

BAB 2

LANDASAN TEORI

2.1 Kanker Pankreas

Kanker pankreas didefinisikan sebagai kondisi keganasan yang ditandai dengan pertumbuhan sel yang tidak terkendali pada jaringan pankreas. Pankreas merupakan organ yang terletak di belakang lambung dan memiliki fungsi penting dalam sistem pencernaan serta pengaturan kadar gula darah melalui produksi enzim pencernaan dan hormon seperti insulin. Kanker pankreas dapat berasal dari bagian eksokrin maupun endokrin pankreas, namun sebagian besar kasus berkembang dari sel-sel eksokrin, terutama sebagai *pancreatic ductal adenocarcinoma* (PDAC), yaitu kanker yang bermula dari sel-sel yang melapisi saluran pankreas [8].

Secara molekuler, kanker pankreas berkaitan dengan akumulasi perubahan genetik dan epigenetik yang mengganggu regulasi pertumbuhan sel, apoptosis, perbaikan DNA, serta jalur pensinyalan seluler. Beberapa gen yang sering terlibat dalam perkembangan kanker pankreas antara lain *KRAS*, *TP53*, *SMAD4*, dan *CDKN2A*. Mutasi pada gen-gen tersebut dapat menyebabkan aktivasi onkogen, inaktivasi gen penekan tumor, serta gangguan mekanisme kontrol siklus sel. Perubahan tersebut memungkinkan sel kanker untuk terus berproliferasi, menghindari kematian sel terprogram, menginvasi jaringan sekitar, dan menyebar ke organ lain melalui proses metastasis [9].

Berdasarkan data GLOBOCAN 2022, kanker pankreas menyumbang sekitar 510.992 kasus baru dan 467.409 kematian di seluruh dunia. Jumlah kematian yang hampir mendekati jumlah kasus baru menunjukkan bahwa kanker pankreas merupakan salah satu kanker dengan tingkat fatalitas yang tinggi [1]. Tingginya angka kematian tersebut berkaitan dengan karakteristik kanker pankreas yang sering berkembang tanpa gejala spesifik pada tahap awal. Akibatnya, sebagian besar pasien baru terdiagnosis ketika penyakit telah berada pada stadium lanjut atau telah mengalami metastasis, sehingga pilihan terapi menjadi lebih terbatas dan prognosis pasien menjadi buruk.

Tingkat kelangsungan hidup pasien kanker pankreas juga masih tergolong rendah dibandingkan dengan banyak jenis kanker lainnya. Beberapa laporan menyebutkan bahwa tingkat kelangsungan hidup lima tahun secara keseluruhan pada kanker pankreas berada pada kisaran sekitar 12–13% [10, 11]. Rendahnya

tingkat kelangsungan hidup ini dipengaruhi oleh beberapa faktor, seperti keterlambatan diagnosis, agresivitas tumor, keterbatasan terapi yang efektif, resistensi terhadap kemoterapi, serta kompleksitas lingkungan mikro tumor. Meskipun perkembangan terapi target dan imunoterapi terus dilakukan, kanker pankreas masih menjadi tantangan besar dalam bidang onkologi.

Oleh karena itu, diperlukan eksplorasi target molekuler dan kandidat senyawa terapeutik baru yang dapat mendukung pengembangan strategi pengobatan kanker pankreas. Pendekatan komputasional seperti analisis jaringan biologis, network centrality, Skyline Query, dan molecular docking dapat digunakan untuk membantu mengidentifikasi gen target potensial serta mengevaluasi interaksi antara protein target dan senyawa kandidat. Dalam konteks penelitian ini, kanker pankreas menjadi fokus utama untuk mengidentifikasi target obat potensial dari tanaman herbal Indonesia melalui pendekatan berbasis data dan analisis *in silico*.

2.2 Tanaman Herbal

Tanaman herbal merupakan salah satu sumber senyawa bioaktif yang banyak digunakan dalam pengobatan tradisional maupun penelitian berbasis komputasional. Dalam konteks penemuan kandidat obat, tanaman herbal dapat dikaji melalui pendekatan *in silico* untuk mengidentifikasi hubungan antara senyawa bioaktif, target protein, dan penyakit tertentu.

