

AI in Drug Discovery Pipelines: Generative Chemistry and Closed-Loop Experimental Validation

DEEPAK SAINI¹, JATIN², ISHITA³, AMAN KUMAR⁴, DEVENDRA PRASAD⁵

^{1, 2, 3, 4, 5}Department of Computer Science and Engineering (AI & ML), Panipat Institute of Engineering and Technology (PIET), Samalkha, Panipat, Haryana, India

Abstract- Bringing a single new drug to market costs on the order of two to three billion dollars and typically takes ten to fifteen years, with roughly nine of every ten candidates that enter human trials ultimately failing on efficacy or safety grounds. Artificial intelligence has been proposed as a remedy at nearly every stage of this pipeline, from target identification through generative molecular design to preclinical validation. This paper examines the technical architecture of modern AI-driven drug discovery, with particular focus on two complementary components: generative chemistry models that propose novel candidate molecules optimized against multiple pharmacological objectives, and closed-loop validation systems that iteratively test, score, and refine those proposals against wet-laboratory or computational feedback. We survey representative generative approaches, including SMILES-based recurrent and transformer models, graph-based generative networks, reinforcement-learning-driven optimization, and diffusion and structure-based design methods exemplified by AlphaFold-derived engines. We then examine the design-make-test-analyze cycle and its AI-augmented variants, including Bayesian active learning, multi-objective optimization, and robotic self-driving laboratories, and we propose a unified closed-loop reference architecture linking generative proposal, in silico triage, and experimental feedback into a single iterative system. Using case studies of Insilico Medicine's Chemistry42 and rentosertib program, Isomorphic Labs' IsoDDE engine, Exscientia's Centaur Chemist, and Recursion's LOWE platform, we assess what has been achieved to date — including hit rates markedly above traditional high-throughput screening and at least one candidate reaching Phase II trials in under three years — against what remains unresolved, including the continued absence, as of this writing, of any AI-originated drug with full regulatory approval. We conclude that generative chemistry substantially compresses early discovery timelines but that the dominant bottleneck has shifted downstream, to the same efficacy and safety uncertainties that have long limited traditional pipelines, and that closing this gap will require tighter integration between generative design and experimental validation loops rather than further scaling of generative models alone.

Keywords— Generative Chemistry, Drug Discovery, Active Learning, Design-Make-Test-Analyze Cycle, Molecular Generation, Reinforcement Learning, Structure-Based Drug Design, Self-Driving Laboratories

I. INTRODUCTION

The economics of pharmaceutical research and development have long been described as unsustainable. Bringing a single approved therapy to market is commonly estimated to cost between two and three billion dollars and to require ten to fifteen years from initial target identification to regulatory approval, and industry-wide analyses have repeatedly found that roughly ninety percent of candidates that enter human clinical trials fail before reaching approval, most often because efficacy or safety problems surface only once the compound is tested in humans rather than in a dish or an animal model. This combination of high cost, long timelines, and low success rates has made drug discovery an obvious target for artificial intelligence, and over the past decade nearly every stage of the discovery pipeline — target identification, virtual screening, molecular generation, property prediction, and preclinical triage — has attracted a dedicated machine learning approach.

Two developments in particular have reshaped the field since the early 2020s. First, generative chemistry models — systems that do not merely score or rank existing molecules but actively propose novel chemical structures optimized against specified property objectives — have matured from academic proof-of-concept into production platforms credited with designing clinical-stage compounds. Insilico Medicine's Chemistry42 platform, for example, generated tens of thousands of candidate inhibitors for a fibrosis-associated kinase target and narrowed that pool to a small synthesized set with a reported hit rate

far exceeding traditional high-throughput screening. Second, structure-based generative design, driven by the accuracy gains that followed DeepMind's AlphaFold protein-structure-prediction models, has extended computational design from ligand-centric molecule generation to explicit modeling of the three-dimensional binding geometry between a candidate molecule and its target protein, an approach Isomorphic Labs has pursued commercially through its IsoDDE engine, released in February 2026 with reported accuracy roughly double that of AlphaFold 3 on the hardest ligand-binding cases.

Generative proposal alone, however, is not discovery. A molecule proposed by a generative model is only a hypothesis until it is synthesized, tested, and its measured properties are fed back into the system that produced it — a cycle traditionally known in medicinal chemistry as the design-make-test-analyze cycle. The central technical claim of this paper is that the quality of this validation loop, not the sophistication of the generative model in isolation, increasingly determines how much a given AI drug discovery pipeline actually compresses real-world development timelines, and that the current bottleneck in the field has shifted from generating plausible candidate molecules, which is now comparatively well solved, to efficiently and reliably validating them against the same efficacy and safety criteria that have always determined clinical success.

It is also important to state plainly what has, and has not, yet been achieved. As of this writing, over two hundred AI-discovered or AI-designed drug candidates are reported to be in some stage of clinical trials, and one program, Insilico Medicine's rentosertib for idiopathic pulmonary fibrosis, is frequently cited as the most advanced fully end-to-end generative-AI-designed compound, having progressed from target identification to human dosing in approximately thirty months, roughly half the traditional industry timeline. Despite this progress, no AI-originated drug candidate has yet received full regulatory approval from a major regulatory authority, and clinical success rates for AI-nominated candidates that have reached human trials remain broadly comparable to, rather than dramatically better than, those observed for traditionally discovered compounds. This paper's analysis is framed explicitly around this tension

between accelerated early-stage timelines and unchanged downstream clinical risk.

The remainder of this paper is organized as follows. Section II surveys related work. Section III introduces the traditional drug discovery pipeline and a taxonomy of where AI methods intervene. Section IV surveys generative chemistry methods in detail. Section V examines validation loops, active learning, and automated experimentation. Section VI proposes a unified closed-loop reference architecture. Section VII presents case studies of deployed platforms. Section VIII discusses evaluation metrics and benchmarking practices. Section IX enumerates open challenges. Section X discusses regulatory considerations, and Section XI addresses ethical and reproducibility concerns, with Section XII concluding.

