

Integrating Artificial Intelligence Across the Drug-Discovery Pipeline: From Targets to Preclinical Proof

Dilhan NAMLI*, Nesrin GÖKHAN KELEKÇİ***

Integrating Artificial Intelligence Across the Drug-Discovery Pipeline: From Targets to Preclinical Proof

İlaç Keşfi Sürecinde Yapay Zekâ Entegrasyonu: Hedeflerden Preklinik Kanıt

SUMMARY

Artificial intelligence (AI)-based approaches have drawn significant attention for their potential to address major challenges in drug discovery processes, such as time constraints, high costs, and low success rates. Specifically, machine learning (ML) and deep learning (DL) algorithms are effectively utilized in various stages of drug development, including target identification, molecular screening, lead compound selection, optimization, and ADMET prediction. In this study, the integration of current AI models into pharmaceutical R&D processes is examined from an interdisciplinary perspective, and their application domains are evaluated through relevant case studies in the literature. It has been observed that ML-based methods can yield successful results even with limited data, while DL architectures offer advantages in modeling complex molecular relationships. Furthermore, the architectural frameworks, training strategies, and diversity of ML and DL algorithms are comprehensively discussed within the scope of the study. It is demonstrated that traditionally experience-based decision processes such as retrosynthetic planning and formulation development can be accelerated and made more sustainable through data-driven systems. Additionally, AI-assisted predictions are shown to reduce the experimental burden and enhance research efficiency in preclinical and clinical stages. These evaluations suggest that AI technologies are not merely supportive tools but also strategic components at the core of innovative drug discovery approaches.

Keywords: Artificial intelligence, drug design, predictive modeling, computational drug discovery, virtual screening.

ÖZ

Yapay zekâ tabanlı yöntemler, ilaç keşfi süreçlerinde karşılaşılan zaman, maliyet ve başarı oranı gibi temel sorunlara çözüm üretme potansiyeliyle dikkat çekmektedir. Özellikle makine öğrenmesi ve derin öğrenme algoritmaları, ilaç geliştirme sürecinin hedef belirleme, molekül tarama, öncü bileşik seçimi, optimizasyon ve ADMET tahmini gibi aşamalarında etkin biçimde kullanılmaktadır. Bu çalışmada, güncel yapay zekâ modellerinin farmasötik Ar-Ge süreçlerine entegrasyonu disiplinlerarası bir yaklaşımla ele alınmış; literatürde yer alan vaka örnekleriyle bu modellerin uygulama alanları değerlendirilmiştir. Makine öğrenmesi tabanlı yöntemlerin sınırlı veriyle dahi başarılı sonuçlar üretebildiği, derin öğrenme mimarilerinin ise karmaşık moleküler ilişkileri tanımlamada avantaj sağladığı gözlemlenmiştir. Ayrıca çalışma kapsamında, makine öğrenmesi ve derin öğrenme algoritmalarının mimari yapıları, eğitim stratejileri ve model çeşitliliği detaylı olarak ele alınmıştır. Retrosentez planlaması ve formülasyon geliştirme gibi geleneksel olarak deneyime dayanan karar süreçlerinin, veri odaklı sistemlerle desteklenerek daha hızlı ve sürdürülebilir hâle getirilebileceği ortaya konmuştur. Preklinik ve klinik aşamalarda yapay zekâ destekli tahminlerin deney yükünü azalttığı ve araştırma verimliliğini artırdığı vurgulanmıştır. Bu kapsamda yapılan değerlendirmeler, yapay zekâ teknolojilerinin yalnızca destekleyici bir araç değil, aynı zamanda yenilikçi ilaç keşif yaklaşımlarının merkezinde yer alan stratejik bir unsur olduğunu göstermektedir.

Anahtar Kelimeler: Yapay zekâ, ilaç tasarımı, öngörülse modelleme, hesaplmalı ilaç keşfi, sanal tarama.

Received: 18.09.2025

Revised: 18.10.2025

Accepted: 23.10.2025

* ORCID: 0009-0004-9682-5556, R&D Formulation Department, GEN İlaç ve Sağlık Ürünleri San. ve Tic. A.Ş., Ankara, Turkey

** ORCID: 0000-0001-7383-8443, Department of Pharmaceutical Chemistry, Faculty of Pharmacy, Hacettepe University, Ankara, Turkey

° Corresponding Author; Nesrin Gökhan Kelekçi
E-mail: onesrin@hacettepe.edu.tr

INTRODUCTION

In recent years, artificial intelligence (AI) has revolutionized many disciplines, particularly the health sciences, by simulating cognitive functions such as learning and decision-making. The integration of this technology into drug discovery processes marks the beginning of a new era by aiming to overcome the limitations of traditional methods and resolve the growing complexity of biomedical research.

Conventional drug discovery methods rely heavily on experimental procedures, which are both time-consuming and costly. Typically, only a few compounds among thousands progress to clinical evaluation, and even fewer reach the market. Developing a single drug may take an average of 10 to 15 years and cost between 1 to 2 billion USD. Furthermore, up to 90% of drug candidates fail during clinical trials. These high failure rates not only cause financial losses but also result in the inefficient use of human resources and experimental infrastructure.

As a result, the need for faster, more flexible, and more accurate systems that can transcend the limitations of traditional approaches has become increasingly urgent. At this point, AI steps in with its data-driven analytical capabilities and provides critical decision support tools from the early phases of drug discovery onward (Parvatikar et al., 2023; Wu et al., 2024b; Yang, Wang, Byrne, Schneider, & Yang, 2019a).

AI-based methods are being integrated into various stages of the drug development pipeline in order to overcome these challenges. By leveraging machine learning (ML) and its subfield, deep learning (DL), accurate predictions can be made on large-scale biomedical datasets. These technologies enable the development of decision support mechanisms across a wide range of processes—from target identification

and molecular design to synthesis planning and pre-clinical or clinical studies.

There exists a hierarchical structure between the concepts of artificial intelligence, machine learning, and deep learning: while AI is the overarching concept, ML represents a subfield that focuses on learning from data. Deep learning, in turn, is a specialized approach within ML that employs multi-layered neural network architectures for high-level pattern recognition (Janiesch, Zschech, & Heinrich, 2021). These interrelated disciplines have the potential to reshape the future of pharmaceutical innovation.

This structural relationship, along with the integration of these approaches into various stages of drug discovery and development, is schematically illustrated in Figure 1.

In this study, the historical development of artificial intelligence and its role in the healthcare sector are first examined. Subsequently, AI-driven drug discovery processes are analyzed in detail under the headings of target identification, lead compound discovery, molecular optimization, and synthesis planning. The following sections explore machine learning and deep learning models, followed by an evaluation of their effectiveness based on recent applications in the literature. Finally, the integration of AI into pharmaceutical development processes is discussed—ranging from preclinical stages to clinical research, and from formulation development to ethical and regulatory considerations. In light of all these sections, the positioning of AI technologies within pharmaceutical research is addressed through an interdisciplinary perspective.

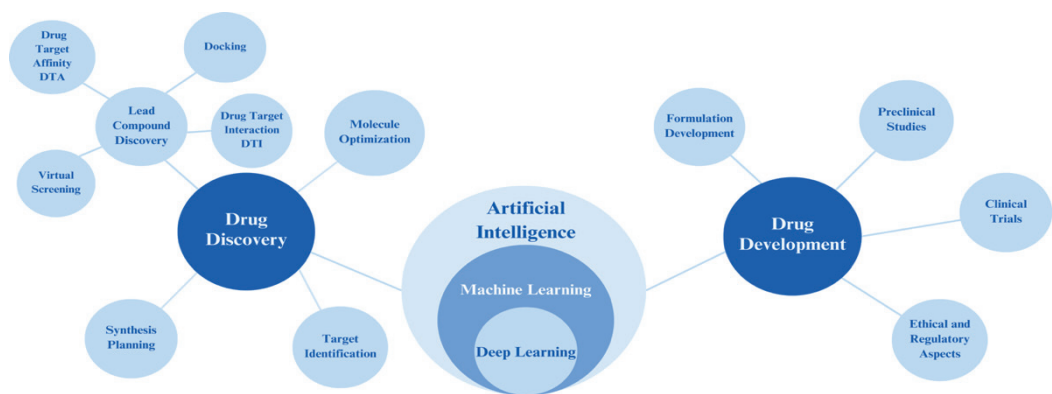


Figure 1. A schematic representation of the structural hierarchy among AI, ML, and DL, and their roles in drug discovery and development processes.

THE HISTORY OF AI & ITS PLACE IN THE HEALTHCARE SECTOR

The historical development of AI technologies provides a crucial foundation for understanding the evolution of their applications in healthcare and drug discovery (Figure 2.). This journey began in 1943 with the mathematical modeling of artificial neural networks by McCulloch and Pitts (McCulloch & Pitts, 1943) and gained conceptual depth with the introduction of the

Turing Test in 1950 (Turing, 1950). In the following decades, the advancement of approaches such as machine learning, deep learning, and big data analytics significantly accelerated progress. Since the 1980s, evolving neural network architectures have laid the groundwork for clinical decision support systems and drug discovery models, especially in areas like visual data processing and natural language processing (Gupta et al., 2021; Kaul, Enslin, & Gross, 2020; Niazi, 2023).

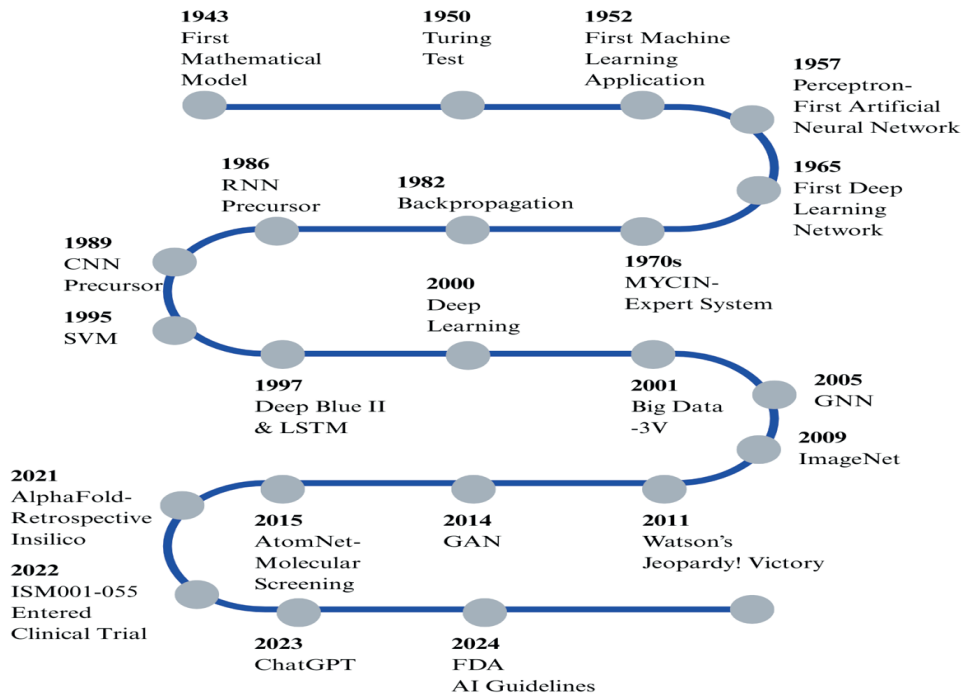


Figure 2. Timeline of historical developments in the field of artificial intelligence (1943–2024).

Visualized through a spiral timeline, these developments present the evolutionary trajectory of AI technologies in drug discovery within a temporal perspective—revealing a multilayered progression from early mathematical models to the widespread use of modern neural network architectures. This timeline not only illustrates a technical chronology, but also highlights scientific breakthroughs, technological leaps, and the stages at which AI integration gained momentum in pharmaceutical R&D processes. In this context, examples such as AtomNet (Wallach, Dzamba, & Heifets, 2015), AlphaFold (Jumper et al., 2021), and Insilico Medicine (Zhavoronkov et al., 2019) demonstrate that artificial intelligence is not only a theoretical concept, but also an effective tool for practical applications (Gupta et al., 2021; Insilico Medicine, 2022; Kaul et al., 2020; Kumar et al., 2022; Li et al., 2024; Niazi, 2023; Rai, Sahu, & Sawant, 2018).

AI-DRIVEN MOLECULAR DISCOVERY, OPTIMIZATION & SYNTHESIS PROCESSES

Target Identification and Lead Compound Discovery

Target Identification

Drug discovery relies on identifying biomolecular targets involved in disease mechanisms and discovering lead compounds that interact with them. In the era of personalized medicine, integrating molecular, environmental, and microbiome data has become

essential for defining clinically relevant targets. The growing regulatory focus on safety and ethics also demands more rigorous target validation studies.

Effectively managing this multidimensional process requires hybrid systems that combine human expertise with AI-based data processing. Human intuition enables recognition of complex biological patterns, whereas AI provides cognitive support in analyzing massive, multidimensional datasets. Such human-AI collaboration-empowered by cheminformatics and knowledge-based technologies-enhances the reliability and efficiency of target identification.

By enabling the integration and interpretation of large-scale biomedical data, systems biology and network-based computational models grounded in omics data have improved the reliability of target identification and promoted more holistic, data-driven strategies in drug discovery (Katsila, Spyrou, Patrinos, & Matsoukas, 2016). Integrating genetic and genomic evidence further improves target validation: clinical success rates can increase by up to 80% when strong genetic support exists for a given target (Wenteler et al., 2024).

Oprea and colleagues (2018) classified the human proteome into four categories -clinically known, chemically known, biologically known, and dark targets- collectively referred to as the “druggable genome,” as illustrated in Figure 3 (Oprea et al., 2018).

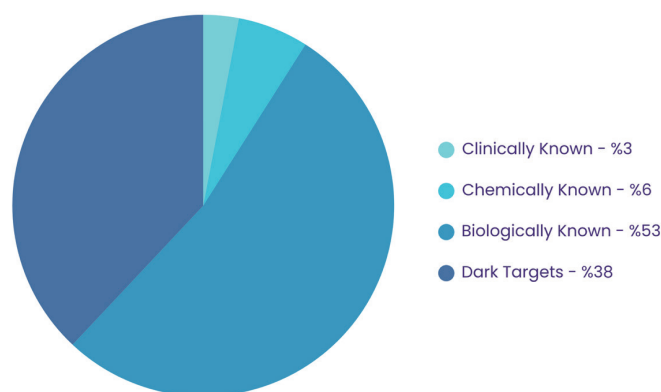


Figure 3. Classification of the human proteome according to Oprea et al., 2018.

This classification clearly underscores the vast potential that remains untapped in the field of drug discovery (Doytchinova, 2022). Accordingly, the Illuminating the Druggable Genome (IDG) initiative, launched by the U.S. National Institutes of Health (NIH) in 2014, aims to systematically elucidate the druggable genome. Open-access platforms such as Pharos and the Target Central Resource Database (TCRD) provide researchers with comprehensive data on over 20,000 human protein targets. These resources have significantly contributed to the understanding of underexplored protein targets and have facilitated the identification of novel therapeutic opportunities (Sheils et al., 2021).

Structural characterization of target macromolecules is essential for rational drug design. Techniques such as nuclear magnetic resonance (NMR) spectroscopy, X-ray crystallography, and more recently, cryo-electron microscopy (Cryo-EM) are used to determine protein structures, which are archived in the

Protein Data Bank (PDB). As of Q1 2025, the PDB contains structural data for 234,785 biomacromolecules, while PDDBind includes 3D structures and binding affinities for 27,385 complexes. This wealth of data on protein-ligand interactions serves as a valuable foundation for AI-integrated models, guiding the identification of new therapeutic opportunities.

In conclusion, the integration of AI-powered *in silico* approaches with genetic data not only enhances the accuracy and clinical success rate of target identification but also enables the development of more time- and cost-efficient drug discovery strategies.

Lead Compound Discovery

Drug discovery has historically relied on various approaches, the earliest of which involved serendipitous findings and the chemical modification of existing bioactive molecules. Over time, these empirical methods evolved into systematic screening strategies based on chemical libraries, molecular targets and AI (Doytchinova, 2022) (Figure 4).

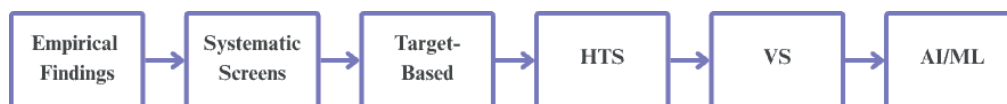


Figure 4. Evolution of lead discovery

With advances in combinatorial synthesis, automation, and genomic sequencing, the field shifted toward more target-oriented and efficient discovery pipelines. However, the exponential growth in compound diversity and target complexity has increasingly challenged conventional experimental methods (Wilkey, Haunso, Tudor, Webb, & Connick, 2017). The demand for systematic screening led to the development of high-throughput screening (HTS), which automates the testing of hundreds of thousands of molecules against biological targets. Despite its util-

ity, HTS is costly and resource-intensive, prompting a transition toward computational strategies such as virtual screening (VS). Integrating AI and ML algorithms has enabled VS to evolve from a computational extension of HTS into a core component of rational drug design (RDD) (Figure 5.). By virtually scanning large chemical libraries *in silico*, VS identifies potential “hit” compounds for experimental validation, thus accelerating early-stage drug discovery (Andricopulo & Ferreira, 2014).

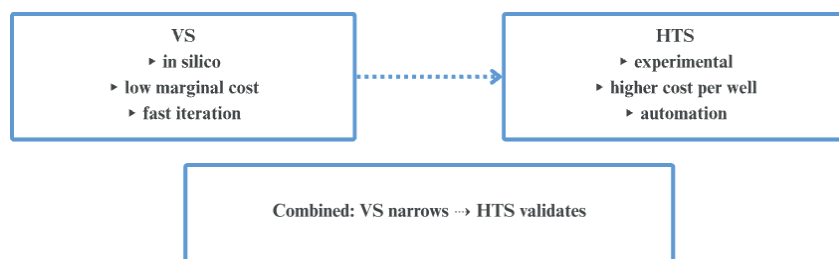


Figure 5. Virtual screening vs HTS

Rational drug design begins with the identification of a biological target associated with a specific disease and proceeds through the discovery and optimization of ligands that exhibit desired pharmacological activity (Doytchinova, 2022). In this context, two primary approaches form the foundation of rational drug design:

- *Ligand-Based Drug Design*

In many cases, the discovery of pharmacologically active new compounds can be achieved without requiring any information about the three-dimensional structure of the target biomolecule. In this context, LBDD strategies are widely used in the design and optimization of new ligands, based on the structural and physicochemical properties derived from previously validated bioactive molecules. In this method, the structural and activity data of inhibitors are analyzed to determine features associated with binding efficacy. Subsequently, large compound databases are screened to identify drug candidates that match these characteristics (Andricopulo & Ferreira, 2014; Wu et al., 2024a).

