

Analysis of Bioinformatics Software Applied in Computer-aided Drug Design

Mohd. Wasiullah¹, Piyush Yadav², Divyansh Singh³, Satish Kumar Yadav^{4,*},
Vinay Kumar Deepak⁵

Abstract

The integration of bioinformatic tools with computational methods has revolutionized the field of Computer-aided Drug Design (CADD), enabling researchers to expedite the discovery and optimization of new therapeutics. This review provides an in-depth analysis of the bioinformatic tools utilized in CADD, encompassing molecular docking, molecular dynamics simulation, virtual screening, homology modelling, and molecular visualization. We go over the tenets, approaches, and uses of these instruments, emphasizing their value in expediting the drug discovery process and tackling challenging biological issues. The convergence of bioinformatics and computational methods has ushered in a new era in drug discovery known as Computer-aided Drug Design (CADD). This review explores the arsenal of bioinformatic tools utilized in CADD, encompassing molecular docking, molecular dynamics simulation, virtual screening, homology modelling, and molecular visualization. By elucidating the principles, methodologies, and applications of these tools, this review aims to provide a comprehensive understanding of their role in accelerating the drug discovery process. From target identification to lead optimization, bioinformatic tools play a pivotal role in expediting the identification and optimization of novel therapeutics, ultimately contributing to the advancement of human health.

Keywords: Computer-aided drug design (CADD), bioinformatic tools, molecular docking, molecular dynamics simulation, virtual screening, homology modelling, molecular visualization, drug discovery, target identification, lead optimization

INTRODUCTION TO COMPUTER-AIDED DRUG DESIGN (CADD)

The multidisciplinary discipline of computer-aided drug design (CADD) uses bioinformatics tools, computational approaches, and molecular modelling techniques to speed up the discovery and development of new medications [1]. It encompasses a wide range of computational techniques and algorithms aimed at understanding molecular interactions, predicting ligand binding modes, and optimizing the pharmacological properties of candidate compounds. CADD is essential to the entire drug discovery process, from lead optimization and preclinical testing to target identification and validation [2–4]. The primary objective of CADD is to leverage computational approaches to streamline the drug discovery process, reduce experimental costs, and accelerate the development of safe and effective therapeutics [5]. By harnessing the power of computational algorithms and high-performance computing resources, researchers can explore vast chemical space, predict the bioactivity of small molecules,

*Author for Correspondence

Satish Kumar Yadav
E-mail: slk.pharma@gmail.com

¹Principal and Professor, Department of Pharmacy, Prasad Institute of Technology, Jaunpur, Uttar, Pradesh, India

²Professor and Academic Head, Department of Pharmacy, Prasad Institute of Technology, Jaunpur, Uttar, Pradesh, India

³Scholar, Department of Pharmacy, Prasad Institute of Technology, Jaunpur, Uttar, Pradesh, India

⁴Associate Professor, Department of Pharmacy, Prasad Institute of Technology, Jaunpur, Uttar, Pradesh, India

⁵Professor, Department of Pharmacy, Prasad Institute of Technology, Jaunpur, Uttar, Pradesh, India

Received Date: May 10, 2024

Accepted Date: May 14, 2024

Published Date: June 21, 2024

Citation: Mohd. Wasiullah, Piyush Yadav, Divyansh Singh, Satish Kumar Yadav, Vinay Kumar Deepak. Analysis of Bioinformatics Software Applied in Computer-aided Drug Design. Research & Reviews: A Journal of Drug Design & Discovery. 2024; 11(2): 1–10p.

