

BAB 2

LANDASAN TEORI

2.1 Tumbuhan Obat Indonesia dan Jamu

Tumbuhan obat Indonesia secara tradisional diolah menjadi jamu, yaitu istilah Jawa untuk ramuan obat tradisional berbahan tumbuhan yang kerap memadukan beberapa jenis tumbuhan, yang telah dipraktikkan masyarakat selama berabad-abad untuk memelihara kesehatan dan mengobati penyakit. Bentuknya beragam, dari minuman segar yang diracik secara rumahan hingga produk industri seperti serbuk, tablet, dan kapsul. Sebagian besar khasiatnya bersandar pada data empiris dari penggunaan turun-temurun, bukan pada uji klinis, sehingga khasiat dan keamanannya tidak dapat dianggap terjamin dengan sendirinya dan masih memerlukan pembuktian ilmiah [1]. Kesenjangan antara penggunaan empiris dan pembuktian ilmiah tersebut menjadi dasar pengaturan tumbuhan obat di Indonesia, baik melalui program Saintifikasi Jamu [2] maupun melalui penggolongan produk menurut tingkat pembuktian khasiatnya.

2.2 Network Pharmacology

Network pharmacology adalah pendekatan farmakologi yang menjelaskan kerja obat sebagai interaksi banyak senyawa dengan banyak target protein di dalam suatu jaringan biologis, bukan sebagai kerja satu senyawa pada satu target [4]. Pendekatan ini muncul sebagai pergeseran paradigma dari model “satu obat, satu target” setelah disadari bahwa sebagian besar obat memengaruhi beberapa target sekaligus, sifat yang disebut polifarmakologi [4]. Dengan memetakan hubungan antara senyawa, target protein, dan penyakit ke dalam satu jaringan, pendekatan ini bertujuan menjelaskan mekanisme kerja obat pada tingkat sistem, yakni menunjukkan senyawa mana yang bekerja pada target mana dan jalur biologis apa yang terpengaruh, bukan menilai satu target yang terpisah [5].

Cara pandang banyak senyawa dengan banyak target ini sejalan dengan sifat tumbuhan obat. Berbeda dengan kebanyakan obat sintetis yang hanya terdiri atas satu senyawa aktif, tumbuhan obat beserta ramuannya mengandung banyak metabolit, yaitu senyawa kimia yang dihasilkan tumbuhan melalui metabolismenya, yang dapat bekerja pada banyak target sekaligus, sehingga mekanismenya sulit dijelaskan bila hanya ditinjau dari satu target [6]. Karena itu, *network pharmacology* banyak dipakai untuk meneliti mekanisme tumbuhan obat, yaitu menyaring senyawa aktifnya serta menelusuri target dan jalur biologis yang diduga mendasari khasiatnya [6, 9].

Secara umum, analisis *network pharmacology* mengikuti rangkaian tahap yang sebagian besar sama di berbagai penelitian. Masukannya umumnya berupa tumbuhan obat

yang diteliti beserta penyakit yang menjadi sasaran [6]. Titik awal analisisnya dapat berupa satu spesies tumbuhan, beberapa spesies, suatu ramuan, maupun senyawa tertentu yang sudah dikenal. Sebagai contoh, analisis senyawa kurkumin (senyawa aktif kunyit, *Curcuma longa*) terhadap diabetes melitus tipe 2 meneliti hanya satu senyawa itu saja [19], sedangkan analisis ramuan Danggui Beimu Kushen Wan, suatu formula obat tradisional Tiongkok, terhadap kanker prostat bertolak dari sebuah ramuan [20]. Tahap awal mengumpulkan senyawa aktif yang terkandung pada masukan tersebut, lalu menyaringnya berdasarkan sifat menyerupai obat (*drug-likeness*), atau disebut juga penyaringan ADME (absorpsi, distribusi, metabolisme, ekskresi) [7]. Senyawa terpilih dipetakan ke target proteinnya, kemudian dicari irisan antara target itu dan target yang berkaitan dengan penyakit yang diteliti. Irisan target dirangkai menjadi jaringan *protein-protein interaction* atau PPI, lalu dianalisis topologinya untuk menemukan target *hub*, yaitu protein yang paling banyak terhubung dalam jaringan, serta diuji pengayaannya (*enrichment analysis*), yaitu uji statistik untuk mengetahui jalur biologis mana yang paling banyak diwakili oleh kumpulan target tersebut [8, 9, 6]. Keluarannya mencakup daftar target *hub*, hasil pengayaan jalur, serta jaringan yang menghubungkan senyawa dengan target dan jalur tersebut atau disebut juga jaringan *compound-target-pathway* (C-T-P) [9, 19].