A further scope clarification is warranted. This paper focuses primarily on small-molecule discovery, where generative chemistry and structure-based design methods are most mature, though the same closed-loop principles increasingly extend to biologics, including antibody and peptide design, where structure-prediction advances such as AlphaFold3's joint modeling of antibody-antigen complexes have begun to enable analogous generative design workflows. We also distinguish, throughout this paper, between AI methods that merely filter or rank an existing compound library and those that generate genuinely novel chemical matter, since the former inherits the diversity limitations of whatever library it screens while the latter is bounded chiefly by the quality of its property-prediction reward signal, a distinction that matters considerably when interpreting reported success metrics across different companies' disclosed methodologies.

II. RELATED WORK

Early computational drug discovery relied primarily on virtual screening: scoring large, pre-existing compound libraries against a target using docking simulations or quantitative structure-activity relationship models, then selecting the top-ranked hits for experimental testing. This approach is fundamentally limited by the size and diversity of the library being screened, since it can only select among

molecules that already exist in a database rather than propose genuinely novel structures. Generative modeling addresses this limitation directly. Gómez-Bombarelli et al. demonstrated that a variational autoencoder could learn a continuous latent representation of molecular structure from SMILES string encodings, enabling gradient-based optimization directly in the learned latent space to discover novel molecules with desired properties [1]. Olivecrona et al. proposed REINVENT, a recurrent-neural-network-based molecular generator fine-tuned with reinforcement learning against a user-specified reward function, establishing reinforcement-learning-driven generative chemistry as a widely adopted paradigm subsequently extended by numerous follow-on systems [2].

Graph-based generative models followed, motivated by the observation that representing a molecule as an atom-bond graph rather than a linear SMILES string avoids the syntactic validity problems that string-based generators can produce, such as unbalanced ring closures or invalid valences. Jin et al. introduced a junction-tree variational autoencoder that constructs molecules hierarchically from valid chemical substructures, substantially improving the validity rate of generated molecules relative to purely string-based approaches [3]. More recently, diffusion-based generative models have been adapted from the image domain to three-dimensional molecular conformation generation, allowing candidate molecules to be generated directly with plausible atomic coordinates rather than as a topological graph alone, which is particularly relevant to structure-based design where the three-dimensional shape of a candidate ligand determines its fit within a target binding pocket.

On the structure-prediction side, Jumper et al.'s AlphaFold2 demonstrated that a deep-learning system could predict protein three-dimensional structure from amino-acid sequence at an accuracy approaching experimental methods for a large fraction of previously unsolved protein families, a result that fundamentally altered the practicality of structure-based drug design for targets lacking an experimentally resolved structure [4]. AlphaFold3 subsequently extended this capability to joint prediction of protein-ligand, protein-nucleic-acid, and antibody-antigen complexes, directly enabling in

silico estimation of how well a candidate molecule binds a target protein without requiring an experimental co-crystal structure [5]. Isomorphic Labs' IsoDDE engine, described in the company's February 2026 technical disclosure, builds on this foundation with reported accuracy improvements specifically on difficult ligand-binding cases with low sequence similarity to the training set [6].

A separate literature addresses the validation side of the pipeline directly. Exscientia's Centaur Chemist platform couples generative molecular design with automated synthesis and testing in an explicitly closed feedback loop, and the company's DSP-1181 candidate, targeting obsessive-compulsive disorder, was widely reported as among the first AI-designed molecules to enter human clinical trials [7]. Insilico Medicine's Pharma.AI platform similarly integrates target identification, generative chemistry through its Chemistry42 module, and property prediction into a single pipeline, with the company reporting that its lead idiopathic-pulmonary-fibrosis candidate progressed from target identification to Phase II dosing in roughly thirty months [8], [9]. Recursion Pharmaceuticals has pursued a complementary strategy centered on large-scale phenotypic screening and learned embeddings of cellular imaging data, using its LOWE embedding system to predict compound activity from correlations with billions of previously observed molecular interactions rather than relying primarily on generative proposal [10].

A further strand of related work addresses ADMET prediction — absorption, distribution, metabolism, excretion, and toxicity — the pharmacokinetic and safety properties that most directly determine whether a potent, selective candidate ultimately survives preclinical and clinical testing. Graph neural networks and, more recently, pretrained molecular transformer models trained on large public and proprietary ADMET datasets have been incorporated as auxiliary reward terms within generative pipelines specifically to bias candidate generation away from structures likely to fail on pharmacokinetic or toxicological grounds later in development, though the sparsity and inconsistency of publicly available ADMET data relative to potency data remains a persistent limitation noted throughout this literature, and one that we revisit as an open problem in Section IX.

III. BACKGROUND AND TAXONOMY

A. The Traditional Drug Discovery Pipeline

Conventional small-molecule drug discovery proceeds through a well-established sequence of stages. Target identification and validation establishes which biological molecule, typically a protein, should be modulated to achieve a therapeutic effect and confirms that modulating it plausibly affects disease biology. Hit identification screens a library of existing compounds, often numbering in the hundreds of thousands to low millions, against the validated target, commonly through high-throughput biochemical or cell-based assays, to identify molecules with any measurable activity. Hit-to-lead optimization iteratively modifies the chemical structure of promising hits to improve potency, selectivity against related off-target proteins, and early drug-like properties such as solubility and metabolic stability. Lead optimization continues this refinement while additionally optimizing pharmacokinetic and toxicological properties in preparation for preclinical animal studies, which in turn precede an Investigational New Drug application and the beginning of human clinical trials across Phase I, II, and III.

Each stage of this pipeline is characterized by a steep attrition rate: only a small fraction of screened hits advance to lead status, only a fraction of leads survive preclinical development, and, as noted in Section I, roughly nine in ten candidates that do reach human trials still fail, most commonly during Phase II or III when efficacy or previously undetected safety issues become apparent in a real patient population rather than a model system. This attrition pattern is central to evaluating any AI intervention in the pipeline: an approach that accelerates early discovery stages but does not address the underlying biological uncertainty that drives clinical-stage attrition will compress timelines without necessarily improving overall success rates.

