



LEARNING THE NEXT EXPERIMENT: A DECISION ARCHITECTURE FOR CLOSED-LOOP DRUG DISCOVERY UNDER COST, UNCERTAINTY, AND SCIENTIFIC VALUE

Daniel Osei^{1*}, Koenraad Janssens², Akua Asare³, Mohammed El-Sayed⁴

1. *Department of Closed-Loop Discovery and Experimental Design, Faculty of Pharmacy, University of Ghana, Accra, Ghana.*
2. *Department of Decision Architecture and Uncertainty Management, Faculty of Pharmaceutical Sciences, KU Leuven, Leuven, Belgium.*
3. *Department of Cost-Effective Experiment Planning, Faculty of Pharmacy, Kwame Nkrumah University of Science and Technology, Kumasi, Ghana.*
4. *Department of Scientific Value Assessment in Drug Discovery, Faculty of Pharmacy, Cairo University, Cairo, Egypt.*

ARTICLE INFO

Received:

07 June 2024

Received in revised form:

15 October 2024

Accepted:

19 October 2024

Available online:

28 October 2024

Keywords: Active learning, Closed-loop drug discovery, Experiment selection, Uncertainty quantification, Scientific value, Information asymmetry

ABSTRACT

Closed-loop drug discovery seeks to connect computational inference, experiment selection, physical execution, measurement, and model updating into an iterative learning process. Yet the central decision in such a loop—what experiment should be performed next—cannot be resolved reliably by maximizing predicted activity, model confidence, expected improvement, or any other single performance measure. Candidate experiments differ not only in their likelihood of producing desirable outcomes but also in their capacity to reduce uncertainty, distinguish competing mechanisms, broaden chemical or biological coverage, reveal model failure, consume scarce resources, and delay other decisions. This article develops a proposed Next-Experiment Decision Architecture for organizing these non-equivalent considerations. The architecture treats experiment selection as a staged scientific decision comprising admissibility assessment, multidimensional value characterization, constrained portfolio assembly, and accountable authorization, execution, and adaptive control. It separates prediction from experimental confirmation, confidence from calibrated uncertainty, structural novelty from mechanistic learning, and data availability from decision suitability. It further introduces an evidence ledger that preserves model, assay, procedural, and human-decision provenance across iterations. The synthesis indicates that closed-loop value depends on the relationship between an experiment and the unresolved decision state, rather than on an intrinsic acquisition score. Important limitations include context dependence, incomplete uncertainty models, difficulty comparing heterogeneous value dimensions, and the absence of prospective validation across discovery stages and automated platforms. The proposed architecture is therefore not a deployment-ready decision tool. It is a conceptual and methodological framework intended to support testable experiment-selection policies, transparent failure handling, and more scientifically bounded implementation of active learning in pharmaceutical discovery.

This is an **open-access** article distributed under the terms of the [Creative Commons Attribution-Non Commercial-Share Alike 4.0 License](https://creativecommons.org/licenses/by-nc-sa/4.0/), which allows others to remix, and build upon the work non commercially.

To Cite This Article: Osei D, Janssens K, Asare A, El-Sayed M. Learning the Next Experiment: A Decision Architecture for Closed-Loop Drug Discovery under Cost, Uncertainty, and Scientific Value. *Pharmacophore*. 2024;15(5):100-10. <https://doi.org/10.51847/m8LTGDHsle>

Introduction

Artificial intelligence and machine learning have expanded the scale at which molecular structures can be generated, prioritized, and evaluated, but their contribution to drug discovery remains conditional on how computational outputs are connected to medicinal-chemistry reasoning, experimental evidence, and consequential project decisions. Predictive performance can be useful without being sufficient: a model may rank compounds accurately within a benchmark while remaining poorly suited to selecting an experiment whose result will alter a real discovery decision. The practical value of computational discovery therefore depends on task definition, data validity, experimental accessibility, and the scientific meaning of the evidence that a model helps obtain [1–3]. This distinction becomes especially important in closed-loop systems,

Corresponding Author: Daniel Osei; Department of Closed-Loop Discovery and Experimental Design, Faculty of Pharmacy, University of Ghana, Accra, Ghana. E-mail: daniel.osei@ug.edu.gh.

where an apparently minor weakness in candidate selection can be propagated through repeated model updates and automated execution.

Generative molecular-design systems illustrate the difference between producing candidates and learning what should be tested. Such systems can generate large numbers of structures that satisfy encoded objectives, but those objectives may only partially represent the requirements of a pharmaceutical project. A proposed molecule can appear favorable under a computational score while remaining difficult to synthesize, chemically unstable, poorly matched to the available assay, redundant with existing evidence, or uninformative about the biological mechanism under investigation. Generative output must therefore be treated as a source of candidate hypotheses rather than as a self-validating discovery result [4]. The selection problem begins after generation: the system must determine which candidate experiment is sufficiently feasible, informative, and decision-relevant to justify consuming laboratory capacity.

This problem is not reducible to predictive accuracy because pharmaceutical research and development proceeds through interacting dimensions of evidence and value. Target rationale, molecular properties, tissue exposure, safety, patient relevance, and operational feasibility can constrain one another, and weakness in one dimension may not be compensated by strength in another. Industrial experience with multidimensional productivity frameworks supports the broader proposition that discovery progress depends on coordinated scientific and decision criteria rather than isolated technical achievements [5]. At the experiment level, the same logic implies that a high predicted response may have limited value when the measurement does not address the project's principal uncertainty, when the assay result cannot discriminate among competing explanations, or when a more modest experiment would resolve a critical feasibility question sooner.

The unresolved conceptual gap is therefore not whether active learning can rank unlabeled candidates, but how a closed-loop discovery system should define and govern the value of its next experimental action. This article proposes a Next-Experiment Decision Architecture that organizes experiment selection into four connected stages: scientific and operational admissibility, multidimensional value characterization, constrained portfolio assembly, and accountable authorization and adaptive control. The proposed contribution does not claim that these stages constitute an empirically validated optimization policy. Instead, it specifies the decision objects, relationships, boundaries, and evaluation needs required to compare future policies without conflating benchmark improvement, information acquisition, mechanistic interpretation, experimental confirmation, and pharmaceutical usefulness.

