



AI-Enabled Neuroprotective Natural Products: Network Biology, Blood–Brain Barrier Prediction, Target Selectivity, and Safety-Aware Discovery Logic

Fatima El Idrissi^{1*}, Samira Bennani¹, Youssef Benali²

¹Department of AI for Argan Oil and Saffron Bioactive Compounds, Faculty of Pharmacy, University of Marrakech, Marrakech, Morocco

²Department of Cheminformatics for Tyrosinase Inhibitors, Faculty of Pharmacy, University of Fez, Fez, Morocco.

ABSTRACT

Natural products provide structurally diverse starting points for neuropharmacology research, but computational evidence of target association or pathway relevance is insufficient to establish neuroprotection, central nervous system exposure, safety, or therapeutic utility. This original computational framework article proposes a structured discovery logic for artificial intelligence-enabled prioritization of natural products and natural product-inspired candidates with potential neuroprotective relevance. The framework begins with verified botanical, phytochemical, structural, stereochemical, bioactivity, assay, and provenance information and then connects neuroprotective hypotheses to disease or phenotype context. Network biology is used to organize compound–target–pathway relationships, while target-confidence, target-selectivity, target-engagement, and off-target analyses constrain mechanistic interpretation. Blood–brain barrier prediction is treated as one component of a wider CNS-relevance assessment that also considers passive permeability, transporter and efflux behavior, metabolic stability, ADME evidence, applicability domain, out-of-distribution status, and predictive uncertainty. Safety-aware discovery incorporates neurotoxicity, systemic toxicity, safety pharmacology, metabolism-related concerns, and herb–drug or drug–drug interaction risk. Advancement occurs only through predefined validation gates and expert neuropharmacology and toxicology review. The principal contribution is an evidence-linked, uncertainty-aware, and claim-bounded framework that supports research prioritization while explicitly preventing computational predictions from being presented as validated neuroprotection, demonstrated brain exposure, established safety, clinical utility, or treatment guidance.

Key Words: Neuroprotective natural products, AI drug discovery, Network biology, Target selectivity, Blood–brain barrier prediction, Safety-aware discovery

eIJPPR 2026; 16(3):117-134

HOW TO CITE THIS ARTICLE: Idrissi FE, Bennani S, Benali Y. AI-Enabled Neuroprotective Natural Products: Network Biology, Blood–Brain Barrier Prediction, Target Selectivity, and Safety-Aware Discovery Logic. Int J Pharm Phytopharmacol Res. 2026;16(3):117-134. <https://doi.org/10.51847/yFeUMz9hbG>

INTRODUCTION

Artificial intelligence-enabled natural product discovery is emerging at the intersection of cheminformatics, pharmacology, bioinformatics, systems biology, and data-intensive molecular science. Natural products occupy broad and structurally distinctive regions of chemical space, but their computational use is complicated by

variable source attribution, incomplete stereochemical records, heterogeneous assay evidence, uncertain constituent identity, and uneven target annotation. Artificial intelligence can help organize these data, identify patterns across chemical and biological evidence, propose target or property hypotheses, and prioritize experiments, provided that predictions remain traceable to curated

Corresponding author: Fatima El Idrissi

Address: Department of AI for Argan Oil and Saffron Bioactive Compounds, Faculty of Pharmacy, University of Marrakech, Marrakech, Morocco.

E-mail: ✉ fatima.idrissi@uca.ma

Received: 14 March 2026; **Revised:** 14 June 2026; **Accepted:** 17 June 2026

This is an **open access** journal, and articles are distributed under the terms of the Creative Commons Attribution-Non Commercial-ShareAlike 4.0 License, which allows others to remix, tweak, and build upon the work non-commercially, as long as appropriate credit is given and the new creations are licensed under the identical terms.



inputs and explicit scientific questions [1]. Contemporary natural product discovery therefore requires more than applying an algorithm to a compound collection; it requires coordinated data curation, model documentation, uncertainty assessment, and development-aware validation planning [2].

These requirements become more stringent in neuropharmacology. Central nervous system drug discovery must contend with biologically heterogeneous disorders, cell-type-specific mechanisms, complex disease progression, limited accessibility of neural tissue, blood–brain barrier constraints, and the possibility that molecular activity outside the brain may not translate into meaningful neural exposure or benefit. Machine learning may assist target discovery, molecular prioritization, pharmacokinetic prediction, and experimental design, but CNS-focused applications remain highly sensitive to endpoint definition, dataset composition, model generalizability, and prospective validation [3]. The convergence of artificial intelligence, natural product research, and neurodegeneration creates a promising hypothesis-generation environment, yet it does not permit predicted pathway association, molecular similarity, docking compatibility, or BBB classification to be interpreted as proof of neuroprotective efficacy [4].