B. Where Artificial Intelligence Intervenes

AI methods have been applied at essentially every stage of the pipeline described above, but with markedly different levels of maturity. Target identification increasingly uses machine learning models trained on multi-omic data — genomic,

transcriptomic, and proteomic profiles — to nominate disease-associated proteins, an approach that remains difficult to validate rigorously because ground-truth labels for successful target validation are sparse and take years to accumulate. Virtual screening and hit identification benefit from learned scoring functions and graph neural networks that predict binding affinity or biochemical activity substantially faster than physics-based docking simulations, allowing much larger virtual libraries to be triaged computationally before any physical synthesis occurs. Generative chemistry, the primary focus of Section IV, intervenes at the hit-to-lead and lead-optimization stages by proposing novel molecular structures directly rather than merely filtering an existing library. Finally, closed-loop validation systems, the focus of Section V, intervene across the hit-to-lead through preclinical stages by coupling generative proposal to automated or semi-automated experimental testing so that measured results continuously retrain and refine the generative model's proposals.

This staged taxonomy is a simplification in one important respect: in a well-designed AI-augmented pipeline, the boundaries between stages become considerably more porous than in the traditional sequential process, since a generative model operating at the hit-to-lead stage may simultaneously incorporate preclinical toxicological signal that would traditionally only be assessed much later, and a structure-based design engine operating on a still-hypothetical target structure may begin proposing candidate scaffolds before a target has been fully experimentally validated. This blurring of traditionally sequential stages is itself one of the principal claimed advantages of AI-augmented discovery, since it allows information that would traditionally arrive too late to influence early design decisions to instead shape candidate generation from the outset, provided the underlying property-prediction models are sufficiently reliable — a caveat that returns throughout this paper, particularly in Sections VI and IX, as the central unresolved condition on which most of the field's remaining promise depends.

IV. GENERATIVE CHEMISTRY METHODS

A. Sequence-Based Molecular Generation

The simplest and historically earliest class of generative chemistry models represents molecules as SMILES strings, a linear text encoding of molecular graphs, and applies sequence models — recurrent neural networks, and more recently transformer architectures — trained to generate syntactically valid and chemically plausible strings token by token. Gómez-Bombarelli et al.'s variational-autoencoder approach maps SMILES strings into a continuous latent space in which nearby points correspond to structurally similar molecules, allowing new candidates to be generated by decoding points sampled near known active compounds, or by performing gradient-based optimization within the latent space toward a target property value predicted by an auxiliary model trained jointly with the autoencoder [1]. A persistent weakness of purely sequence-based generation is that a meaningful fraction of sampled strings correspond to chemically invalid structures — for example, an unbalanced ring-closure digit or an atom with an impossible valence — requiring a validity-filtering step that discards or repairs invalid outputs before any downstream property evaluation.

Reinforcement-learning-based fine-tuning substantially improves the practical utility of sequence-based generators by directly optimizing for a downstream reward rather than purely for the likelihood of the training distribution. REINVENT, introduced by Olivecrona et al., initializes a recurrent generator by standard likelihood training on a large corpus of known drug-like molecules, then fine-tunes the same network using policy-gradient reinforcement learning against a scalar reward function that can combine predicted binding affinity, synthetic accessibility, and similarity constraints relative to a known active scaffold [2]. This reward-driven fine-tuning stage is precisely the mechanism through which sequence-based generative chemistry connects to the validation loops discussed in Section V, since the reward function itself is frequently derived from experimental or computational feedback collected during earlier iterations of a design-make-test-analyze cycle.

B. Graph-Based and Fragment-Based Generation

Representing a molecule directly as a graph of atoms and bonds, rather than as a linearized string, allows a generative model to enforce chemical validity constraints structurally rather than relying on the model to learn valid syntax implicitly from data. Jin et al.'s junction-tree variational autoencoder constructs candidate molecules hierarchically, first generating a tree of valid chemical substructures such as rings and functional groups, then assembling those substructures into a complete molecular graph, which substantially improves the fraction of generated molecules that are chemically valid relative to purely sequential string generation [3]. Fragment-based approaches extend this idea further by constraining generation to recombinations of a curated library of known drug-like fragments, trading some generative diversity for a higher guarantee of synthetic accessibility, since fragments drawn from a curated library are more likely to be composable using established synthetic routes than an arbitrary graph structure a purely unconstrained generator might produce.

C. Structure-Based and Diffusion Generative Design

Structure-based generative design conditions molecular generation explicitly on the three-dimensional geometry of a target protein's binding pocket, rather than relying solely on a learned mapping from molecular structure to a scalar activity prediction. This approach became substantially more practical following the accuracy gains of DeepMind's AlphaFold2 and AlphaFold3 systems, which predict protein structure, and in the case of AlphaFold3, joint protein-ligand complex structure, directly from sequence at accuracy levels that in many cases approach experimentally resolved crystal structures [4], [5]. Diffusion-based generative models have been adapted to this setting by learning to generate three-dimensional atomic coordinates for a candidate ligand conditioned on the geometry of a fixed or co-generated protein pocket, iteratively denoising an initial random atomic arrangement toward a chemically valid, well-fitting ligand conformation in a manner directly analogous to image-diffusion denoising.

Isomorphic Labs' IsoDDE engine, disclosed in a February 2026 technical whitepaper, integrates structure prediction, ligand-binding affinity estimation, and antibody-interaction modeling into a

single unified pipeline built on AlphaFold foundations, and the company reports that IsoDDE roughly doubles AlphaFold3's accuracy on the hardest ligand-binding cases, defined as targets with less than twenty percent sequence similarity to the training data [6]. This class of method is particularly valuable for target classes historically considered difficult for purely ligand-centric generative approaches, including protein-protein interaction interfaces and antibody-antigen binding, where the geometry of the interaction surface is central to activity in a way that is difficult to capture from ligand structure alone.