Keluaran tersebut masih berupa prediksi komputasi sehingga, sebagai tahap validasi, analisis *network pharmacology* modern umumnya dilanjutkan dengan pengujian bertingkat. Mayoritas penelitian melakukan hingga tahap pertama, yakni *in-silico* atau bersifat komputasi: penambatan molekuler (*molecular docking*) memperkirakan kekuatan ikatan antara senyawa dan target yang dihasilkan analisis, lalu sebagian studi meneruskannya dengan simulasi dinamika molekuler (*molecular dynamics*) untuk menguji kestabilan ikatan tersebut dari waktu ke waktu [12, 21]. Sebagian melanjutkan penelitian hingga pembuktian biologisnya kemudian dilakukan di laboratorium melalui pengujian *in vitro*, yaitu pada sel atau jaringan, dan *in vivo*, yaitu pada organisme hidup [6, 9, 21]. Dengan demikian, target *hub*, hasil pengayaan jalur, dan jaringan C-T-P yang dihasilkan *network pharmacology* berperan sebagai hipotesis mekanisme yang mengarahkan rangkaian pengujian tersebut, bukan sebagai bukti akhir [22].

2.3 Tahapan Alur Kerja Network Pharmacology

Bagian ini merinci metode, algoritme, dan rumus yang umum dipakai pada setiap tahap alur kerja *network pharmacology*. Pembahasannya dikelompokkan menjadi lima subbab menurut keterkaitan metode antartahap, yaitu penyaringan ADME, identifikasi target dan irisan target, konstruksi dan analisis jaringan PPI, analisis pengayaan, serta validasi lanjutan.

2.3.1 Penyaringan ADME

Penyaringan ADME (absorpsi, distribusi, metabolisme, ekskresi) adalah salah satu tahap awal analisis *network pharmacology* yang menyaring senyawa dari tumbuhan obat terlebih dahulu berdasarkan sifat menyerupai obat, yaitu kemiripan sifat fisikokimia (sifat fisik dan kimia) suatu senyawa dengan sifat obat yang diberikan secara oral [7]. Penyaringan ini menyingkirkan senyawa yang kecil kemungkinannya terserap apabila diberikan secara oral sehingga kandidat yang diteruskan ke pemetaan target lebih terfokus.

Pendekatan yang umum dipakai terbagi atas dua jenis, yaitu aturan berbasis batas sifat yang diwakili Aturan Lima Lipinski dan aturan Veber, serta skor gabungan yang diwakili skor ketersediaan hayati oral dan skor sifat menyerupai obat, dengan SwissADME sebagai alat yang paling sering digunakan untuk penyaringan ADME [21].

Aturan Lima Lipinski (*Lipinski's Rule of Five*) memperkirakan bahwa penyerapan suatu senyawa cenderung buruk apabila senyawa tersebut melanggar dua atau lebih dari empat batas berikut [23].

1. massa molekul, yaitu ukuran besar molekul, tidak lebih dari 500 dalton (Da);
2. koefisien partisi oktanol–air ($\log P$), yaitu ukuran kecenderungan senyawa larut dalam lemak dibanding air, tidak lebih dari 5;
3. jumlah donor ikatan hidrogen, yaitu gugus yang dapat menyumbang atom hidrogen pada ikatan hidrogen seperti gugus N–H dan O–H, tidak lebih dari 5;
4. jumlah akseptor ikatan hidrogen, yaitu atom yang dapat menerima ikatan hidrogen seperti atom nitrogen dan oksigen, tidak lebih dari 10.

Aturan ini disebut “aturan lima” karena keempat batasnya merupakan kelipatan lima [23]. Senyawa dengan paling banyak satu pelanggaran masih dianggap lolos, sedangkan dua pelanggaran atau lebih menandai penyerapan oral yang mungkin buruk.

Aturan Veber melengkapi Lipinski dengan dua sifat lain yang turut menentukan ketersediaan hayati oral (*oral bioavailability*), yaitu seberapa besar bagian senyawa yang mencapai peredaran darah ketika diberikan melalui mulut. Suatu senyawa memenuhi aturan Veber apabila kedua batas berikut dipenuhi sekaligus [24].

1. jumlah ikatan terputar (*rotatable bond*), yaitu ikatan tunggal yang memungkinkan bagian molekul berputar sehingga menentukan kelenturan molekul, tidak lebih dari 10;
2. luas permukaan polar topologis (*topological polar surface area*, TPSA), yaitu luas bagian permukaan molekul yang ditempati atom polar seperti nitrogen dan oksigen, tidak lebih dari $140 \text{ \AA}^2$, dengan $\text{\AA}$ (angstrom) sebagai satuan panjang seukuran atom sehingga $\text{\AA}^2$ merupakan satuan luas.

Berbeda dengan Lipinski yang memberi kelonggaran satu pelanggaran, kedua syarat Veber bersifat wajib sehingga keduanya harus dipenuhi sekaligus [24].