Pada penelitian ini, tanaman herbal yang digunakan berasal dari Indonesia, yaitu jahe, bawang merah, dan kunyit. Ketiga tanaman tersebut dipilih karena memiliki kandungan senyawa bioaktif yang dapat ditelusuri melalui basis data tanaman dan senyawa. Data tanaman dan senyawa dalam penelitian ini diperoleh dari IJAH Analytics, yaitu basis data tanaman herbal Indonesia, serta KNApSACk, yaitu basis data metabolit tanaman secara global [12, 13]. Melalui kedua basis data tersebut, senyawa-senyawa bioaktif dari tanaman herbal dapat diidentifikasi dan digunakan sebagai dasar untuk prediksi target molekuler.

Jahe (*Zingiber officinale*) merupakan tanaman herbal yang banyak digunakan dalam pengobatan tradisional dan dikenal memiliki berbagai senyawa bioaktif. Dalam penelitian ini, jahe digunakan sebagai salah satu sumber kandidat senyawa yang berpotensi berinteraksi dengan target molekuler yang berkaitan dengan kanker pankreas. Senyawa-senyawa dari jahe kemudian dianalisis lebih lanjut untuk memperoleh target gen atau protein yang relevan.

Bawang merah (*Allium cepa*) juga digunakan sebagai salah satu tanaman

herbal dalam penelitian ini. Bawang merah diketahui memiliki beragam senyawa bioaktif yang dapat dikaji dalam konteks farmakologi komputasional. Target-target molekuler dari senyawa bawang merah dikumpulkan dan digabungkan dengan target dari tanaman lainnya untuk membentuk dataset target tanaman herbal.

Kunyit (*Curcuma longa*) merupakan tanaman herbal yang banyak diteliti karena kandungan senyawa bioaktifnya, terutama kurkuminoid. Dalam penelitian ini, kunyit digunakan sebagai sumber kandidat senyawa yang kemudian dipetakan terhadap target molekuler. Data target dari kunyit digabungkan dengan data target jahe dan bawang merah untuk memperoleh dataset gabungan tanaman herbal.

2.3 Pendekatan Network Farmakologi

Network pharmacology merupakan pendekatan komputasional yang mengintegrasikan biologi sistem, bioinformatika, dan farmakologi untuk memahami hubungan kompleks antara senyawa bioaktif, target protein, penyakit, serta jalur biologis yang terlibat di dalamnya [14]. Pendekatan ini relevan untuk digunakan dalam penelitian berbasis tanaman herbal karena tanaman herbal umumnya memiliki banyak senyawa bioaktif yang dapat bekerja pada lebih dari satu target molekuler. Berbeda dengan paradigma konvensional yang berfokus pada satu obat dan satu target, network pharmacology memungkinkan analisis multi-senyawa, multi-target, dan multi-jalur secara lebih sistematis.

Dalam konteks kanker pankreas, pendekatan network pharmacology dapat digunakan untuk mengidentifikasi hubungan antara senyawa bioaktif dari tanaman herbal Indonesia dengan gen atau protein yang berkaitan dengan perkembangan penyakit. Kanker pankreas merupakan penyakit kompleks yang melibatkan banyak perubahan molekuler, interaksi genetik, serta jalur pensinyalan seluler. Oleh karena itu, analisis berbasis jaringan diperlukan untuk memahami keterlibatan berbagai target molekuler dan menentukan target yang memiliki peran penting dalam jaringan biologis.

Pada penelitian ini, pendekatan network pharmacology digunakan untuk mengkaji potensi target dari tiga tanaman herbal Indonesia, yaitu jahe (*Zingiber officinale*), bawang merah (*Allium cepa*), dan kunyit (*Curcuma longa*). Data senyawa bioaktif diperoleh dari basis data IJAH Analytics dan KNApSACK. Selanjutnya, senyawa tersebut diseleksi menggunakan parameter *Oral Bioavailability* dan *Drug-likeness* melalui SwissADME dan Molsoft [15, 16, 17]. Senyawa yang memenuhi kriteria kemudian digunakan untuk prediksi target protein

menggunakan SwissTargetPrediction [18, 19].

Data target penyakit kanker pankreas dikumpulkan dari basis data publik yang relevan, kemudian dilakukan proses penghapusan duplikasi menggunakan Venny [20]. Setelah data target tanaman dan data gen penyakit diperoleh, dilakukan analisis irisan untuk mendapatkan gen yang sama antara target tanaman herbal dan gen terkait kanker pankreas. Gen hasil irisan tersebut menjadi dasar dalam pembentukan jaringan interaksi protein dan analisis lanjutan.