One of the prominent approaches within LBDD is the modeling of Quantitative Structure-Activity Relationships (QSAR). QSAR approaches aim to describe and quantitatively model the relationships between molecular features and biological activity. These relationships are modeled through mathematical formulas that allow the prediction of the potential activity of compounds that have not yet been synthesized. Another frequently employed method within LBDD strategies is pharmacophore modeling. This approach seeks to identify structural motifs that are commonly

found in a set of ligands and are assumed to play a critical role in binding to the target protein. The process involves analyzing the conformational space of molecules and aligning common features, ultimately leading to the generation of a three-dimensional pharmacophore hypothesis. This hypothesis can then be used to screen compound databases containing molecules with unknown activity (Andricopulo & Ferreira, 2014).

Another strategy within LBDD is ligand-based virtual screening (LBVS), which enables the identification of novel candidate molecules by relying solely on the structural and physicochemical properties of active compounds, without directly utilizing the three-dimensional structural information of the macromolecular target. In recent years, the integration of ML and deep learning (DL) algorithms into LBVS workflows has gained significant attention. These algorithms represent drug and target protein data using various encoding methods in vector or graph formats and extract meaningful features by processing this information. In this context, predictions of drug–target interaction (DTI) and drug–target affinity (DTA) have emerged as key steps within LBVS. This workflow is illustrated in Figure 6.

DTI prediction aims to computationally identify potential interactions between drug molecules and numerous possible targets in the organism. Compared to virtual screening, DTI has a narrower focus and can be considered a subcomponent of it. In this process, existing experimental data and known drug–target interactions are leveraged using ML and data mining techniques to predict novel potential associations. In

ML-based integrated models, DTI prediction is generally approached as a classification problem, here DTI tools produce binary outputs indicating the presence or absence of interaction between a molecule

and a target. In contrast, DTA prediction employs regression-based models to estimate quantitative values such as K_d , K_i , or IC_{50} that reflect the strength of the interaction (Wu et al., 2024a).

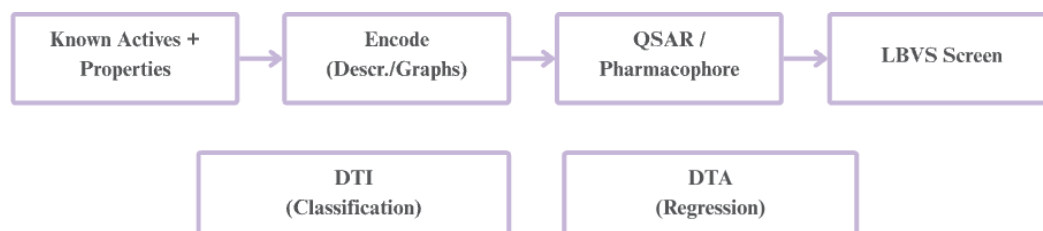


Figure 6. LBDD

• Structure-Based Drug Design – SBDD

Structure-based drug design is an approach that utilizes the known 3D structure of a target protein. More specifically, it is referred to as structure-based virtual screening (Ghislat, Rahman, & Ballester, 2021). Widely employed in solving various drug design challenges, Once the structure is resolved, an appropriate binding site on the target protein is identified, followed either by the optimization of existing ligands or the discovery of new ligands through *in silico* screening processes. Promising compounds resulting from this process may be synthesized or procured commercially and tested *in vitro* against the target protein.

Structure-Based Virtual Screening (SBVS) methods rely on the computational docking of large libraries of small molecules into the binding site of the target protein. Scoring functions (SFs) are employed to predict the binding affinity of ligands positioned within the binding site. these functions enable the ranking of compounds in chemical libraries based on their predicted affinity for the target protein. Compounds estimated to have higher binding affinity (i.e., lower K_d , K_i , or $\Delta G_{\text{binding}}$ values) are ranked at the top. These compounds are generally assumed to exhibit lower IC_{50} or EC_{50} values. Scoring functions are not limited to regression-based models; some also estimate class probabilities or directly classify molecules as active/inactive. Molecular docking is a widely used technique aimed at predicting the potential conformations of ligands that can fit into the

binding pocket of a biological receptor (Andricopulo & Ferreira, 2014; Salgin-Goksen et al., 2021; Tuncel et al., 2025). The identification of promising lead compounds is a complex process that requires the integrative application of various strategies. In this context, considering that target biomolecules may undergo conformational changes at different levels, SBVS strategies can be supported by molecular dynamics (MD) simulations. MD simulations play a fundamental role in predicting molecular motions and structural changes.

In certain cases, structural alterations in the target protein may be limited, allowing ligands to bind readily to well-defined binding pockets. However, it is also known that some proteins undergo significant conformational changes during the molecular recognition process. In such cases, a representative ensemble of multiple conformations of the target protein is generated, and incorporating these conformational models into SBVS analyses contributes to obtaining more meaningful and reliable results.

MD simulations have been successfully employed in elucidating molecular mechanisms supported by experimental data. Although they present some limitations in terms of system size and computational cost, when combined with medicinal chemistry approaches, they offer multifaceted contributions to drug design efforts (Andricopulo & Ferreira, 2014).

However, in cases where experimental structural data for the target protein are unavailable, struc-

ture-based screening processes can be supported by AI-based prediction methods. High-accuracy protein structure predictions can be achieved using advanced deep learning models such as AlphaFold2, while generative chemistry platforms like Chemistry42 can be used to design new molecules based on these pre-

dictions. Consequently, even for biomolecules with limited structural information, virtual screening and lead discovery can be effectively carried out (Niazi, Zamara, & Magoola, 2024). The AI-based SBDD workflow is illustrated in Figure 7.

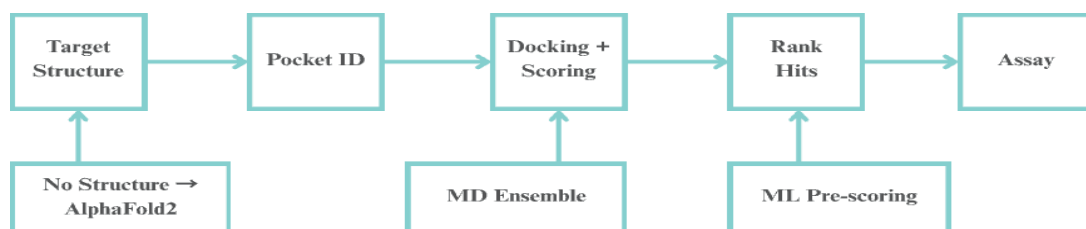


Figure 7. SBDD

ML, one of the most advanced fields of artificial intelligence, plays a significant role in drug design processes. In scenarios where target-specific ML-based scoring functions can be trained, various studies have reported higher SBVS accuracy compared to classical scoring functions. ML-based scoring functions developed within this context have been shown to process available target data effectively. Especially in the re-scoring of crystal structures, ML-based models have demonstrated high accuracy in predicting the binding affinity of ligands to target proteins. Similarly, successful results have been obtained for redocked li-

gand poses (Ghislat et al., 2021).

In recent years, studies that integrate MD simulations with ML approaches to improve the accuracy and efficiency of binding affinity predictions have drawn significant attention. These studies incorporate not only classical 3D structural descriptors but also dynamic information that reflects structural changes of the system over time, thereby generating more comprehensive 4D descriptors (Figure 8.). MD simulations are capable of modeling dynamic processes in biological systems across timescales ranging from picoseconds to milliseconds.

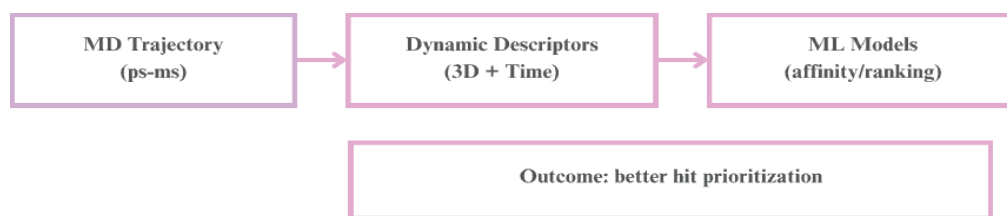


Figure 8. MD + ML Integration

In this context, a study by Jamal et al. developed ML-based models to predict biologically active compounds for the Caspase-8 target. During the training phase of these models, descriptors derived from MD simulations were employed. The study reported that models constructed with descriptors derived from longer simulations achieved the highest classification performance. These findings demonstrate that such models can be effectively utilized in the prioritization

and optimization of lead compounds (Jamal, Grover, & Grover, 2019).

Lead Optimization

Lead optimization is the systematic refinement of early hit/lead compounds to improve potency, selectivity, and pharmacokinetic/tox profiles (ADME), bridging hit discovery to preclinical/clinical development. Building on candidates from HTS and VS (see: Target Identification and Lead Discovery), this stage

enhances target interactions while improving properties such as solubility, stability, and bioavailability a critical inflection point in drug design.

Traditionally, teams iterate through design synthesis test cycles guided by empirical rules and expert intuition; this is effective but time-consuming. AI-driven methods accelerate the loop by learning from large structure–activity and experimental datasets to predict binding affinity and ADMET, prioritize modifications, and reduce the number of wet-lab iterations.

Recent platforms BIOVIA GTD, Query-based Molecule Optimization (QMO), Chemistry42, and ZairaChem apply generative models (e.g., genetic algorithms, neural networks) to propose potent, selective, and structurally novel variants (Bleicher et al., 2022). For example, in SYK inhibitor programs (entospletinib, lanraplenib), GTD re-identified and refined candidates by constraining chemical space and recombining features across series, improving lead quality (Bleicher et al., 2022; Loos et al., 2024).

Looking ahead, tighter integration of deep learning and reinforcement learning with established medicinal chemistry workflows and hybrid models that blend expert knowledge with data-centric systems should further improve multi-parameter optimization. Progress will depend on robust data infrastructures and close interdisciplinary collaboration (Niazi et al., 2024).

AI-Assisted Synthesis Planning and Synthetic Accessibility

The 1828 Wöhler urea synthesis launched total synthesis as a field. Today, discovery pipelines use ligand-based and structure-based virtual screening to

find binders, but these methods are limited to pre-existing molecules, constraining chemical space, patentability, and diversity. Hence the rise of *de novo* design yet synthetic accessibility remains the key bottleneck. Classical retrosynthesis plans routes by working backward from the target to purchasable precursors. Effective but manual, it scales poorly for complex structures and depends heavily on expert intuition motivating AI integration. AI-driven retrosynthesis couples large reaction databases with ML to propose feasible routes. Typical systems comprise:

1. Representations (SMILES, fingerprints, graph encodings),
2. Single-step prediction (template-based rules vs template-free bond disconnections),
3. Multi-step planning, and
4. Route evaluation.

Quality control includes round-trip evaluation (forward models regenerate the product), reaction diversity metrics (e.g., Jensen–Shannon divergence), and feasibility scores (SCScore) alongside ML predicted yields. Multi-step planning often mirrors reinforcement-learning-like decision sequences optimized for cost, step count, and overall yield (Jiang, Yu, Kong, Mei, & Luo, 2023). Data-driven approaches now often surpass rule-based systems. Wang et al. introduced RetroExplainer, combining the Retro*10 algorithm with commercial reagent databases to ensure synthesizability and avoid manual reactant selection; across 101 test cases, 86.9% of single-step predictions matched literature routes. The method also offered interpretable energy scores and reproduced a four-step route to protokylol, as shown in Figure 9, consistent with reported chemistry (Wang et al., 2023).

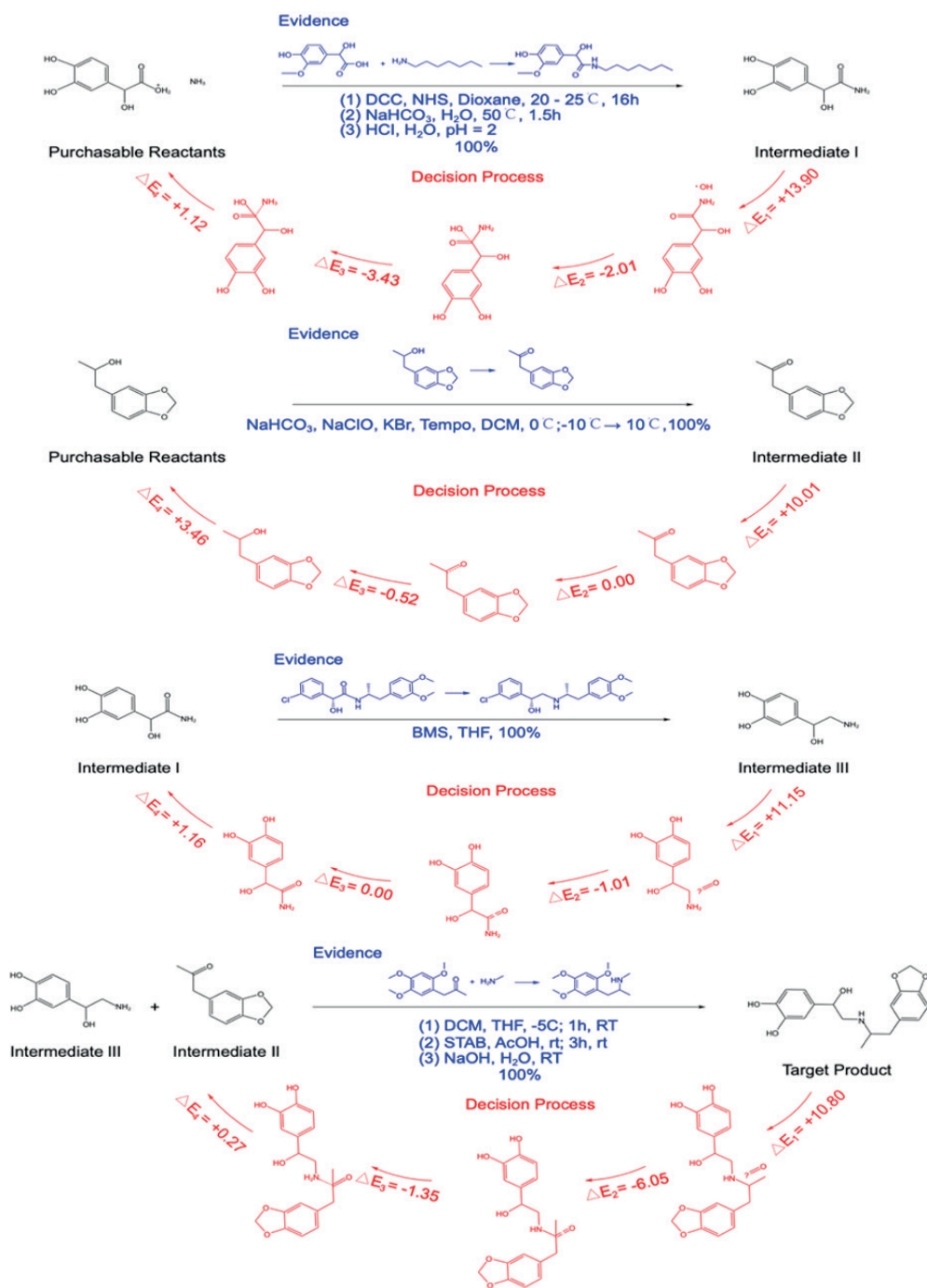


Figure 9. Retrosynthetic route for protokylol generated by RetroExplainer. The structures shown in grey represent intermediate compounds and the final target molecule along the synthetic pathway. At each retrosynthetic step, the blue-highlighted regions indicate literature or dataset examples provided by the model as supporting evidence. The red-highlighted structures and arrows represent the model's decision-making process, where the ΔE value displayed on each arrow corresponds to the energy score of the respective synthetic step. Lower ΔE scores indicate chemically more plausible transformations.

MACHINE LEARNING

AI is a field of science and engineering that aims to enable computers to imitate human behavior, replicate decision-making processes, and solve complex tasks autonomously. In its early stages, AI systems were built upon predefined rules and formal representations of knowledge, where decision-making mechanisms relied on these fixed rules. However, the dynamic nature of real-world problems has revealed the limitations of such rigid systems, highlighting the need for more flexible, data-driven learning capabilities.

In response to this need, ML methods have emerged, allowing computer programs to improve their performance on specific tasks through experience. ML automates analytical modeling processes and enables cognitive functions -such as classification, regression, and clustering- to be learned from data without explicit programming. With the increasing volume of data, advancements in computational

infrastructure, and the development of novel algorithms, ML has found widespread application across numerous industries.

Progress in ML has also led to the evolution of artificial neural networks (ANNs) into deeper and more complex architectures. DL, which utilizes deep neural networks (DNNs) comprising multiple layers, enables the learning of high-level representations from data. DL models can directly extract intricate patterns from raw data and deliver superior performance on large-scale datasets. Nevertheless, in scenarios with low data dimensionality or limited training data, traditional -or shallow- ML algorithms often outperform deep learning models, offering not only better accuracy but also higher interpretability of the results (Janiesch, Zschech, & Heinrich, 2021). A visual representation of this conceptual distinction is provided in Figure 10, which broadly illustrates the positions of shallow and deep learning models within the broader machine learning framework.

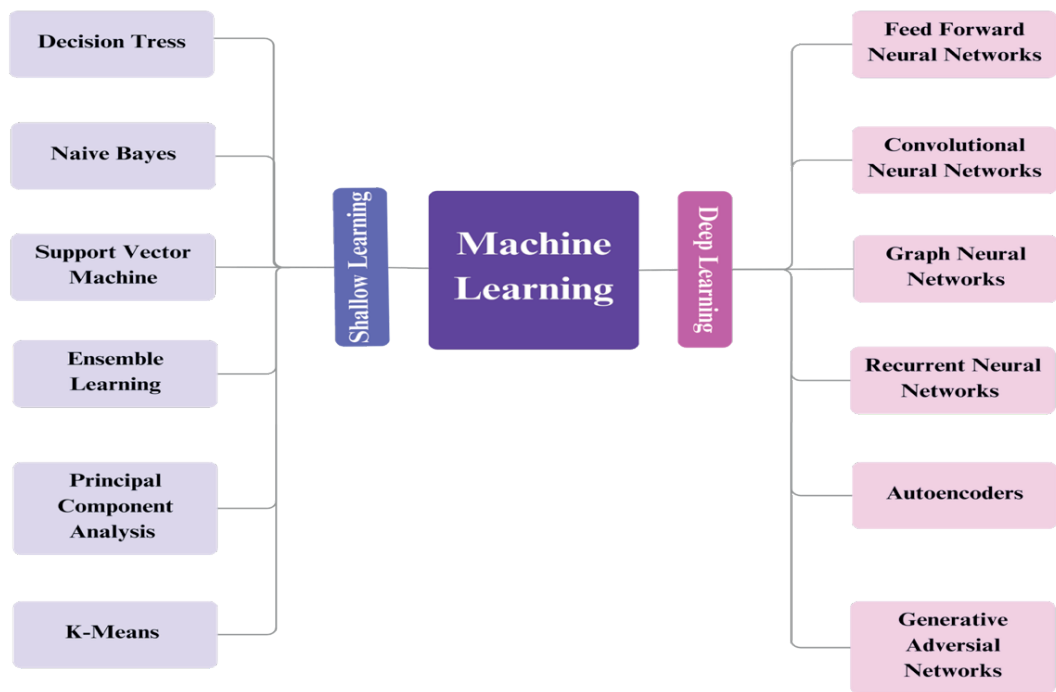


Figure 10. A schematic classification illustrating the positions of shallow and deep learning algorithms within the broader framework of machine learning.