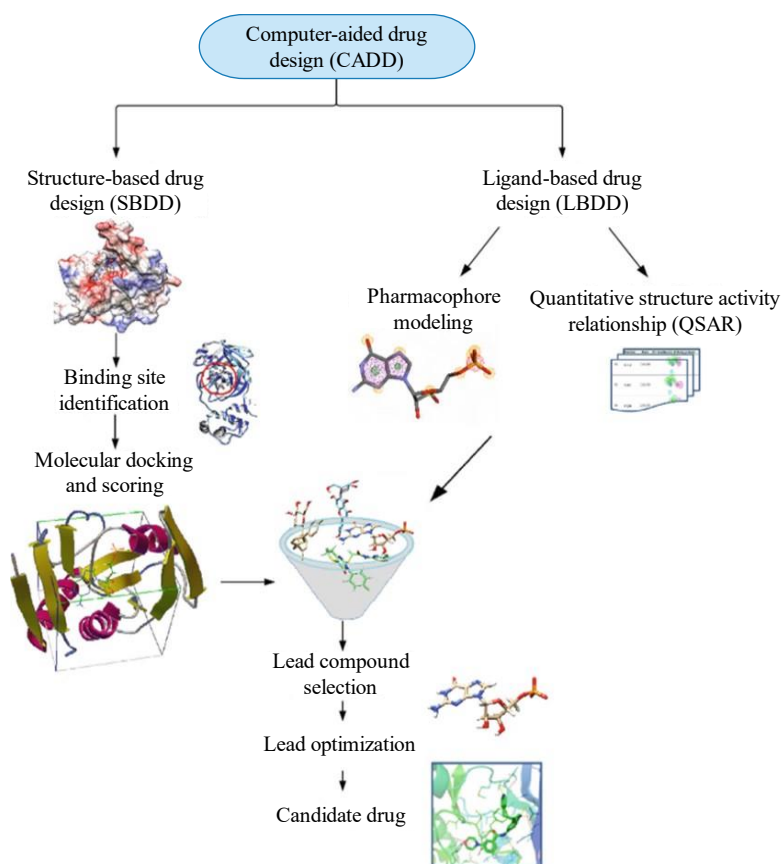


Figure 1. Computer aided drug design process.

and optimize lead compounds with improved potency, selectivity, and pharmacokinetic properties. CADD techniques are used to rationalize structure-activity relationships, identify potential drug targets, and guide the design of novel therapeutics for various diseases and conditions [3]. In recent years, advances in computational chemistry, machine learning, and structural biology have revolutionized the field of CADD, enabling researchers to tackle complex drug discovery challenges with unprecedented precision and efficiency [6]. With the exponential growth of biological and chemical data, as well as the increasing computational power and algorithmic sophistication, CADD continues to evolve as a critical component of modern drug discovery pipelines, offering new opportunities for innovation and discovery in pharmaceutical research. In order to analyze, understand, and manage biological data, the multidisciplinary area of bioinformatics combines computer science, mathematics, statistics, and biology [7, 8]. It includes a broad spectrum of software tools and computational methods for storing, retrieving, and analyzing biological data, such as gene expression patterns, protein structures, DNA sequences, and metabolic pathways [9–12]. Bioinformatics plays a crucial role in various areas of biological research, including genomics, proteomics, structural biology, and systems biology [13]. The primary goal of bioinformatics is to extract meaningful insights from large-scale biological data, elucidate the structure and function of biomolecules, and understand the underlying mechanisms of biological processes [12]. Bioinformatics tools and algorithms are used to analyze DNA sequences, predict protein structures, annotate genes and genomes, and identify genetic variations associated with diseases and traits. By integrating computational approaches with experimental techniques, bioinformatics enables researchers to address fundamental questions in biology and medicine, ranging from understanding the genetic basis of diseases to designing personalized therapies (Figure 1) [14].

Overview of CADD

The multidisciplinary field of computer-aided drug design (CADD) lies at the interface of computational science, biology, and chemistry [3]. It speeds up the search for and development of new

medications by utilizing bioinformatics tools and computational techniques [7]. To anticipate how tiny compounds would interact with target proteins, CADD includes a broad range of methods such as molecular docking, molecular dynamics modelling, virtual screening, and structure-based drug design [1].

Computational Methods' Significance in Drug Discovery

Conventional drug discovery techniques entail expensive and time-consuming experimental procedures [3]. Computational approaches in CADD offer a cost-effective and efficient alternative by enabling researchers to screen large compound libraries, predict ligand binding modes, and optimize lead compounds *in silico* before experimental validation [4]. These computational methods complement experimental techniques, allowing researchers to prioritize candidate compounds and focus resources on the most promising drug candidates [2]. Bioinformatic tools play a crucial role in CADD by providing the computational infrastructure and analytical capabilities needed to process, analyze, and interpret biological data [8]. These tools enable researchers to extract meaningful insights from large-scale datasets, such as protein structures, ligand databases, and genomic sequences [4]. By integrating bioinformatic tools with computational methods, researchers can gain a deeper understanding of the molecular basis of diseases and identify novel therapeutic targets for drug discovery [9].