Sebagai alternatif aturan yang memberi keputusan lolos atau tidak lolos, sebagian besar penelitian juga memakai skor gabungan, terutama skor ketersediaan hayati oral (OB) dan skor sifat menyerupai obat (DL) yang dipopulerkan basis data TCMSP [7, 21]. Tidak seperti batas Lipinski dan Veber yang dihitung langsung dari sifat fisikokimia, OB dan DL merupakan skor hasil model komputasi tersendiri dalam basis data tersebut. Senyawa kemudian diteruskan bila skornya melampaui batas yang ditetapkan. Sebagai contoh, penelusuran senyawa aktif suatu ramuan dengan TCMSP memakai batas OB minimal 40% dan DL minimal 0,18 [7]. Pendekatan skor gabungan ini banyak dipakai pada penelitian *network pharmacology* obat tradisional Tiongkok [9, 21].

2.3.2 Identifikasi Target dan Irisan Target

Pemetaan senyawa ke target proteinnya, yaitu protein yang menjadi sasaran kerja senyawa di dalam tubuh, dilakukan melalui dua pendekatan. Pendekatan prediksi memperkirakan target suatu senyawa dari kemiripan strukturnya dengan senyawa yang targetnya sudah diketahui. Alat yang umum dipakai untuk ini adalah SwissTargetPrediction [21], yang kerap dipadukan dengan alat sejenis seperti STITCH dan SuperPred. Sebagai contoh, analisis kurkumin terhadap diabetes melitus tipe 2 memprediksi target senyawanya dengan memadukan SwissTargetPrediction, STITCH, dan beberapa alat lain [19]. Sebaliknya, pendekatan eksperimental mengambil target dari pengukuran bioaktivitas yang sudah tercatat pada basis data, misalnya ChEMBL dan PubChem BioAssay [25, 26].

Secara terpisah, protein yang berkaitan dengan penyakit yang diteliti dikumpulkan dari basis data target–penyakit. Sumber yang umum dipakai antara lain GeneCards, OMIM, dan DisGeNET [21].

Target yang menjadi pusat perhatian adalah irisan kedua kelompok tersebut, yaitu protein yang sekaligus dikenai senyawa dan berkaitan dengan penyakit. Irisan ini divisualisasikan dengan diagram Venn dan menjadi kumpulan target kandidat yang dibawa ke tahap berikutnya [6].

2.3.3 Konstruksi dan Analisis Jaringan PPI

Analisis pada tahap ini berlandaskan graf, yaitu kumpulan simpul (*node*) yang dihubungkan oleh sisi (*edge*) [27]. Banyaknya sisi yang terhubung pada sebuah simpul disebut derajat (*degree*) simpul tersebut, sedangkan lintasan terpendek antara sepasang simpul adalah rangkaian sisi paling singkat yang menghubungkan keduanya [28]. Keterhubungan antarsimpul dapat dinyatakan sebagai matriks ketetanggaan (*adjacency*

matrix), yaitu matriks yang unsurnya bernilai 1 bila dua simpul terhubung dan bernilai 0 bila tidak [29]. Graf yang sisinya tidak memiliki arah disebut graf tak berarah (*undirected*), dan pada graf demikian hubungan antara dua simpul bersifat timbal balik [27].

Target kandidat yang diperoleh pada tahap sebelumnya dirangkai menjadi jaringan interaksi protein–protein, disebut juga jaringan PPI (*Protein–Protein Interaction*), agar hubungan antartarget dapat dianalisis secara menyeluruh [6]. Pada konteks ini, simpul mewakili entitas seperti senyawa, protein, atau penyakit, sedangkan sisi mewakili hubungan di antaranya, misalnya interaksi antarprotein [22]. Interaksi antarprotein tidak memiliki arah sehingga jaringan PPI berupa graf tak berarah yang simpulnya berupa protein dan sisinya berupa interaksi antarprotein [27]. Struktur keterhubungan inilah yang dianalisis untuk menemukan target *hub* pada tahap analisis topologi.

Data interaksi tersebut diambil dari basis data interaksi protein, dan yang paling umum dipakai adalah STRING, yaitu basis data yang menghimpun interaksi protein–protein hasil percobaan maupun prediksi komputasi [30, 21]. Jaringan yang terbentuk lalu divisualisasikan dan dianalisis menggunakan perangkat lunak analisis dan visualisasi jaringan, yang paling umum adalah Cytoscape [9, 21]. Sebagai contoh, analisis kurkumin terhadap diabetes melitus tipe 2 membangun jaringan dari target bersamanya melalui STRING dan memvisualisasikannya pada Cytoscape [19].

Setelah jaringan terbentuk, analisis topologi mengukur kepentingan tiap simpul untuk menemukan target *hub*, yaitu protein sentral dalam jaringan dengan konektivitas [6]. Kepentingan ini dinilai melalui beberapa ukuran sentralitas (*centrality*); empat yang umum dipakai adalah *degree*, *betweenness*, *closeness*, dan *eigenvector* [31].