Ultimately, a hierarchical relationship exists among the concepts of AI, ML, and DL: AI represents the broadest umbrella term, encompassing the general field of intelligent systems, while ML constitutes a subfield of AI that focuses on data-driven learning. Within ML, deep learning emerges as a specialized approach that leverages multi-layered neural network architectures to enable advanced pattern recognition and representation learning.

ML approaches are generally categorized into three main types: supervised learning, unsupervised learning, and semi-supervised learning. Each of these paradigms offers specialized techniques tailored to distinct data structures and learning objectives (Janiesch et al., 2021).

Supervised Learning

Supervised learning is based on training a model using labeled datasets. In this approach, each input example is presented to the algorithm along with its corresponding correct output label. The objective of the model is to learn the patterns within these input-output pairs and make accurate predictions on new, unseen data with similar structure.

The training process typically involves two main phases. In the first phase, the model is trained on a labeled training dataset. In the second phase, the model's generalization ability is evaluated using a test dataset composed of previously unseen data. During this process, the model's predictions are compared to the true outputs, and the discrepancy is quantified using a loss function. Optimization techniques are then employed to minimize this error (Mahesh, 2020; Nasteski, 2017).

Supervised learning methods primarily address two types of problems: classification and regression. Classification tasks aim to assign input data to predefined categories. For instance, classifying lung nodules as benign or malignant based on imaging data exemplifies a classification problem. In contrast, regression tasks focus on predicting continuous numerical variables. An example of this is estimating the op-

timal drug dosage for patients based on demographic and clinical data using regression-based models (Ryan et al., 2023).

This approach is widely applied across both traditional (shallow) machine learning algorithms and deep learning architectures. Traditional methods include algorithms such as Decision Trees, Naive Bayes, and Support Vector Machines, while deep learning models such as Convolutional Neural Networks (CNNs) and Recurrent Neural Networks (RNNs) are also trained within the supervised learning paradigm.

Unsupervised Learning

Unsupervised learning is a subfield of machine learning that operates on unlabeled data. In this approach, the algorithm is trained without any predefined labels or known output values associated with the input data. The primary goal of the model is to uncover similarities, patterns, or structural relationships within the data and to organize the input in a meaningful way.

Unlike supervised learning, unsupervised learning does not involve a "correct answer." The algorithms freely analyze the internal structure of the data and attempt to identify hidden patterns or relationships. This process is typically associated with tasks such as clustering and dimensionality reduction. Clustering aims to group data samples with similar characteristics, while dimensionality reduction techniques seek to reduce the number of variables in high-dimensional data while preserving critical information, thereby enabling more concise and interpretable data representations.

Because unsupervised learning attempts to reveal inherent structures in data without external guidance, the learning process is generally more ambiguous and complex than in supervised approaches. Nevertheless, it offers significant advantages, particularly in large and complex datasets, by enabling automatic grouping of data or the discovery of previously unrecognized patterns (Greene, Cunningham, & Mayer, 2008; Mahesh, 2020).

This approach is implemented across both shallow ML algorithms and deep learning architectures. For instance, K-means clustering and Principal Component Analysis (PCA) are widely used conventional methods, while Autoencoders and Generative Adversarial Networks (GANs) represent deep learning models commonly employed in unsupervised learning frameworks.

Semi-Supervised Learning

Situated between supervised and unsupervised learning, *semi-supervised learning* seeks to improve model performance by leveraging both labeled and unlabeled data. Particularly in tasks such as classification and clustering, the incorporation of unlabeled data can help the model make more robust generalizations. For example, in classification problems, unlabeled instances complementing a limited number of labeled examples may enable more accurate definition of decision boundaries. Similarly, in clustering applications, prior knowledge indicating that certain data points belong to the same group can facilitate the formation of more coherent clusters.

Semi-supervised learning becomes especially valuable when labeled data is scarce or difficult and costly to obtain—scenarios commonly encountered in domains such as computer-aided diagnosis or drug discovery. In these contexts, insights derived from large volumes of unlabeled data can enhance model accuracy and lead to better overall performance compared to models trained solely on labeled datasets (Van Engelen & Hoos, 2020).

Within the realm of deep learning, semi-supervised learning is implemented through models trained on partially labeled datasets. In certain instances, models such as GANs are also employed under this category. Additionally, RNN architectures like Gated Recurrent Units (GRUs) and Long Short-Term Memory (LSTM) networks are occasionally used within semi-supervised learning paradigms. One of the key advantages of this approach lies in its ability to significantly reduce the reliance on large amounts of labeled data (Alzubaidi et al., 2021).

Data Resources

In machine learning-based drug discovery studies, the availability of high-quality and structured data resources is of critical importance.

PubChem is the world's largest open-access chemical information database, providing comprehensive data on chemical compounds and their biological activities. It includes chemical structures, physicochemical properties, toxicity data, and 2D/3D structural information, as well as protein interaction data derived from biochemical assays. The data are organized into three main categories: *Substance* (raw chemical records), *Compound* (unique structures), and *BioAssay* (results of biological experiments).

ChEMBL is a freely accessible database specifically developed for drug discovery purposes, containing biologically active molecules. Established in 2002 by EMBL-EBI, it compiles compounds and their associated bioassay results from medicinal chemistry literature, approved drugs, and clinical trial data. Its latest version includes information on over 1.9 million compounds, more than 10,000 drugs, and upwards of 12,000 target proteins.

DrugBank is one of the most widely used and reliable open-access drug reference databases. Launched in 2006, it offers extensive bioinformatics and cheminformatics data on drugs and their corresponding target proteins. Drug-target interaction data in DrugBank are curated from scientific publications, textbooks, and other major databases.

Protein Data Bank (PDB) serves as the primary source for three-dimensional structural information on proteins and protein-ligand complexes. As one of the earliest open-access biological data repositories, it provides 3D structures of proteins, nucleic acids, and their complexes, mostly derived through X-ray crystallography and, to a lesser extent, nuclear magnetic resonance (NMR) spectroscopy.

UniProt is among the most comprehensive protein-centric databases, based on protein sequences obtained from genome sequencing projects. It pro-

vides detailed functional annotations of proteins, and its *Swiss-Prot* subset contains manually curated information on more than 560,000 verified protein entries.

BRENDA, is an extensive enzyme database established in 1987. It contains data on approximately 84,000 enzymes and nearly 205,000 enzyme–ligand relationships. These records are manually curated from over 140,000 scientific publications. The ligands include substrates, products, activators, inhibitors, and cofactors. Users have full access to the entire dataset.

BindingDB, is an open-access database that focuses on binding affinities between drug-like small molecules and protein targets. Its content is derived from scientific articles and patents and is integrated with other databases such as ChEMBL and PubChem.

PDBbind, is a specialized database that aggregates experimentally determined binding affinity data for biomolecular complexes from the PDB. Introduced in 2004, it aims to bridge the gap between protein structural information and binding energetics. The dataset is based on complex structures available in the PDB and their associated binding affinity data reported in the scientific literature (Niazi et al., 2024).

Data Representations

The successful application of machine learning and deep learning algorithms in drug discovery critically depends on the ability to represent both drug candidates and target proteins in numerically interpretable formats. These representations are regarded as a fundamental step in computational modeling, as they allow for the encoding of the structural, chemical, and biological properties of molecules and proteins. In general, such representations are categorized into three main types: string-based, fingerprint-based, and graph-based (2D/3D) formats.

One of the most common approaches for ligand representation is the Simplified Molecular Input Line Entry System (SMILES), which expresses a compound as a sequence of characters. This method linearizes the molecular graph to generate a textual representation. However, due to its sequential nature, SMILES

often fails to capture spatial relationships within the molecule, potentially leading to the loss of critical structural information.

To overcome these limitations, fingerprint-based representations are employed. These representations encode molecules as binary vectors based on the presence or absence of specific substructures-assigning “1” to indicate presence and “0” for absence. Among the most widely used are Extended Connectivity Fingerprints (ECFP) or Morgan fingerprints. These representations have shown strong performance in molecular similarity analysis and property prediction tasks.

In 2D graph-based molecular representations, atoms are treated as nodes and bonds as edges. Each node is described by a feature vector capturing its chemical or structural characteristics, while an adjacency matrix encodes the connectivity between atoms. Some representations are further enriched with edge features, offering a more detailed view of the molecular structure. This format aligns seamlessly with Graph Neural Network (GNN)-based models, enabling more accurate analysis of molecular topology and interactions.

3D graph-based representations consider the spatial coordinates of atoms in a molecule and incorporate stereochemical information encoded through x, y, and z coordinates. These representations are especially effective in quantum-level property prediction tasks and in identifying binding pockets.

Similar strategies are employed for protein representations. Protein sequences are typically transformed into binary vectors using one-hot encoding. However, such sequence-based representations do not inherently capture structural information. To address this, contact maps -a form of 2D graphical representation- are commonly used. These maps encode pairwise relationships between amino acid residues in matrix format, allowing for the integration of both sequence and structural information.

Finally, 3D graph-based protein representations are constructed using structures obtained from data-

bases such as PDB and AlphaFold. In these graphs, nodes may represent either amino acid residues or atoms, and edges denote the physical or chemical interactions between them. These 3D representations are particularly valuable in advanced tasks such as identifying protein–ligand interaction sites, as they enhance model accuracy and generalizability (Wu et al., 2024a).

Traditional (Shallow) Learning (SL)

Traditional machine learning, also referred to as shallow learning (SL), encompasses the majority of machine learning models proposed prior to 2006. Shallow learning aims to capture patterns in data using relatively simple and direct approaches. These methods typically involve a limited number of parameters and shallow model architectures, focusing on identifying explicit relationships within the data without learning complex representations. Due to their simplicity, shallow learning algorithms offer faster training processes and higher interpretability. However, when it comes to modeling multi-layered or complex structural relationships, their performance is often inferior to that of deep learning techniques (Xu, Zhou, Sekula, & Ding, 2021).

Within the framework of shallow learning, a wide range of algorithms has been developed to address different types of problems. This section provides a detailed overview of core models including Decision Trees, Naive Bayes, Support Vector Machines (SVM), Ensemble Learning methods, PCA, and K-means clustering.

Decision Tree

Decision trees are a machine learning method primarily used for classification tasks. This technique is based on representing decisions and their possible consequences in a tree-like structure, enabling systematic modeling of the decision-making process. The construction of a decision tree begins with the entire dataset and involves recursively splitting the data into two subsets based on the value of a specific feature. This splitting continues at each step until the resulting subsets contain instances belonging to a single class.

To determine the decision nodes, the concept of

entropy is often employed. For each feature, the entropy of the resulting subsets is calculated, and the feature-value pair that results in the lowest total entropy is selected as the decision node. Accordingly, each node in the decision tree represents a feature and a threshold value associated with that feature.

During classification, the data instance is evaluated starting from the root node, and at each decision node, it is compared with the corresponding feature and directed to the appropriate branch. This process continues until a leaf node is reached. The leaf node indicates the predicted class of the instance, thereby completing the classification process.

In a decision tree, each node represents a condition or decision, while branches denote the possible outcomes of those decisions. In other words, nodes reflect the attributes used for classification, and the branches correspond to the possible values of those attributes (Mahesh, 2020; Nasteski, 2017).

Naive Bayes

Naive Bayes is a simple yet effective classification technique that falls under the category of supervised learning methods and is fundamentally based on probability theory. The algorithm aims to estimate the conditional probability of each class $\mathcal{Y}$ given an object $\mathcal{X}$, expressed as $P(\mathcal{Y} | \mathcal{X})$. Through conditional probability computations, Naive Bayes is widely utilized in classification tasks and decision support systems.

Thanks to its straightforward computational requirements, the Naive Bayes algorithm is capable of delivering fast results and exhibits reliable performance even on small datasets, with low variance. Moreover, it allows for incremental learning by seamlessly incorporating new data into the system, enabling continuous model updates. The algorithm is also notably robust in the presence of errors and missing values within the dataset.

The flexible and robust nature of the Naive Bayes algorithm has facilitated its widespread and effective application across various domains (Mahesh, 2020; Nasteski, 2017; Webb, Keogh, & Miikkulainen, 2010).

Support Vector Machines (SVM)

SVMs are a widely used and effective technique in the field of machine learning, primarily employed for classification and regression tasks within the supervised learning paradigm. SVM aims to represent each input in a high-dimensional feature space and to determine the optimal hyperplane that separates the data samples. This hyperplane is constructed to maximize the margin between classes, thereby minimizing classification error.

The SVM algorithm is not limited to linearly separable problems; it can also be applied to nonlinear classification tasks through the use of a technique known as the kernel trick. This approach maps the original data into a higher-dimensional feature space, where a linear separation becomes possible.

SVMs are particularly effective in domains characterized by complex data structures and have shown strong performance on high-dimensional datasets, especially where the number of features exceeds the number of samples. However, without proper regularization, SVM models are prone to overfitting, particularly in noisy or sparse data environments.

On the downside, SVM is often considered a “black-box” algorithm, meaning that the underlying decision-making process can be difficult to interpret. Specifically, the rationale behind the selection and optimization of the separating hyperplane is not readily transparent, which limits the explainability of model outputs (Jiang, Gradus, & Rosellini, 2020; Mahesh, 2020).

Ensemble Learning (EL)

EL is a technique in supervised machine learning that combines multiple models to produce a collective decision. In this approach, each model functions as a base learner, trained on labeled data to generate predictions for new, unlabeled instances. EL may incorporate various machine learning algorithms, including decision trees, artificial neural networks, or linear regression. The central rationale behind this method is that the individual errors made by different models can offset each other, thereby improving overall pre-

dictive accuracy. As a result, ensemble methods often achieve higher performance compared to any single model acting alone.

This approach reflects the principle of “wisdom of the crowd” in machine learning. A classic demonstration of this idea was provided by British philosopher and statistician Sir Francis Galton. In a contest held at a livestock fair, Galton asked participants to estimate the weight of an ox. Although none of the individual guesses were perfectly accurate, the average of all the predictions closely approximated the actual weight. This experiment highlighted the potential of aggregated predictions to yield more accurate results than individual estimates. Inspired by this concept, ensemble models aim to combine the outputs of multiple learners to generate more reliable and accurate predictions.

Moreover, ensemble models are particularly efficient when the base learners have low computational costs. This enables both more accurate and faster predictions to be achieved.

Ensemble learning methods are generally categorized into two main frameworks: dependent and independent.

In the dependent framework, the output of each model influences the training of subsequent models. Here, misclassified instances in previous iterations are given greater importance during the training of new models. A well-known example of this approach is the AdaBoost algorithm.

In contrast, the independent framework trains each model separately, without influence from others. The outputs are then aggregated using techniques such as majority voting. In this framework, a bootstrapping technique is often employed to generate subsets of the data—where some samples may be used multiple times and others not at all. The Random Forest algorithm, which applies the bagging method to decision trees, is a representative example of an independent ensemble learning strategy (Arifa, Aditsania, & Kurniawan, 2022).

Principal Component Analysis (PCA)

PCA is a widely used technique for reducing the dimensionality of large datasets. By projecting data onto a lower-dimensional space, PCA facilitates faster and more efficient analysis. The method aims to represent numerous variables using smaller groups while preserving the essential information in the dataset and minimizing noise, thus revealing the most meaningful underlying structure.

Also known as the Karhunen–Loève transform, PCA is utilized in various applications such as dimensionality reduction, data compression, and visualization. Extended versions of this method, including probabilistic PCA and kernel PCA, have also been developed. These approaches contribute to highlighting key features and uncovering hidden structural relationships, particularly in large and complex datasets.

The computational process of PCA is grounded in linear algebra and can be efficiently executed by computers. By isolating only the most informative components, PCA enhances the performance and speed of subsequent machine learning algorithms. Additionally, it helps mitigate the challenges associated with high-dimensional data and reduces the risk of overfitting in regression-based models (Kurita, 2019; Naeem, Ali, Anam, & Ahmed, 2023).

K-Means

The K-Means algorithm is one of the most widely used techniques within unsupervised learning. Its primary objective is to partition a dataset into K predefined groups, or clusters, based on data similarity. The algorithm operates by assigning data points to clusters in such a way that the distance between the data points and the centroids of their respective clusters is minimized.

Initially, K centroids are randomly selected. Each data point is then assigned to the nearest centroid based on a distance metric, typically Euclidean distance. Subsequently, the centroids are updated by calculating the mean position of all data points assigned to each cluster. These steps are iteratively repeated un-

til the cluster assignments stabilize and the centroids no longer change significantly.

K-Means is particularly effective for large-scale and numerical datasets due to its computational efficiency and simplicity. However, since the initial selection of centroids can significantly influence the final clustering outcome, it is often recommended to perform multiple runs with different initializations to enhance robustness and accuracy (Mahesh, 2020; Naeem et al., 2023).

DEEP LEARNING (DL)

DL, a more advanced subfield of machine learning, aims to capture complex patterns in data through multilayered ANNs. These networks are inspired by the information processing mechanisms of the biological nervous system and are composed of interconnected artificial neurons designed to process inputs and generate outputs. Each neuron performs basic mathematical operations by weighting and summing the incoming signals, which are then transformed into outputs through activation functions. This layered organization of numerous neurons lays the foundation for learning more intricate representations and decision-making processes (Montesinos López, Montesinos López, & Crossa, 2022; Singh & Banerjee, 2019).

Artificial neurons, inspired by their biological counterparts, operate by aggregating incoming signals, applying weights to them, and transforming the results into outputs through activation functions once a predefined threshold is reached. In ANNs, the input layer is responsible for receiving data from the external environment, hidden layers process this data through nonlinear transformations, and the output layer generates the final prediction or decision. The connections between neurons are mediated by adjustable weights that govern the flow of information and are optimized during the learning process. In certain neural architectures, feedback connections where the output of a layer is fed back into a previous layer enable more dynamic and interactive learning mechanisms.

Deep learning overcomes the limited representational capacity of traditional neural networks, which are often constrained to a single hidden layer, by introducing multiple hidden layers to learn more complex and hierarchical patterns. These architectures, referred to as DNNs, are capable of mapping intricate relationships between inputs and outputs, particularly when trained on large-scale datasets. As the number of hidden layers increases, so does the model's learning capacity, enabling it to extract progressively higher-level abstractions from the data (Montesinos López et al., 2022).

The topology of an artificial neural network defines the overall structure of the network and the pattern of connections among neurons. Different types of connections such as inter-layer, intra-layer, and self-connections play a crucial role in determining the network's information processing capacity. The weights assigned to each connection govern the strength and direction of information flow, serving as fundamental parameters during the learning process (Abiodun et al., 2018).