MOLECULAR DOCKING

Molecular docking serves as a cornerstone of structure-based drug design, facilitating the rational exploration of ligand-receptor interactions and the identification of potential drug candidates [3]. This subsection provides an overview of the principles underlying molecular docking, including scoring functions, search algorithms, and binding site characterization [7]. Additionally, it emphasizes how important molecular docking is to speeding up lead optimization and identification in the drug discovery process [11].

Principles and Methods of Molecular Docking

A computer method called molecular docking is used to forecast how small compounds, or ligands, will connect to target proteins, or receptors, as well as their affinity [15]. The process involves the flexible docking of ligands into the binding site of a protein, followed by scoring and ranking of the docked poses based on their binding energies [3]. To search the conformational space of ligands and find the best binding orientations, molecular docking systems use a variety of search techniques and scoring functions [7].

Popular Docking Software

Several software tools are available for molecular docking, each offering unique features and capabilities [9]. Auto Dock is a widely used docking program that employs a Lamarckian genetic algorithm to explore ligand conformational space and optimize ligand-protein interactions [10]. Auto Dock Vina is a successor to Auto Dock, offering improved performance and efficiency in docking calculations [16]. Glide, developed by Schrödinger, is a commercial docking program that utilizes a hierarchical search algorithm and advanced scoring functions for accurate ligand docking and scoring [13].

Applications in Drug Design

Molecular docking has numerous applications in drug discovery, including lead identification, lead optimization, and off-target prediction [17]. In lead identification, docking is used to screen compound libraries and identify potential ligands with high affinity for the target protein [14]. In lead optimization, docking helps optimize the chemical structure of lead compounds to improve their binding affinity and selectivity. Additionally, docking can be used to predict the off-target interactions of lead compounds, thereby minimizing the risk of adverse effects and improving drug safety profiles [4].

MOLECULAR DYNAMICS SIMULATION

Molecular dynamics simulation offers valuable insights into the dynamic behavior and conformational changes of biological macromolecules, including proteins, nucleic acids, and membranes [7]. This

subsection introduces the fundamental principles of molecular dynamics simulation, including Newton's equations of motion, force fields, and integration algorithms [13]. It underscores the importance of molecular dynamics in elucidating protein-ligand interactions, understanding drug binding kinetics, and predicting protein flexibility [16].

Basics of Molecular Dynamics

A computational method for studying the motions and interactions of atoms and molecules over time is called molecular dynamics (MD) simulation [17]. It involves the numerical integration of Newton's equations of motion to simulate the behavior of a molecular system under specified conditions. Protein folding, ligand binding, and conformational changes are only a few examples of the dynamic behavior of biomolecular systems that can be understood by MD simulations [14].

Software for Molecular Dynamics: GROMACS, AMBER

Several software packages are available for performing molecular dynamics simulations, each offering specialized features and capabilities [16]. One popular MD simulation program that is well-known for its great performance and scalability is called GROMACS (Groningen Machine for Chemical Simulations) [18]. AMBER (Assisted Model Building with Energy Refinement) is another popular MD simulation package that offers a comprehensive suite of tools for simulating biomolecular systems with accuracy and efficiency [19].

Role in Drug Design and Optimization

Molecular dynamics simulations are essential for understanding the dynamic behavior of biomolecular systems, which is important for drug design and optimization [20]. By simulating the interactions between ligands and target proteins at the atomic level, MD simulations can elucidate the mechanisms of ligand binding, identify key protein-ligand interactions, and predict ligand binding modes with high accuracy [17]. Additionally, MD simulations can reveal conformational changes in target proteins induced by ligand binding, aiding in the design of allosteric modulators and structure-based drug design strategies [15].

STRUCTURE-BASED VIRTUAL SCREENING

Structure-based virtual screening leverages the three-dimensional structure of target proteins to identify small molecule ligands with high binding affinity and selectivity [14]. This subsection provides an overview of the rationale behind structure-based virtual screening, emphasizing its role in accelerating lead identification and optimization [2]. It also discusses the advantages of structure-based approaches over ligand-based methods and highlights key challenges in virtual screening [1].