D. Multi-Objective Property Optimization

A candidate molecule intended for clinical development must simultaneously satisfy many independent objectives — potency against the intended target, selectivity against related off-target proteins, aqueous solubility, metabolic stability, synthetic accessibility, and the absence of structural alerts associated with toxicity — and optimizing for any single objective in isolation typically produces

molecules that fail on one or more of the others. Practical generative chemistry systems therefore combine multiple predicted properties into a single multi-objective reward, either through a weighted scalarization or through Pareto-based selection that maintains a diverse frontier of candidates trading off different objectives rather than collapsing them into a single score prematurely. Insilico Medicine's Chemistry42 platform, for example, is reported to have generated approximately seventy-eight thousand virtual candidate inhibitors against a fibrosis-associated kinase target, filtered this pool through multi-objective optimization across potency, selectivity, and drug-like property criteria, and ultimately synthesized sixty top-ranked compounds, of which a reported 16.7 percent showed meaningful activity — a hit rate roughly two orders of magnitude higher than the approximately 0.1 percent typically reported for traditional high-throughput screening campaigns [8], [9].

TABLE I. Comparison of representative generative chemistry method families used across the discovery pipeline.

Method Family	Molecular Representation	Optimization Mechanism	Key Strength	Representative System
Sequence-based VAE	SMILES string	Latent-space gradient search	Continuous, differentiable latent space	Gómez-Bombarelli et al.
RL-fine-tuned generator	SMILES string	Policy-gradient reward optimization	Directly optimizes arbitrary reward	REINVENT
Graph / junction-tree VAE	Molecular graph	Hierarchical substructure assembly	High chemical validity rate	Jin et al. JT-VAE
3D diffusion design	Atomic coordinates	Iterative denoising	Explicit 3D pocket fit	Pocket-conditioned diffusion models
Structure-based unified engine	Protein-ligand complex	Joint structure + affinity prediction	Handles difficult low-similarity targets	Isomorphic IsoDDE / AlphaFold3

E. Retrosynthesis Planning and Synthesizability

A generated molecule is only useful to a discovery program if it can actually be synthesized within a reasonable number of steps using accessible starting materials and established reaction chemistry, and retrosynthesis-planning models address this by learning to decompose a target molecule backward

into a sequence of simpler precursor compounds and known reaction templates, mirroring the manual retrosynthetic analysis long practiced by synthetic chemists. Modern retrosynthesis planners combine a learned single-step reaction predictor, which proposes plausible precursor sets for a given target molecule, with a tree-search algorithm that explores multi-step

synthetic routes toward commercially available starting materials, and increasingly these planners are integrated directly into the generative loop as a filtering or reward-shaping component, penalizing or discarding candidate molecules for which no reasonably short synthetic route can be found, rather than treating synthesizability as a post-hoc concern addressed only after a candidate has already been selected on potency and selectivity grounds alone.

V. VALIDATION LOOPS AND CLOSED-LOOP EXPERIMENTATION

A. The Design-Make-Test-Analyze Cycle

The design-make-test-analyze cycle, long central to medicinal chemistry practice, structures iterative compound optimization as a repeated loop: chemists design a set of candidate structures, synthesize the most promising subset, test the synthesized compounds against relevant biochemical and cellular assays, and analyze the resulting data to inform the next round of design. AI-augmented discovery pipelines automate or accelerate each stage of this cycle without eliminating its fundamentally iterative structure: the design stage is handled by the generative chemistry methods described in Section IV, the make stage increasingly involves automated or robotic synthesis platforms, the test stage may combine physical assays with fast computational property predictors as a pre-filter, and the analyze stage feeds measured results back as additional training data or reward signal for the next generative round. The central insight motivating this section is that the cycle time and information efficiency of this loop, not the raw generative capacity of the design stage alone, ultimately governs how quickly a pipeline converges on a viable clinical candidate.

B. Active Learning and Bayesian Optimization

Because physical synthesis and experimental testing remain orders of magnitude more expensive and time-consuming than computational property prediction, an effective validation loop must select which of many generated candidates to actually synthesize and test using an information-efficient strategy rather than testing candidates in an arbitrary or purely rank-based order. Active learning frameworks address this by explicitly modeling predictive uncertainty alongside the point estimate of a property prediction, and

preferentially selecting candidates for experimental testing that are expected to be maximally informative — either because the model is highly uncertain about their properties, or because they lie near a decision boundary relevant to the optimization objective, such as the threshold separating active from inactive compounds. Bayesian optimization formalizes this selection using an acquisition function, commonly expected improvement or upper confidence bound, that balances exploration of uncertain regions of chemical space against exploitation of regions already known to contain promising candidates, and has been widely adopted in both academic and industrial discovery pipelines specifically because it provably minimizes, under standard modeling assumptions, the number of expensive experimental evaluations required to reach a target level of optimization performance.

C. Automated Synthesis and Self-Driving Laboratories

A parallel line of engineering effort has focused on automating the make and test stages of the cycle directly, using robotic liquid-handling systems, automated flow-chemistry synthesis platforms, and high-throughput assay robotics to reduce the wall-clock time between a generative model proposing a candidate and experimental data on that candidate becoming available to retrain the model. Exscientia's Centaur Chemist platform is explicitly described by the company as combining AI molecular design with automated synthesis and testing in a tight feedback loop, an architecture credited with contributing to DSP-1181 becoming among the first AI-designed molecules to enter Phase I clinical trials [7]. Fully autonomous self-driving laboratories extend this concept further, closing the loop with minimal human intervention between rounds: a control policy, sometimes itself learned, selects the next batch of experiments, robotic hardware executes the synthesis and assay, and the resulting data is automatically incorporated into the next round of model retraining and candidate selection, in principle allowing many more design-make-test-analyze cycles to be executed per unit time than a conventional, human-paced laboratory workflow permits.