The Closed-Loop Discovery Cycle and the Value of the Next Experiment

A closed-loop discovery cycle links the current state of knowledge to a candidate action, executes that action, evaluates the resulting evidence, and uses the evidence to revise subsequent decisions. Self-driving laboratory concepts formalize this iterative relation among experiment design, execution, measurement, analysis, and updating, while robotic reactivity searches and AI-informed synthesis systems demonstrate that portions of the cycle can be coupled within bounded chemical domains [6–8]. These systems establish that experimental order need not be fixed in advance. They do not, however, establish that every computable acquisition objective represents pharmaceutical value. The loop can execute efficiently while pursuing a poorly specified endpoint, repeatedly sampling a biased region, or producing precise answers to questions that no longer matter to the project.

The “next experiment” should therefore be defined as a structured scientific action rather than as a molecular item at the top of a ranked list. It comprises an intervention or material, an experimental context, a measurement endpoint, appropriate controls, execution requirements, and an anticipated role in a decision. Bayesian reaction optimization shows how accumulated observations and predictive uncertainty can be used to select subsequent experimental conditions when evaluations are limited [9]. Within drug discovery, however, the selected action may concern synthesis, binding, phenotype, exposure, selectivity, toxicity, formulation, or mechanism, and each domain has different sources of noise and different consequences for being wrong. A transferable architecture must represent these differences explicitly rather than assuming that one acquisition function can compare all possible experiments on a common scientific scale.

The value of a next experiment is relational: it depends on the current evidence state, the unresolved decision, and the alternatives that could be performed instead. An experiment may have exploitation value because it tests a candidate expected to perform well; exploration value because it probes a poorly represented region; diversity value because it reduces redundancy within a batch; or mechanistic value because it distinguishes competing explanations. It may also have diagnostic value when it tests model calibration, assay reliability, synthesis feasibility, or domain transfer. Reviews of active learning in drug discovery indicate that practical performance is sensitive to data regimes, learning algorithms, acquisition policies, batch construction, and validation protocols [10]. Accordingly, experiment value cannot be assigned independently of the model, laboratory, decision stage, and evidence already accumulated.

The proposed architecture represents next-experiment value as a multidimensional profile rather than an intrinsic scalar. This profile should preserve at least the expected decision-relevant outcome, anticipated information gain, contribution to evidence diversity, capacity for mechanistic discrimination, uncertainty, cost, delay, execution risk, and downstream usability of the result. Scalarization may still be used within a defined decision context, but weights and trade-offs should remain inspectable and should not allow a high predicted benefit to compensate automatically for scientific invalidity or infeasibility. **Table 1** organizes the evidence, constructs, and interpretation boundaries needed to develop the closed-loop discovery cycle and the value of the next experiment within the article's central argument.

Table 1. Evidence domains, core questions, scientific requirements, and interpretation boundaries for The closed-loop discovery cycle and the value of the next experiment

Construct or evidence domain	Core question	Scientific basis	Role in the article	Failure or overclaiming risk	Boundary statement
Discovery-state evidence ledger	What is currently known, inferred, missing, failed, or contested?	Versioned molecular data, assay results, computational outputs, protocol records, negative findings, and human decisions	Defines the state from which the next experiment derives its value	Treating an incomplete or biased record as a complete representation of the discovery problem	Recorded evidence may remain unsuitable, inconsistent, or affected by selection bias
Candidate experiment object	What exactly would be performed, measured, and interpreted?	Intervention, material, context, endpoint, controls, procedure, and anticipated decision use	Prevents the architecture from equating a molecule with a fully specified experiment	Ranking compounds without defining the assay, controls, or interpretation pathway	A molecular candidate is not equivalent to an executable and interpretable experiment
Current decision state	Which scientific or project decision remains unresolved?	Explicit progression question, competing hypotheses, uncertainty, and available alternatives	Anchors experimental value to a consequential question	Optimizing an endpoint that does not alter the relevant decision	A statistically informative result may remain decision-irrelevant
Expected outcome value	Could the experiment identify a desirable or actionable result?	Predictive models, prior observations, domain knowledge, and endpoint relevance	Represents the exploitation-oriented component of value	Treating a high predicted result as experimental confirmation or therapeutic promise	Prediction supports prioritization but does not establish activity, safety, or utility
Information value	How much could the result reduce decision-relevant uncertainty?	Predictive distributions, model disagreement, alternative explanations, and expected evidence change	Represents learning beyond immediate candidate performance	Maximizing abstract uncertainty reduction without resolving a pharmaceutical question	Information gain is valuable only relative to a specified decision or knowledge gap
Diversity contribution	Does the experiment reduce redundancy and expand evidential coverage?	Structural, physicochemical, biological, assay, and mechanistic relationships	Supports batch complementarity and protection against local overexploitation	Equating structural dissimilarity with biological or mechanistic diversity	Structural novelty alone does not establish scientific novelty
Mechanistic discrimination	Can the experiment distinguish among plausible explanations?	Competing hypotheses, perturbations, orthogonal endpoints, temporal evidence, and controls	Separates mechanism-directed learning from predictive improvement	Presenting association, attribution, or model explanation as causal mechanism	Mechanistic interpretation requires appropriately designed experimental evidence
Cost, delay, and opportunity cost	What resources are consumed, and which alternatives are displaced?	Synthesis burden, assay duration, material requirements, instrument availability, and project timing	Constrains the feasible and preferable next action	Using generic cost assumptions as if they were locally valid	Resource estimates are context-dependent and must be updated as execution conditions change
Executability and quality	Can the experiment be performed reproducibly and interpreted reliably?	Protocol maturity, material availability, equipment state, controls, and quality criteria	Establishes non-compensatory admissibility before optimization	Allowing high predicted value to override poor controls or infeasible execution	Eligibility for execution does not establish expected scientific value
Returned evidence	Can the result legitimately update the discovery state?	Quality control, provenance, protocol adherence, censoring, and result interpretation	Closes the loop and determines whether learning should occur	Automatically training on failed, deviated, or uninterpretable experiments	Every result should be quality-checked before it modifies the model or decision state
Human authorization	Who may approve, defer, override, escalate, or stop the experiment?	Scientific responsibility, expertise, documented reasoning, and governance	Preserves accountability for consequential selections	Treating human involvement as automatic protection against error or bias	Oversight improves contestability but does not guarantee correctness
Continuation value	Is another iteration preferable to stopping, redirecting, or escalating?	Remaining uncertainty, marginal information, resource state, and unresolved conflicts	Connects experiment selection to explicit continuation policies	Continuing because automation can proceed rather than because the next iteration has scientific value	A closed loop should not be assumed to require indefinite continuation