Today, deep learning has become an indispensable approach in domains that rely heavily on large-scale datasets, particularly for tasks such as pattern recognition, feature extraction, data-driven prediction, and complex decision-making. While previous sections have explored how AI-based techniques are integrated into the drug discovery pipeline -from target identification to lead compound optimization- this section introduces the core principles of deep learning and outlines selected neural network architectures that have demonstrated particular relevance in pharmaceutical research.

Feedforward Networks (FFNNs)

FFNNs represent the most fundamental architecture of artificial neural networks. In these networks, information flows in a single direction—from the input layer, through the hidden layers, to the output layer. Connections exist only between successive layers; there are no intra-layer or skip connections. This

unidirectional flow of information contributes to the relatively simple and interpretable structure of feedforward neural networks, facilitating their analytical tractability (Abiodun et al., 2018).

The architecture of feedforward neural networks consists of three fundamental components: the input layer, hidden layers, and the output layer. The input layer serves as the initial interface where external data enters the network. Each input neuron represents an independent variable in the model and directly influences the output depending on the training conditions. Data are received at the input layer, passed through one or more hidden layers via feedforward propagation, and finally reach the output layer.

The output layer transforms the processed data into the final output of the network. The number of neurons in this layer is determined by the specific task the network is designed to perform. For instance, in classification tasks, one output neuron may be assigned to each category, whereas in tasks like noise reduction, the input and output layers might need to contain the same number of neurons.

Hidden layers, located between the input and output layers, enhance the network's learning capacity. These layers, which are not directly connected to the external environment, extract pattern-based features from the input data and contribute significantly to the generation of the final output. The number of neurons in the hidden layers plays a critical role in determining model performance. If too few neurons are used, the network may fail to learn complex patterns, resulting in underfitting. Conversely, an excessive number of neurons may increase the risk of overfitting and unnecessarily prolong the training process, thereby raising computational costs.

Therefore, when designing a feedforward neural network, both the number of layers and the number of neurons per layer must be carefully determined. These parameters directly affect the model's learning capacity and should be optimized based on the nature of the problem and the structure of the data (Ben-Bright et al., 2017).

In training feedforward neural networks, trial-and-error-based strategies are commonly employed to determine the optimal architecture. One such strategy is the Forward Selection Method, which begins by constructing a model with a small number of hidden neurons and gradually increases the number of neurons based on training and test performance. Typically, the process starts with two hidden neurons, and additional neurons are incrementally added as performance improves.

In contrast, the Backward Selection Method starts with a model containing a relatively large number of hidden neurons and gradually reduces the number of neurons based on performance evaluations. Neurons are removed one by one until a noticeable decline in performance is observed. Both approaches aim to balance model capacity and optimize the learning process.

In addition to these strategies, pruning techniques are also employed to reduce unnecessary complexity and accelerate the training process. Pruning involves analyzing the connection weights within the network and eliminating neurons associated with connections whose weights are close to zero. This approach not only reduces computational cost but also facilitates the development of a more parsimonious and generalizable model (Ben-Bright et al., 2017).

Convolutional Neural Networks (CNNs)

CNNs are among the core deep learning approaches, distinguished by their multilayered archi-

tectures and high accuracy, particularly in visual data processing tasks.

The architecture of CNNs is generally composed of three main components: convolutional layers, pooling layers, and fully connected layers.

Convolutional layers process input data and intermediate feature maps using a set of kernels to generate new feature maps.

Following the convolutional layers, pooling layers are applied to reduce the dimensionality of the feature maps and to decrease the number of model parameters.

Finally, fully connected layers transform the two-dimensional feature maps into one-dimensional vectors and are primarily responsible for tasks such as classification. These layers often account for the majority of a CNN's total parameters, which significantly increases computational cost.

The training of a CNN model typically consists of two main phases: forward propagation and backpropagation. During forward propagation, the input data are processed using the current weights and biases, and the predicted output is compared with the true labels to compute the loss. Subsequently, in the backpropagation phase, gradients are calculated using the chain rule, and the model parameters are updated accordingly. After a sufficient number of iterations, the learning process is completed. The workflow of the general CNN architecture is shown in Figure 11 (Guo et al., 2016).

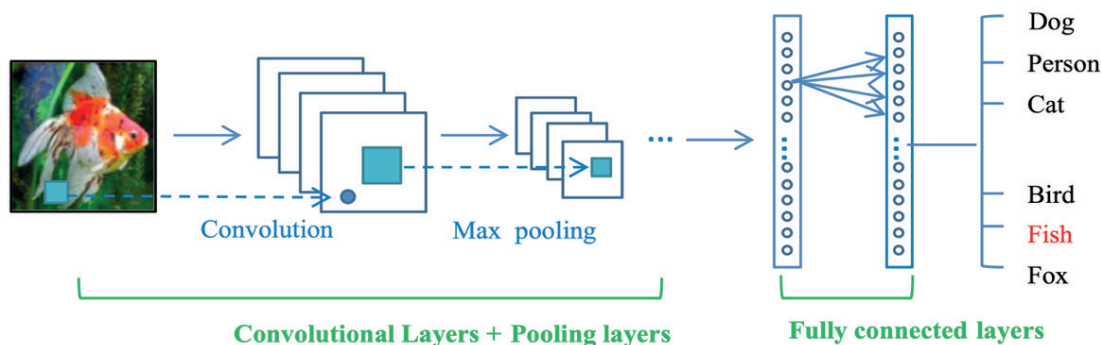


Figure 11. Schematic workflow of a general CNN.

The large number of parameters involved in deep neural networks often leads to undesirable issues such as overfitting. To address this challenge, various regularization strategies have been developed. One such strategy is the Dropout technique, which randomly deactivates a subset of feature detectors during each training iteration to prevent the model from overfitting to the training data. This technique enhances the model's generalization capability and reduces the risk of overfitting. Dropout is also considered a form of ensemble learning, as it effectively trains multiple sub-models simultaneously.

In CNNs, which are widely used in visual object recognition tasks, data augmentation techniques are frequently employed to expand the training dataset. These techniques enhance data diversity without additional labeling costs, thereby supporting better generalization. AlexNet stands as an early and notable example of leveraging data augmentation to improve performance.

Another effective strategy is pre-training, which involves initializing a network with previously learned parameters instead of random weights. This approach accelerates the learning process and improves generalization performance. For instance, models such as AlexNet, trained on the ImageNet dataset, are often fine-tuned for new tasks and datasets, allowing for faster adaptation and improved results.

Another strategy, fine-tuning, is a crucial step for adapting models to specific tasks and datasets. For instance, models such as AlexNet, trained on the ImageNet dataset, are often fine-tuned for new tasks and datasets, enabling faster adaptation and improved performance. During this process, class labels are typically required to compute loss functions for model optimization. However, in certain scenarios, these labels may be unavailable. To address this limitation, similarity learning methods have been proposed, allowing the model to adapt without the need for explicit class annotations.

All of these regularization strategies can be applied individually or in combination to further enhance the

overall performance and robustness of CNN models (Guo et al., 2016).

CNNs Models

Among the prominent CNN architectures, AlexNet stands out as one of the earliest deep learning models that brought significant attention to the field. Comprising five convolutional layers and three fully connected layers -eight layers in total- AlexNet processes fixed-size input images through successive convolution and pooling operations, ultimately producing the final output via fully connected layers. Trained on the ImageNet dataset, AlexNet integrated regularization techniques such as data augmentation and Dropout, and its victory at the 2012 ImageNet Large Scale Visual Recognition Challenge (ILSVRC) marked a turning point for deep learning research and applications (Guo et al., 2016) .

Aiming to increase network depth by reducing the size of convolutional kernels, the Visual Geometry Group (VGG) model focuses on building deeper and more structured networks using small convolutional filters. VGG achieved second place in the 2014 ILSVRC competition with a Top-5 error rate of 7.3%, demonstrating the positive correlation between network depth and performance. In particular, VGG-16 has become widely used in image processing tasks due to its simple architecture and strong compatibility with transfer learning.

The winner of the same competition, GoogLeNet, introduced the Inception module, which allows multiple convolution operations with different kernel sizes to be executed simultaneously. Despite its depth, GoogLeNet effectively reduced the number of parameters and offered a more computationally efficient architecture. By incorporating sparse connections and global average pooling, the model significantly lowered computational costs while maintaining a lightweight design.

Following the success of GoogLeNet, the Inception architecture was developed with the goal of increasing network depth while maintaining computational effi-

ciency. Evolving from Inception V1 to V2 and V3, this architecture optimizes multi-scale processing capabilities, enabling deep learning models to be effectively applied to larger and more complex datasets.

The Residual Network (ResNet) model was developed to address the degradation problem caused by increasing the number of layers in deep neural networks. ResNet adopts the residual learning approach, which facilitates the training of very deep networks by incorporating sequential residual blocks. This architecture enables the model to learn identity mappings more effectively, thereby improving gradient flow and reducing the risk of performance degradation as the network depth increases. For instance, a ResNet model with 152 layers possesses a deeper architecture compared to shallower networks like VGG, while achieving lower error rates and higher accuracy (Zhao et al., 2024).

Graph Neural Networks (GNNs)

Graph structures consist of nodes, which represent entities, and edges, which denote the relationships or connections between these nodes. Such structures play a crucial role in modeling complex relational systems, including social networks, citation networks, and molecular structures. GNN-based approaches have found significant applications in pharmaceutical research, particularly in modeling molecular structures and in the discovery of novel drug candidates.

GNNs have been developed to process graph-structured data characterized by irregular and dynamic topologies, drawing inspiration from traditional CNNs. While CNNs excel at capturing spatial dependencies in grid-like data and RNNs are effective at learning sequential correlations, GNNs demonstrate superior performance on graph-based data by modeling complex dependencies through nodes and edges (Khemani, Patil, Kotecha, & Tanwar, 2024).

A graph is a structure composed of nodes interconnected by edges. Edges represent the connections between nodes. Graphs are employed in relevant domains to model and analyze the relationships between objects or entities.

GNNs are neural networks that operate on data structures consisting of nodes and edges. The fundamental components of a GNN include nodes, edges, layers, activation functions, pooling, aggregation, and other common neural network components.

A node is a point or vertex in a graph that can be connected to other nodes. In GNNs, each node is associated with a feature vector that contains the properties of the corresponding entity and primarily represents the node's attributes. Node classification is a powerful technique in classification problems.

Nodes play a critical role in enabling the network to learn from the graph data structure and its connections. During GNN training, information between nodes is propagated via the edges that connect them. This process allows the network to learn from the relationships among nodes and to make predictions for previously unseen nodes in the graph.

Edges are the connections between nodes in a graph. Edge features can provide significant information regarding the relationships between nodes. This facilitates understanding the types of connections shared by nodes in one-to-one, one-to-many, or many-to-many relationship contexts. Edge prediction plays an important role in link prediction, node classification, and graph-level classification.

Edge features can be represented as a vector and are typically used together with node features, enabling the network to learn from both the attributes of individual nodes and the relationships between nodes.

GNN layers are the fundamental components that enable the network to learn from the connections within the graph data structure. Each layer aggregates information from the feature vectors of neighboring nodes and combines it with the current node's feature vector to generate new representations.

In GNNs, activation functions are employed to introduce non-linearities to the outputs of each layer, allowing the network to learn complex patterns in the data. The choice of activation function in a GNN de-

depends on the nature of the input data and the specific application. Different activation functions can have varying effects on the training process and the performance of the network.

Pooling and aggregation are used in GNNs to reduce the dimensionality of the feature space and to enable the network to handle graphs of varying sizes. Pooling refers to the process of consolidating information from multiple nodes or subgraphs into a sin-

gle representation. Aggregation, on the other hand, involves combining information from neighboring nodes into a single representation for each node. Aggregation is often used in conjunction with pooling to create more compact representations of the graph data. By pooling and aggregating information from multiple nodes and subgraphs, the network can learn more effectively and produce scalable representations of graph data.

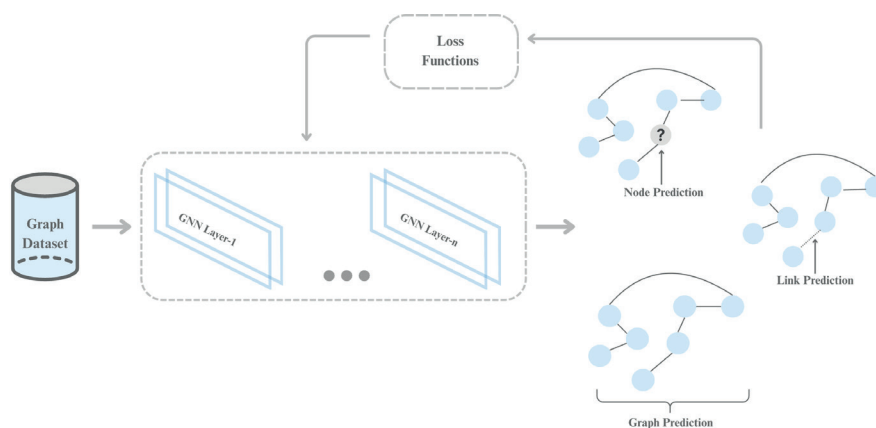


Figure 12. Schematic workflow of a general GNN.

Figure 12 illustrates the workflow of a GNN schematically. The graph dataset is input to the network and passed through GNN layers that incorporate the fundamental components described above. The loss function is applied as in other deep learning systems, and the network is trained until a predefined error threshold or number of iterations is reached. The task is typically a classification at the node, edge, or graph level. Finally, the trained model is evaluated by making predictions on the test data (Sharma, Singh, & Ratna, 2024).

Node-level tasks focus on determining the identity or function of each node within a graph. They are typically used when working with unlabeled data; for instance, predicting whether a specific individual is a smoker.

Edge-level tasks (link prediction) concentrate on analyzing the relationships between pairs of nodes in a graph. An example of such a task is assessing the likelihood or compatibility of a connection between two en-

tities. On platforms such as Netflix, an edge-level task may involve predicting the next video to recommend based on a user's viewing history and preferences.

Graph-level tasks aim to predict an overall property or behavior that encompasses the entire graph. For example, in the evaluation of a newly synthesized chemical compound, a graph-level task may seek to determine whether the molecule has the potential to serve as an effective drug (Khemani et al., 2024).

GNN Models

GNNs have been developed in various architectures based on their application domains and can be broadly categorized into four main types: RecGNNs, ConvGNNs, Graph Autoencoders (GAEs), and Spatio-Temporal Graph Neural Networks (STGNNs) (Park, Yi, & Ji, 2020).

RecGNNs are among the pioneering architectures of GNNs and aim to learn node representations through an iterative information exchange mecha-

nism. This approach refines the concept of message passing and serves as the foundation for the development of ConvGNNs.

ConvGNNs generalize traditional convolution operations from grid-structured data to graph-structured data and generate high-level node representations by stacking multiple graph convolution layers. ConvGNNs are widely used in a variety of tasks, including node classification and graph classification.

GAEs are based on an unsupervised learning approach and aim to encode nodes or entire graphs into a latent vector space from which the original graph structure can be reconstructed. They are effectively used in applications such as network embedding and graph generation.

STGNNs aim to learn latent patterns from dynamic data by simultaneously modeling spatial and temporal dependencies within graph-structured information. These models are widely applied in time-sensitive tasks, particularly in traffic speed prediction and human activity recognition (Wu et al., 2021).

Recurrent Neural Networks (RNNs)

RNNs are deep learning models specifically designed to process sequential data. Unlike traditional feedforward neural networks (FFNs), RNNs can re-

tain information from previous inputs in their memory and utilize this information when processing new data. This capability enables the network to effectively learn temporal dependencies and sequential patterns (Lipton, Berkowitz, & Elkan, 2015; Park, Yi, & Ji, 2020; Yang et al., 2019a).

RNNs are an extension of FFNs; however, unlike FFNs, RNNs incorporate recurrent connections between hidden units, thereby providing the model with temporal context information. This structure allows RNNs to effectively process sequential dependencies in time-series data. The network is designed with a hidden state mechanism that retains information from previous inputs, and its core architecture consists of an input layer, a hidden layer, and an output layer. At each time step, the hidden units receive information both from the current data point and from the hidden state of the previous step; consequently, the output is computed using both the current and accumulated past information, facilitating a cyclic flow of information within the network. This recurrent mechanism enables the learning of long-term dependencies (Lipton et al., 2015; Mienye, Swart, & Obaido, 2024). Figure 13 presents a structural comparison between the general FFN and RNN architectures.

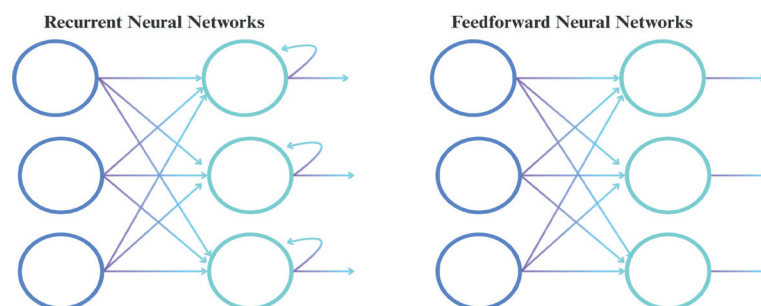


Figure 13. Structural Comparison of general FFNs and RNNs.

Training RNNs is a complex process, particularly due to the challenges encountered when learning long-term dependencies. During training, the Backpropagation Through Time (BPTT) algorithm is applied (Lipton et al., 2015; Park et al., 2020; Yang et al., 2019a).

During backpropagation, the propagation of errors through time often leads to the vanishing gradient and exploding gradient problems. The exploding gradient problem arises when the gradient norms grow excessively during training, causing unstable learning.

To address this issue, the Truncated Backpropagation Through Time (TBPTT) method has been proposed, which limits the span of error propagation, enabling more stable and balanced learning.

On the other hand, the vanishing gradient problem occurs when gradients exponentially decay over time, hindering the model's ability to learn long-term dependencies (Pascanu, Mikolov, & Bengio, 2013). To overcome this limitation, the Long Short-Term Memory (LSTM) network has been proposed, enabling effective learning of long-term dependencies.

Although structurally similar to conventional RNNs, LSTMs employ memory cells instead of hidden units, allowing them to learn long-term dependencies. Each memory cell maintains its own internal state, which helps prevent information loss over time. LSTMs utilize input and output gates to control when information is written to or read from memory and employ a constant error carousel mechanism to mitigate vanishing or exploding gradients. Thanks to these capabilities, LSTMs are widely applied in fields such as natural language processing, speech recognition, and handwriting recognition (Lipton et al., 2015).

As an alternative to LSTMs, the Gated Recurrent Unit (GRU) architecture has been proposed, offering comparable performance while featuring a simpler structure and fewer parameters, thereby accelerating the training process (Mienye et al., 2024).