Predicted activity alone is insufficient because a computationally attractive compound may have uncertain identity, weak disease-context relevance, poor target selectivity, substantial off-target liability, limited metabolic stability, active efflux, inadequate unbound brain exposure, or unrecognized neurotoxicity. Conversely, a compound with modest predicted affinity toward one target may remain scientifically interesting when supported by reproducible phenotypic evidence, coherent network positioning, relevant pathway perturbation, acceptable exposure plausibility, and a clearly defined validation strategy. Candidate assessment must therefore distinguish molecular association from target engagement, network proximity from causal mechanism, BBB prediction from measured permeability, permeability from brain exposure, and the absence of a predicted toxicity alert from demonstrated safety.

The objective of this article is to propose an original computational framework for AI-enabled discovery of neuroprotective natural products that integrates natural product identity, neuroprotective hypothesis classification, network biology, target annotation, target confidence, target selectivity, off-target risk, BBB prediction, CNS relevance, ADME and ADMET assessment, neurotoxicity and systemic toxicity prediction, interaction-risk evaluation, model reliability, validation gates, and expert oversight. The framework is intended to organize evidence and prioritize research rather than generate treatment

recommendations. Its central principle is that a candidate should advance only when its chemical identity, biological context, mechanistic plausibility, CNS-exposure hypothesis, safety profile, uncertainty, and next validation step can be examined together without converting computational outputs into unsupported therapeutic claims.

Neuroprotective natural products

Natural products are valuable sources of molecular diversity and biological hypotheses, but the label “neuroprotective natural product” should be used cautiously. In an early computational discovery context, the more defensible description is a natural product or natural product-inspired candidate associated with a neuroprotective hypothesis. Such a hypothesis may concern attenuation of neuroinflammatory signaling, modulation of oxidative stress responses, preservation of mitochondrial function, alteration of protein homeostasis, support of neuronal survival, or maintenance of synaptic function. Reviews of natural products in degenerative brain disorders illustrate the breadth of these proposed mechanisms while also showing that the underlying evidence varies substantially across compounds, assays, experimental systems, and disease contexts [5]. Translational interpretation remains constrained by uncertain pharmacokinetics, bioavailability, exposure, toxicity, standardization, and the frequent absence of sufficiently discriminating mechanistic validation [6].

A credible evidence foundation begins with botanical source, taxonomic identification, extraction or preparation context, constituent assignment, chemical structure, stereochemistry, purity, and assay metadata. These elements are not peripheral documentation; they determine whether the biological evidence can be attributed to a defined molecular entity. Neuroprotective hypothesis generation should then specify the disease or phenotype context, biological system, measured endpoint, direction of effect, concentration or exposure conditions where reported, and whether the evidence derives from biochemical, cellular, organoid, animal, omics, or other experimental observations. For example, evidence that a constituent modulates Nrf2-associated signaling may support an oxidative-stress or neuroinflammation hypothesis, but pathway association alone does not establish neuronal protection, target engagement, or therapeutic efficacy [7]. Likewise, reports of plant-derived bioactives with potential relevance to Alzheimer’s- or Parkinson’s-related processes remain predominantly sources of preclinical hypotheses rather than validated preventive or therapeutic conclusions [8].

Mechanistic organization should preserve distinctions among broad stress-response pathways.

Neuroinflammation relevance may involve microglial activation, innate immune signaling, cytokine regulation, inflammasome-related processes, or interactions between neural and immune cell types. Oxidative stress relevance may involve reactive oxygen species, endogenous antioxidant systems, redox-sensitive transcriptional programs, or damage-associated signaling. Mitochondrial relevance may include respiration, membrane potential, mitophagy, calcium handling, metabolic adaptation, or apoptosis-related processes. Protein aggregation relevance should be included only when supported by disease-specific aggregation, clearance, proteostasis, or toxicity evidence. Similarly, neuronal-survival or synaptic relevance requires neuron-specific functional evidence rather than generic cytoprotection in unrelated cell types. Safety caution must be introduced at the same stage as efficacy-related hypothesis generation. Natural origin does not imply low toxicity, predictable composition, adequate brain exposure, absence of pharmacological interactions, or compatibility with chronic administration. A constituent may produce desirable pathway modulation and still present off-target, metabolic, neuroactive, genotoxic, cardiovascular, hepatic, or interaction-related concerns. The framework therefore treats each proposed neuroprotective effect as one side of an evidence balance that must also include identity confidence, assay quality, exposure plausibility, toxicity signals, model uncertainty, and validation feasibility. This approach prevents attractive mechanistic narratives from advancing independently of chemical, pharmacokinetic, and safety constraints.

Network biology and target selectivity

Network biology provides a structured means of relating natural product constituents to disease-associated proteins, pathways, phenotypes, and molecular modules. Its value depends on the quality of the chemical and biological entities entering the network. Analytical phytochemical characterization can help ensure that compound–target–pathway relationships are linked to constituents that were actually detected or otherwise credibly assigned, rather than to an unrestricted list of database-retrieved molecules [9]. Network-medicine approaches can then assess whether herb- or compound-associated targets occupy relevant positions within disease or symptom modules, but network proximity remains conditional on interactome completeness, target annotation accuracy, phenotype definition, and the resolution at which mixtures are represented [10].