Jaringan interaksi protein dibentuk menggunakan STRING v12.0 untuk memperoleh hubungan fungsional antarprotein yang terlibat dalam dataset overlap [21]. Jaringan ini kemudian dianalisis menggunakan beberapa ukuran *network centrality*, yaitu degree centrality, betweenness centrality, closeness centrality, dan eigenvector centrality.

Setelah nilai centrality diperoleh, penelitian ini menerapkan algoritma Skyline Query dengan metode Block Nested Loop untuk melakukan seleksi multi-kriteria terhadap kandidat gen target. Kriteria yang digunakan dalam proses Skyline Query adalah empat nilai centrality, yaitu degree, betweenness, closeness, dan eigenvector centrality.

Secara keseluruhan, pendekatan network pharmacology dalam penelitian ini digunakan sebagai kerangka kerja untuk menghubungkan senyawa bioaktif tanaman herbal, target protein, dan kanker pankreas. Integrasi antara seleksi senyawa, prediksi target, analisis jaringan protein, perhitungan centrality, Skyline Query, dan molecular docking diharapkan dapat menghasilkan proses identifikasi target obat yang lebih sistematis dan objektif.

2.4 Sentralitas Jaringan

Sentralitas jaringan merupakan salah satu pendekatan analisis dalam teori graf yang digunakan untuk mengukur tingkat kepentingan suatu node berdasarkan posisinya di dalam jaringan. Dalam konteks network pharmacology, node dapat merepresentasikan gen atau protein, sedangkan edge merepresentasikan hubungan atau interaksi antarprotein. Analisis sentralitas menjadi penting karena tidak semua protein dalam jaringan memiliki peran yang sama. Beberapa protein dapat memiliki banyak koneksi langsung, menjadi penghubung antarbagian jaringan, memiliki akses yang cepat ke node lain, atau terhubung dengan protein lain yang juga memiliki pengaruh tinggi [22].

Penggunaan sentralitas jaringan dapat membantu mengidentifikasi protein

atau gen yang memiliki peran kunci dalam struktur jaringan biologis. Protein dengan nilai sentralitas tinggi sering dianggap sebagai kandidat penting karena berpotensi memengaruhi stabilitas jaringan, jalur pensinyalan, maupun proses biologis yang berkaitan dengan penyakit. Dalam penelitian ini, analisis sentralitas digunakan untuk mengevaluasi gen hasil irisan antara dataset gen kanker pankreas dan target gen dari tanaman herbal Indonesia. Empat ukuran sentralitas yang digunakan adalah degree centrality, betweenness centrality, closeness centrality, dan eigenvector centrality.

1. Degree Centrality

Degree centrality mengukur jumlah koneksi langsung yang dimiliki oleh suatu node dalam jaringan [23]. Semakin tinggi nilai degree centrality suatu node, semakin banyak hubungan langsung yang dimiliki node tersebut dengan node lainnya. Dalam konteks jaringan protein, protein dengan nilai degree tinggi dapat dianggap sebagai *hub protein* karena memiliki banyak interaksi langsung dengan protein lain.

Secara matematis, degree centrality dapat dirumuskan sebagai berikut:

$$C_D(v) = k(v) \quad (2.1)$$

Keterangan:

- $C_D(v)$ adalah nilai degree centrality dari node v .
- v adalah node yang dianalisis dalam jaringan.
- $k(v)$ adalah jumlah koneksi langsung atau edge yang dimiliki oleh node v .

2. Betweenness Centrality

Betweenness centrality mengukur seberapa sering suatu node berada pada jalur terpendek antara dua node lain dalam jaringan [24]. Node dengan nilai betweenness centrality tinggi memiliki peran penting sebagai penghubung atau mediator dalam jaringan. Dalam jaringan biologis, protein dengan nilai betweenness tinggi dapat berperan sebagai pengatur utama dalam aliran informasi atau interaksi antarprotein.

Secara matematis, betweenness centrality dapat dirumuskan sebagai berikut:

$$C_B(v) = \sum_{s \neq v \neq t} \frac{\sigma_{st}(v)}{\sigma_{st}} \quad (2.2)$$

Keterangan:

- $C_B(v)$ adalah nilai betweenness centrality dari node v .
- σ_{st} adalah jumlah jalur terpendek dari node s ke node t .
- $\sigma_{st}(v)$ adalah jumlah jalur terpendek dari node s ke node t yang melewati node v .
- s dan t adalah pasangan node berbeda dalam jaringan.