Sources of Uncertainty, Cost, Delay, and Information Asymmetry

Uncertainty in closed-loop drug discovery is not a single quantity. It can arise from limited or biased data, model specification, parameter estimation, assay variability, measurement censoring, domain shift, unmodeled biology, uncertain synthesis outcomes, or execution deviations. Comparative studies of molecular uncertainty quantification show that estimation methods differ in calibration, scalability, and behavior across tasks, while uncertainty-based error-control approaches demonstrate that uncertainty can help identify predictions at greater risk of failure [11–13]. These findings support using uncertainty as a decision input, but not treating it as self-interpreting evidence. A numerical confidence value may combine non-equivalent sources of uncertainty, remain poorly calibrated outside the training domain, or omit uncertainties that were never represented by the model.

Information asymmetry arises when the decision system observes only a selective and imperfect record of the relevant scientific space. Tested compounds are shaped by historical project choices, synthesis accessibility, assay availability, publication practices, and prior model preferences. The resulting dataset may contain extensive information about familiar

chemical series while providing little evidence about structurally or biologically distinct regions. Ligand-based benchmarks can reward memorization of similarity patterns rather than genuine generalization, illustrating how apparently strong predictive results may conceal dependence on information unavailable in prospective settings [14]. In a closed loop, this asymmetry can become self-reinforcing: the model selects candidates similar to those it already understands, receives confirming labels from that region, and becomes increasingly confident without resolving its lack of coverage elsewhere.

Cost and delay further distinguish experiment selection from ordinary ranking. An experiment consumes reagents, compounds, biological materials, equipment time, specialist labor, analytical capacity, and organizational attention. It may also delay a competing experiment whose result would have greater decision value. Estimates of medicine-development investment demonstrate the broader resource burden and risk surrounding pharmaceutical research [15], although such aggregate analyses should not be converted into invented local experiment costs. The proposed architecture therefore treats cost and delay as context-specific attributes that must be measured within the relevant laboratory and project. A low-cost experiment may be preferable when it rapidly resolves feasibility, but it may be wasteful if it yields evidence too weak to alter the decision. Conversely, a costly experiment may be justified when it discriminates among high-consequence alternatives that cheaper proxies cannot resolve.

These factors interact rather than operating as independent penalties. High uncertainty may increase the potential information value of an experiment, but only when the uncertainty is credible, reducible, and connected to a relevant decision. Information asymmetry may justify exploration, yet exploration into an infeasible or uninterpretable domain does not create useful evidence. Delay may reduce the value of an otherwise informative experiment when a project decision is time-sensitive, while rapid execution may be misleading if weak quality control produces an unusable result. The proposed architecture consequently requires separate representations of uncertainty, resource burden, evidential coverage, and expected decision effect. It does not assume that these dimensions can always be combined objectively, and it treats unresolved trade-offs as reasons for sensitivity analysis, expert review, escalation, or redesign rather than as invitations to conceal them within a single acquisition score.

Proposed Decision Architecture for Experiment Selection

The proposed Next-Experiment Decision Architecture begins from the premise that acquisition rules are components of a decision system rather than complete decision policies. Comparative evaluation of active-learning protocols for ligand-binding affinity prediction shows that outcomes depend on the learner, acquisition strategy, batch design, data regime, and evaluation procedure [16]. The architecture therefore requires every selection policy to declare its model assumptions, candidate pool, optimization objective, available evidence, resource constraints, and intended decision use. These declarations form part of the experiment-selection record and prevent a favorable outcome under one protocol from being generalized automatically to another target, assay, or discovery stage.

The first architectural layer is an admissibility gate that determines whether a candidate experiment is sufficiently specified, interpretable, feasible, and quality-controllable to enter value comparison. The second layer characterizes eligible experiments without immediately collapsing their attributes into a single score. Pool-based active learning can concentrate computational evaluation on selected regions of very large molecular spaces [17], but efficient narrowing alone does not determine whether the resulting candidates are experimentally accessible or scientifically informative. The proposed architecture therefore separates scalable computational prioritization from the stronger decision to allocate synthesis, assay, or expert resources.

The third layer constructs an experimental portfolio from candidates with different value profiles. Active-learning-guided lead optimization illustrates how surrogate models can allocate expensive binding free-energy calculations among candidate compounds [18]. Active learning has also been adapted to plasma-exposure prediction, demonstrating that experiment selection can be tailored to a noisy pharmaceutical endpoint and its chemical domain [19]. These examples support domain-specific acquisition, but they do not establish a universal selection rule. Within the proposed architecture, each experiment retains a reason-coded role, such as exploitation, uncertainty reduction, evidence diversification, feasibility testing, or discrimination between competing mechanistic hypotheses.

The fourth layer governs authorization, execution, evidence return, and adaptive control. Surrogate-assisted platforms can reduce the fraction of an ultra-large molecular library receiving an expensive docking evaluation [20], but docking scores remain computational proxies rather than confirmation of binding, efficacy, developability, or therapeutic value. The architecture therefore records which evidence tier an experiment can produce and which downstream actions that evidence may legitimately support. **Table 2** organizes the evidence, constructs, and interpretation boundaries needed to develop proposed decision architecture for experiment selection within the article's central argument. **Figure 1** presents the closed-loop experimental value compass, showing how the article's key components and boundaries are connected within proposed decision architecture for experiment selection.