Sentralitas *degree* menilai kepentingan simpul dari banyaknya sisi yang terhubung padanya, sehingga protein yang berinteraksi langsung dengan lebih banyak target lain memperoleh nilai lebih tinggi [28]. Nilai ini dinormalkan terhadap jumlah keterhubungan maksimum yang mungkin, seperti ditunjukkan Rumus 2.1.

$$C_D(v) = \frac{k_v}{N-1} \quad (2.1)$$

dengan $C_D(v)$ sentralitas *degree* simpul v , k_v banyaknya sisi yang terhubung pada simpul v , dan N jumlah simpul dalam jaringan.

Sentralitas *betweenness* mengukur seberapa sering suatu simpul berada pada lintasan terpendek antarpasangan simpul lain, sehingga simpul dengan nilai tinggi berperan sebagai penghubung lalu lintas informasi dalam jaringan [28]. Nilainya dihitung seperti pada Rumus 2.2.

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

dengan $C_B(v)$ sentralitas *betweenness* simpul v , σ_{st} banyaknya lintasan terpendek dari

simpul s ke simpul t , dan $\sigma_{st}(v)$ banyaknya lintasan tersebut yang melewati v ; penjumlahan dilakukan atas semua pasangan simpul s dan t yang berbeda dari v .

Sentralitas *closeness* mengukur seberapa dekat suatu simpul dengan seluruh simpul lain, dihitung sebagai kebalikan dari total jarak lintasan terpendeknya ke semua simpul lain, sehingga simpul yang dekat dengan banyak simpul memperoleh nilai lebih tinggi [28]. Perhitungannya ditunjukkan Rumus 2.3.

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

dengan $C_C(v)$ sentralitas *closeness* simpul v , $d(v, u)$ panjang lintasan terpendek antara simpul v dan simpul u , serta N jumlah simpul seperti sebelumnya.

Sentralitas *eigenvector* menilai kepentingan simpul tidak hanya dari jumlah tetangganya, tetapi juga dari kepentingan tetangga tersebut, sehingga simpul yang terhubung ke simpul-simpul penting memperoleh nilai lebih tinggi [29]. Penilaian tersebut bersandar pada vektor eigen (*eigenvector*) beserta nilai eigen (*eigenvalue*) dari matriks ketetanggaan jaringan. Sebagian besar vektor berubah arah ketika dikalikan dengan sebuah matriks, tetapi terdapat vektor istimewa yang arahnya tetap sama dan hanya berubah panjangnya; vektor demikian disebut vektor eigen dari matriks tersebut, sedangkan besar perubahan panjangnya disebut nilai eigen [32]. Perhitungannya ditunjukkan Rumus 2.4.

$$C_E(v) = \frac{1}{\lambda} \sum_u a_{vu} C_E(u) \quad (2.4)$$

dengan $C_E(v)$ sentralitas *eigenvector* simpul v , a_{vu} unsur matriks ketetanggaan (*adjacency matrix*) yang bernilai 1 bila simpul v dan u terhubung serta 0 bila tidak, dan λ nilai eigen terbesar dari matriks tersebut.

Berdasarkan ukuran-ukuran tersebut, simpul dengan sentralitas *degree* tinggi disebut *hub*, sedangkan simpul dengan sentralitas *betweenness* tinggi disebut *bottleneck*, yaitu titik penghubung yang dilalui banyak lintasan terpendek sehingga ditentukan langsung oleh sentralitas keantaraan tanpa rumus tersendiri [33]. Protein yang sekaligus menjadi *hub* dan *bottleneck* cenderung paling berpengaruh dan paling sering merupakan protein esensial, yaitu protein yang penting bagi kelangsungan hidup sel [33].

Analisis topologi jaringan PPI umumnya dilakukan melalui Cytoscape [21]. CytoHubba, sebuah *plugin* Cytoscape untuk analisis topologi, menyediakan sebelas metode pemeringkatan target *hub*, dan metode unggulannya, *Maximal Clique Centrality* (MCC), terbukti paling akurat memprediksi protein esensial dibandingkan sepuluh metode lainnya [34]. MCC suatu simpul dihitung dari klik maksimal yang memuatnya. Klik adalah kelompok simpul yang setiap pasangannya saling terhubung, sedangkan klik maksimal adalah klik yang tidak dapat ditambah simpul lain tanpa kehilangan sifat tersebut. Nilainya dirumuskan pada Rumus 2.5.

$$MCC(v) = \sum_{C \in S(v)} (|C| - 1)! \quad (2.5)$$

dengan $S(v)$ kumpulan klik maksimal yang memuat simpul v dan $|C|$ jumlah simpul dalam klik C ; bila tetangga simpul v tidak saling terhubung, nilai MCC-nya sama dengan derajatnya [34].

Dari hasil pemaduan tersebut, sejumlah simpul teratas (*top-N*) diambil sebagai *hub*, misalnya sepuluh besar [34]. Sebagai contoh, analisis kurkumin terhadap penyakit diabetes melitus tipe 2 memberikan peringkat pada target *hub* menggunakan MCC [19].