3. Closeness Centrality

Closeness centrality mengukur seberapa dekat suatu node terhadap seluruh node lain dalam jaringan [25]. Node dengan nilai closeness centrality tinggi memiliki jarak rata-rata yang lebih pendek terhadap node lain, sehingga dapat menjangkau bagian lain dalam jaringan secara lebih efisien. Dalam konteks protein interaction network, nilai closeness yang tinggi dapat menunjukkan bahwa suatu protein memiliki akses yang cepat terhadap protein lain dalam sistem biologis.

Secara matematis, closeness centrality dapat dirumuskan sebagai berikut:

$$C_C(v) = \frac{n-1}{\sum_{u \neq v} d(v,u)} \quad (2.3)$$

Keterangan:

- $C_C(v)$ adalah nilai closeness centrality dari node v .
- n adalah jumlah seluruh node dalam jaringan.
- $d(v,u)$ adalah jarak terpendek antara node v dan node u .
- u adalah node lain dalam jaringan selain node v .

4. Eigenvector Centrality

Eigenvector centrality mengukur tingkat kepentingan suatu node berdasarkan koneksinya dengan node lain yang juga memiliki tingkat kepentingan tinggi [26]. Berbeda dengan degree centrality yang hanya menghitung jumlah

koneksi langsung, eigenvector centrality mempertimbangkan kualitas koneksi suatu node. Dengan demikian, suatu node dapat memiliki nilai eigenvector centrality tinggi apabila terhubung dengan node-node yang juga penting dalam jaringan.

Secara matematis, eigenvector centrality dapat dirumuskan sebagai berikut:

$$C_E(v) = \frac{1}{\lambda} \sum_{w \in V} A_{vw} \cdot C_E(w) \quad (2.4)$$

Keterangan:

- $C_E(v)$ adalah nilai eigenvector centrality dari node v .
- λ adalah eigenvalue yang digunakan sebagai faktor normalisasi.
- V adalah himpunan seluruh node dalam jaringan.
- w adalah node lain dalam jaringan yang memiliki kemungkinan hubungan dengan node v .
- $\sum_{w \in V}$ adalah penjumlahan terhadap seluruh node w dalam himpunan node V .
- A_{vw} adalah elemen adjacency matrix yang menunjukkan hubungan antara node v dan node w . Jika node v terhubung dengan node w , maka nilainya adalah 1; jika tidak terhubung, maka nilainya adalah 0.
- $C_E(w)$ adalah nilai eigenvector centrality dari node w .

Kombinasi keempat ukuran sentralitas tersebut digunakan untuk memperoleh pemahaman yang lebih menyeluruh mengenai peran setiap gen dalam jaringan. Degree centrality menunjukkan jumlah interaksi langsung, betweenness centrality menunjukkan peran sebagai penghubung, closeness centrality menunjukkan efisiensi akses terhadap node lain, dan eigenvector centrality menunjukkan keterhubungan dengan node penting lainnya. Dalam penelitian ini, nilai dari keempat centrality tersebut digunakan sebagai kriteria dalam proses prioritisasi gen target kanker pankreas menggunakan algoritma Skyline Query.

2.5 Algoritma Skyline Query

Skyline Query merupakan metode seleksi data berbasis *multi-criteria decision analysis* yang digunakan untuk memilih objek terbaik berdasarkan beberapa kriteria secara bersamaan [6]. Dalam Skyline Query, suatu objek dikatakan mendominasi objek lain apabila objek tersebut memiliki nilai yang sama atau lebih baik pada seluruh kriteria, serta memiliki nilai yang lebih baik pada minimal satu kriteria. Dengan demikian, hasil dari Skyline Query adalah sekumpulan objek yang tidak terdominasi oleh objek lain.

Secara umum, misalkan terdapat dua objek p dan q dengan sejumlah kriteria d . Objek p dikatakan mendominasi objek q apabila memenuhi kondisi berikut:

$$p \succ q \iff \forall i \in d, p_i \geq q_i \wedge \exists j \in d, p_j > q_j \quad (2.5)$$

Keterangan:

- $p \succ q$ menunjukkan bahwa objek p mendominasi objek q .
- p_i dan q_i merupakan nilai objek p dan q pada kriteria ke- i .
- $\forall i \in d$ menunjukkan bahwa kondisi berlaku untuk semua kriteria.
- $\exists j \in d$ menunjukkan bahwa terdapat minimal satu kriteria dengan nilai yang lebih baik.