Table 2. Components, relationships, uncertainties, and validation needs within Proposed decision architecture for experiment selection

Component or layer	Inputs	Core function	Expected output	Uncertainty or limitation	Validation requirement
Discovery-state evidence ledger	Experimental observations, predictions, failed studies, protocol	Establish the current state of knowledge and unresolved questions	Versioned and auditable decision state	Records may be incomplete, biased, inconsistent, or outdated	Audit completeness, provenance quality, and update fidelity

versions, provenance, and human decisions					
Candidate experiment representation	Candidate material, intervention, assay, controls, endpoint, procedure, and resource requirements	Convert a molecular proposal into a fully specified experimental action	Comparable candidate-experiment object	Some execution conditions or outcomes may remain unknown	Test whether representations are sufficient for reproducible execution and interpretation
Admissibility gate	Scientific rationale, assay validity, synthesis feasibility, safety constraints, and equipment state	Exclude, defer, or escalate candidates that should not enter compensatory ranking	Set of eligible experiments	Eligibility criteria may be context-dependent or incompletely observed	Prospective review of false exclusion and inappropriate admission
Uncertainty decomposition	Model distributions, calibration evidence, assay variability, domain-shift indicators, and execution risk	Separate non-equivalent sources of uncertainty	Qualified uncertainty profile	Methods may be miscalibrated or omit unknown sources	Temporal, scaffold, assay, and external calibration studies
Cost–delay model	Material use, labor, instrument availability, queue time, synthesis burden, and opportunity cost	Represent local resource consequences	Context-specific resource profile	Estimates can change after selection	Compare predicted and realized cost and delay
Scientific-value characterization	Expected outcome, information gain, diversity, mechanistic discrimination, and decision relevance	Preserve distinct reasons for conducting an experiment	Multidimensional value profile	Attributes may use incompatible scales and uncertain assumptions	Sensitivity analysis and prospective decision-utility evaluation
Portfolio assembly	Eligible experiments and value profiles	Select a complementary batch or sequence under constraints	Reason-coded experimental portfolio	Trade-offs may be underdetermined or stakeholder-dependent	Compare against random, expert, exploitation-only, and alternative active-learning policies
Human authorization	Objectives, assumptions, limitations, conflicts, and proposed reasons	Authorize, defer, modify, escalate, or stop consequential actions	Accountable decision record	Human review can introduce bias, inconsistency, or automation dependence	Human-factors studies and override audits
Execution and provenance layer	Authorized protocol, materials, equipment state, controls, and stopping conditions	Perform the experiment and preserve deviations and lineage	Quality-controlled observations and execution record	Hardware or procedural failure may mimic scientific failure	Reproducibility, protocol-fidelity, and failure-injection testing
Adaptive controller	Returned evidence, calibration status, marginal value, quality signals, and resource state	Continue, redirect, escalate, recover, or stop the loop	Updated decision state and next authorized action	Counterfactual value of unperformed experiments is difficult to observe	Prospective trajectory and policy-comparison studies

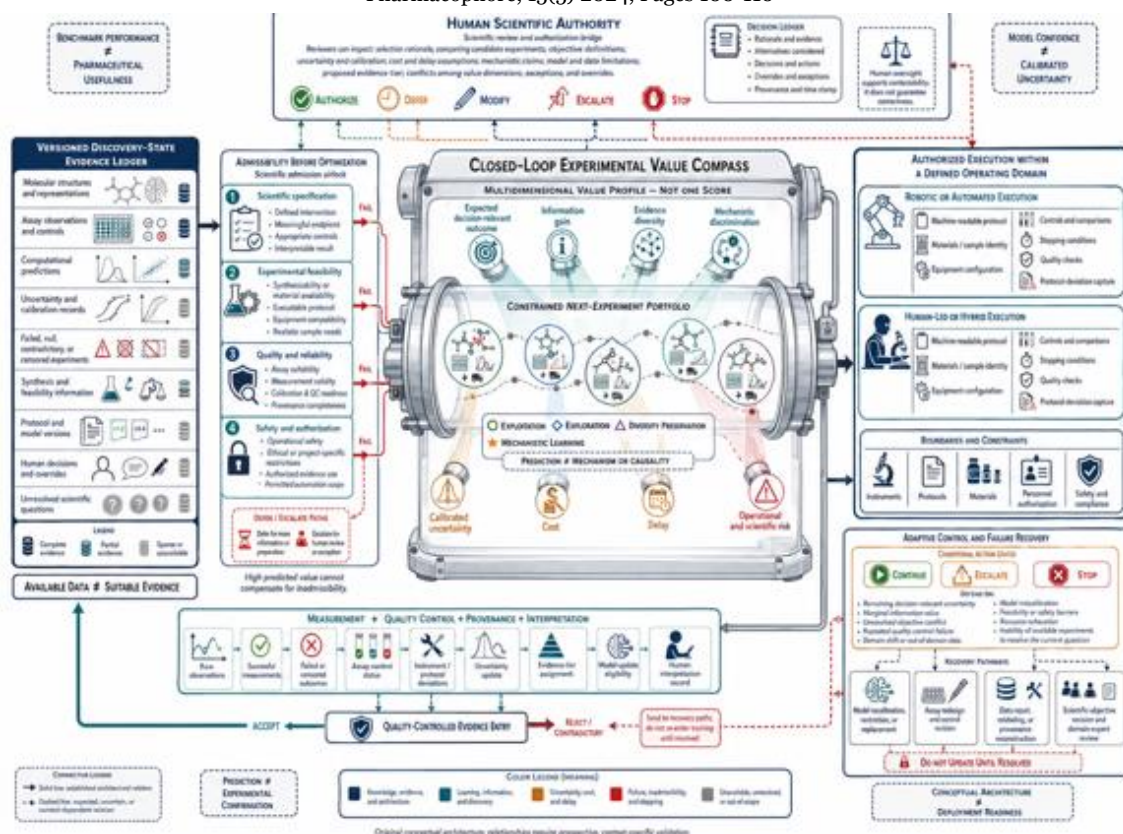


Figure 1. Closed-Loop Experimental Value Compass.