Another successful RNN architecture is the Bidirectional Recurrent Neural Network (BRNN). BRNNs can learn both past and future context by processing input sequences in both forward and backward directions. The BRNN architecture consists of two separate hidden layers: one processes the inputs in the forward temporal direction, while the other operates in reverse. This bidirectional processing allows the model to make predictions at each time step based on both previous and subsequent inputs. However, the dependence on future information may limit the applicability of BRNNs in online learning or real-time data processing scenarios. Nevertheless, for fixed-length

sequences, this bidirectional structure offers significant performance advantages.

LSTM and BRNN architectures possess complementary characteristics, and their combination—Bidirectional LSTM (BLSTM)—has demonstrated superior performance in sequential data tasks such as phoneme classification and handwriting recognition (Lipton et al., 2015).

Autoencoders (AEs)

Autoencoders (AEs) serve as fundamental building blocks that enable the hierarchical structuring of deep neural network models. These networks are designed to learn and reconstruct input data. By organizing, compressing, and extracting high-level features from the data, they facilitate unsupervised learning and the discovery of nonlinear features. In unsupervised learning, the primary objective is to transform raw data into more meaningful and information-rich representations.

AEs are feedforward neural networks that transmit information unidirectionally and are particularly prominent in tasks such as feature learning and dimensionality reduction. Consequently, they are widely applicable across different data types and domains.

The most basic autoencoder, which does not incorporate additional complexity or architectural modifications, is referred to as the “vanilla autoencoder.” A basic autoencoder typically consists of an input layer (encoder), one or more hidden layers (latent space), and an output layer (decoder). This structure compresses the input data into intermediate representations and subsequently reconstructs the original input from these representations (Berahmand, Nasiri, & Karimi, 2024). Figure 14 provides a visual depiction of the architecture of a general autoencoder neural network.

Conventional autoencoders typically employ a single-layer encoder. However, this limitation can impede the learning of deep and abstract features. To address this issue, deeper network architectures and layer-wise learning approaches have been developed. In this context, Stacked Autoencoders (SAEs) are con-

structured by hierarchically stacking multiple basic autoencoders layer by layer.

SAEs adopt the principle of layer-wise learning, wherein each autoencoder compresses the data into a lower-dimensional representation to extract high-

er-level features. The final output is obtained from the combined outputs of all individual autoencoders. This approach enables the learning of deeper and more abstract representations and allows for effective modeling of complex data structures.

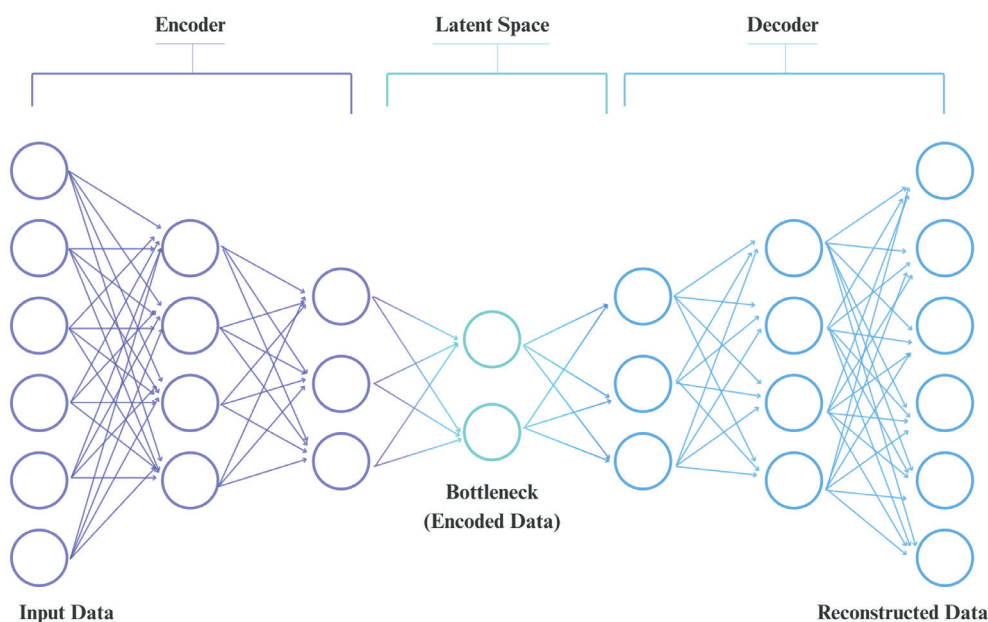


Figure 14. Schematic depiction of the architecture of a general AE.

The training of autoencoders involves hyperparameter and architectural design choices that directly affect model performance. These decisions are typically optimized through experimental evaluation and validation procedures to achieve task-specific performance objectives.

The depth of the network and its capacity to capture complex data patterns are primarily determined by the number of hidden layers. Increasing the number of layers can enhance representational capacity but also introduces optimization challenges and raises the risk of overfitting.

The bottleneck structure refers to the latent layer that compresses the most essential features of the input data. The dimensionality of this layer governs the trade-off between information retention and compression. An excessively small bottleneck may lead to significant information loss.

Activation functions enhance the network's ability to model nonlinear relationships, thereby improving its representational power. Functions such as Sigmoid, Tanh, ReLU, and SELU are commonly employed in autoencoder architectures.

Optimization algorithms update the weights and biases to minimize the loss function. Techniques such as Stochastic Gradient Descent (SGD), Adam, and Adagrad are selected based on dataset size and model complexity.

The learning rate determines the step size of weight updates; excessively high or low values can adversely affect the training process.

The batch size refers to the number of samples used in each optimization step. Smaller batches provide faster but noisier updates, whereas larger batches offer more stable gradients at the cost of higher memory usage. The optimal batch size should be chosen

according to the scale of the dataset and the available computational resources (Berahmand et al., 2024).

AE Models

Autoencoders can be adapted for different data types and tasks through architectural modifications and objective function adjustments, offering a broad range of applications.

Regularized Autoencoders (RAEs) aim to represent data in a compressed latent space while applying regularization constraints to ensure that the learned representations exhibit meaningful and structured characteristics. These regularization techniques promote the learning of sparse representations, manifold structures, or orthogonal features, enhancing interpretability and robustness.

Robust Autoencoders are designed to handle noisy or corrupted inputs and improve model resilience against common real-world data issues such as outliers and missing values. Robust AE models are commonly classified into three main categories: Denoising Autoencoders (DAE), Marginalized Denoising Autoencoders (mDAE), and $L_{2,1}$ -norm Autoencoders, each offering varying degrees of robustness and reconstruction fidelity.

GAEs go beyond traditional data compression by learning the underlying probability distribution of the data, thereby enabling the generation of new samples resembling the training data. Among these, Variational Autoencoders (VAEs) have attracted significant attention in unsupervised learning due to their ability to represent and compress complex data distributions. In pharmaceutical research, VAEs have been effectively utilized in the discovery and design of novel drug molecules, with a notable example being the pioneering work conducted by Gómez-Bombarelli et al. (2018) (Gangwal et al., 2024).

Convolutional Autoencoders (CAEs) employ convolutional layers in both encoder and decoder modules instead of fully connected layers, allowing them to capture spatial dependencies within the data. CAEs have demonstrated high performance in image-relat-

ed tasks, including image denoising, inpainting, segmentation, and super-resolution.

Recurrent Autoencoders (RAEs) are tailored for sequential data (e.g., time series) and incorporate architectures such as Long Short-Term Memory (LSTM) or Gated Recurrent Units (GRU) in their encoder and decoder modules. This allows them to effectively preserve and learn from temporal dependencies within the sequence data.

GAEs aim to enhance the efficiency of graph analysis by transforming graph-structured data into lower-dimensional embeddings. In this architecture, the encoder compresses the input graph into a vector representation, and the decoder reconstructs the original graph structure from this latent space. When combined with models like Graph Convolutional Networks (GCNs), GAEs can effectively capture node attributes and topological relationships, offering significant advantages in graph-based learning tasks (Berahmand et al., 2024).

Generative Adversarial Networks (GAN)

Developed in 2014, GANs represent a significant milestone in generative artificial intelligence (GAI)-based modeling. These models are capable of learning the fundamental distribution of real-world data and generating various types of synthetic data, such as images, videos, and even molecular structures. By employing an adversarial training strategy, GANs enable the creation of new compounds with desired properties. In particular, GAN-based approaches have been effectively utilized in drug discovery to expand the chemical space and identify novel chemical structures (Gangwal et al., 2024).

GANs consist of two adversarial neural networks: a generator (G) and a discriminator (D). Typically constructed using convolutional and/or fully connected layers, these two models are trained in opposition to one another, continuously improving in response to each other's performance (Creswell et al., 2018).

In the GAN architecture, the generator network produces synthetic data samples and presents them

to the discriminator. The generator aims to map from a latent variable space to the data space, generating samples that closely resemble the true data distribu-

tion. Conversely, the discriminator is trained to determine whether a given input originates from the real dataset or has been generated by the generator.

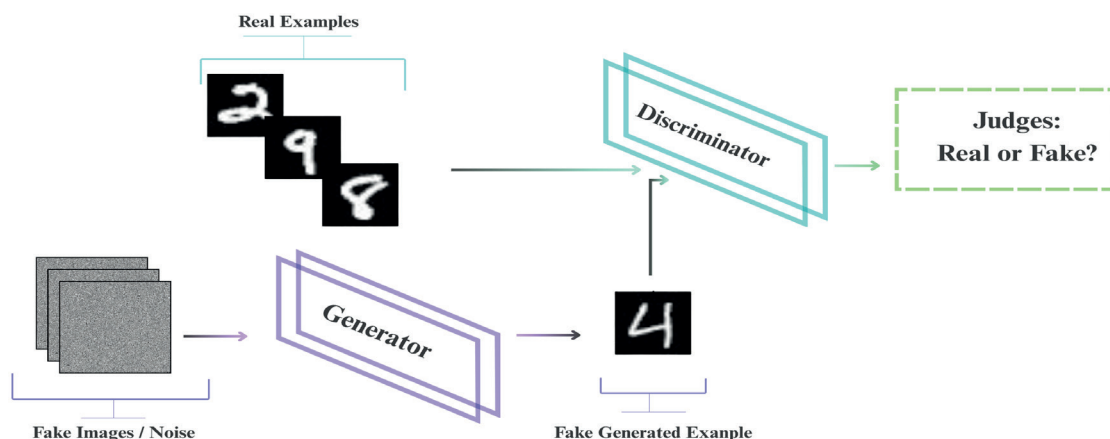


Figure 15. Schematic workflow of a general GAN.

The discriminator is optimized to assign high probabilities to real data and low probabilities to synthetic (fake) data. Meanwhile, the generator seeks to deceive the discriminator into classifying its outputs as real, thereby maximizing the likelihood that its generated data will be perceived as genuine. This adversarial process is grounded in zero-sum game theory: while the generator minimizes its loss the discriminator simultaneously strives to maximize its classification accuracy. Figure 15 illustrates the schematic workflow of the GAN architecture.

The ultimate objective of GANs is to establish a Nash equilibrium between the generator and discriminator networks. Theoretically, this equilibrium is reached when the distribution of data produced by the generator aligns with the true data distribution. At this point, the generator has effectively learned the real data distribution and can produce novel, highly realistic samples that are indistinguishable from genuine inputs (Gui, Sun, Wen, Tao, & Ye, 2021; Oussidi & Elhassouny, 2018). However, achieving this equilibrium is challenging, as real and generated data occupy only limited regions within the same space.

When the discriminator rapidly attains high accuracy, the gradients propagated to the generator

approach zero, causing the learning process to stall. Additionally, the alternating parameter updates of the two networks can introduce instability during training; under such conditions, the generator tends to produce highly similar outputs for different inputs—a phenomenon known as mode collapse.

One of the first major improvements addressing these issues was introduced through the Deep Convolutional GAN (DCGAN) architecture. The reduction of fully connected layers and the use of batch normalization enhanced the stability and efficiency of training deeper networks. Moreover, employing leaky ReLU activation functions in the intermediate layers of the discriminator yielded superior performance compared to conventional activations.

Subsequently, several heuristic methods were proposed to further stabilize GAN training. Among these, feature matching allows the generator to learn by imitating the intermediate activations of the discriminator rather than solely attempting to deceive it. The mini-batch discrimination technique introduces an additional input feature to the discriminator, preventing the generator from producing identical or overly similar outputs. Additionally, label smoothing sets the target label for real samples to 0.9 instead of

1.0, thereby softening the discriminator's decision boundary, reducing overconfidence, and providing the generator with more balanced gradients.

To mitigate the vanishing gradient problem—a key challenge—f-GAN was proposed by generalizing the loss function used in classical GANs. This approach provides greater flexibility to the training process.

Subsequently, the Wasserstein GAN (WGAN) model was developed to address the vanishing gradient problem inherent in classical GANs. WGAN provides the generator with more meaningful and stable gradients, thereby enhancing the overall stability of the training process. In this model, the discriminator—referred to as the “critic”—does not perform probabilistic classification; instead, it measures the distance between the real and generated data distributions.

In general, GANs are highly powerful models capable of generating new data from random noise. However, challenges such as vanishing gradients, convergence difficulties, and mode collapse make the training process highly complex. Successful training of GANs depends on maintaining a balanced adversarial dynamic between the generator and the discriminator. Consequently, through various developed approaches over time, the training strategies of GANs have become more stable, balanced, and efficient (Creswell et al., 2018).

GAN Models

The original GAN (Vanilla GAN) architecture paved the way for the emergence of various structural and methodological variants. One of the major challenges in GAN training is the vanishing gradient problem, which can severely impede model convergence. To address this issue, the Wasserstein GAN (W-GAN) model was developed, measuring the difference between real and generated data distributions using the Wasserstein distance and enforcing a Lipschitz continuity constraint on the discriminator; this results in a more stable training process.

Nevertheless, W-GAN may still experience subop-

timal sample generation or convergence issues under certain conditions. To overcome these limitations, the Loss-Sensitive GAN (LS-GAN) model was proposed, aiming to achieve a more balanced learning process by controlling the discriminative capacity of the discriminator. Both approaches preserve the core structure of GANs while employing different strategies to enhance training stability (Sengar, Hasan, Kumar, & Carroll, 2024).

Since the Vanilla GAN architecture does not have direct access to label information during data generation, its ability to perform controlled data synthesis is limited. To address this, the Conditional GAN (CGAN) was developed. CGAN integrates auxiliary information, such as class labels, into both the generator and discriminator, enabling data generation conditioned on specific attributes (Navidan, Dehghan-tanha, & Parizi, 2021).

Extending the CGAN approach, InfoGAN was developed to strengthen the relationship between observed data and latent variables. This model maximizes the generator's loss while minimizing the discriminator's loss, thereby increasing the mutual dependency between generated data and latent codes. Consequently, it enhances the interpretability of the generative process and provides a more structured and disentangled generation mechanism (Sengar et al., 2024; Wang, She, & Ward, 2021).

Designed for semi-supervised learning, the Auxiliary Classifier GAN (AC-GAN) allows the generated samples to be classified according to both real/fake status and class labels, enabling more targeted and label-informed data generation (Navidan et al., 2021).

A significant architectural advancement in GANs was achieved with the Deep Convolutional GAN (DCGAN), which employs deconvolution operations in the generator to better model spatial patterns and utilizes convolutional layers for more realistic data generation.

The Laplacian Pyramid Adversarial Network (LAPGAN) aims to produce high-resolution outputs

from low-resolution inputs. Using a multi-scale learning approach based on Laplacian pyramids, it generates sharper and more detailed images through progressive resolution enhancement (Wang et al., 2021).

While classical GANs focus solely on generating data from latent variables, the Bidirectional GAN (BiGAN) introduces a mechanism to map real data into the latent space, providing a unified framework for both data generation and representation learning.

Adversarial Autoencoders combine traditional autoencoder architectures with the adversarial learning framework to learn more structured and semantically rich latent representations. Building on this concept, the Adversarial Variational Bayes (AVB) model was developed, enhancing the Variational Autoencoder (VAE) framework with adversarial learning to obtain more expressive data representations (Creswell et al., 2018).

REAL-WORLD APPLICATIONS

In recent years, the integration of AI technologies into drug discovery processes has garnered significant attention due to their potential to reduce development costs and accelerate timelines. AI-based approaches are increasingly being adopted, particularly with the aim of identifying novel therapeutic candidates more efficiently and at lower cost.

AI-Assisted Molecular Discovery Approaches in Scientific Literature

In a study conducted by Chen et al. in 2020, an artificial intelligence-based approach was developed to identify potential dual inhibitors targeting fibroblast growth factor receptor 4 (FGFR4) and epidermal growth factor receptor (EGFR). In this work, four different machine learning models including SVM and Random Forest (RF) were trained to predict the biological activities of both targets. The IC_{50} values of 843 compounds for FGFR4 and 5,088 compounds for EGFR were collected from BindingDB to support the modeling process.

The results demonstrated that the SVM model achieved the highest prediction accuracy for both targets. This model was subsequently applied to predict the biological activities of a set of in house compounds previously synthesized and archived by the research group. To evaluate the accuracy of the predictions, the kinase inhibition potentials of selected compounds were experimentally tested. Among these, compound 1 exhibited significant inhibitory activity against both FGFR4 ($IC_{50} = 86.2$ nM) and EGFR ($IC_{50} = 83.9$ nM) kinases.

In conclusion, this study demonstrates the feasibility of in silico modeling for predicting dual-target inhibition against FGFR4 and EGFR kinases, offering a valuable methodological framework for assessing the biological activity of compounds in early-stage drug discovery. Compound 1, selected based on predictive modeling, showed potent inhibitory activity against both targets and emerged as a promising candidate for further investigation. Compound 1 is illustrated in Figure 16 (Chen et al., 2020).

In the previously discussed study, classical machine learning algorithms were prominent; however, the following example demonstrates the modeling capacity of these technologies in drug discovery through a deep learning-based approach enriched with graph neural network architectures.

Zhi, Zhao, Lee, & Chen (2021) developed an AI-assisted multilayered approach for the discovery of novel dihydroorotate dehydrogenase (DHODH) inhibitors that may be effective in the treatment of small-cell lung cancer (SCLC). In this study, both GNN architectures and traditional machine learning algorithms namely, RF and SVR were utilized to evaluate the biological activity of candidate compounds through various modeling techniques. The workflow employed in this process is presented below in Figure 17.

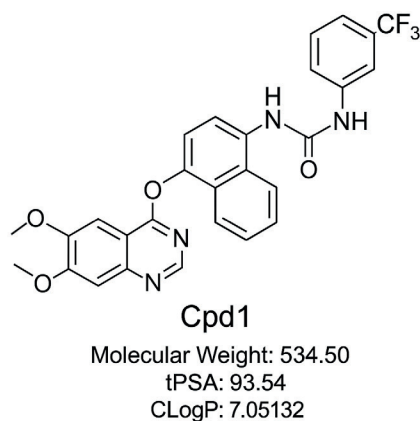


Figure 16. Structure of Compound 1 identified by Chen et al. As a dual FGFR4/EGRF inhibitor.