D. Multi-Fidelity and Simulation-in-the-Loop Validation

Because physical experimentation remains the ultimate bottleneck even in highly automated pipelines, many practical systems introduce one or more intermediate, lower-cost validation stages between generative proposal and full experimental testing, an approach often termed multi-fidelity optimization. A typical cascade might first filter tens of thousands of generated candidates using fast, approximate computational scoring, such as a learned binding-affinity predictor or a simplified docking simulation; then apply a smaller number of more expensive but more accurate simulations, such as free-energy perturbation calculations, to a shortlist of several hundred candidates; and finally synthesize and experimentally test only the highest-scoring tens of compounds that survive both filtering stages. This cascaded structure is precisely what allows platforms such as Chemistry42 to generate tens of thousands of virtual candidates while synthesizing only a small fraction of that number, concentrating expensive experimental resources on the candidates most likely to succeed based on progressively more accurate, and progressively more expensive, computational triage [8].

VI. PROPOSED UNIFIED CLOSED-LOOP FRAMEWORK

A. Architecture

We propose a reference architecture, termed the Generative-Validation Loop, that formalizes the interaction between the generative chemistry and experimental validation components described in Sections IV and V as three coupled stages operating on a shared candidate pool. The proposal stage uses a generative model, of any of the families surveyed in Section IV, to produce a batch of candidate molecules conditioned on the current state of a multi-objective reward model and, where applicable, a target protein structure. The triage stage applies the multi-fidelity cascade described in Section V.D to rank the proposed batch and select a small, information-maximizing subset for physical testing, using an acquisition function that explicitly accounts for the current model's predictive uncertainty rather than its point estimate alone. The feedback stage incorporates the resulting experimental measurements — potency,

selectivity, solubility, and any observed toxicity signal — as additional training data for both the property-prediction models used in triage and, where the generative model itself is reward-driven, as an updated reward signal for the next generation round.

Critically, this architecture treats the generative model and the property-prediction models used for triage as two distinct components trained on potentially different, though overlapping, data distributions, rather than as a single monolithic system. This separation matters because the generative model's training objective is to produce chemically diverse, synthetically accessible candidates, while the triage model's objective is calibrated property prediction; conflating the two risks a generative model that is confidently wrong about its own proposed candidates' properties, since it would be evaluating its own outputs using the same biases that produced them in the first place.

B. Convergence Behavior and Cycle Efficiency

The efficiency of the Generative-Validation Loop is governed primarily by two quantities: the information gain per experimental cycle, determined by how well the triage stage's acquisition function targets genuinely uncertain or promising candidates, and the wall-clock latency per cycle, determined by how quickly physical synthesis and assay results become available to retrain the system. Improving the generative model's raw sample quality, holding these two quantities fixed, yields diminishing returns beyond a certain point, because a pipeline limited primarily by experimental cycle latency cannot exploit an improved generative model any faster than its physical validation infrastructure can process the resulting candidates. This observation motivates our Section I claim that recent gains in generative chemistry model quality, while real, have shifted the binding constraint on overall discovery timelines toward the validation loop's throughput and information efficiency rather than the generative model's sample quality per se.

C. Threat to Validity: Distributional Shift Between Design and Test

A persistent failure mode in closed-loop generative discovery is distributional shift between the chemical space the generative model is trained to propose and the chemical space in which the triage models'

property predictions remain reliable. Property-prediction models are typically trained on historical assay data concentrated around previously studied chemical scaffolds, and a generative model optimizing aggressively against such a model's predictions can drift toward regions of chemical space where the underlying property predictor is poorly calibrated, producing candidates that score well on the (miscalibrated) predicted objective but perform poorly when actually synthesized and tested — a phenomenon closely related to reward hacking in the broader reinforcement learning literature. Mitigating this failure mode requires the triage stage to explicitly penalize candidates far from the property predictor's training distribution, or to route such candidates preferentially toward experimental testing specifically because their predicted properties are least trustworthy, reinforcing the importance of uncertainty-aware acquisition functions described in Section V.B.

VII. CASE STUDIES

A. Insilico Medicine: Chemistry42 and Rentosertib

Insilico Medicine's Pharma.AI platform integrates target identification, generative chemistry through its Chemistry42 module, and property prediction into a single pipeline, and the company's rentosertib program for idiopathic pulmonary fibrosis is widely cited as the most advanced fully end-to-end example of the approach to date. The platform identified TNIK, a protein kinase implicated in fibrosis pathways, as a novel target using omics-data-driven target discovery, then used Chemistry42 to generate approximately seventy-eight thousand virtual candidate TNIK inhibitors, filtered this pool through multi-objective optimization, and synthesized sixty top-ranked compounds, achieving a reported 16.7 percent hit rate compared with roughly 0.1 percent for traditional high-throughput screening [8], [9]. The company reports that the resulting lead compound progressed from target identification to human Phase I dosing in approximately thirty months, roughly half the traditional industry average timeline, and that the compound subsequently advanced into Phase II trials.

B. Isomorphic Labs: IsoDDE

Isomorphic Labs, a DeepMind-originated, Alphabet-backed company, released its IsoDDE engine on

February 10, 2026, describing it as a unified drug-design system combining protein structure prediction, ligand-binding-affinity estimation, and antibody-interaction modeling within a single pipeline built on AlphaFold foundations [6]. The company reports that IsoDDE roughly doubles AlphaFold3's accuracy on the hardest ligand-binding cases, defined by low sequence similarity to the training data, and has stated an intention to advance internally designed candidates toward first-in-human trials by the end of 2026. Unlike Insilico's approach, which centers on de novo generative chemistry guided primarily by predicted activity and drug-like property models, Isomorphic's approach centers explicitly on structural prediction of the three-dimensional binding interaction, reflecting the company's origin in structure-prediction research rather than generative chemistry research; industry commentary has accordingly described the company's approach as distinguished less by generative novelty and more by binding-geometry prediction accuracy [6].