Techniques and Algorithms

A computational method called “structure-based virtual screening” is employed to find putative ligands that bind to a target protein on the basis of their structural complementarity [19]. This methodology incorporates multiple techniques, such as quantitative structure-activity relationship (QSAR) analysis, pharmacophore modelling, and molecular docking [20]. Molecular docking involves the computational docking of ligands into the binding site of a target protein, followed by scoring and ranking of the docked poses based on their binding energies [14]. Pharmacophore modeling identifies common chemical features (e.g., hydrogen bond donors, acceptors, hydrophobic regions) essential for ligand binding and uses them to screen compound libraries for potential ligands [21]. The prediction of ligand potency and selectivity is made possible by QSAR analysis, which creates quantitative models linking the chemical structure of ligands with their biological function [22].

Tools for Structure-based Virtual Screening

Several software tools are available for performing structure-based virtual screening, each offering unique features and capabilities [13]. DOCK is a widely used docking program that employs a geometric matching algorithm to identify potential ligands that complement the binding site of a target

protein [20]. The Schrödinger Suite, developed by Schrödinger, offers a comprehensive suite of tools for molecular modeling, including Glide for molecular docking, Phase for pharmacophore modeling, and Qik Prop for QSAR analysis [19].

Examples of Successful Applications

Structure-based virtual screening has been successfully applied in various drug discovery projects to identify lead compounds, optimize their chemical structures, and analyze structure-activity relationships [14]. For example, virtual screening campaigns have led to the discovery of novel inhibitors targeting protein kinases, G protein-coupled receptors, and viral proteases [17]. Structure-based methods have also been used to optimize the potency and selectivity of lead compounds through rational design and virtual screening against focused compound libraries [15]. Additionally, structure-activity relationship analysis has provided insights into the molecular determinants of ligand binding and guided the design of next-generation therapeutics with improved pharmacological properties [18].

LIGAND-BASED VIRTUAL SCREENING

Ligand-based virtual screening relies on the similarity principle, where active compounds are used as queries to search for structurally or chemically related molecules in large compound libraries [23]. This subsection provides an overview of the rationale behind ligand-based virtual screening, highlighting its versatility and applicability to diverse target classes [24–26]. It also discusses the advantages of ligand-based approaches, such as the ability to identify novel scaffolds and the independence from target structure information [11].

Methods: Similarity Searching, Pharmacophore Modelling, QSAR Analysis

Ligand-based virtual screening is a computational approach used to identify potential ligands based on their similarity to known active compounds [27]. This methodology comprises multiple techniques, including as pharmacophore modelling, similarity searching, and quantitative structure-activity relationship (QSAR) research [14]. Similarity searching compares the chemical structures of candidate compounds to those of known active compounds in a reference database, identifying structurally similar molecules with potential biological activity [23]. Pharmacophore modeling generates three-dimensional models of the essential chemical features (e.g., hydrogen bond donors, acceptors, hydrophobic regions) required for ligand binding, allowing for the screening of compound libraries based on their pharmacophore fit [16]. By using molecular descriptors to predict ligand potency and selectivity, QSAR analysis creates quantitative models that link ligand chemical structure to biological activity [25].

Software for Ligand-based Virtual Screening

Several software tools are available for performing ligand-based virtual screening, each offering unique features and capabilities [23]. Discovery Studio, developed by BIOVIA, provides a comprehensive suite of tools for molecular modeling and virtual screening, including tools for similarity searching, pharmacophore modeling, and QSAR analysis [12]. MOE (Molecular Operating Environment), developed by Chemical Computing Group, offers a wide range of tools for ligand-based virtual screening, including tools for similarity searching, pharmacophore modeling, and QSAR analysis [19]. The OpenEye Suite, developed by OpenEye Scientific Software, provides a collection of tools for cheminformatics and virtual screening, including tools for similarity searching, shape-based screening, and pharmacophore modelling [26].