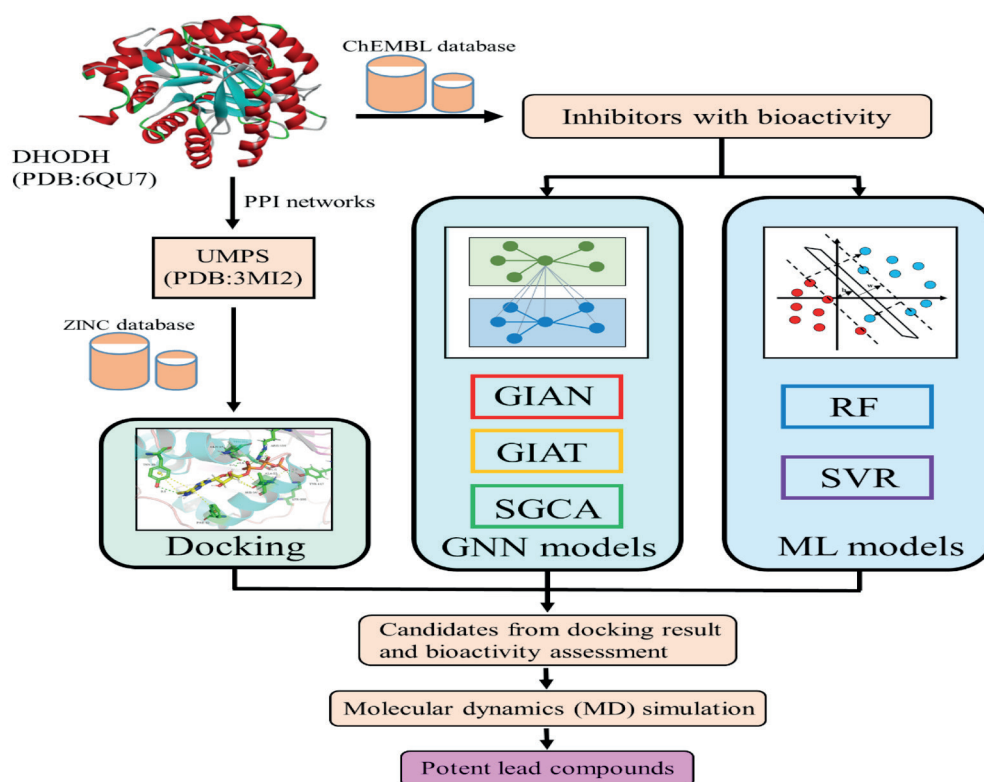


Figure 17. Workflow diagram of the study conducted for the discovery of DHODH inhibitors.

The modeling process was based on two primary data sources: data on known inhibitors with biological activity were retrieved from the ChEMBL database, while potential novel compound candidates were obtained from the ZINC database. The retrieved structures were subjected to molecular docking analysis against DHODH (PDB: 6QU7) and the related pro-

tein UMPS (PDB: 3MI2). Binding scores and interaction profiles were evaluated to select lead compounds.

Subsequently, the selected compounds were analyzed using three different GNN architectures, as well as RF and SVR models. The results demonstrated that graph neural networks achieved higher accuracy in predicting biological activity compared to traditional

methods. The top performing compounds were then subjected to MD simulations, in which the binding stability and interaction dynamics of the ligand–protein complexes were thoroughly assessed.

As a result of these analyses, three compounds ZINC8577218, ZINC95618747, and ZINC4261765 were identified as exhibiting high binding stability

and potential inhibitory effects on the DHODH target. These findings highlight the potential of deep learning-assisted modeling as a powerful tool for identifying effective compounds in early-stage drug discovery (Zhi et al., 2021). The structural representations of these three lead candidate compounds are presented in Figure 18.

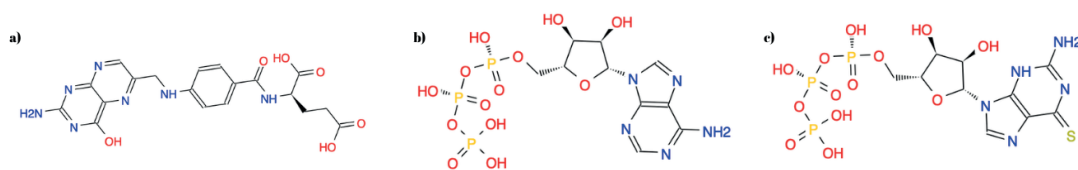


Figure 18. 2D structural representations of the compounds identified as potential DHODH inhibitors: a) ZINC8577218, b) ZINC4261765, and c) ZINC95618747.

Applications of AI-Based Molecule Discovery in The Pharmaceutical Industry

In January 2020, the British pharmaceutical company Exscientia announced the initiation of Phase I clinical trials for DSP-1181, a compound developed for the treatment of obsessive-compulsive disorder (OCD). DSP-1181 was designed using an artificial intelligence platform that screens chemical libraries to identify the most promising candidates, making it the first AI-designed drug candidate to enter human clinical trials. Exscientia developed this molecule in collaboration with Japan-based Sumitomo Dainippon Pharma and completed the process from initial screening to preclinical testing in just 12 months. In contrast, this stage typically takes an average of 4–6 years in the pharmaceutical industry, and only about 1 in 1000 screened molecules generally advance to clinical phases (Burki, 2020).

Similarly, Insilico Medicine has made notable progress in AI-driven drug discovery. In February 2022, the company initiated a Phase I clinical trial for ISM001-055, a small-molecule inhibitor developed for the treatment of idiopathic pulmonary fibrosis (IPF). The discovery process of ISM001-055 was entirely AI-driven, from target selection via the PandaOmics™ platform to molecule design using the Chemistry42™

engine. The company emphasized that the total time from target identification to Phase I trial initiation was less than 30 months (Insilico Medicine, 2022).

In January 2022, within the framework of a collaboration between Insilico Medicine and Fosun Pharma, ISM004-1057D was developed as a modulator of the QPCTL enzyme, which plays a regulatory role in the CD47–SIRPα axis—an important signaling pathway in immuno-oncology. This molecule has been advanced to the preclinical stage as a potential first-in-class small-molecule inhibitor targeting this pathway (Kirkpatrick, 2022).

For the treatment of anemia associated with chronic kidney disease, a PHD1/2 inhibitor was discovered using Insilico Medicine's Pharma.AI platform. Leveraging structural information of target proteins, a structure-based virtual screening approach was implemented through the Generative Chemistry application to design candidate molecules targeting the PHD1/2 enzymes. These candidates were optimized over several iterations in terms of physicochemical properties as well as *in vitro* and *in vivo* ADME profiles, and subsequently evaluated in preclinical studies. In June 2023, an Investigational New Drug (IND) application was submitted in China, and Phase I clinical trials are currently ongoing (Insilico Medicine, 2025).

In the field of rare diseases, HLX-1502, developed by Healx, stands out as a promising candidate. Discovered through Healx's AI-driven platform, HLX-1502 has been granted Fast Track, Orphan Drug, and Rare Pediatric Disease designations by the U.S. Food and Drug Administration (FDA). The compound completed its Phase I clinical trials in 2024 and is currently being evaluated in Phase II as of 2025 (Healx, 2024).

Exscientia and Recursion are two additional key players attracting attention for their work in AI-driven drug discovery. These companies aim to accelerate the drug development process and deliver novel therapeutic options for rare and serious diseases through their technology-based platforms.

Within this context, Exscientia's CDK7 inhibitor REC-617, intended for the treatment of advanced solid tumors, is expected to have Phase I monotherapy safety and pharmacokinetic/pharmacodynamic data disclosed by the end of 2024. Additionally, updated Phase I dose-escalation data for the RBM39-targeting compound REC-1245 are anticipated in the first half of 2026. Initial patient dosing for the MALT1 inhibitor REC-3565 and the LSD1 inhibitor REC-4539 is projected to occur in Q1 and the first half of 2025, respectively. For REC-4881, a MEK1/2 inhibitor being developed for familial adenomatous polyposis, Phase 1b/2 safety and efficacy data are expected in the first half of 2025. As for REC-2282, an HDAC inhibitor targeting neurofibromatosis type 2, results from the PFS6 (6-month progression-free survival) analysis are scheduled to be released in 2025.

Furthermore, the development candidate for REV-1025, an ENPP1 inhibitor for the treatment of hypophosphatasia, is expected to be identified in Q4 of 2024. An update on the Phase II progress of REC-3964, developed to prevent recurrent *Clostridium difficile* infections, is planned for Q1 2026. Finally, IND-enabling studies are ongoing for REC-4209, an investigational compound targeting idiopathic pulmonary fibrosis, whose molecular target has not yet been disclosed (Exscientia, 2024).

In the field of antibiotic discovery, a study published in *Cell* in February 2020 reported the identification of a novel antibiotic using a deep learning-based model. The researchers trained a neural network on a dataset comprising 2,335 compounds to identify molecules capable of inhibiting *E. coli* growth. Subsequently, the model was applied to a library of 6,111 molecules. Based on the model's predictions, the compound coded as SU3327 was identified for its antibacterial activity and renamed "Halicin"—a reference to the artificial intelligence system HAL 9000 from Stanley Kubrick's film 2001: A Space Odyssey. Halicin demonstrated efficacy not only against *E. coli* but also against multidrug-resistant pathogens such as *Clostridium difficile* and *Acinetobacter baumannii*. However, this compound has not yet progressed to clinical trials (Burki, 2020).

Overall, AI-driven drug discovery studies have demonstrated significant success in accelerating development processes and identifying novel molecular entities. AI technologies, particularly in the stages of target identification, molecule design, and prioritization, have overcome limitations inherent in traditional approaches and introduced a new paradigm. Real-world applications confirm that AI serves not only as a theoretical promise in drug discovery but also as a transformative tool in practice.

ARTIFICIAL INTELLIGENCE-ENABLED DRUG DEVELOPMENT & EVALUATION PROCESSES

Preclinical Stages

In preclinical stages, evaluating drug exposure in humans and optimizing the pharmacokinetic (PK) profile represent key objectives of drug discovery and development efforts (Healx, 2024). Since the late 1990s, it has become increasingly evident that inadequate pharmacokinetic properties of candidate molecules are a leading cause of clinical failures, prompting significant paradigm shifts within the pharmaceutical industry (Exscientia, 2024). Even today, poor ADME characteristics of small-molecule drug candidates

remain one of the primary reasons for development setbacks. However, recent years have witnessed considerable advancements in addressing these issues (Obrezanova, 2023).

During the drug discovery process, prior to first-in-human (FIH) studies, dose estimation and exposure assessment are typically carried out using *in vivo* animal models and *in vitro* systems derived from human sources. Physiologically Based Pharmacokinetic (PBPK) models simulate the time-dependent ADME behavior of drugs through mathematical formulations and contribute significantly to predicting biopharmaceutical profiles in humans, especially in advanced development stages. However, due to their high cost and limited scalability, PBPK models often fall short in enabling the rapid screening of large numbers of candidate compounds (Obrezanova, 2023).

One of the main challenges in PBPK modeling lies in the lack of compound-specific pharmacokinetic parameters for newly developed drugs. The absence of experimentally validated parameter data for most novel molecules limits the predictive accuracy of these models. To overcome this gap, several researchers have proposed integrated systems that combine simplified PBPK frameworks with ML algorithms to predict various PK parameters. This emerging approach is regarded as a multi-step process for the characterization and optimization of key ADME properties (Chou & Lin, 2023).

In the initial phase, databases containing *in vivo* time-concentration profiles and pharmacokinetic parameters are either constructed or derived from existing resources. Within this framework, the structural and physicochemical properties of selected drugs, along with *in vitro* ADME experimental data, must also be compiled.

In the second phase, ML/AI-based computational models are developed using these datasets, and the models are then employed to predict various ADME parameters—such as the tissue-plasma partition coefficient (K_p), clearance (Cl), and unbound fraction

(fu)—based on the structural and physicochemical attributes of the compounds. The predicted values are subsequently integrated into a general PBPK model.

In the third phase, time-concentration profiles in plasma and tissues are simulated using the integrated model, allowing for the estimation of key PK parameters such as the area under the curve (AUC) and maximum concentration (C_{max}). The resulting simulation data can either enhance existing databases or contribute to the development of new datasets for future modeling cycles (Chou & Lin, 2023).

In pharmacokinetic simulations specifically developed for animal models, a broad spectrum of methodologies has been employed—ranging from traditional ML techniques to advanced approaches such as deep learning. For example, Obrezanova and colleagues utilized a dataset comprising over 3,000 compounds to model time-concentration curves for nine different PK parameters following both intravenous and oral administration in rats. In their study, high-accuracy curve predictions were achieved for intravenous applications; however, only limited success was obtained for oral routes (Obrezanova, 2023).

There are certain structural challenges associated with modeling human pharmacokinetics using ML and AI techniques. One major limitation is the relatively small amount of available data for human PK modeling compared to preclinical animal data. While datasets for humans are typically limited to around 1,400 compounds, animal-based datasets encompass 5,000 to 60,000 compounds. Furthermore, due to the inherent constraints of clinical trials—such as high costs, long durations, and ethical or regulatory barriers—the expansion of human PK data occurs at a much slower pace.

In addition, compounds entering clinical trials are often selected based on “desirable” pharmacokinetic profiles, such as low clearance and high bioavailability. As a result, ML training datasets tend to underrepresent compounds with “suboptimal” properties, such as low bioavailability or high clearance. This under-

representation hinders the models' ability to learn the full spectrum of pharmacokinetic behavior and limits their generalizability.

Moreover, human data involve substantially greater complexity compared to controlled animal studies, due to inter-individual variability, disease states, dietary factors, and variations in measurement techniques. Consequently, a rigorous data curation process is essential when working with human-derived datasets.

In this context, Miljkovic et al. utilized a curated dataset of 1,000 clinical compounds to predict 12 different intravenous and oral human PK parameters. They employed random forest algorithms along with 2D chemical descriptors and selected *in vivo* rat PK parameters. The models developed for predicting oral C_{max}, AUC, and volume of distribution (V_d) demonstrated sufficient accuracy to be considered usable during the design stage and were validated using clinical data provided by AstraZeneca (Obrezanova, 2023).

Accurate and rapid prediction of pharmacokinetic parameters in the preclinical phase is a critical requirement for enhancing the efficiency of drug

discovery processes. In this regard, *in silico* ADMET modeling, which enables pharmacokinetic predictions based solely on chemical structures and fundamental molecular properties, has emerged as a valuable scientific tool. These approaches improve the effectiveness of project teams in designing and selecting molecules with favorable ADMET profiles and aim to reduce the number of compounds that need to be synthesized and experimentally tested by directing resources toward more promising candidates (Chou & Lin, 2023).

In recent years, the growing availability of experimental ADMET data, combined with advances in artificial intelligence algorithms, has led to the development of numerous AI-based tools for predicting ADMET properties. These tools significantly contribute to the efficient evaluation of candidate molecules in drug discovery and offer researchers substantial advantages in terms of both time and cost. Among these tools, five publicly available ADMET prediction platforms developed in the early 2020s have undergone extensive benchmarking and stand out as particularly notable. A summary of key information regarding these models is presented in Table 1 (Tran, Tayara, & Chong, 2023).

Table 1. AI-based ADMET prediction models and key features

Model	Latest Version	Number of Prediction	ML Algorithm	Website / Source
ADMETlab 3.0	3.0 (2024)	119	GNN	https://admetlab3.scbdd.com/
FP-ADMET	1.0 (2021)	≈50	RF	gitlab.com/vishsoft/fpadmet
AdmetSAR 3.0	3.0 (2023)	107	RF, SVM, k-NN	https://lmmd.ecust.edu.cn/admetSar3/index.php
Interpretable-ADMET	1.0 (2022)	59	GCNN, GAT	cadd.pharmacy.nankai.edu.cn/interpretableadmet
HelixADMET	1.0 (2022)	52	RF, GNN	paddlehelix.baidu.com/app/drug/admet/train

AI-assisted Pharmaceutical Formulation Development

Pharmaceutical formulation development is a multi-stage process aimed at designing safe, effective, and patient-compliant products by combining active pharmaceutical ingredients (APIs) with appropriate excipients (Afrin & Gupta, 2025). Traditional formu-

lation strategies predominantly rely on experimental trial-and-error approaches and often utilize statistical tools such as Design of Experiments (DoE) or Quality by Design (QbD) throughout the process. However, the multidimensional nature of formulation parameters frequently leads to prolonged timelines, high costs, and intensive labor requirements (Bao et al., 2023).

For instance, approximately 40% of newly developed compounds face solubility/bioavailability issues during early-stage screening. A large proportion of next-generation molecules exhibit poor aqueous solubility, which reduces bioavailability and increases the risk of formulation failure (Bhalani, Nutan, Kumar, & Singh Chandel, 2022). These challenges can result in reduced therapeutic efficacy or compromised product stability, further complicating the development process. Several reports emphasize that conventional formulation practices are labor-intensive, time-consuming, and costly due to their empirical basis (Afrin & Gupta, 2025). Consequently, there is a growing demand for more efficient, faster, and predictive methods in pharmaceutical formulation development.

In this context, AI, ML, and DL techniques have initiated a significant transformation within pharmaceutical technology in recent years. With a computational pharmaceuticals approach, drug formulations are being designed, optimized, and evaluated using big data and simulation techniques, enabling the extraction of meaningful relationships from experimental datasets. ML models allow for the prediction of drug–excipient interactions and compatibilities, while DL networks can forecast critical parameters such as drug release profiles. As a result, formulation design has moved beyond sole reliance on laboratory experimentation, evolving into a more systematic, rapid, and high-accuracy process (Dong, Wu, Xu, & Ouyang, 2023; Yang et al., 2019b).

The integration of AI into pharmaceutical formulation development has led to substantial paradigm shifts compared to conventional approaches. Primarily, AI-powered methods reduce both the number of experimental trials and overall development costs and timelines (Yang et al., 2019b). Leveraging historical datasets, predictive models can estimate critical properties such as solubility, bioavailability, and release kinetics in advance, thereby enabling more accurate formulation decisions. Furthermore, AI-based systems can uncover structure–property relationships that may be overlooked by traditional techniques,

thereby facilitating the identification of optimal formulation combinations and the development of innovative strategies (Challener, 2024).

Recent projects and platforms have demonstrated significant progress in AI-assisted pharmaceutical formulation development. For instance, the FormulationAI platform provides pre-trained AI models capable of predicting formulation properties across a variety of systems, from cyclodextrin complexes to liposomal delivery systems. Users can input information about active pharmaceutical ingredients and excipients to obtain predictions for multiple formulation characteristics (Dong et al., 2023).