2.3.4 Analisis Pengayaan

Analisis pengayaan (*enrichment analysis*) digunakan untuk menjelaskan makna biologis dari sekumpulan gen atau protein dengan membandingkan daftar tersebut terhadap himpunan gen yang telah didefinisikan dalam basis data anotasi fungsi atau jalur biologis [35]. Pada konteks *network pharmacology*, daftar tersebut berasal dari target yang beririsan, gen hub, atau protein yang dianggap relevan terhadap hubungan senyawa dan penyakit [35].

Dua basis data anotasi yang paling umum dipakai adalah GO dan KEGG [9, 21]. GO (Gene Ontology) adalah sistem klasifikasi baku yang menggolongkan fungsi gen beserta produknya ke dalam tiga aspek, yaitu proses biologis (*biological process*), komponen seluler (*cellular component*), dan fungsi molekuler (*molecular function*) [36]. Ketiga aspek tersebut berturut-turut menerangkan rangkaian kegiatan biologis yang diikuti gen, lokasi gen bekerja di dalam sel, dan aktivitas biokimia yang dijalankannya. Sementara itu, KEGG (Kyoto Encyclopedia of Genes and Genomes) memetakan gen ke jalur biologis, yaitu rangkaian reaksi atau interaksi molekul yang menjalankan suatu fungsi seluler, seperti jalur metabolisme dan jalur pensinyalan [37]. Anotasi GO menjelaskan fungsi tiap target secara terpisah, sedangkan jalur KEGG menempatkan target pada peta interaksi antarmolekul yang lebih luas.

Keterwakilan berlebih pada tiap kategori, baik satu istilah GO maupun satu jalur KEGG, diuji melalui analisis keterwakilan berlebih (*over-representation analysis*, ORA). Uji ini membandingkan banyaknya target yang teranotasi pada suatu kategori dengan banyaknya yang diharapkan muncul secara kebetulan, dan peluangnya dihitung memakai uji hipergeometrik seperti pada Rumus 2.6 [38].

$$P(X \geq k) = \sum_{i=k}^{\min(n,K)} \frac{\binom{K}{i} \binom{N-K}{n-i}}{\binom{N}{n}} \quad (2.6)$$

dengan N jumlah seluruh gen pada himpunan latar (*background*), K jumlah gen latar yang teranotasi pada kategori yang diuji, n jumlah target yang dianalisis, k jumlah target yang

teranotasi pada kategori tersebut, dan $P(X \geq k)$ peluang memperoleh paling sedikit k target teranotasi secara kebetulan.

Peluang inilah yang menjadi nilai-p (*p-value*) kategori tersebut: makin kecil nilainya, makin kecil kemungkinan keterwakilan berlebih itu terjadi secara acak sehingga kategori yang bersangkutan makin layak dianggap signifikan.

Karena satu analisis menguji ribuan kategori sekaligus, sebagian kategori dapat tampak signifikan semata-mata akibat banyaknya pengujian. Persoalan pengujian berganda (*multiple testing*) ini diatasi dengan mengoreksi nilai-p berdasarkan laju penemuan palsu (*false discovery rate*, FDR), yaitu proporsi temuan signifikan yang sebenarnya merupakan positif palsu. Prosedur koreksi FDR yang banyak dipakai adalah Benjamini–Hochberg [9], yang mengurutkan seluruh nilai-p secara menaik lalu menetapkan batas signifikansi seperti pada Rumus 2.7 [39].

$$k = \max \left\{ i : p_{(i)} \leq \frac{i}{m} q \right\} \quad (2.7)$$

dengan $p_{(i)}$ nilai-p ke- i setelah seluruh nilai-p diurutkan menaik, m jumlah seluruh kategori yang diuji, dan q laju penemuan palsu yang ditargetkan; seluruh kategori berperingkat sampai k dinyatakan signifikan.

Analisis pengayaan dijalankan menggunakan alat khusus yang memuat anotasi GO dan KEGG. DAVID (Database for Annotation, Visualization and Integrated Discovery) menyediakan anotasi fungsi sekaligus pengayaan GO dan KEGG bagi sekumpulan gen [40]. Metascape memadukan banyak basis data anotasi dan menyajikan pengayaan GO serta KEGG dalam satu antarmuka [41]. g:Profiler menjalankan pengayaan fungsi sekaligus mengonversi penanda gen antarbasis data [42]. Sebagai contoh, analisis kurkumin di atas menjalankan pengayaan GO dan KEGG atas targetnya menggunakan DAVID dengan batas FDR di bawah 0,01, lalu memperoleh 615 istilah GO dan 168 jalur KEGG yang signifikan [19].

Secara keseluruhan, analisis pengayaan dapat dikelompokkan ke dalam beberapa pendekatan [35]. Perbedaan utama antarmetode terletak pada bentuk data masukan dan informasi tambahan yang digunakan. Metode berbasis tumpang tindih hanya menguji irisan antara daftar gen dan himpunan gen rujukan, sedangkan pendekatan berbasis topologi atau jaringan mempertimbangkan hubungan antar gen atau struktur jalur. Oleh karena itu, hasil pengayaan perlu dibaca sebagai petunjuk biologis yang bergantung pada metode, basis data, dan parameter yang digunakan, bukan sebagai bukti final bahwa suatu jalur pasti menjadi mekanisme utama [35].