Examples of Successful Applications

Ligand-based virtual screening has been successfully applied in various drug discovery projects to identify hit compounds, optimize their chemical structures, and prioritize them for further experimental validation [13]. For example, virtual screening campaigns have led to the discovery of novel inhibitors targeting protein kinases, proteases, and ion channels [15]. Ligand-based methods have also been used to optimize the potency and selectivity of hit compounds through scaffold hopping, analog searching, and structure-activity relationship analysis [17]. Additionally, ligand-based virtual screening has been

used to prioritize hit compounds for experimental testing based on their predicted pharmacological properties and ADME (absorption, distribution, metabolism, excretion) profiles [17].

HOMOLOGY MODELLING AND PROTEIN STRUCTURE PREDICTION

The foundation of homology modelling, sometimes referred to as comparative modelling, is the idea that proteins with similar sequences have comparable structures and activities [22]. This subsection provides an overview of homology modelling and protein structure prediction, emphasizing their importance when experimental structures of target proteins are unavailable or difficult to obtain [17]. It discusses the challenges associated with homology modelling, such as template selection and model validation, and highlights the impact of these techniques on structure-based drug design [23]. The foundation of homology modelling, sometimes referred to as comparative modelling, is the idea that proteins with comparable sequences have comparable structures and functions [24]. Finding appropriate template structures that have a high degree of sequence similarity to the target protein, aligning the target sequence with the template structure, creating a target protein model based on the template structure, and honing the model to increase accuracy and dependability are the steps in the process [24]. Techniques such as loop modelling, side-chain optimization, and model validation are used to enhance the quality of the homology model and ensure its compatibility with experimental data [27].

Fundamentals and Methods of Homology Modelling

A computational method called homology modelling, sometimes referred to as comparison modelling, is used to estimate a protein's three-dimensional structure based on how close its sequence is to a known template structure [23]. Finding a good template structure with high sequence identity to the target protein, aligning the target sequence with the template structure, and creating a target protein model based on the template structure are the steps in the procedure [26]. Homology modeling relies on the principle of evolutionary conservation, assuming that proteins with similar sequences have similar structures and functions [11].

Software for Homology Modelling

Several software tools are available for performing homology modeling, each offering unique features and capabilities [24]. MODELLER is a widely used software package for protein structure modeling that employs a satisfaction of spatial restraints approach to generate models with optimized stereochemistry and energetics [13]. SWISS-MODEL is an automated homology modeling server that provides user-friendly access to state-of-the-art modeling methods and template libraries. Phyre2 is an advanced protein structure prediction server that combines homology modeling, *ab initio* modeling, and fold recognition methods to generate accurate models of protein structures [24].

Impact on Drug Design

Homology modeling plays a crucial role in drug design by providing structural insights into target proteins and facilitating rational ligand design. Homology modelling allows for the rational design of ligands with enhanced binding affinity and selectivity, as well as the identification of ligand binding sites and the characterization of protein-ligand interactions by forecasting the three-dimensional structure of target proteins. Additionally, homology modeling can be used to validate drug targets, assess the impact of genetic variations on protein structure and function, and guide the design of structure-based drug discovery campaigns [27].

MOLECULAR VISUALIZATION

Importance of Visualization in Drug Design

Molecular visualization plays a crucial role in drug design by providing insights into the three-dimensional structure and interactions of biomolecules. Visualization tools allow researchers to explore complex molecular structures, analyze protein-ligand interactions, and visualize the dynamic behavior of biomolecular systems [23]. Researchers can better comprehend the molecular principles underlying drug action and develop more potent therapies by visualizing molecular structures and interactions [18].

Software for Molecular Visualization

Several software tools are available for molecular visualization, each offering unique features and capabilities. PyMOL is a popular molecular visualization program that provides advanced rendering and animation capabilities for visualizing protein structures, ligand binding sites, and molecular dynamics trajectories. VMD (Visual Molecular Dynamics) is a powerful visualization and analysis tool designed for studying large biomolecular systems, such as protein complexes, lipid membranes, and nucleic acids [15]. Chimera is a versatile molecular modeling program that integrates molecular visualization, analysis, and modeling tools for studying complex biological systems [3].