C. Exscientia's Centaur Chemist and Recursion's LOWE

Exscientia's Centaur Chemist platform, which explicitly combines reinforcement-learning-driven generative molecular design with automated synthesis and testing in a tight closed feedback loop, was associated with DSP-1181, widely reported as among the earliest AI-designed molecules to enter human clinical trials, targeting obsessive-compulsive disorder [7]. Recursion Pharmaceuticals has pursued a related but architecturally distinct strategy centered on its LOWE embedding system, which learns representations from large-scale cellular imaging and molecular interaction data spanning billions of observations, using these learned embeddings to predict compound activity computationally without necessarily relying on de novo generative proposal as the primary discovery mechanism; the company has advanced multiple AI-curated oncology candidates into clinical trials using this phenotypic-screening-centered approach.

D. Synthesis: What the Case Studies Show

Taken together, these case studies illustrate both the genuine progress and the persistent limitations of AI-driven drug discovery as deployed in practice by mid-2026. Reported hit rates from generative triage

substantially exceed traditional high-throughput screening, and at least one fully end-to-end AI-designed candidate has progressed from target identification through Phase II trials in a compressed timeline. At the same time, over two hundred AI-discovered or AI-designed candidates are reported across the industry to be in some stage of clinical trials, yet no AI-originated therapy had received full regulatory approval from a major authority as of this writing, and industry analyses estimate roughly a coin-flip probability that a first such approval occurs within the current or following year. This pattern is consistent with our central claim in Section I: generative chemistry and structure-based design have measurably compressed the early stages of discovery, but the clinical-stage attrition driven by efficacy and safety uncertainty — the same forces that have limited traditional drug discovery for decades — has not yet been comparably reduced.

It is also worth noting the methodological diversity across these four programs, since it illustrates that

there is not yet a single dominant architecture for AI-driven drug discovery in the way that, for instance, transformer architectures have become dominant in natural language processing. Insilico's approach centers on de novo generative chemistry guided by omics-derived target discovery; Isomorphic's centers on structural prediction of binding geometry; Exscientia's centers on tight integration between generative design and automated wet-laboratory execution; and Recursion's centers on large-scale phenotypic screening with learned embeddings rather than explicit molecular generation. This diversity suggests the field has not yet converged on which combination of methods most reliably translates into clinical success, and that the coming several years of clinical readouts across these differently architected programs may provide unusually informative comparative evidence about which design choices matter most for eventual regulatory approval.

TABLE II. Comparison of representative AI drug discovery platforms, primary methodology, and reported clinical milestones as of mid-2026.

Platform Company	Primary Method	Flagship Candidate	Reported Milestone
Insilico Medicine	Generative chemistry (Chemistry42) + omics target ID	Rentosertib (TNIK inhibitor, IPF)	Target ID to Phase I dosing in ~30 months; Phase II reached
Isomorphic Labs	Structure-based design (IsoDDE, AlphaFold-derived)	Undisclosed oncology/immunology leads	IsoDDE released Feb. 2026; first-in-human targeted end of 2026
Exscientia	RL-driven generative design + automated synthesis (Centaur Chemist)	DSP-1181 (OCD)	Among first AI-designed molecules to enter Phase I
Recursion Pharmaceuticals	Phenotypic screening + learned embeddings (LOWE)	Multiple oncology candidates	Several AI-curated candidates advanced to clinical trials

VIII. EVALUATION METRICS AND BENCHMARKING

Evaluating a generative chemistry system in isolation typically relies on a standard set of distributional metrics: validity, the fraction of generated molecules that correspond to chemically well-formed structures; uniqueness, the fraction of valid generated molecules that are distinct from one another; novelty, the fraction of valid, unique molecules absent from the training data; and synthetic accessibility, an estimate of how difficult a proposed molecule would be to synthesize using established chemical reactions. These metrics are useful for comparing generative architectures but are, on their own, insufficient to evaluate a full discovery pipeline, since a generative model can score well on all four while producing candidates that fail on the pharmacological properties that actually determine clinical success.

A more meaningful evaluation for the closed-loop systems described in Section VI reports pipeline-level outcomes directly: the experimental hit rate achieved among synthesized candidates relative to a traditional high-throughput screening baseline, the number of design-make-test-analyze cycles required to reach a target potency or selectivity threshold, and, where available, the wall-clock time and total cost from initial target identification to a defined clinical milestone such as Investigational New Drug filing or first human dosing. Reported figures along these dimensions — such as Insilico's approximately thirty-month target-to-Phase-I timeline and its 16.7 percent synthesized-candidate hit rate — provide considerably stronger evidence of practical utility than generative-distribution metrics alone, though cross-company comparison remains difficult because companies do not uniformly disclose comparable baselines, and because a single favorable case study does not establish a generalizable success rate across a company's full portfolio of programs [8], [9]. We accordingly recommend that future benchmarking efforts prioritize portfolio-level statistics, such as the fraction of a company's AI-nominated programs that survive each clinical phase relative to industry-wide historical baselines, over single flagship case studies, which are subject to obvious selection and reporting bias.

IX. CHALLENGES AND OPEN RESEARCH PROBLEMS

The most consequential open problem is the persistence of clinical-stage attrition despite accelerated early discovery. Because efficacy and safety failures typically emerge only during human testing, and because the biological complexity underlying these failures — off-target effects, unanticipated toxicity pathways, and disease heterogeneity across patient populations — is not fully captured by any current preclinical model, computational or otherwise, it remains unclear how much further generative chemistry and structure-based design alone can reduce this failure rate without corresponding advances in preclinical predictive biology, an area only partially addressed by the methods surveyed in this paper.

A second open problem concerns synthetic feasibility and retrosynthesis. A generative model can readily propose a molecule with excellent predicted potency and selectivity that is, in practice, extremely difficult or prohibitively expensive to synthesize using known chemical reactions, and while synthetic-accessibility scoring functions and retrosynthesis-planning models have improved substantially, they remain imperfect predictors of true synthetic difficulty, occasionally passing candidates that later prove impractical to make at the make stage of the design-make-test-analyze cycle. A third problem concerns data scarcity for rare or novel target classes: generative and property-prediction models trained predominantly on historical data for well-studied target families transfer poorly to genuinely novel biology, precisely the setting — such as the low-sequence-similarity ligand-binding cases that motivated Isomorphic's IsoDDE development — where AI-driven discovery would be most valuable if it generalized reliably [6].