2.3.5 Validasi Lanjutan

Target *hub* beserta jalur yang diperkaya pada tahap-tahap sebelumnya masih berupa hipotesis mekanisme hasil prediksi komputasi sehingga perlu diuji sebelum dijadikan simpulan [12]. Pengujian ini disusun bertingkat, dari pemeriksaan komputasi yang cepat dan murah menuju pembuktian laboratorium yang lebih mahal tetapi lebih meyakinkan, dengan tiap tingkat menguji aspek yang berbeda dari hipotesis tersebut [12].

Pertama, secara *in silico*, umumnya dilakukan *molecular docking*, yang menguji apakah senyawa dapat terikat pada targetnya secara struktural [12, 6]. *Molecular docking* hanya memberi gambaran statis pada satu saat, sehingga sebagian penelitian meneruskannya dengan simulasi *molecular dynamics* [21]. Setelah pemeriksaan komputasi tersebut, pembuktian biologis dilakukan di laboratorium melalui pengujian *in vitro* pada sel atau jaringan dan *in vivo* pada organisme hidup, yang menguji apakah efek yang diprediksi benar-benar terjadi dan bukan sekadar mungkin secara struktural [6, 9]. Sebagai contoh, analisis kurkumin di atas memvalidasi ikatan kurkumin pada target *hub*-nya melalui *molecular docking*, lalu mengujinya secara *in vivo* [19].

2.4 Prinsip FAIR

Prinsip FAIR adalah pedoman penatakelolaan data ilmiah agar data dapat ditemukan (*Findable*), diakses (*Accessible*), dioperasikan bersama (*Interoperable*), dan digunakan ulang (*Reusable*) [13]. Aspek *Findable* mensyaratkan setiap data memiliki pengidentifikasi unik dan metadata yang memadai; aspek *Accessible* mensyaratkan data dapat diambil melalui protokol baku yang terbuka; aspek *Interoperable* memerlukan keseragaman representasi dan pengidentifikasi agar data dapat digabungkan lintas sumber; dan aspek *Reusable* memerlukan kelengkapan deskripsi beserta asal-usul data agar dapat digunakan ulang [13].

2.5 Validasi terhadap Studi Pembandingan

Ketepatan keluaran suatu platform komputasi perlu dibuktikan sebelum hasilnya dipakai lebih lanjut. Panduan evaluasi metode *network pharmacology* menyatakan bahwa penelitian yang membangun algoritme, basis data, maupun platform analisis komputasi perlu memverifikasi hasil prediksi yang dikeluarkannya, dan menganjurkan evaluasi kesesuaian terhadap studi lain sebagai salah satu caranya [43]. Cara ini sejalan dengan penilaian metode komputasi secara umum, yaitu memeriksa apakah suatu metode mampu memunculkan kembali temuan yang telah dilaporkan sebelumnya, meskipun penilaian semacam itu bergantung pada keabsahan hasil penelitian yang dijadikan acuan [44].

Keluaran platform untuk masukan yang sama dibandingkan terhadap keluaran studi acuan, lalu tingkat kesesuaian di antara keduanya diukur.

Ukuran yang dipakai untuk menilai kesesuaian tersebut adalah *recall*, yang juga disebut sensitivitas, yaitu rasio antara jumlah informasi yang terdeteksi dengan benar dan jumlah seluruh informasi relevan yang ada pada acuan [43]. Panduan yang sama menempatkan *recall* sebagai indikator bagi keutuhan data, yakni sejauh mana data yang berhasil dihimpun mencakup keseluruhan data relevan yang tersedia [43]. Perhitungannya ditunjukkan Rumus 2.8.

$$R = \frac{|H \cap A|}{|A|} \times 100\% \quad (2.8)$$

dengan H himpunan keluaran sistem, A himpunan keluaran acuan, $|H \cap A|$ banyaknya anggota acuan yang berhasil dimunculkan kembali oleh sistem, dan $|A|$ banyaknya seluruh anggota acuan.

Nilai R yang tinggi berarti sistem berhasil memunculkan kembali sebagian besar keluaran acuan, sedangkan nilai yang rendah menandakan sebagian keluaran acuan tidak terjangkau oleh sistem. Karena penyebutnya adalah jumlah anggota acuan, ukuran ini menilai kelengkapan keluaran sistem terhadap acuan, bukan ketepatan setiap anggota keluaran yang dihasilkan sistem tersebut.

2.6 System Usability Scale

System Usability Scale (SUS) adalah instrumen kuesioner yang digunakan untuk memperoleh penilaian subjektif terhadap kegunaan suatu sistem secara menyeluruh. Instrumen ini terdiri dari sepuluh pernyataan dan dirancang sebagai metode evaluasi yang ringkas untuk menghasilkan satu skor komposit kegunaan sistem [45].