Dalam penelitian ini, seluruh kriteria yang digunakan bersifat *maximize*, yaitu semakin tinggi nilai kriteria maka semakin baik kandidat tersebut. Kriteria yang digunakan dalam proses Skyline Query adalah *degree centrality*, *betweenness centrality*, *closeness centrality*, dan *eigenvector centrality*. Oleh karena itu, gen yang memiliki nilai centrality lebih tinggi pada beberapa aspek dan tidak lebih rendah pada aspek lainnya dapat dianggap lebih dominan dibandingkan gen lain.

Algoritma yang digunakan untuk menjalankan Skyline Query dalam penelitian ini adalah *block nested loop*. Algoritma *block nested loop* bekerja dengan menyimpan kandidat sementara dalam sebuah daftar atau *window*. Setiap data baru akan dibandingkan dengan kandidat yang sudah ada di dalam *window*. Apabila data baru didominasi oleh kandidat yang telah ada, maka data tersebut tidak dimasukkan ke dalam hasil skyline. Sebaliknya, apabila data baru mendominasi kandidat yang telah ada, maka kandidat yang terdominasi akan dihapus dari *window*. Jika

data baru tidak terdominasi oleh kandidat mana pun, maka data tersebut akan dimasukkan sebagai kandidat skyline.

Langkah-langkah umum penerapan Skyline Query dengan metode Block Nested Loop dalam penelitian ini adalah sebagai berikut:

1. Mengumpulkan data gen hasil irisan antara dataset target tanaman herbal Indonesia dan dataset gen kanker pankreas.
2. Membentuk jaringan interaksi protein berdasarkan gen hasil irisan.
3. Menghitung nilai degree centrality, betweenness centrality, closeness centrality, dan eigenvector centrality untuk setiap gen dalam jaringan.
4. Menentukan keempat nilai centrality tersebut sebagai kriteria dalam proses Skyline Query.
5. Membandingkan setiap gen berdasarkan konsep dominasi menggunakan metode Block Nested Loop.
6. Menghapus gen yang terdominasi oleh gen lain berdasarkan kombinasi nilai centrality.
7. Menghasilkan gen skyline, yaitu gen yang tidak terdominasi dan memiliki potensi lebih kuat sebagai kandidat target obat kanker pankreas.

Dalam konteks penelitian ini, Skyline Query digunakan untuk memprioritaskan kandidat gen target kanker pankreas berdasarkan beberapa ukuran sentralitas jaringan secara bersamaan. Pendekatan ini dipilih karena penggunaan satu parameter centrality saja dapat menghasilkan penilaian yang terbatas. Sebagai contoh, suatu gen dapat memiliki degree centrality tinggi, tetapi belum tentu memiliki betweenness atau eigenvector centrality yang tinggi. Dengan menggunakan Skyline Query, proses prioritisasi dapat mempertimbangkan beberapa karakteristik jaringan secara bersamaan tanpa harus mengabaikan salah satu kriteria.

Integrasi antara network pharmacology, analisis sentralitas jaringan, dan Skyline Query memungkinkan proses identifikasi kandidat target obat menjadi lebih sistematis dan objektif. Hasil dari proses Skyline Query digunakan sebagai dasar untuk menentukan gen atau protein prioritas yang kemudian dipertimbangkan dalam tahap molecular docking. Dengan demikian, Skyline Query berperan sebagai metode seleksi multi-kriteria yang mendukung proses identifikasi target obat kanker pankreas dari tanaman herbal Indonesia.

2.6 Protein dan Compound

Protein dan ligan yang didapat nantinya akan dipersiapkan untuk dilakukan tahap lebih lanjut. Protein akan dipersiapkan menjadi reseptor atau induknya. Compound akan dipersiapkan menjadi ligan dan akan dilakukan *docking* dengan reseptor.

2.7 Enrichment Analysis

Enrichment Analysis dilakukan untuk menginterpretasikan lebih lanjut relevansi biologis dari gen-gen yang telah diprioritaskan. Pada penelitian ini, analisis pengayaan difokuskan pada analisis *Gene Ontology* (GO) dan jalur *Kyoto Encyclopedia of Genes and Genomes* (KEGG). *Gene Ontology* menyediakan informasi terstruktur mengenai fungsi gen dan produk gen. GO dibagi menjadi tiga kategori utama, yaitu *Molecular Function* (MF), *Biological Process* (BP), dan *Cellular Component* (CC). *Molecular Function* menjelaskan aktivitas pada tingkat molekuler yang dilakukan oleh produk gen, *Biological Process* menjelaskan program atau proses biologis yang lebih luas yang melibatkan berbagai aktivitas molekuler, sedangkan *Cellular Component* menjelaskan lokasi atau struktur seluler tempat produk gen menjalankan fungsinya [27].