Similarly, the DE-INTERACT project has developed an AI-based system capable of predicting drug–excipient compatibility with 98% accuracy by utilizing molecular descriptors and model ensemble techniques. Platforms like Intrepid Labs, which operate as “self-driving” laboratories, combine AI and robotics to automate experimental optimization workflows, significantly shortening formulation development timelines (Intrepid Labs, 2024). Additionally, the AI-driven autonomous lab established through the collaboration between Atinary Technologies and Takeda has markedly improved formulation efficiency (Atinary Technologies, 2023) (Hang et al., 2024).

Collectively, these advances illustrate how the integration of AI, ML, and DL technologies is transforming pharmaceutical formulation development into a more rapid, efficient, and intelligent process. As these technologies become more widely adopted, it is expected that the timelines for drug discovery and development will be significantly shortened, while simultaneously reducing research and development costs.

AI Applications in Clinical Trials

The development and commercialization of a new drug typically span a lengthy period of 10 to 15 years and require high costs ranging between approximately 1.5 to 2.0 billion USD. Nearly half of this time and

financial investment is consumed during drug discovery, optimization, and preclinical research stages, while the remaining 50% is allocated to clinical research processes. The high failure rate observed during clinical trials stands out as one of the major obstacles in drug development. Only one-third of compounds that reach Phase II proceed to Phase III, and more than one-third of those that reach Phase III fail during the approval process. Considering that Phase III studies alone account for nearly 60% of total R&D costs, a failed clinical trial can result in a financial loss ranging from 0.8 to 1.4 billion USD, representing a significant portion of R&D investment.

At this point, AI technologies emerge as promising tools to make clinical research processes more efficient, faster, and effective. ML and DL techniques analyze patterns within large and diverse data sources—such as electronic health records (EHR), medical literature, and clinical databases—to optimize patient-trial matching processes and enable a more reliable assessment of clinical endpoints (Harrer, Shah, Antony, & Hu, 2019).

In 2024, the Research Affairs Committee of the American College of Clinical Pharmacy (ACCP) conducted a comprehensive review of the role of AI in clinical pharmacy research and scientific publishing. Their evaluation highlighted how AI technologies can be effectively utilized at various stages of scientific research—including research question development, study design, data analysis, and results reporting (Chan et al., 2025).

Patient recruitment is one of the major challenges encountered in clinical trials. In Phase I clinical studies, approximately 80% of trials experience delays in patient enrollment. Traditional recruitment methods require healthcare professionals to manually review large volumes of medical records, a process that is time-consuming, costly, and prone to human error. Artificial intelligence offers significant advantages in this regard by leveraging electronic health records, social media platforms, and other digital data sources to

rapidly identify potential participants who meet specific study criteria. For instance, Deep 6 AI technology is capable of screening millions of patient records to swiftly detect eligible candidates for clinical trials (Wu et al., 2024b).

Wout Brusselaers, founder and CEO of Deep 6 AI, has emphasized that their system is not limited to traditional insurance records or structured EHR data. Instead, it harnesses the power of AI to mine deeper and less structured data sources, enabling the identification of target patient populations with high precision and speed.

Traditional data collection and evaluation methods often suffer from limitations such as low efficiency, limited capacity, susceptibility to error, and a lack of real-time monitoring. Artificial intelligence offers promising solutions to these challenges. ML techniques can manage and analyze large-scale clinical datasets swiftly and effectively, thereby contributing to the early identification of overlooked issues and hidden risks. Moreover, the analysis of real-time health data collected from wearable devices enables continuous and dynamic monitoring of participants' health status, providing researchers with accurate and up-to-date insights.

Artificial intelligence also plays a critical role in the clinical trial design phase. AI-powered predictive analytics can analyze historical clinical trial data to forecast potential outcomes, supporting researchers in making more informed decisions regarding optimal dosage selection, appropriate patient cohort identification, and the early detection of possible adverse events (Wu et al., 2024b).

In addition, AI holds great potential for enhancing operational efficiency and patient management in clinical research. AI applications in study design and patient stratification processes contribute to improved trial performance and support the development of targeted therapeutic strategies. Machine learning algorithms can analyze prior clinical data, electronic health records, and genomic information to facilitate

the creation of more efficient study protocols and accurate identification of suitable patient populations.

One study reported a 25% reduction in protocol amendments and a 15% increase in patient enroll-

ment rates through the application of ML techniques. These tangible improvements in clinical trial efficiency are summarized in Table 2 (Huang, Yang, Wen, Xia, & Yuan, 2024).

Table 2. The impact of artificial intelligence on clinical trial efficiency

Criterion	Traditional Approach	AI-Assisted Approach	Improvement Rate
Protocol Amendments	4.2 per study	3.15 per study	25% reduction
Patient Enrollment Rate	60%	69%	15% increase
Study Completion Time	3.5 years	2.8 years	20% reduction
Cost Reduction	-	-	18% reduction

Integration of artificial intelligence into the healthcare and pharmaceutical sectors: ethical and regulatory perspectives

The integration of AI technologies into the healthcare and pharmaceutical sectors has sparked a range of ethical debates. One of the key concerns pertains to the potential impact of these technologies on the labor market; for instance, it has been suggested that AI-driven systems may reduce employment opportunities by automating tasks traditionally carried out by healthcare professionals. In addition, questions surrounding the ownership of discoveries generated by machine learning algorithms—and whether such findings can be patented—represent unresolved ethical and legal challenges.

The degree to which physicians should rely on health profiles generated by AI systems, the extent to which these systems should influence diagnostic processes, and the level of trust patients place in such recommendations are all matters of growing concern. The absence of assurance regarding the impartiality and completeness of the data used to train these algorithms undermines trust in AI applications. Furthermore, there is a significant risk that existing social inequalities may be exacerbated by biased models trained on non-representative datasets. In particular, insufficient representation of disadvantaged patient groups in training data may lead to discriminatory practices, thereby violating the principle of equity in healthcare delivery.

The prediction of an individual’s predisposition to certain diseases based on pharmacogenomic data poses a significant ethical dilemma. While such predictive information can facilitate preventive medical interventions and offer potential health benefits, it may also exert adverse effects on an individual’s psychological well-being and emotional stability. Moreover, if health insurance providers were to gain access to such data, they might begin to demand future risk reports from individuals, potentially leading to new ethical and privacy-related concerns.

In response to these and similar ethical challenges, various international platforms have been developed to foster global discourse on the ethical implications of artificial intelligence. One notable initiative is the *Moral Machine*, which provides a simulation environment to explore ethical dilemmas encountered by autonomous vehicles. With input from over 40 million participants, the platform has allowed researchers to examine cross-cultural variations in moral decision-making. Findings have suggested that certain ethical principles may be shared across cultures, indicating the presence of some universal moral values.

Milena Pribić, a representative of IBM, has emphasized that insufficient prioritization of AI ethics by institutions may result in adverse long-term consequences. Therefore, it is imperative that AI systems be designed and implemented in alignment with fundamental ethical principles. In clinical practice, protocols such as CONSORT-AI and SPIRIT-AI have been

introduced to enhance transparency and accountability by ensuring that AI-driven interventions are reported in a comprehensive and standardized manner within the context of clinical trials (Arabi, 2021).

In a study conducted by Karimian, Petelos, and Evers in 2021 (Karimian, Petelos, & Evers, 2022), the ethical challenges potentially arising from the integration of artificial intelligence into the healthcare sector were systematically examined. The review was based on the Ethics Guidelines for Trustworthy AI, published by the European Commission, which outlines seven fundamental ethical principles: human agency and oversight, technical robustness and safety, privacy and data governance, transparency, diversity and non-discrimination, societal and environmental well-being, and accountability.

Among the most frequently discussed ethical issues in the reviewed literature were fairness, the preservation of human autonomy, explainability, and patient privacy. Notably, the principle of non-maleficence—"do no harm"—was found to be under-represented in the existing academic discourse. Furthermore, the study highlighted a substantial gap in the development of practical evaluation frameworks aimed at assessing the compliance of AI-based systems with these ethical principles. Only a limited number of publications offered actionable solutions for protecting patient privacy in healthcare settings, and empirical evidence supporting the operationalization of other ethical principles was reported to be scarce (Karimian et al., 2022).

In conclusion, the development and implementation of artificial intelligence technologies within ethical boundaries are of paramount importance to maintain a balance between technological progress and human values. AI solutions that operate in collaboration with human expertise and complement ethical gaps such as intuition and emotion have the potential to enhance societal benefit when properly managed. However, the continuous ethical evaluation and responsible governance of these technologies are essen-

tial for building a more equitable and trustworthy AI ecosystem (Arabi, 2021).

Pharmacogenomics, Bioinformatics, and Data Privacy

Pharmacogenomics (PGx), a cornerstone of personalized medicine, aims to maximize drug efficacy and minimize adverse drug reactions by integrating individuals' genetic characteristics into therapeutic decision-making processes.

The integration of AI with pharmacogenomic data introduces an innovative approach to precision medicine and facilitates the development of individualized treatment strategies. For instance, models utilizing ML algorithms can predict drug responses based on patients' genetic profiles, thereby enabling safer treatment options and optimized dosage regimens. In domains such as oncology—where therapeutic options may be limited—these models have contributed to the identification of novel genetic biomarkers associated with drug response through the analysis of transcriptomic data, thereby advancing targeted drug discovery processes.

However, the ethical, legal, and social implications associated with technologies developed using PGx data are receiving increasing attention. Among these concerns, genetic data privacy remains one of the most pressing issues. Such data not only reveal personal information but also encode familial biological relationships and ancestral origins. This inherent sensitivity renders full anonymization challenging and increases the risk of re-identification through cross-referencing even when data are presumed to be anonymized. Consequently, the use of genetic data for training AI models necessitates the highest standards of data security and a strong commitment to maintaining patient trust.

To address these concerns, the development of ethical guidelines for the use of AI in healthcare plays a pivotal role. For example, Char (2020) proposed a four-phase framework for assessing ethical challenges in AI-based healthcare applications. This framework

includes: (i) development and deployment of AI tools, (ii) establishment of regulatory and oversight mechanisms, (iii) ethical evaluation of AI-driven decisions, and (iv) consideration of broader ethical concerns. Such a comprehensive approach promotes the responsible and transparent implementation of AI systems in healthcare settings. Additionally, techniques such as federated learning—categorized under decentralized data processing strategies—help reduce the need for inter-institutional data sharing, thereby contributing to the protection of patient privacy. These methods allow AI models to be trained on decentralized datasets while preserving the confidentiality of sensitive health information.

In conclusion, the development of AI-based pharmacogenomic models must be grounded in ethical principles, ensuring the careful handling of genetic data and the protection of patient confidentiality.

Employing fair and representative datasets that reflect genetic diversity and actively addressing algorithmic biases are essential for building trustworthy and ethically sound AI systems. Such an approach not only enhances clinical efficacy but also reinforces public confidence in AI-driven healthcare technologies (Haga, 2024).

RESULTS AND DISCUSSION

This study has comprehensively demonstrated how AI technologies can be integrated across all stages of the drug-discovery process—from target identification and validation to lead optimization, from synthesis planning to ADME/Tox prediction and preclinical assessment. Beyond structure- and ligand-based virtual screening, the capacity of state-of-the-art models—particularly graph neural networks—to capture complex molecular relationships delivers clear gains in candidate selection, target fit, and optimization speed. In addition to these modeling advances, we outline data-representation choices, algorithmic architectures, and training strategies. Synthesis planning and synthetic-accessibility prediction enhance resource efficiency and sustainability,

while early computational ADME/Tox filters enable a more controlled and predictable triage of candidates with the potential to advance to the clinic. Industrial and academic case studies surveyed under Real-World Applications show that AI is not merely a theoretical promise but a transformative tool with practical impact. Nevertheless, realizing AI's full potential requires continuous improvements in data volume and quality, explainability, reproducibility, and governance (data security, ethics, standards). AI technologies are emerging not only as supportive tools in pharmaceutical research but as strategic components at the center of innovative drug-discovery paradigms. To this end, building large and balanced datasets; strengthening labeling and curation; defining performance metrics transparently and comparably; making model decision logic auditable; and corroborating results through independent validation are critical. Regulatory compliance and transparent reporting should be supported by end-to-end traceability, versioning, and audit trails.

At the preclinical stage, we recommend systematically aligning *in silico* predictions with *in vitro/in vivo* readouts and regularly monitoring out-of-model errors (dataset shift, bias). From a sustainability perspective, adopting low-energy, low-resource workflows; prioritizing synthesis routes consistent with green-chemistry principles; and establishing automation and active-learning cycles that improve the cost-time balance are valuable.

In summary, the assessments presented here indicate that AI is a complementary accelerator in pharmaceutical R&D; when robust data governance, explainability, reproducibility, and regulatory compliance are ensured, it reduces risk while increasing accuracy and efficiency. To further strengthen AI integration into pharmaceutical R&D, we recommend enhancing interdisciplinary collaboration, prioritizing explainable-AI approaches, and restructuring health-specific ethical and regulatory standards accordingly.

ACKNOWLEDGEMENTS

This work is not funded by any organization.

AUTHOR CONTRIBUTION STATEMENT

Conception (NGK), Literature research and writing (DN), Reviewing the text (NGK, DN), Supervision (NGK).

CONFLICT OF INTEREST

The authors declare that there is no conflict of interest.

REFERENCES

- Abiodun, O. I., Jantan, A., Omolara, A. E., Dada, K. V., Mohamed, N. A., & Arshad, H. (2018). State-of-the-art in artificial neural network applications: A survey. *Heliyon*, 4(11), e00938. doi:10.1016/j.heliyon.2018.e00938
- Afrin, S., & Gupta, V. (2025). Pharmaceutical formulation. In *StatPearls*. Treasure Island, FL: StatPearls Publishing. Retrieved from <https://www.ncbi.nlm.nih.gov/books/NBK562239/>
- Alzubaidi, L., Zhang, J., Humaidi, A. J., Al-Dujaili, A., Duan, Y., Al-Shamma, O.,...Farhan, L. (2021). Review of deep learning: Concepts, CNN architectures, challenges, applications, and future directions. *Journal of Big Data*, 8(1), 53. doi:10.1186/s40537-021-00444-8
- Andricopulo, A. D., & Ferreira, L. G. (2014). Medicinal chemistry approaches to neglected-diseases drug discovery. *Journal of Modern Medicinal Chemistry*, 2(1), 20–31. doi:10.12970/2308-8044.2014.02.01.4
- Arabi, A. A. (2021). Artificial Intelligence in Drug Design: Algorithms, Applications, Challenges and Ethics. *Future Drug Discovery*, 3(2). <https://doi.org/10.4155/fdd-2020-0028>
- Arifa, I., Aditsania, A., & Kurniawan, I. (2023). The implementation of genetic algorithm-ensemble learning on QSAR study of diacylglycerol acyltransferase-1 (DGAT1) inhibitors as anti-diabetes. In Y. B. Wah, M. W. Berry, A. Mohamed, & D. Al-Jumeily (Eds.), *Data science and emerging technologies (DaSET 2022), Lecture notes on data engineering and communications technologies* (Vol. 165, pp. 282–292). Singapore: Springer. doi:10.1007/978-981-99-0741-0_20
- Bao, Z., Bufton, J., Hickman, R. J., Aspuru-Guzik, A., Bannigan, P., & Allen, C. (2023). Revolutionizing drug formulation development: The increasing impact of machine learning. *Advanced Drug Delivery Reviews*, 202, 115108. doi:10.1016/j.addr.2023.115108
- Ben-Bright, B., Zhan, Y., Ghansah, B., Amankwah, R., Wornyo, D. K., & Ansah, E. (2017). Taxonomy and a theoretical model for feedforward neural networks. *International Journal of Computer Applications*, 162(4), 1–7. Retrieved from <https://www.ijcaonline.org/archives/volume162/number4/3274-2017913274>
- Berahmand, K., Nasiri, M., & Karimi, M. (2024). Autoencoders and their applications in machine learning: A survey. *Artificial Intelligence Review*, 57(2), 28. doi:10.1007/s10462-023-10662-6
- Bhalani, D. V., Nutan, B., Kumar, A., & Singh Chandel, A. K. (2022). Bioavailability enhancement techniques for poorly aqueous soluble drugs and therapeutics. *Biomedicines*, 10(9), 2055. doi:10.3390/biomedicines10092055
- Bleicher, L. S., Van Daelen, T., Honeycutt, J. D., Hassan, M., Chandrasekhar, J., Shirley, W., Tsui, V., & Schmitz, U. (2022). Enhanced utility of AI/ML methods during lead optimization by inclusion of 3D ligand information. *Frontiers in Drug Discovery*, 2, 1074797. doi:10.3389/fddis.2022.1074797

- Burki, T. (2020). A new paradigm for drug development. *The Lancet Digital Health*, 2(5), e220. doi:10.1016/s2589-7500(20)30088-1
- Challener, C. A. (2024). Using advanced algorithms to solve formulation challenges. *Pharmaceutical Technology*, 48(6), 14–17. Retrieved from <https://www.pharmtech.com/view/using-advanced-algorithms-to-solve-formulation-challenges>
- Chan, A., Baker, W. L., Abazia, D., Bauman, J., DeVane, C. L., Goodlet, K. J.,...Zhou, C. (2025). Impact of artificial intelligence on future clinical pharmacy research and scholarship. *Journal of the American College of Clinical Pharmacy*, 8(4), 311–316. doi:10.1002/jac5.70003
- Chen, X., Xie, W., Yang, Y., Hua, Y., Xing, G., Liang, L.,...Zhang, Y. (2020). Discovery of dual FGFR4 and EGFR inhibitors by machine learning and biological evaluation. *Journal of Chemical Information and Modeling*, 60(10), 4640–4652. doi:10.1021/acs.jcim.0c00652
- Chou, W. C., & Lin, Z. (2023). Machine learning and artificial intelligence in physiologically based pharmacokinetic modeling. *Toxicological Sciences*, 191(1), 1–14. doi:10.1093/toxsci/kfac101
- Creswell, A., White, T., Dumoulin, V., Arulkumaran, K., Sengupta, B., & Bharath, A. A. (2018). Generative adversarial networks: An overview. *IEEE Signal Processing Magazine*, 35(1), 53–65. doi:10.1109/MSP.2017.2765202
- Doytchinova, I. (2022). Drug design—past, present, future. *Molecules*, 27(5), 1492. <https://doi.org/10.3390/molecules27051496>
- Dong, J., Wu, Z., Xu, H., & Ouyang, D. (2024). FormulationAI: A novel web-based platform for drug formulation design driven by artificial intelligence. *Briefings in Bioinformatics*, 25(1), bbad419. doi:10.1093/bib/bbad419
- Exscientia. (2024). Recursion and Exscientia, two leaders in the AI drug discovery space, have officially combined to advance the industrialization of drug discovery. Retrieved April 26, 2025, from <https://investors.exscientia.ai/press-releases/press-release-details/2024/Recursion-and-Exscientia-two-leaders-in-the-AI-drug-discovery-space-have-officially-combined-to-advance-the-industrialization-of-drug-discovery/default.aspx>
- Gangwal, A., Ansari, A., Ahmad, I., Azad, A. K., Kumarasamy, V., Subramaniyan, V., & Wong, L. S. (2024). Generative artificial intelligence in drug discovery: Basic framework, recent advances, challenges, and opportunities. *Frontiers in Pharmacology*, 15, 1331062. doi:10.3389/fphar.2024.1331062
- Ghislat, G., Rahman, T., & Ballester, P. J. (2021). Recent progress on the prospective application of machine learning to structure-based virtual screening. *Current Opinion in Chemical Biology*, 65, 28–34. doi:10.1016/j.cbpa.2021.04.009
- Greene, D., Cunningham, P., & Mayer, R. (2008). Unsupervised learning and clustering. In M. Kantardzic (Ed.), *Machine learning techniques for multimedia: Case studies on organization and retrieval* (pp. 51–90). Berlin, Germany: Springer.
- Gómez-Bombarelli, R., Wei, J. N., Duvenaud, D., Hernández-Lobato, J. M., Sánchez-Lengeling, B., Sheberla, D.,...Aspuru-Guzik, A. (2018). Automatic chemical design using a data-driven continuous representation of molecules. *ACS Central Science*, 4, 268–276. doi:10.1021/acscentsci.7b00572
- Gui, J., Sun, Z., Wen, Y., Tao, D., & Ye, J. (2023). A review on generative adversarial networks: Algorithms, theory, and applications. *IEEE Transactions on Knowledge and Data Engineering*, 35(4), 3313–3332. doi:10.1109/TKDE.2021.3076029