SUS merupakan bentuk penerapan skala Likert. Skala Likert merupakan metode psikometrik untuk mengukur sikap, persepsi, atau preferensi responden terhadap suatu objek melalui tingkat persetujuan terhadap serangkaian pernyataan [46]. Pada SUS, format tersebut diterapkan melalui lima pilihan jawaban yang bergerak dari sangat tidak setuju sampai sangat setuju. Makna setiap poin skala ditunjukkan pada Tabel 2.1.

Tabel 2.1. Makna skala lima poin pada *System Usability Scale* (SUS)

Skor	Makna
1	Sangat tidak setuju
2	Tidak setuju
3	Netral
4	Setuju
5	Sangat setuju

Pernyataan dalam SUS disusun dengan nada positif dan negatif secara berselang-seling. Penyusunan ini bertujuan mengurangi kecenderungan responden menjawab secara mekanis tanpa membaca setiap pernyataan dengan cermat [45]. Kesepuluh pernyataan SUS ditunjukkan secara verbatim pada Tabel 2.2.

Tabel 2.2. Sepuluh pernyataan pada kuesioner *System Usability Scale* (SUS)

No.	Pernyataan
1	I think that I would like to use this system frequently.
2	I found the system unnecessarily complex.
3	I thought the system was easy to use.
4	I think that I would need the support of a technical person to be able to use this system.
5	I found the various functions in this system were well integrated.
6	I thought there was too much inconsistency in this system.
7	I would imagine that most people would learn to use this system very quickly.
8	I found the system very cumbersome to use.
9	I felt very confident using the system.
10	I needed to learn a lot of things before I could get going with this system.

Pada pelaksanaan evaluasi, responden mengisi SUS setelah menggunakan sistem yang dinilai dan sebelum sesi diskusi atau pengarahan lanjutan. Responden diminta memberi jawaban berdasarkan kesan langsung terhadap setiap pernyataan. Apabila responden tidak dapat menentukan jawaban untuk suatu pernyataan, titik tengah skala dapat dipilih [45].

Karena SUS bernada positif–negatif secara berselang-seling, skor akhir SUS dihitung dari kontribusi skor setiap pernyataan, bukan dari rata-rata langsung seluruh

jawaban. Untuk pernyataan bernomor ganjil, kontribusi dihitung dengan mengurangi skor jawaban sebesar satu. Untuk pernyataan bernomor genap, kontribusi dihitung dengan mengurangkan skor jawaban dari lima. Setiap kontribusi berada pada rentang 0 sampai 4. Jumlah kontribusi dari sepuluh pernyataan kemudian dikalikan dengan 2,5 sehingga menghasilkan skor akhir pada rentang 0 sampai 100, seperti ditunjukkan pada Rumus 2.9 [45].

$$SUS = 2,5 \times \left[\sum_{i \text{ ganjil}} (x_i - 1) + \sum_{i \text{ genap}} (5 - x_i) \right] \quad (2.9)$$

dengan x_i menyatakan skor jawaban pada pernyataan ke- i . Skor SUS yang lebih tinggi menunjukkan persepsi kegunaan yang lebih baik. Namun, skor tersebut perlu dipahami sebagai ukuran komposit terhadap kegunaan sistem secara keseluruhan. Oleh karena itu, skor masing-masing butir tidak ditafsirkan secara terpisah sebagai ukuran kegunaan yang berdiri sendiri [45].

Penggunaan SUS sebagai instrumen evaluasi kegunaan juga didukung oleh evaluasi empiris yang menghimpun lebih dari 2.300 kuesioner SUS dari lebih dari 200 studi. Penelitian tersebut melaporkan reliabilitas internal yang tinggi dan menunjukkan bahwa skor SUS peka terhadap perbedaan jenis antarmuka, seperti web, antarmuka grafis, dan sistem interaktif berbasis suara [47].

Jumlah responden pada pengukuran SUS berkaitan dengan seberapa stabil skor yang dihasilkannya. Perbandingan lima kuesioner kegunaan terhadap 123 partisipan menunjukkan bahwa kesimpulan SUS mulai stabil pada jumlah sampel sekitar 12, yaitu ketika SUS memberikan kesimpulan yang sama dengan data penuh pada sekurang-kurangnya 90% pengambilan sampel [48]. Sebagai gambaran praktiknya, meta-analisis terhadap 117 skor SUS dari aplikasi kesehatan digital mencatat rata-rata jumlah responden sebanyak 6 orang dengan rentang 2 sampai 31 [49]. Rata-rata tersebut berada di bawah jumlah sampel yang membuat kesimpulan SUS mulai stabil, sehingga sampel kecil merupakan praktik yang umum ditemukan pada studi SUS yang dipublikasikan.

