. Carayon P, Wooldridge AR, Hoonakker P, Hundt AS, Kelly MM. SEIPS 3.0: Human-centered design of the patient journey for patient safety. *Appl Ergon*. 2020;84:103033. doi:10.1016/j.apergo.2019.103033
12. Lebovitz S, Lifshitz-Assaf H, Levina N. To engage or not to engage with AI for critical judgments: How professionals deal with opacity when using AI for medical diagnosis. *Organ Sci*. 2022;33(1):126-48. doi:10.1287/orsc.2021.1549
13. Tschandl P, Rinner C, Apalla Z, Argenziano G, Codella NCF, Halpern AC, et al. Human-computer collaboration for skin cancer recognition. *Nat Med*. 2020;26(8):1229-34. doi:10.1038/s41591-020-0942-0
14. Gaube S, Suresh H, Raue M, Merritt A, Berkowitz SJ, Lerner E, et al. Do as AI say: Susceptibility in deployment of clinical decision-aids. *NPJ Digit Med*. 2021;4(1):31. doi:10.1038/s41746-021-00385-9
15. Cabitza F, Rasoini R, Gensini GF. Unintended consequences of machine learning in medicine. *JAMA*. 2017;318(6):517-8. doi:10.1001/jama.2017.7797
16. Lyell D, Coiera E. Automation bias and verification complexity: A systematic review. *J Am Med Inform Assoc*. 2017;24(2):423-31. doi:10.1093/jamia/ocw105
17. Endsley MR. From here to autonomy: Lessons learned from human-automation research. *Hum Factors*. 2017;59(1):5-27. doi:10.1177/0018720816681350
18. Kwan JL, Lo L, Ferguson J, Goldberg H, Diaz-Martinez JP, Tomlinson G, et al. Computerised clinical decision support systems and absolute improvements in care: Meta-analysis of controlled clinical trials. *BMJ*. 2020;370:m3216. doi:10.1136/bmj.m3216
19. Wong A, Otles E, Donnelly JP, Krumm A, McCullough J, DeTroyer-Cooley O, et al. External validation of a widely implemented proprietary sepsis prediction model in hospitalized patients. *JAMA Intern Med*. 2021;181(8):1065-70. doi:10.1001/jamainternmed.2021.2626
20. Liu Y, Carrero ZI, Jiang X, Ferber D, Wölflein G, Zhang L, et al. Benchmarking large language model-based agent systems for clinical decision tasks. *NPJ Digit Med*. 2026;9(1):259. doi:10.1038/s41746-026-02443-6
21. Reddy S, Allan S, Coghlan S, Cooper P. A governance model for the application of AI in health care. *J Am Med Inform Assoc*. 2020;27(3):491-7. doi:10.1093/jamia/ocz192
22. Char DS, Shah NH, Magnus D. Implementing machine learning in health care—addressing ethical challenges. *N Engl J Med*. 2018;378(11):981-3. doi:10.1056/NEJMp1714229
23. Vayena E, Blasimme A, Cohen IG. Machine learning in medicine: Addressing ethical challenges. *PLoS Med*. 2018;15(11):e1002689. doi:10.1371/journal.pmed.1002689
24. Price WN, Gerke S, Cohen IG. Potential liability for physicians using artificial intelligence. *JAMA*. 2019;322(18):1765-6. doi:10.1001/jama.2019.15064
25. McDougall RJ. Computer knows best? The need for value-flexibility in medical AI. *J Med Ethics*. 2019;45(3):156-60. doi:10.1136/medethics-2018-105118
26. Liu X, Cruz Rivera S, Moher D, Calvert MJ, Denniston AK, SPIRIT-AI and CONSORT-AI Working Group. Reporting guidelines for clinical trial reports for interventions involving artificial intelligence: The CONSORT-AI extension. *Nat Med*. 2020;26(9):1364-74. doi:10.1038/s41591-020-1034-x
27. Cruz Rivera S, Liu X, Chan AW, Denniston AK, Calvert MJ, SPIRIT-AI and CONSORT-AI Working Group. Guidelines for clinical trial protocols for interventions involving artificial intelligence: The SPIRIT-AI extension. *Nat Med*. 2020;26(9):1351-63. doi:10.1038/s41591-020-1037-7
28. Vasey B, Nagendran M, Campbell B, Clifton DA, Collins GS, Denaxas S, et al. Reporting guideline for the early-stage clinical evaluation of decision support systems driven by artificial intelligence: DECIDE-AI. *Nat Med*. 2022;28(5):924-33. doi:10.1038/s41591-022-01772-9
29. Labkoff S, Oladimeji B, Kannry J, Solomonides A, Leftwich R, Koski E, et al. Toward a responsible future: Recommendations for AI-enabled clinical decision support. *J Am Med Inform Assoc*. 2024;31(11):2730-9. doi:10.1093/jamia/ocae209
30. Hatzimanolis J, Riley B, El-Den S, Aslani P, Zhou J, Chaar BB. Applications of artificial intelligence in current pharmacy practice: A scoping review. *Res Social Adm Pharm*. 2025;21(3):134-41. doi:10.1016/j.sapharm.2024.12.007
31. Graafsma J, Murphy RM, van de Garde EMW, Karapinar-Çarkit F, Derijks HJ, Hoge RHL, et al. The use of artificial intelligence to optimize medication alerts generated by clinical decision support systems: A scoping review. *J Am Med Inform Assoc*. 2024;31(6):1411-22. doi:10.1093/jamia/ocae076
32. Ryan DK, Maclean RH, Balston A, Scourfield A, Shah AD, Ross J. Artificial intelligence and machine learning for clinical pharmacology. *Br J Clin Pharmacol*. 2024;90(3):629-39. doi:10.1111/bcp.15930
33. Alqahtani SS, Menachery SJ, Alshahrani A, Albalkhi B, Alshayban D, Iqbal MZ. Artificial intelligence in clinical pharmacy—a systematic review of current scenario and future perspectives. *Digit Health*. 2025;11:20552076251388145. doi:10.1177/20552076251388145
34. Obermeyer Z, Powers B, Vogeli C, Mullainathan S. Dissecting racial bias in an algorithm used to manage the health of populations. *Science*. 2019;366(6464):447-53. doi:10.1126/science.aax2342
35. Finlayson SG, Bowers JD, Ito J, Zittrain JL, Beam AL, Kohane IS. Adversarial attacks on medical machine learning. *Science*. 2019;363(6433):1287-9. doi:10.1126/science.aaw4399
36. Ghassemi M, Oakden-Rayner L, Beam AL. The false hope of current approaches to explainable artificial intelligence in health care. *Lancet Digit Health*. 2021;3(11):e745-50. doi:10.1016/S2589-7500(21)00208-9
37. Futoma J, Simons M, Panch T, Doshi-Velez F, Celi LA. The myth of generalisability in clinical research and machine learning in health care. *Lancet Digit Health*. 2020;2(9):e489-92. doi:10.1016/S2589-7500(20)30186-2