Visualization in Drug Design

Molecular visualization tools are used extensively in drug design to visualize protein-ligand interactions, analyze protein structures, and validate computational models. Researchers can determine important interactions between ligands and target proteins, such as hydrogen bonds, hydrophobic contacts, and electrostatic interactions, by visualizing protein-ligand complexes. Visualization tools also allow researchers to analyze protein structures, identify structural features relevant to ligand binding, and assess the impact of mutations or modifications on protein function [23]. Additionally, molecular visualization is used to validate computational models generated by molecular docking, molecular dynamics simulation, and homology modeling, ensuring their accuracy and reliability for drug discovery applications [12].

DATABASE RESOURCES FOR DRUG DESIGN

Overview of Databases

Databases play a crucial role in drug design by providing access to a wealth of biological, chemical, and pharmacological data. Proteins, nucleic acids, and complex assemblies are among the biological macromolecules whose three-dimensional structures have been experimentally identified and are compiled in the Protein Data Bank (PDB [23]). PubChem is a publicly available database maintained by the National Institutes of Health (NIH), containing chemical structures, biological activities, and experimental data for millions of compounds. DrugBank is a comprehensive database of drug and drug target information, offering detailed annotations on drug properties, mechanisms of action, and therapeutic indications [16].

Utilization of Databases in CADD

Databases are extensively utilized in computer-aided drug design (CADD) for data mining, compound screening, and target identification. Researchers can query databases to retrieve relevant biological and chemical data for target proteins, ligands, and drug-like compounds. Databases provide valuable resources for virtual screening, allowing researchers to screen compound libraries against target proteins and identify potential drug candidates. Additionally, databases facilitate target identification and validation by providing information on protein structures, functions, and interactions. By integrating data from multiple databases, researchers can gain insights into the relationships between chemical structures, biological activities, and drug targets, guiding the design of novel therapeutics [17].

Challenges and Considerations

While databases offer valuable resources for drug design, several challenges and considerations must be addressed. Data quality is a critical issue, as inaccuracies or inconsistencies in database entries can lead to erroneous conclusions in drug discovery research. Ensuring data integrity and reliability requires rigorous curation and validation processes.

Ethical and regulatory considerations, such as data privacy, intellectual property rights, and compliance with ethical guidelines, must also be taken into account when using databases for drug design. By addressing these challenges and considerations, researchers can harness the full potential of databases to accelerate drug discovery and development.

INTEGRATION OF BIOINFORMATIC TOOLS IN DRUG DISCOVERY PIPELINES

Workflow of CADD

A multi-step procedure that starts with target selection and validation and continues with hit identification, lead optimization, and preclinical testing is involved in the integration of bioinformatic tools into drug discovery pipelines [16]. Bioinformatic tools are utilized at each stage of the drug discovery pipeline to facilitate data analysis, model building, and decision-making. In the target identification phase, bioinformatic tools are used to analyze biological data, predict protein structures, and prioritize potential drug targets based on their biological relevance and drug ability [20]. In the hit identification phase, bioinformatic tools are employed to screen compound libraries, analyze structure-activity relationships, and identify lead compounds with desired pharmacological properties. In the lead optimization phase, bioinformatic tools are utilized to refine lead compounds, predict their ADMET (absorption, distribution, metabolism, excretion, and toxicity) properties, and optimize their pharmacokinetic profiles [22].

Challenges and Considerations

The integration of bioinformatic tools in drug discovery pipelines presents several challenges and considerations that must be addressed to ensure successful outcomes. Data integration is a key challenge, as drug discovery pipelines rely on diverse datasets from various sources, including biological, chemical, and pharmacological data [23]. To maintain data quality, consistency, and accessibility, integrating and harmonizing different datasets requires strong data management and integration procedures [24]. Algorithm validation is another challenge, as bioinformatic tools often rely on complex algorithms and computational models that require validation against experimental data [13]. Ensuring the accuracy, reliability, and reproducibility of these algorithms is essential for making informed decisions in drug discovery. Regulatory compliance is also a consideration, as drug discovery pipelines must adhere to ethical, legal, and regulatory requirements governing the use of biological and chemical data, as well as the development and testing of therapeutic agents [24].