The figure is an original conceptual synthesis that organizes the article's central contribution across closed-loop discovery cycle and the value of the next experiment, sources of uncertainty, cost, delay, and information asymmetry, sources of uncertainty, information asymmetry, proposed decision architecture for experiment selection, balancing exploitation, exploration, diversity, and mechanistic learning. Arrows and grouping indicate proposed or evidence-supported relationships rather than measured effect sizes. The figure does not represent a validated predictive model, regulatory determination, clinical recommendation, or deployment-ready system.

Balancing Exploitation, Exploration, Diversity, and Mechanistic Learning

Exploitation prioritizes experiments expected to produce favorable outcomes under the current model. It is appropriate when the objective is well specified, the model is sufficiently supported in the relevant domain, and the project benefits from concentrating resources around promising candidates. Inverse molecular design demonstrates that generated candidates are inseparable from the objectives and constraints used to define desirability [21]. The architecture therefore treats exploitation as conditional on objective validity. A compound that optimizes the wrong endpoint, operates outside the model's supported domain, or duplicates evidence already available may have little next-experiment value despite a favorable score.

Exploration prioritizes experiments that expand knowledge beyond well-characterized regions. Diversity preservation is related but distinct: it limits redundancy within a selected batch and can operate across structural, physicochemical, biological, assay, or mechanistic dimensions. Benchmarking work in de novo molecular design separates goal-directed success from distribution learning, novelty, and diversity, showing that these dimensions describe different aspects of model behavior [22]. In the proposed architecture, exploration is justified by an unresolved knowledge gap, whereas diversity is justified by the need for complementary evidence. Neither should be represented through structural analysis alone when the scientific question concerns biological or mechanistic coverage.

Mechanistic learning requires experiments that can discriminate among explicit alternative explanations. It should not be equated with model interpretability, feature attribution, chemical novelty, or improvement in predictive error. The synthesizability of generated molecules illustrates how success under a computational objective can diverge from practical experimental feasibility [23]. An analogous divergence occurs when a candidate appears predictive but cannot distinguish causal from correlational explanations. Mechanistic experiments generally require appropriate perturbations, controls, orthogonal measurements, temporal structure, or counterfactual comparisons. Their immediate performance value may be lower than that of exploitation candidates, while their long-term scientific value may be greater.

The proposed portfolio layer balances these modes without assuming that they are commensurable in every context. Multi-objective molecular optimization provides methods for representing competing property goals and Pareto trade-offs [24], but a Pareto-efficient candidate is not necessarily the scientifically preferred experiment. Selection depends on the current discovery stage, unresolved risks, evidence quality, laboratory constraints, and tolerance for delayed benefit. The architecture

therefore permits explicit weighting, constrained optimization, or Pareto-set selection, provided that assumptions remain visible. When reasonable weighting schemes produce materially different portfolios, instability should be reported and escalated rather than hidden through arbitrary scalarization.

Stopping, Escalation, and Failure-Recovery Rules

A closed loop requires an explicit basis for deciding whether another iteration remains worthwhile. Bayesian optimization methods such as Phoenix illustrate how posterior beliefs can guide experiment selection when evaluations are scarce [25]. The proposed architecture extends this logic by distinguishing predictive improvement from scientific continuation value. A loop may stop when the current decision has been resolved, when expected marginal learning is too small relative to cost and delay, when remaining uncertainty cannot be reduced by available experiments, or when another experimental strategy has greater opportunity value. Stopping does not imply that the scientific space is exhausted; it means that continuation is no longer justified under the declared objective and constraints.

Escalation is required when the system encounters a conflict it is not authorized or evidentially equipped to resolve. Multi-objective active learning can preserve competing objectives during sequential optimization [26], but it cannot determine which trade-off an organization or scientific team should prefer without a defensible decision rule. Escalation triggers may include unstable rankings across plausible weights, conflict between efficacy and safety-related objectives, insufficient evidence for a consequential action, unsupported domain transfer, or disagreement between computational and experimental indicators. Escalation may request stronger assays, orthogonal evidence, reformulation of the objective, or accountable expert judgment. Failure recovery begins by distinguishing scientific failure from data, model, assay, and execution failure. Historical experimental datasets can encode human selection biases that restrict the observed search space and thereby hinder exploration [27]. Repeatedly selecting within such a record may intensify the bias while appearing to improve internal performance. A recovery controller should therefore detect concentration of sampling, loss of calibration, contradictory measurements, missing provenance, control failure, and procedural deviation. Suspect results should be quarantined until their evidential status is resolved rather than incorporated automatically into model retraining.

Recovery actions must correspond to the suspected failure source. Uncertainty-calibrated semi-supervised learning may help acquisition under sparse molecular labels [28], but additional learning cannot repair every problem. Model failure may require recalibration, replacement, or restriction of the supported domain. Data failure may require relabeling, provenance reconstruction, or recollection. Assay failure may require revised controls or a different endpoint, while execution failure may require instrument inspection and protocol repetition. Governance failure—including unauthorized objective changes or undocumented overrides—requires suspension and decision reconstruction. Recovery is therefore a controlled scientific response, not merely another optimization iteration.

Human Oversight and Implementation in Automated Discovery Platforms

Implementation requires translating an authorized scientific decision into reproducible physical or computational operations. Modular robotic synthesis systems demonstrate that machine-readable procedures can coordinate hardware and chemical operations within a supported environment [29]. For the proposed architecture, this means that the experiment-selection record must connect directly to an executable protocol, material specification, control structure, quality criteria, and provenance model. The system should preserve the distinction between selecting an experiment, authorizing it, executing it, and interpreting its result. Automation of one stage does not transfer authority automatically to the others.