Salah satu acuan interpretasi skor akhir SUS adalah Sauro–Lewis *Curved Grading Scale* (CGS) [50]. Skala ini memetakan skor SUS 0 sampai 100 ke dalam nilai huruf dan rentang persentil berdasarkan kumpulan data empiris dari ratusan studi. Dalam skala tersebut, skor 68 berada pada pusat kategori C dan digunakan sebagai acuan rata-rata empiris, sedangkan skor yang lebih tinggi menunjukkan posisi yang lebih baik terhadap data pembandingan. Rincian interpretasi Sauro-Lewis CGS ditunjukkan pada Tabel 2.3.

Tabel 2.3. Interpretasi skor SUS berdasarkan Sauro–Lewis *Curved Grading Scale* (CGS)

Rentang Skor SUS	Nilai	Rentang Persentil
84,1–100	A+	96–100
80,8–84,0	A	90–95
78,9–80,7	A-	85–89
77,2–78,8	B+	80–84
74,1–77,1	B	70–79
72,6–74,0	B-	65–69
71,1–72,5	C+	60–64
65,0–71,0	C	41–59
62,7–64,9	C-	35–40
51,7–62,6	D	15–34
0,0–51,6	F	0–14

Berdasarkan acuan tersebut, hasil SUS ditafsirkan dengan melihat posisi skor rata-rata responden terhadap rentang nilai pada Tabel 2.3. Misalnya, skor pada kategori C menunjukkan hasil yang berada di sekitar rata-rata empiris, sedangkan skor pada kategori B atau A menunjukkan hasil di atas rata-rata. Sebaliknya, skor pada kategori D atau F menunjukkan bahwa persepsi kegunaan sistem masih berada di bawah acuan umum dan perlu diperhatikan pada pembahasan hasil. Dengan demikian, hasil SUS tidak hanya dilaporkan sebagai angka, tetapi juga dijelaskan sebagai posisi relatif terhadap acuan empiris.

2.7 Platform Terkait

Sejumlah platform *network pharmacology* telah berupaya mengotomatiskan seluruh tahapan analisisnya. Platform-platform tersebut dapat dikelompokkan menurut cakupan dan fokusnya.

2.7.1 Platform Berbasis TCM

TCMSP, BATMAN-TCM, dan TCMNPAS adalah beberapa platform *network pharmacology* yang terfokus pada cakupan obat tradisional Tiongkok atau Traditional Chinese Medicine (TCM). TCMSP (TCM Systems Pharmacology) menyediakan basis data sistem farmakologi berbasis web yang memuat kandungan kimia, sifat ADME, target, penyakit, dan jaringannya untuk bahan TCM [7]. BATMAN-TCM (Bioinformatics Analysis Tool for Molecular Mechanism of TCM) memprediksi target ramuan TCM serta

menyediakan analisis pengayaan dan visualisasi jaringan yang menautkan ramuan, target, dan jalur biologis [14, 51]. TCMNPAS (TCM Network Formulaology and Pharmacology Analysis System) memadukan *network formulaology* dan *network pharmacology* dengan cakupan paling lengkap di antara ketiganya, mulai dari penyaringan sifat menyerupai obat (QED, Lipinski, dan Veber), penambatan molekuler, hingga pengayaan jalur [10].

2.7.2 Platform Analisis Jaringan

Platform lain berfokus pada tahap analisis jaringan dan mengolah data yang sudah disediakan pengguna, seperti NeXus dan SmartGraph. NeXus adalah aplikasi yang tersedia dalam bentuk aplikasi *command-line* berbasis Python serta web yang membangun jaringan, menghitung sentralitas, mendeteksi komunitas dengan algoritme Louvain, serta menjalankan pengayaan multimetode (ORA, GSEA, dan GSVA) [11]. SmartGraph menyediakan aplikasi web berbasis basis data graf untuk menelusuri relasi senyawa–target dan memprediksi bioaktivitas berbasis kerangka molekul [15].

2.7.3 Platform di Indonesia

Di Indonesia, sudah ada beberapa platform *network pharmacology*, bahkan yang dikhususkan untuk tumbuhan obat Indonesia. IJAH (Indonesia Jamu–Herbs) Analytics adalah aplikasi web yang menelusuri keterhubungan antara tumbuhan, senyawa, protein, dan penyakit [16]. IJAH mencakup tahapan dari tumbuhan sampai penyakit, tetapi seluruh penelusurannya hanya menggunakan basis data bawaan yang sudah dikurasi sebelumnya dan dijalankan satu kali hingga keluaran akhir, tanpa penyaringan ADME maupun peninjauan antar-tahap. iPrimeTarget, dari kelompok yang sama, adalah platform analisis jaringan protein yang memprioritaskan target dari penyakit [17, 18]; platform ini bekerja pada lapis protein sehingga melengkapi tahap analisis dan tidak mencakup pengumpulan data di awal.

U N I V E R S I T A S
M U L T I M E D I A
N U S A N T A R A