- Guo, Y. M., Liu, Y., Oerlemans, A., Lao, S., Wu, S., & Lew, M. S. (2016). Deep learning for visual understanding: A review. *Neurocomputing*, 187, 27–48. doi:10.1016/j.neucom.2015.09.116
- Gupta, R., Srivastava, D., Sahu, M., Tiwari, S., Ambasta, R. K., & Kumar, P. (2021). Artificial intelligence to deep learning: Machine intelligence approach for drug discovery. *Molecular Diversity*, 25(3), 1315–1360. doi:10.1007/s11030-021-10217-3
- Haga, S. B. (2024). Artificial intelligence, medications, pharmacogenomics, and ethics. *Pharmacogenomics*, 25(14–15), 611–622. <https://doi.org/10.1080/14622416.2024.2428587>
- Hang, N. T., Long, N. T., Duy, N. D., Chien, N. N., & Phuong, N. V. (2024). Towards safer and efficient formulations: Machine learning approaches to predict drug–excipient compatibility. *International Journal of Pharmaceutics*, 653, 123884. doi:10.1016/j.ijpharm.2024.123884
- Harrer, S., Shah, P., Antony, B., & Hu, J. (2019). Artificial intelligence for clinical trial design. *Trends in Pharmacological Sciences*, 40(8), 577–591. doi:10.1016/j.tips.2019.05.005
- Healx. (2024). Healx announces first patient dosed in Phase 2 trial evaluating HLX-1502 for the treatment of neurofibromatosis type 1. Retrieved April 26, 2025, from <https://healx.ai/healx-announces-first-patient-dosed-in-phase-2-trial-evaluating-hlx-1502-for-the-treatment-of-neurofibromatosis-type-1/>
- Huang, D., Yang, M., Wen, X., Xia, S., & Yuan, B. (2024). AI-driven drug discovery: Accelerating the development of novel therapeutics in biopharmaceutics. *Journal of Knowledge Learning and Science Technology*, 3(3), 206–224. doi:10.60087/jklst.vol3.n3.p.206-224
- Insilico Medicine. (2022). Insilico announces first-in-human trial of AI-discovered and AI-designed drug for idiopathic pulmonary fibrosis. Retrieved April 26, 2025, from <https://insilico.com/phase1>
- Insilico Medicine. (n.d.). Pipeline target: PHD1/2. Retrieved April 26, 2025, from https://insilico.com/pipeline_target_phd
- Jamal, S., Grover, A., & Grover, S. (2019). Machine learning from molecular dynamics trajectories to predict Caspase-8 inhibitors against Alzheimer’s disease. *Frontiers in Pharmacology*, 10, 780. <https://doi.org/10.3389/fphar.2019.00780>
- Janiesch, C., Zschech, P., & Heinrich, K. (2021). Machine learning and deep learning. *Electronic Markets*, 31(3), 685–695. doi:10.1007/s12525-021-00475-2
- Jiang, T., Gradus, J. L., & Rosellini, A. J. (2020). Supervised machine learning: A brief primer. *Behavior Therapy*, 51(5), 675–687. doi:10.1016/j.beth.2020.05.002
- Jiang, Y., Yu, Y., Kong, M., Mei, Y., & Luo, T. (2023). Artificial intelligence for retrosynthesis prediction. *Engineering*, 25, 32–50. doi:10.1016/j.eng.2022.04.021
- Jumper, J., Evans, R., Pritzel, A., Green, T., Figurnov, M., Ronneberger, O.,...Hassabis, D. (2021). Highly accurate protein structure prediction with AlphaFold. *Nature*, 596(7873), 583–589. doi:10.1038/s41586-021-03819-2
- Karimian, G., Petelos, E., & Evers, S. M. (2022). The ethical issues of the application of artificial intelligence in healthcare: A systematic scoping review. *AI and Ethics*, 2(4), 539–551. <http://dx.doi.org/10.1007/s43681-021-00131-7>

- Katsila, T., Spyrou, S., Patrinos, G. P., & Matsoukas, M. T. (2016). Computational approaches in target identification and drug discovery. *Computational and Structural Biotechnology Journal*, 14, 177–184. doi:10.1016/j.csbj.2016.04.004
- Kaul, V., Enslin, S., & Gross, S. A. (2020). History of artificial intelligence in medicine. *Gastrointestinal Endoscopy*, 92(4), 807–812. doi:10.1016/j.gie.2020.06.040
- Khemani, B., Patil, S., Kotecha, K., & Tanwar, S. (2024). A review of graph neural networks: Concepts, architectures, techniques, challenges, datasets, applications, and future directions. *Journal of Big Data*, 11, 18. doi:10.1186/s40537-023-00876-4
- Kirkpatrick, P. (2022). Artificial intelligence makes a splash in small-molecule drug discovery. *Nature Biotechnology*, 40(7), 937. <https://doi.org/10.1038/d43747-022-00104-7>.
- Kumar, A., Tejaswini, P., Nayak, O., Kujur, A. D., Gupta, R., Rajanand, A., & Sahu, M. (2022, May). A survey on IBM Watson and its services. *Journal of Physics: Conference Series*, 2273(1), 012022. doi:10.1088/1742-6596/2273/1/012022
- Kurita, T. (2019). Principal component analysis (PCA). In K. Ikeuchi (Ed.), *Computer vision: A reference guide* (pp. 1–4). Cham, Switzerland: Springer.
- Li, B., Tan, K., Lao, A. R., Wang, H., Zheng, H., & Zhang, L. (2024). A comprehensive review of artificial intelligence for pharmacology research. *Frontiers in Genetics*, 15, 1450529. doi:10.3389/fgene.2024.1450529
- Lipton, Z. C., Berkowitz, J., & Elkan, C. (2015). A critical review of recurrent neural networks for sequence learning. *arXiv preprint arXiv:1506.00019*. Retrieved from <https://arxiv.org/abs/1506.00019>
- Loos, N. H. C., Sparidans, R. W., Heydari, P., Bui, V., Lebre, M. C., Beijnen, J. H., & Schinkel, A. H. (2024). The ABCB1 and ABCG2 efflux transporters limit brain disposition of the SYK inhibitors entospletinib and lanraplenib. *Toxicology and Applied Pharmacology*, 485, 116911. doi:10.1016/j.taap.2024.116911
- Mahesh, B. (2020). Machine learning algorithms: A review. *International Journal of Science and Research (IJSR)*, 9(1), 381–386. doi:10.21275/ART20203995
- McCulloch, W. S., & Pitts, W. (1943). A logical calculus of the ideas immanent in nervous activity. *The Bulletin of Mathematical Biophysics*, 5(4), 115–133. doi:10.1007/BF02478259
- Mienye, I. D., Swart, T. G., & Obaido, G. (2024). Recurrent neural networks: A comprehensive review of architectures, variants, and applications. *Information*, 15(9), 395. doi:10.3390/info15090517
- Montesinos López, O. A., Montesinos López, A., & Crossa, J. (2022). Fundamentals of artificial neural networks and deep learning. In O. A. Montesinos López, A. Montesinos López, & J. Crossa (Eds.), *Multivariate statistical machine learning methods for genomic prediction* (pp. 379–425). Cham, Switzerland: Springer. doi:10.1007/978-3-030-89010-0_10
- Naeem, S., Ali, A., Anam, S., & Ahmed, M. M. (2023). An unsupervised machine learning algorithms: Comprehensive review. *International Journal of Computing and Digital Systems*, 13(1), 911–921. doi:10.12785/ijcds/130172
- Nasteski, V. (2017). An overview of the supervised machine learning methods. *Horizons International Scientific Journal. Series B*, 21(5), 51–62. doi:10.20544/HORIZONS.B.04.1.17.P05

- Navidan, H., Dehghantanha, A., & Parizi, R. M. (2021). Generative adversarial networks (GANs) in networking: A comprehensive survey and evaluation. *Computer Networks*, 194, 108149. doi:10.1016/j.comnet.2021.108149
- Niazi, S. K. (2023). The coming of age of AI/ML in drug discovery, development, clinical testing, and manufacturing: The FDA perspectives. *Drug Design, Development and Therapy*, 17, 2691–2725. doi:10.2147/DDDT.S424991
- Niazi, S. K., Zamara, M., & Magoola, M. (2024). Optimizing lead compounds: The role of artificial intelligence in drug discovery. *Preprints*. doi:10.20944/preprints202404.0055.v1
- Obrezanova, O. (2023). Artificial intelligence for compound pharmacokinetics prediction. *Current Opinion in Structural Biology*, 79, 102546. <https://doi.org/10.1016/j.sbi.2023.102546>
- Oprea, T. I., Bologa, C. G., Brunak, S., Campbell, A., Gan, G. N., Gaulton, A.,...Zahoránszky-Köhalmi, G. (2018). Unexplored therapeutic opportunities in the human genome. *Nature Reviews Drug Discovery*, 17(5), 317–332. <https://doi.org/10.1038/nrd.2018.14>
- Oussidi, A., & Elhassouny, A. (2018, April). Deep generative models: Survey. In *2018 International Conference on Intelligent Systems and Computer Vision (ISCV)* (pp. 1–8). Piscataway, NJ: IEEE. doi:10.1109/ISACV.2018.8354080
- Park, J., Yi, D., & Ji, S. (2020). Analysis of recurrent neural network and predictions. *Symmetry*, 12(4), 615. <https://doi.org/10.3390/sym12040615>
- Parvatikar, P. P., Patil, S., Khaparkhantkar, K., Patil, S., Singh, P. K., Sahana, R.,...Raghu, A. V. (2023). Artificial intelligence: Machine learning approach for screening large database and drug discovery. *Antiviral Research*, 220, 105740. doi:10.1016/j.antiviral.2023.105740
- Pascanu, R., Mikolov, T., & Bengio, Y. (2013). On the difficulty of training recurrent neural networks. In *Proceedings of the 30th International Conference on Machine Learning* (pp. 1310–1318). Atlanta, GA: PMLR.
- Rai, R. B., Sahu, S., & Sawant, S. (2018). An analysis on MYCIN expert system. *International Journal of Research in Engineering, Science and Management*, 1(1), 1–4. Retrieved from <https://www.ijresm.com>
- Ryan, D. K., Maclean, R. H., Balston, A., Scourfield, A., Shah, A. D., & Ross, J. (2024). Artificial intelligence and machine learning for clinical pharmacology. *British Journal of Clinical Pharmacology*, 90(3), 629–639. doi:10.1111/bcp.15930
- Salgin-Goksen, U., Telli, G., Erikci, A., Dedecengiz, E., Tel, B. C., Kaynak, F. B.,...Gokhan-Kelekci, N. (2021). New 2-pyrazoline and hydrazone derivatives as potent and selective monoamine oxidase A inhibitors. *Journal of Medicinal Chemistry*, 64(4), 1989–2009. doi: 10.1021/acs.jmedchem.0c01504.
- Sengar, S. S., Hasan, A. B., Kumar, S., & Carroll, F. (2025). Generative artificial intelligence: A systematic review and applications. *Multimedia Tools and Applications*, 84, 23661–23700. doi:10.1007/s11042-024-20016-1
- Sharma, A., Singh, S., & Ratna, S. (2024). Graph neural network operators: A review. *Multimedia Tools and Applications*, 83(8), 23413–23436. <https://doi.org/10.1007/s11042-023-16440-4>.
- Sheils, T. K., Mathias, S. L., Kelleher, K. J., Siramshetty, V. B., Nguyen, D.-T., Bologa, C. G.,...Oprea, T. I. (2021). TCRD and Pharos 2021: Mining the human proteome for disease biology. *Nucleic Acids Research*, 49(D1), D1334–D1346. doi:10.1093/nar/gkaa993

- Singh, J., & Banerjee, R. (2019, March). A study on single and multi-layer perceptron neural network. In *Proceedings of the 3rd International Conference on Computing Methodologies and Communication (ICCMC)* (pp. 35–40). Piscataway, NJ: IEEE. doi:10.1109/ICCMC.2019.8819780
- Tran, T. T. V., Tayara, H., & Chong, K. T. (2023). Recent studies of artificial intelligence on in silico drug distribution prediction. *International Journal of Molecular Sciences*, 24(3), 2112. <https://doi.org/10.3390/ijms24031815>
- Tuncel, S. T., Gunal, S. E., Demir, I., Baysal, I., Erdem, S. S., Telli, G.,...Kelekci, N. G. (2025). Multitarget-directed carbamate and carbamothioate derivatives: Cholinesterase and monoamine oxidase inhibition, anti- β -amyloid (A β) aggregation, antioxidant and blood–brain barrier permeation properties against Alzheimer’s disease. *European Journal of Medicinal Chemistry*, 272, 117986. <https://doi.org/10.1016/j.ejmech.2025.117986>
- Turing, A. M. (1950). Computing machinery and intelligence. *Mind*, 59(236), 433–460. doi:10.1093/mind/LIX.236.433
- Van Engelen, J. E., & Hoos, H. H. (2020). A survey on semi-supervised learning. *Machine Learning*, 109(2), 373–440. <https://doi.org/10.1007/s10994-019-05855-6>
- Vijayan, R. S. K., Kihlberg, J., Cross, J. B., & Poon-gavanam, V. (2022). Enhancing preclinical drug discovery with artificial intelligence. *Drug Discovery Today*, 27(4), 967–985. doi:10.1016/j.dru-dis.2021.11.023
- Wallach, I., Dzamba, M., & Heifets, A. (2015). Atom-Net: A deep convolutional neural network for bioactivity prediction in structure-based drug discovery. *arXiv preprint arXiv:1510.02855*. Retrieved from <https://arxiv.org/abs/1510.02855>
- Wang, Y., Pang, C., Wang, Y., Jin, J., Zhang, J., Zeng, X.,...Wei, L. (2023). Retrosynthesis prediction with an interpretable deep-learning framework based on molecular assembly tasks. *Nature Communications*, 14(1), 6155. doi:10.1038/s41467-023-41698-5
- Wang, Z., She, Q., & Ward, T. E. (2021). Generative adversarial networks in computer vision: A survey and taxonomy. *ACM Computing Surveys*, 54(2), 1–38. <https://doi.org/10.1145/3439723>
- Webb, G. I., Keogh, E., & Miikkulainen, R. (2010). Naïve Bayes. In C. Sammut & G. I. Webb (Eds.), *Encyclopedia of machine learning* (pp. 713–714). Boston, MA: Springer. doi:10.1007/978-0-387-30164-8_576
- Wenteler, A., Cabrera, C. P., Wei, W., Neduva, V., & Barnes, M. R. (2024). AI approaches for the discovery and validation of drug targets. *Precision Clinical Medicine*, 7(2), e7. doi:10.1017/pcm.2024.4
- Wildey, M. J., Haunso, A., Tudor, M., Webb, M., & Connick, J. H. (2017). High-throughput screening. In M. L. Neidle (Ed.), *Annual Reports in Medicinal Chemistry* (Vol. 50, pp. 149–195). Amsterdam: Elsevier. doi:10.1016/bs.armc.2017.08.004
- Wu, H., Liu, J., Zhang, R., Lu, Y., Cui, G., Cui, Z., & Ding, Y. (2024a). A review of deep learning methods for ligand-based drug virtual screening. *Fundamental Research*, 4(4), 715–737. doi:10.1016/j.fmre.2024.02.011
- Wu, Y., Li, H., Zhang, X., Chen, J., Zhao, L., & Wang, Q. (2024b). The role of artificial intelligence in drug screening, drug design, and clinical trials. *Frontiers in Pharmacology*, 15, 1459954. doi:10.3389/fphar.2024.1459954
- Wu, Z., Pan, S., Chen, F., Long, G., Zhang, C., & Yu, P. S. (2021). A comprehensive survey on graph neural networks. *IEEE Transactions on Neural Networks and Learning Systems*, 32(1), 4–24. doi:10.1109/TNNLS.2020.2978386

- Xu, Y., Zhou, Y., Sekula, P., & Ding, L. (2021). Machine learning in construction: From shallow to deep learning. *Developments in the Built Environment*, 6, 100045. doi:10.1016/j.dibe.2021.100045
- Yang, X., Wang, Y., Byrne, R., Schneider, G., & Yang, S. (2019a). Concepts of artificial intelligence for computer-assisted drug discovery. *Chemical Reviews*, 119(18), 10520–10594. doi:10.1021/acs.chemrev.8b00728
- Yang, Y., Ye, Z., Su, Y., Zhao, Q., Li, X., & Ouyang, D. (2019b). Deep learning for *in vitro* prediction of pharmaceutical formulations. *Acta Pharmaceutica Sinica B*, 9(1), 177–185. doi:10.1016/j.apsb.2018.09.010
- Zhao, X., Wang, L., Zhang, Y., Han, X., Deveci, M., & Parmar, M. (2024). A review of convolutional neural networks in computer vision. *Artificial Intelligence Review*, 57(4), 99. doi:10.1007/s10462-024-10721-6
- Zhavoronkov, A., Ivanenkov, Y. A., Aliper, A., Veselov, M. S., Aladinskiy, V. A., Aladinskaya, A.,... Aspuru-Guzik, A. (2019). Deep learning enables rapid identification of potent DDR1 kinase inhibitors. *Nature Biotechnology*, 37(9), 1038–1040. doi:10.1038/s41587-019-0224-x
- Zhi, H. Y., Zhao, L., Lee, C. C., & Chen, C. Y. C. (2021). A novel graph neural network methodology to investigate dihydroorotate dehydrogenase inhibitors in small cell lung cancer. *Biomolecules*, 11(3), 380. doi:10.3390/biom11030477