Robotic platforms can navigate, conduct experiments, collect measurements, and update subsequent actions within a bounded domain [30]. Such demonstrations support the feasibility of closed-loop execution but not unrestricted scientific autonomy. The limits of the experimental environment, available procedures, instruments, materials, and model assumptions must remain explicit. Human operators should be able to pause execution, inspect the basis of selection, intervene during anomalous conditions, and prevent results affected by procedural deviations from entering the evidence ledger. A bounded platform may be technically autonomous while remaining scientifically and organizationally supervised.

Reliable implementation also depends on standardized procedural representations. Systems that digitize and automatically execute synthesis literature show the importance of converting natural-language procedures into explicit, machine-readable operations [31]. Yet procedural executability does not determine whether an experiment has sufficient scientific value to be selected. The proposed architecture therefore places procedural representation downstream of admissibility and value characterization but upstream of physical execution. Versioned protocols, instrument states, material identities, deviations, and quality-control outcomes should accompany every returned observation so that later model updates can be traced to the conditions under which the evidence was produced.

Human oversight should focus on consequential uncertainty, objective conflict, exceptions, and evidence boundaries rather than merely confirming an algorithmic recommendation. Explainable AI can support scrutiny and communication in drug discovery, but an explanation does not establish model correctness, causal mechanism, or pharmaceutical validity [32]. Oversight should therefore include access to objective definitions, uncertainty qualification, cost and delay assumptions, candidate alternatives, model limitations, selection reasons, and the evidence tier expected from the experiment. Reviewers should be able to authorize, defer, modify, escalate, stop, or initiate recovery, with reasons preserved for later audit. **Table 3** organizes the evidence, constructs, and interpretation boundaries needed to develop human oversight and implementation in automated discovery platforms within the article's central argument.

Table 3. Evaluation, implementation, and research priorities arising from Learning the Next Experiment: A Decision Architecture for Closed-Loop Drug Discovery under Cost, Uncertainty, and Scientific Value

Evaluation dimension	What must be tested	Suitable evidence or assessment	Failure signal	Interpretive limitation	Decision relevance
Prospective experiment-selection utility	Whether selected experiments resolve predefined pharmaceutical decisions more effectively than credible alternatives	Prospective comparisons with random, expert, exploitation-only, and alternative active-learning policies	Retrospective gains fail to improve prospective decisions	One project cannot establish general utility	Determines whether the architecture contributes beyond ranking performance
Uncertainty calibration	Whether uncertainty estimates correspond to observed error and failure under relevant shifts	Reliability analysis, coverage assessment, selective-risk evaluation, and temporal or scaffold splits	Confident errors or unstable uncertainty rankings	Calibration is conditional on model, dataset, and shift type	Determines whether uncertainty can guide selection or escalation
Cost and delay representation	Whether predicted resource requirements correspond to realized execution burden	Local workflow records, queue times, material consumption, repeat rates, and expert review	Systematically underestimated burden or displaced higher-value work	Aggregate industry estimates cannot substitute for local measurements	Supports context-specific allocation of scarce resources
Exploration and diversity	Whether selected portfolios expand useful coverage without excessive redundancy	Structural, biological, mechanistic, and evidence-level coverage analyses	Repeated selection around familiar series or one hypothesis	Distance measures may not capture biological novelty	Protects against self-reinforcing exploitation
Mechanistic discrimination	Whether selected experiments distinguish competing explanations	Predefined hypotheses, orthogonal endpoints, perturbations, controls, and causal analysis	Prediction improves without clarifying the underlying explanation	Mechanistic conclusions require appropriately designed experiments	Determines whether learning advances scientific understanding
Stopping-policy validity	Whether continuation and termination decisions avoid premature stopping and wasteful persistence	Prospective trajectories, simulations, delayed-outcome review, and opportunity-cost analysis	Loops stop before resolving key questions or continue without marginal value	Counterfactual value of unperformed experiments is partly unobservable	Governs responsible use of experimental capacity
Failure detection and recovery	Whether model, data, assay, execution, and governance failures are identified and repaired	Failure-injection tests, incident logs, quarantine records, and recovery audits	Repeated or silent propagation of the same failure	Recovery pathways are platform- and context-specific	Prevents corrupted evidence from shaping future selections
Human oversight quality	Whether decisions are understandable, contestable, and accountable	Reason-code completeness, override review, disagreement analysis, and human-factors studies	Automation bias, ceremonial review, or undocumented intervention	Human agreement does not establish scientific correctness	Allocates decision authority and escalation responsibility
Procedural interoperability	Whether selected experiments can be transferred across supported instruments and laboratories	Cross-platform execution, protocol translation, and provenance comparison	Loss of meaning or reproducibility during transfer	Portable syntax does not ensure identical scientific conditions	Determines implementation scope and infrastructure needs
Pharmaceutical usefulness	Whether the architecture improves decisions relevant to discovery progression	Longitudinal evaluation linked to predefined project questions	Optimization of proxies without improved progression decisions	Downstream outcomes depend on many factors beyond experiment selection	Separates computational efficiency from pharmaceutical value
Governance integrity	Whether objectives, permissions, exceptions, and overrides remain controlled and auditable	Authorization logs, objective-version histories, access controls, and incident review	Objective drift or unauthorized execution	Auditability alone does not establish decision validity	Protects accountability in automated environments
Readiness assessment	Whether evidence supports progression from conceptual evaluation to bounded implementation	Stage-gated validation, external testing, safety review, and operational qualification	Deployment claims precede prospective validation	Readiness is task-, laboratory-, and platform-specific	Prevents conceptual maturity from being confused with operational readiness

Limitations and Boundary Conditions

The proposed architecture is a conceptual contribution and has not been prospectively validated as a unified experiment-selection policy. Its components are supported by related work in molecular prediction, active learning, optimization,

uncertainty quantification, robotic chemistry, and pharmaceutical decision-making, but evidence for individual components does not validate their integration. Performance is likely to depend on target class, modality, discovery stage, assay system, candidate-pool construction, laboratory infrastructure, organizational priorities, and the consequences of false selection. The architecture should therefore be evaluated as a family of context-dependent policies rather than as a universal algorithm.