A fourth problem concerns reproducibility and independent validation. Much of the most widely cited evidence for AI drug discovery's effectiveness, including the case studies in Section VII, originates from company press releases and investor disclosures rather than from independent, peer-reviewed replication, and the field currently lacks standardized, company-agnostic benchmarks analogous to those that have driven progress in other machine learning

domains, making it difficult for outside researchers to assess whether reported hit rates and timeline compressions generalize across programs, targets, and companies or reflect favorable, selectively reported outcomes.

A fifth, cross-cutting problem concerns the interpretability gap between what a generative or predictive model has learned and what a medicinal chemist needs to trust that learning enough to act on it. Even when a property-prediction model achieves strong held-out accuracy, medicinal chemists and regulatory reviewers alike frequently want a mechanistic rationale for why a given candidate is predicted to be potent, selective, or safe, and current deep generative and predictive models generally do not produce such rationales natively, requiring post-hoc interpretability techniques of uneven reliability. Bridging this gap without sacrificing the predictive performance gains that motivated adopting deep learning in the first place remains an active and largely unresolved area of methodological research, and one with direct consequences for the regulatory acceptance discussed in Section X.

X. REGULATORY CONSIDERATIONS

Regulatory authorities have begun to issue specific guidance addressing the use of artificial intelligence and machine learning within drug development, recognizing that models used to nominate candidates, predict properties, or support regulatory submissions raise validation questions distinct from those posed by conventional computational chemistry tools. Regulatory guidance in this area generally emphasizes the need for a documented, auditable rationale connecting a model's predictions to an eventual regulatory decision, rather than treating a model's output as self-validating simply because it was produced by a sophisticated system, echoing broader regulatory principles applied to software-as-a-medical-device and computational modeling in other therapeutic contexts. Because no AI-originated candidate has yet completed the regulatory approval process, as discussed in Section VII.D, much of this guidance remains prospective rather than tested against an actual completed approval, and how regulators will evaluate an efficacy or safety dossier

substantially shaped by generative chemistry and closed-loop validation methods remains, to a meaningful degree, an open institutional question that the industry expects to be substantially clarified by the outcome of the first such approval decision.

A related regulatory consideration concerns model transparency and explainability. Because generative chemistry models and the property predictors that guide triage are frequently large, opaque deep-learning systems, regulators and internal quality-assurance functions face a genuine tension between the practical performance benefits of these models and the interpretability that has traditionally been expected of methods supporting a regulatory submission; industry practice has increasingly favored supplementing model predictions with orthogonal, more interpretable validation — for instance, confirming a predicted binding mode with an independent physics-based simulation or, where feasible, an experimental structure — rather than relying on a single model's output as the sole basis for an advancement decision.

XI. ETHICAL AND REPRODUCIBILITY CONSIDERATIONS

The concentration of the most capable generative chemistry and structure-based design systems within a small number of well-resourced commercial organizations — chiefly large, well-funded biotechnology companies and Alphabet-backed Isomorphic Labs — raises equity concerns about which diseases receive AI-accelerated research attention. Commercial incentives naturally favor conditions with large addressable markets, and there is a meaningful risk that the throughput gains described in this paper concentrate disproportionately on well-funded therapeutic areas rather than on neglected diseases predominantly affecting lower-income populations, a pattern already observed in traditional pharmaceutical research and development and not obviously corrected by AI-driven acceleration absent deliberate policy intervention.

A second consideration concerns the reproducibility gap noted in Section IX: because leading case studies are predominantly disclosed through company communications rather than independent peer review,

the scientific community's ability to verify reported hit rates, timeline compressions, and success probabilities independently remains limited, and we join other observers in this field in calling for greater standardized, third-party benchmarking as the technology matures and as increasing numbers of AI-nominated candidates progress through clinical trials, providing a growing body of outcome data against which claims of AI-driven acceleration can eventually be tested at a portfolio level rather than through individual flagship examples.

A third consideration concerns intellectual property and data-sharing incentives within the field. Because proprietary training data, particularly internally generated ADMET and binding-affinity measurements, provides a substantial competitive advantage to companies with large historical assay libraries, there is limited commercial incentive for firms to share the negative or failed results that would be most valuable for training property-prediction models capable of avoiding previously discovered dead ends, a dynamic that likely slows field-wide progress relative to a counterfactual in which such data were pooled, even though it is individually rational for any single company. Public-private data-sharing consortia and pre-competitive collaborations have been proposed as partial remedies, though participation remains voluntary and uneven across the industry as of this writing.

XII. CONCLUSION AND FUTURE WORK

This paper examined the technical architecture of modern AI-driven drug discovery, distinguishing generative chemistry methods that propose novel candidate molecules from the closed-loop validation systems that test and refine those proposals against experimental feedback. We surveyed sequence-based, graph-based, and structure-based generative approaches, examined the design-make-test-analyze cycle and its AI-augmented variants including active learning and self-driving laboratories, and proposed a unified Generative-Validation Loop architecture that makes explicit the coupling between generative proposal and experimental feedback. Case studies of Insilico Medicine, Isomorphic Labs, Exscientia, and Recursion Pharmaceuticals illustrate genuine,

measurable progress — hit rates far exceeding traditional screening, and at least one candidate compressing target-to-Phase-I timelines by roughly half — set against the continued absence, as of this writing, of any fully approved AI-originated therapy.

We identify closing the gap between accelerated early discovery and unchanged clinical-stage attrition as the central open problem for the field going forward, and we argue that further scaling of generative model capacity alone is unlikely to resolve this gap absent corresponding advances in the fidelity of preclinical validation and in the information efficiency of experimental feedback loops. Promising directions for future work include tighter integration of self-driving laboratory infrastructure directly into generative training loops so that experimental feedback arrives with substantially lower latency; improved uncertainty quantification in property-prediction models specifically to counter the distributional-shift failure mode described in Section VI.C; standardized, company-agnostic benchmarks reporting portfolio-level clinical success rates rather than individual flagship case studies; and continued development of structure-based methods, following the trajectory from AlphaFold2 through AlphaFold3 to Isomorphic's IsoDDE, toward reliable generalization on the low-similarity, poorly characterized target classes where computational discovery methods currently perform least reliably. The outcome of the first regulatory decision on a fully AI-originated therapy, expected within the next one to two years according to current industry estimates, is likely to substantially shape the field's subsequent direction regardless of whether that decision is favorable.