Selain analisis GO, analisis jalur KEGG dilakukan untuk mengidentifikasi jalur biologis yang berkaitan dengan gen-gen terpilih. Analisis jalur KEGG menyediakan informasi mengenai jaringan interaksi molekuler, proses seluler, serta jalur yang berkaitan dengan penyakit yang mungkin relevan dengan kanker pankreas. Pada penelitian ini, analisis pengayaan GO dan jalur KEGG dilakukan menggunakan Enrichr, yaitu platform analisis pengayaan set gen berbasis web yang dikembangkan oleh Ma'ayan Laboratory [28]. Melalui analisis ini, gen-gen terpilih dapat diinterpretasikan lebih lanjut berdasarkan fungsi biologis, lokasi seluler, aktivitas molekuler, serta keterlibatannya dalam jalur biologis yang relevan.

2.8 Molecular Docking dan Visualisasi

Molecular docking merupakan tahap yang digunakan untuk menganalisis interaksi antara protein target dan ligan kandidat dari tanaman herbal Indonesia. Pada penelitian ini, molecular docking dilakukan setelah proses prioritisasi gen target menggunakan analisis *network centrality* dan algoritma *Skyline Query*. Gen atau protein yang terpilih dari proses tersebut digunakan sebagai dasar dalam

pemilihan reseptor untuk proses docking. Tujuan dari tahap ini adalah untuk mengevaluasi potensi interaksi antara reseptor dan ligan berdasarkan nilai *binding affinity* yang dihasilkan.

Sebelum proses docking dilakukan, reseptor dan ligan perlu dipersiapkan terlebih dahulu. Struktur reseptor diperoleh dari RCSB Protein Data Bank, sedangkan struktur ligan diperoleh dari basis data senyawa yang relevan. Persiapan reseptor dilakukan menggunakan AutoDockTools, yang digunakan untuk membersihkan struktur protein, menambahkan atom hidrogen, serta mengatur format file agar sesuai dengan kebutuhan docking [29]. Sementara itu, Open Babel digunakan untuk membantu proses konversi format file ligan, misalnya dari format SDF menjadi format PDB atau format lain yang diperlukan dalam proses persiapan docking [30].

Proses molecular docking pada penelitian ini dilakukan menggunakan AutoDock Vina versi 1.2.7. AutoDock Vina digunakan untuk memprediksi posisi ikatan ligan terhadap reseptor serta menghitung nilai *binding affinity* dari setiap pasangan protein–ligan [31]. Nilai *binding affinity* yang lebih rendah atau lebih negatif menunjukkan prediksi interaksi yang lebih kuat antara ligan dan reseptor. Pada penelitian ini, proses docking dilakukan terhadap tiga reseptor, yaitu 6YNE, 8RCI, dan 8UW9. Setiap reseptor kemudian dianalisis berdasarkan tiga ligan terbaik yang diperoleh dari hasil docking.

Hasil molecular docking selanjutnya divisualisasikan untuk memahami pola interaksi antara ligan dan reseptor. Visualisasi tiga dimensi dilakukan menggunakan PyMOL untuk menampilkan posisi ligan pada sisi ikatan reseptor serta bentuk interaksi secara spasial [32]. Selain itu, visualisasi dua dimensi dilakukan menggunakan LigPlot+ untuk mengidentifikasi jenis interaksi antara ligan dan residu asam amino pada reseptor, seperti ikatan hidrogen dan interaksi hidrofobik [33]. Visualisasi ini digunakan untuk mendukung interpretasi hasil docking dan memperjelas interaksi protein–ligan yang terbentuk.

Secara umum, tahapan molecular docking dan visualisasi dalam penelitian ini meliputi persiapan reseptor, persiapan ligan, konversi format file, pelaksanaan docking menggunakan AutoDock Vina, pemilihan hasil docking berdasarkan nilai *binding affinity*, serta visualisasi hasil dalam bentuk dua dimensi dan tiga dimensi. Hasil dari tahap ini digunakan untuk mengevaluasi kandidat protein–ligan yang berpotensi sebagai dasar dalam identifikasi target obat kanker pankreas.