Several constructs remain difficult to measure. Expected information gain depends on the adequacy of the model and the set of hypotheses represented. Mechanistic value may be underestimated when important explanations are absent, while cost and delay estimates may change after synthesis or execution begins. Uncertainty estimates can remain poorly calibrated under chemical, biological, temporal, or procedural shift [11]. Portfolio weights may also reflect institutional preference rather than scientific fact. The architecture makes these assumptions visible, but transparency does not eliminate their effects or create objective commensurability among heterogeneous values.

Misuse could arise if the architecture were treated as an autonomous drug-discovery authority, a validated predictor of project success, or a mechanism for converting weak proxy evidence into high-consequence decisions. Docking, molecular-property prediction, generative design, and model explanations cannot alone establish experimental activity, causal mechanism, clinical utility, safety, developability, or regulatory acceptability. Automated execution likewise does not establish that the selected objective was appropriate. The proposed framework is limited to organizing experiment-selection reasoning and specifying how its claims should be tested; it does not replace medicinal chemistry, pharmacology, toxicology, formulation science, clinical evaluation, or accountable scientific judgment.

Research Agenda and Implementation Priorities

The first research priority is prospective comparison of next-experiment policies under predefined pharmaceutical questions. Studies should compare random selection, expert selection, exploitation-oriented acquisition, uncertainty-oriented acquisition, diversity-aware selection, mechanistic discrimination, and the proposed portfolio logic under common budgets and stopping conditions. Evaluations should record not only predictive improvement but also decision resolution, experimental usability, realized cost and delay, calibration, redundancy, and failure recovery. Active-learning benchmarks already demonstrate sensitivity to protocol design [16]; future studies must determine whether such differences persist when candidate synthesis, assay execution, and delayed outcomes are included.

The second priority is infrastructure for evidence provenance, executable experiment representations, and failure-aware model updating. Shared schemas should distinguish proposed candidates, authorized experiments, executed procedures, raw observations, quality-control outcomes, deviations, censored results, and interpretation decisions. Machine-readable synthesis procedures can support execution [31], but comparable representations are also needed for biological assays, simulations, formulation experiments, and hybrid workflows. Validation should include cross-platform transfer, deliberate failure injection, domain shift, incomplete metadata, and recovery from corrupted evidence rather than demonstrating only nominal execution. The third priority is staged implementation with explicit readiness boundaries. Early studies should operate in retrospective or simulated environments, followed by prospective advisory use in which humans retain direct experiment-selection authority. Bounded semi-automated trials may proceed only after uncertainty monitoring, provenance, stopping controls, and recovery pathways have been tested. Greater automation should require task-specific evidence that the platform can recognize unsupported domains, control objective changes, and escalate consequential ambiguity. Human oversight should itself be evaluated for automation bias, workload, disagreement, and override quality rather than assumed to be effective merely because a reviewer remains nominally present.

Conclusion

Learning the next experiment requires more than selecting the candidate with the highest predicted performance or uncertainty score. The proposed Next-Experiment Decision Architecture reframes closed-loop drug discovery as a staged scientific decision involving admissibility, multidimensional value characterization, portfolio construction, accountable authorization, evidence-preserving execution, and explicit continuation, escalation, stopping, and recovery controls. Its central contribution is to keep exploitation, exploration, diversity, mechanistic learning, uncertainty, cost, delay, and information asymmetry visible as distinct and contestable dimensions. The architecture does not establish an optimal acquisition policy, autonomous scientific authority, pharmaceutical success, clinical utility, regulatory acceptability, or deployment readiness. Its value lies in providing a bounded structure through which future experiment-selection policies can be compared prospectively, failures can be handled transparently, and automated discovery systems can remain connected to scientific questions rather than to isolated computational objectives.