Future Directions

Despite the challenges, the integration of bioinformatic tools in drug discovery pipelines offers tremendous opportunities for innovation and discovery. Future advancements in artificial intelligence (AI), machine learning, and deep learning are expected to revolutionize drug discovery by enabling predictive modeling, virtual screening, and *de novo* drug design. Cloud computing technologies are also poised to transform drug discovery pipelines by providing scalable, cost-effective computational resources for data analysis, modeling, and simulation. Additionally, open science initiatives aimed at sharing data, tools, and resources are fostering collaboration and accelerating the pace of drug discovery research. By embracing these advancements and overcoming the challenges, the integration of bioinformatic tools in drug discovery pipelines holds great promise for advancing our understanding of disease mechanisms and developing novel therapeutics to improve human health [27].

CONCLUSION

Summary of Key Points

In conclusion, this review has provided an extensive overview of the bioinformatic tools used in Computer-aided Drug Design (CADD). We have discussed the principles, methodologies, and applications of various bioinformatic techniques, including molecular docking, molecular dynamics simulation, virtual screening, homology modeling, and molecular visualization. These tools play a critical role in accelerating the drug discovery process, from target identification and hit identification to lead optimization and preclinical testing.

Implications and Outlook

The integration of bioinformatic tools with computational methods has revolutionized drug discovery, enabling researchers to tackle complex biological questions and identify novel therapeutic targets. As technology continues to advance, we can expect further innovations in CADD and bioinformatics, including the development of more accurate and efficient algorithms, the integration of

multi-omics data, and the adoption of cloud-based computing platforms. These advancements hold the potential to transform drug discovery by facilitating the rapid identification and optimization of new drugs for a wide range of diseases and conditions.

Closing Remarks

In closing, successful drug discovery requires collaboration, innovation, and knowledge sharing across disciplines. By bringing together expertise from biology, chemistry, computational science, and bioinformatics, we can overcome the challenges of drug discovery and develop novel therapeutics to address unmet medical needs. Through continued collaboration and innovation, we can harness the power of bioinformatic tools to accelerate drug discovery and improve human health.

REFERENCES

1. Rognan D. The impact of in silico screening in the discovery of novel and safer drug candidates. *Pharmacol Ther.* 2017 Jul 1; 175: 47–66. <https://doi.org/10.1016/j.phar>
2. Robinson BS, Riccardi KA, Gong YF, Guo Q, Stock DA, Blair WS, Terry BJ, Deminie CA, Djang F, Colonna RJ, Lin PF. BMS-232632, a highly potent human immunodeficiency virus protease inhibitor that can be used in combination with other available antiretroviral agents. *Antimicrob Agents Chemother.* 2000 Aug 1; 44(8): 2093–9.
3. Anderson AC. The process of structure-based drug design. *Chem Biol.* 2003 Sep 1; 10(9): 787–97.
4. Rutenber EE, Stroud RM. Binding of the anticancer drug ZD1694 to *E. coli* thymidylate synthase: assessing specificity and affinity. *Structure.* 1996 Nov 15; 4(11): 1317–24.
5. Jhoti H, Leach AR. *Structure-based drug discovery.* Dordrecht (Netherlands): Springer; 2007; 250.
6. Vidal D, Garcia-Serna R, Mestres J. Ligand-based approaches to in silico pharmacology. *Cheminformatics and computational chemical biology.* 2011; 672: 489–502.
7. Daina A, Michielin O, Zoete V. SwissADME: a free web tool to evaluate pharmacokinetics, drug-likeness and medicinal chemistry friendliness of small molecules. *Sci Rep.* 2017; 7: 42717.
8. Wishart DS. Introduction to cheminformatics. *Curr Protoc Bioinformatics.* 2007 Jun; 18(1): 14–21.
9. Begam BF, Kumar JS. A study on cheminformatics and its applications on modern drug discovery. *Procedia Eng.* 2012 Jan 1; 38: 1264–75.
10. Manzoni C, Kia DA, Vandrovcova J, Hardy J, Wood NW, Lewis PA, Ferrari R. Genome, transcriptome and proteome: the rise of omics data and their integration in biomedical sciences. *Brief Bioinformatics.* 2018 Mar; 19(2): 286–302.
11. Lesk AM. (2023 Sep 12). Bioinformatics. *Encyclopedia Britannica.* [Online]. URL: <https://www.britannica.com/science/bioinformatics>.
12. Xia X. Bioinformatics and drug discovery. *Curr Top Med Chem.* 2017 Jun 1; 17(15): 1709–26.
13. Wishart DS, Knox C, Guo AC, Cheng D, Shrivastava S, Tzur D, Gautam B, Hassanali M. DrugBank: a knowledgebase for drugs, drug actions and drug targets. *Nucleic Acids Res.* 2008 Jan 1; 36(suppl_1): D901–6.
14. Irwin JJ, Tang KG, Young J, Dandarchuluun C, Wong BR, Khurelbaatar M, Moroz YS, Mayfield J, Sayle RA. ZINC20—a free ultra large-scale chemical database for ligand discovery. *J Chem Inf Model.* 2020 Oct 29; 60(12): 6065–73.
15. Pence HE, Williams A. ChemSpider: an Online Chemical Information Resource. *J Chem Educ.* 2010; 87(11): 1123–1124.
16. Kim S, Chen J, Cheng T, Gindulyte A, He J, He S, Li Q, Shoemaker BA, Thiessen PA, Yu B, Zaslavsky L. PubChem in 2021: new data content and improved web interfaces. *Nucl Acids Res.* 2021 Jan 8; 49(D1): D1388–95.
17. Zdrazil B, Felix E, Hunter F, Manners EJ, Blackshaw J, Corbett S, de Veij M, Ioannidis H, Lopez DM, Mosquera JF, Magarinos MP. The ChEMBL Database in 2023: a drug discovery platform spanning multiple bioactivity data types and time periods. *Nucl Acids Res.* 2024 Jan 5; 52(D1): D1180–92.
18. Wishart DS, Feunang YD, Guo AC, Lo EJ, Marcu A, Grant JR, Sajed T, Johnson D, Li C, Sayeeda Z, Assempour N. DrugBank 5.0: a major update to the DrugBank database for 2018. *Nucl Acids Res.* 2018 Jan 4; 46(D1): D1074–82.