More broadly, the history surveyed in this paper suggests a useful corrective to overly simple narratives about AI's role in drug discovery. Neither the pessimistic view that AI contributes only incremental screening efficiency, nor the optimistic view that generative chemistry alone will soon resolve the field's decades-old attrition problem, is well supported by the evidence assembled in Sections VII through IX. The more accurate picture is one of genuine, substantial acceleration concentrated specifically at the stages of the pipeline — target-to-hit and hit-to-lead — where computational methods can be validated relatively quickly and cheaply against in vitro and in silico data,

coupled with essentially unchanged uncertainty at the stages — Phase II and III efficacy and safety evaluation in real patient populations — where the fundamental bottleneck has always been biological rather than computational. Recognizing this distinction clearly is, in our view, a prerequisite for allocating future research investment toward the parts of the pipeline, particularly preclinical predictive biology and closed-loop experimental validation infrastructure, where continued progress is most likely to translate compressed early timelines into correspondingly higher clinical success rates rather than simply moving failures earlier without reducing their overall frequency.

REFERENCES

- [1] R. Gómez-Bombarelli et al., "Automatic chemical design using a data-driven continuous representation of molecules," *ACS Cent. Sci.*, vol. 4, no. 2, pp. 268–276, 2018.
- [2] M. Olivecrona, T. Blaschke, O. Engkvist, and H. Chen, "Molecular de-novo design through deep reinforcement learning," *J. Cheminform.*, vol. 9, no. 48, 2017.
- [3] W. Jin, R. Barzilay, and T. Jaakkola, "Junction tree variational autoencoder for molecular graph generation," in *Proc. Int. Conf. Mach. Learn. (ICML)*, 2018.
- [4] J. Jumper et al., "Highly accurate protein structure prediction with AlphaFold," *Nature*, vol. 596, pp. 583–589, 2021.
- [5] J. Abramson et al., "Accurate structure prediction of biomolecular interactions with AlphaFold3," *Nature*, vol. 630, pp. 493–500, 2024.
- [6] Presenc AI, "AI drug discovery: IsoDDE and AlphaFold in 2026," *presenc.ai*, May 2026.
- [7] IntuitionLabs, "Isomorphic Labs and AlphaFold: AI drug discovery in trials," *intuitionlabs.ai*, Apr. 2026.
- [8] AIM Media House, "2026 is the year AI drug discovery meets clinical reality," *aimmediahouse.com*, 2026.
- [9] Humai Blog, "AI drugs reach the clinic: 2026 is the year of the first large-scale test," *humai.blog*, Mar. 2026.
- [10] AI Magicx, "How AI is compressing 10-year drug discovery timelines to 18 months: The 2026 biotech revolution," *aimagicx.com*, Mar. 2026.
- [11] BuildMVPFast, "AI drug discovery: What's working, what's failing in 2026," *buildmvpfast.com*, Mar. 2026.
- [12] IntuitionLabs, "Isomorphic Labs \$2.1B Series B: AI drug design analysis," *intuitionlabs.ai*, 2026.
- [13] A. Zhavoronkov et al., "Deep learning enables rapid identification of potent DDR1 kinase inhibitors," *Nat. Biotechnol.*, vol. 37, pp. 1038–1040, 2019.
- [14] Insilico Medicine, "Chemistry42: AI-powered generative chemistry platform," *insilico.com*, 2023–2026.
- [15] Exscientia, "Centaur Chemist: AI-driven design with automated synthesis and testing," *exscientia.ai*, 2020–2026.
- [16] Recursion Pharmaceuticals, "LOWE: Library of Optimized Weighted Embeddings," *recursion.com*, 2024–2026.
- [17] D. C. Swinney and J. Anthony, "How were new medicines discovered?," *Nat. Rev. Drug Discov.*, vol. 10, pp. 507–519, 2011.
- [18] J. Arrowsmith, "Trial watch: Phase II failures: 2008–2010," *Nat. Rev. Drug Discov.*, vol. 10, pp. 328–329, 2011.
- [19] B. Sanchez-Lengeling and A. Aspuru-Guzik, "Inverse molecular design using machine learning: Generative models for matter engineering," *Science*, vol. 361, no. 6400, pp. 360–365, 2018.
- [20] G. Schneider, "Automating drug discovery," *Nat. Rev. Drug Discov.*, vol. 17, pp. 97–113, 2018.
- [21] C. M. Dobson, "Chemical space and biology," *Nature*, vol. 432, pp. 824–828, 2004.
- [22] P. B. Jørgensen, "Fragment-based drug discovery and generative chemistry: Bridging validity and synthesizability," *J. Med. Chem.*, 2022.
- [23] B. Settles, "Active learning literature survey," *Univ. Wisconsin-Madison Comput. Sci. Tech. Rep.* 1648, 2009.
- [24] J. Snoek, H. Larochelle, and R. P. Adams, "Practical Bayesian optimization of machine learning algorithms," in *Proc. Adv. Neural Inf. Process. Syst. (NeurIPS)*, 2012.

- [25] B. Burger et al., "A mobile robotic chemist," Nature, vol. 583, pp. 237–241, 2020.
- [26] U.S. Food and Drug Administration, "Artificial intelligence and machine learning in drug development: Discussion paper," FDA.gov, 2023–2025.
- [27] Fierce Biotech, "Isomorphic Labs raises \$2.1 billion Series B to advance AI-designed drug candidates," fiercebiotech.com, May 2026.
- [28] Endpoints News, "AI-discovered drug pipeline surpasses 200 candidates in clinical trials," endpts.com, 2026.