Acknowledgments: None

Conflict of interest: None

Financial support: None

Ethics statement: None

References

1. Schneider P, Walters WP, Plowright AT, Sieroka N, Listgarten J, Goodnow RA Jr, et al. Rethinking drug design in the artificial intelligence era. *Nat Rev Drug Discov.* 2020;19(5):353-64. doi:10.1038/s41573-019-0050-3
2. Vamathevan J, Clark D, Czodrowski P, Dunham I, Ferran E, Lee G, et al. Applications of machine learning in drug discovery and development. *Nat Rev Drug Discov.* 2019;18(6):463-77. doi:10.1038/s41573-019-0024-5
3. Bender A, Cortés-Ciriano I. Artificial intelligence in drug discovery: What is realistic, what are illusions? Part 1: Ways to make an impact, and why we are not there yet. *Drug Discov Today.* 2021;26(2):511-24. doi:10.1016/j.drudis.2020.12.009
4. Meyers J, Fabian B, Brown N. De novo molecular design and generative models. *Drug Discov Today.* 2021;26(11):2707-15. doi:10.1016/j.drudis.2021.05.019
5. Morgan P, Brown DG, Lennard S, Anderton MJ, Barrett JC, Eriksson U, et al. Impact of a five-dimensional framework on R&D productivity at AstraZeneca. *Nat Rev Drug Discov.* 2018;17(3):167-81. doi:10.1038/nrd.2017.244
6. Häse F, Roch LM, Aspuru-Guzik A. Next-generation experimentation with self-driving laboratories. *Trends Chem.* 2019;1(3):282-91. doi:10.1016/j.trechm.2019.02.007
7. Granda JM, Donina L, Dragone V, Long DL, Cronin L. Controlling an organic synthesis robot with machine learning to search for new reactivity. *Nature.* 2018;559(7714):377-81. doi:10.1038/s41586-018-0307-8
8. Coley CW, Thomas DA, Lummiss JAM, Jaworski JN, Breen CP, Schultz V, et al. A robotic platform for flow synthesis of organic compounds informed by AI planning. *Science.* 2019;365(6453). doi:10.1126/science.aax1566
9. Shields BJ, Stevens J, Li J, Parasram M, Damani F, Martinez Alvarado JI, et al. Bayesian reaction optimization as a tool for chemical synthesis. *Nature.* 2021;590(7844):89-96. doi:10.1038/s41586-021-03213-y
10. Wang L, Zhou Z, Yang X, Shi S, Zeng X, Cao D. The present state and challenges of active learning in drug discovery. *Drug Discov Today.* 2024;29(6):103985. doi:10.1016/j.drudis.2024.103985
11. Scalia G, Grambow CA, Pernici B, Li YP, Green WH. Evaluating scalable uncertainty estimation methods for deep learning-based molecular property prediction. *J Chem Inf Model.* 2020;60(6):2697-717. doi:10.1021/acs.jcim.9b00975
12. Hirschfeld L, Swanson K, Yang K, Barzilay R, Coley CW. Uncertainty quantification using neural networks for molecular property prediction. *J Chem Inf Model.* 2020;60(8):3770-80. doi:10.1021/acs.jcim.0c00502
13. Janet JP, Duan C, Yang T, Nandy A, Kulik HJ. A quantitative uncertainty metric controls error in neural network-driven chemical discovery. *Chem Sci.* 2019;10(34):7913-22. doi:10.1039/C9SC02298H
14. Wallach I, Heifets A. Most ligand-based classification benchmarks reward memorization rather than generalization. *J Chem Inf Model.* 2018;58(5):916-32. doi:10.1021/acs.jcim.7b00403
15. Wouters OJ, McKee M, Luyten J. Estimated research and development investment needed to bring a new medicine to market, 2009–2018. *JAMA.* 2020;323(9):844-53. doi:10.1001/jama.2020.1166
16. Gorantla R, Kubincová A, Suutari B, Cossins BP, Mey ASJS. Benchmarking active learning protocols for ligand-binding affinity prediction. *J Chem Inf Model.* 2024;64(6):1955-65. doi:10.1021/acs.jcim.4c00220
17. Graff DE, Shakhnovich EI, Coley CW. Accelerating high-throughput virtual screening through molecular pool-based active learning. *Chem Sci.* 2021;12(22):7866-81. doi:10.1039/D0SC06805E
18. Gusev F, Gutkin E, Kurnikova MG, Isayev O. Active learning guided drug design lead optimization based on relative binding free energy modeling. *J Chem Inf Model.* 2023;63(2):583-94. doi:10.1021/acs.jcim.2c01052
19. Ding X, Cui R, Yu J, Liu T, Zhu T, Wang D, et al. Active learning for drug design: A case study on the plasma exposure of orally administered drugs. *J Med Chem.* 2021;64(22):16838-53. doi:10.1021/acs.jmedchem.1c01683
20. Gentile F, Agrawal V, Hsing M, Ton AT, Ban F, Norinder U, et al. Deep Docking: A deep learning platform for augmentation of structure-based drug discovery. *ACS Cent Sci.* 2020;6(6):939-49. doi:10.1021/acscentsci.0c00229
21. Sánchez-Lengeling B, Aspuru-Guzik A. Inverse molecular design using machine learning: Generative models for matter engineering. *Science.* 2018;361(6400):360-5. doi:10.1126/science.aat2663
22. Brown N, Fiscato M, Segler MHS, Vaucher AC. GuacaMol: Benchmarking models for de novo molecular design. *J Chem Inf Model.* 2019;59(3):1096-108. doi:10.1021/acs.jcim.8b00839
23. Gao W, Coley CW. The synthesizability of molecules proposed by generative models. *J Chem Inf Model.* 2020;60(12):5714-23. doi:10.1021/acs.jcim.0c00174
24. Fromer JC, Coley CW. Computer-aided multi-objective optimization in small molecule discovery. *Patterns (N Y).* 2023;4(2):100678. doi:10.1016/j.patter.2023.100678
25. Häse F, Roch LM, Kreisbeck C, Aspuru-Guzik A. Phoenix: A Bayesian optimizer for chemistry. *ACS Cent Sci.* 2018;4(9):1134-45. doi:10.1021/acscentsci.8b00307
26. Torres JAG, Lau SH, Anchuri P, Stevens JM, Tabora JE, Li J, et al. A multi-objective active learning platform and web app for reaction optimization. *J Am Chem Soc.* 2022;144(43):19999-20007. doi:10.1021/jacs.2c08592
27. Jia X, Lynch A, Huang Y, Danielson M, Lang'at I, Milder A, et al. Anthropogenic biases in chemical reaction data hinder exploratory inorganic synthesis. *Nature.* 2019;573(7773):251-5. doi:10.1038/s41586-019-1540-5
28. Zhang Y, Lee AA. Bayesian semi-supervised learning for uncertainty-calibrated prediction of molecular properties and active learning. *Chem Sci.* 2019;10(35):8154-63. doi:10.1039/C9SC00616H

29. Steiner S, Wolf J, Glatzel S, Andreou A, Granda JM, Keenan G, et al. Organic synthesis in a modular robotic system driven by a chemical programming language. *Science*. 2019;363(6423). doi:10.1126/science.aav2211
30. Burger B, Maffettone PM, Gusev VV, Aitchison CM, Bai Y, Wang X, et al. A mobile robotic chemist. *Nature*. 2020;583(7815):237-41. doi:10.1038/s41586-020-2442-2
31. Mehr SHM, Craven M, Leonov AI, Keenan G, Cronin L. A universal system for digitization and automatic execution of the chemical synthesis literature. *Science*. 2020;370(6512):101-8. doi:10.1126/science.abc2986
32. Jiménez-Luna J, Grisoni F, Schneider G. Drug discovery with explainable artificial intelligence. *Nat Mach Intell*. 2020;2(10):573-84. doi:10.1038/s42256-020-00236-4