19. Sayers EW, Cavanaugh M, Clark K, Pruitt KD, Sherry ST, Yankie L, Karsch-Mizrachi I. GenBank 2023 update. *Nucl Acids Res.* 2023 Jan 6; 51(D1): D141–4.
20. Hecker N, Ahmed J, Von Eichborn J, Dunkel M, Macha K, Eckert A, Gilson MK, Bourne PE, Preissner R. SuperTarget goes quantitative: update on drug–target interactions. *Nucl Acids Res.* 2012 Jan 1; 40(D1): D1113–7.
21. Chang A, Jeske L, Ulbrich S, Hofmann J, Koblitz J, Schomburg I, Neumann-Schaal M, Jahn D, Schomburg D. BRENDA, the ELIXIR core data resource in 2021: new developments and updates. *Nucl Acids Res.* 2021 Jan 8; 49(D1): D498–508.
22. Andreeva A, Howorth D, Brenner SE, Hubbard TJ, Chothia C, Murzin AG. SCOP database in 2004: refinements integrate structure and sequence family data. *Nucl Acids Res.* 2004 Jan 1; 32(suppl_1): D226–9.
23. Gillespie M, Jassal B, Stephan R, Milacic M, Rothfels K, Senff-Ribeiro A, Griss J, Sevilla C, Matthews L, Gong C, Deng C. The reactome pathway knowledgebase 2022. *Nucl Acids Res.* 2022 Jan 7; 50(D1): D687–92.
24. Raslan MA, Raslan SA, Shehata EM, Mahmoud AS, Sabri NA. Advances in the Applications of Bioinformatics and Chemoinformatics. *Pharmaceuticals.* 2023 Jul 24; 16(7): 1050.
25. Lin X, Li X, Lin X. A review on applications of computational methods in drug screening and design. *Molecules.* 2020 Mar 18; 25(6): 1375.
26. Kassab MM. Development of Novel Antimicrobial Tetracycline Analog B (Iodocycline) By Chemo-Informatics. *Ain Shams Med J.* 2022 Dec 1; 73(4): 969–81.
27. Dawod EF, Mahmoud N, Elsisi A. Hybrid approach for COVID-19 detection from chest radiography. *Inter J Computers and Information (IJCI).* 2021; 8: 71–76.