Target annotation should combine measured bioactivity where available with carefully evaluated computational predictions. Transfer-learning strategies may be useful for natural products because their chemical structures and

target labels are often underrepresented in conventional training sets, yet model output should remain accompanied by confidence, chemical-domain, and evidence-provenance information [11]. Ensemble target-fishing methods can provide ranked target hypotheses and may reduce dependence on a single algorithmic representation, but they do not establish physical binding, functional modulation, intracellular access, or target engagement in a disease-relevant neural system [12]. The framework therefore separates target annotation, target confidence, target engagement hypotheses, and target validation as distinct evidentiary stages.

Disease-network confidence can be strengthened by integrating protein–protein interaction information with transcriptomic, proteomic, metabolomic, genetic, or other omics evidence. Machine-learning and multiomics approaches can help identify proteins or modules associated with neurodegenerative disease biology, although prioritized proteins remain influenced by data availability, tissue composition, batch effects, disease-stage heterogeneity, and annotation bias [13]. Network-derived targets and pathways should consequently be viewed as experimentally testable hypotheses requiring orthogonal molecular, cellular, and disease-relevant validation [14]. Text-mining evidence may broaden literature retrieval and reveal underrecognized compound–target–phenotype relationships, but pivotal extracted relations must be manually verified because publication frequency, terminology ambiguity, and automated relation errors can distort apparent support.

Target selectivity provides an essential counterweight to indiscriminate multi-target narratives. Polypharmacology may be scientifically relevant when several targets converge coherently on a defined phenotype, but a large predicted target set can also indicate nonspecificity, high lipophilicity, assay interference, or increased safety liability. Selectivity assessment should compare intended targets with alternative pharmacological and safety-relevant targets, consider the quality of underlying activity evidence, and distinguish predicted preference from experimentally measured selectivity. Iterative computational design linked to biochemical and cellular testing demonstrates how natural product-inspired structures can be refined toward experimentally assessed target selectivity rather than relying on target prediction alone [15]. Within the proposed framework, network coherence and predicted selectivity support mechanistic plausibility, but neither is permitted to function as mechanism proof or evidence of therapeutic benefit.

Table 1 summarises network biology and target selectivity requirements for AI-enabled neuroprotective natural product discovery.

Table 1. Network Biology and Target Selectivity Requirements for AI-Enabled Neuroprotective Natural Product Discovery: Disease Context, Target Annotation, Target Confidence, Target Engagement Hypotheses, Off-Target Risk, Neuroinflammation, Oxidative Stress, Mitochondrial Dysfunction, Protein Aggregation, Omics Evidence, Text Evidence, Mechanistic Plausibility, and Validation Boundaries

Discovery component	Core neuropharmacology question	Required evidence or input	Computational function	Potential output	Risk of overinterpretation	Validation requirements	Framework role
Neuroprotective hypothesis	Which neuronal process or phenotype is proposed to be preserved or restored?	Explicit phenotype, disease context, assay system, endpoint definition, and provenance	Classify the hypothesis by mechanism, phenotype, and evidence tier	Context-qualified neuroprotective hypothesis	Treating a broad label as demonstrated efficacy	Disease-relevant functional confirmation	Defines the primary research question
Disease or phenotype context	In which biological setting could the candidate be relevant?	Disease module, phenotype, cell type, experimental model, and stage information	Contextualize target and pathway evidence	Disease- or phenotype-specific evidence profile	Generalizing across unrelated disorders or stages	Context-matched experimental models	Limits interpretation and model transfer
Network biology	How are proposed targets connected to disease-associated molecular systems?	Curated interaction networks, disease genes, regulatory relationships, and pathway data	Module identification, proximity analysis, topology, and centrality assessment	Network-relevance profile	Equating connectivity with causal mechanism	Perturbation and functional studies	Organizes systems-level evidence
Network pharmacology	Could multiple constituents or targets converge on relevant pathways?	Verified constituents, target evidence, pathways, and disease modules	Construct compound–target–pathway relationships	Multi-target pathway hypothesis	Inflated networks based on weak database predictions	Chemical verification and pathway-specific testing	Links phytochemistry to systems hypotheses
Target annotation	Which proteins are plausibly associated with the candidate?	Measured activities, curated target records, structures, and validated prediction models	Rank and annotate candidate targets	Traceable target list	Presenting predicted targets as measured binders	Orthogonal binding or functional assays	Supplies candidate target hypotheses
Target confidence	How reliable is each target assignment?	Experimental support, model confidence, source agreement, and domain status	Evidence weighting and calibration	Confidence-tiered target set	Combining correlated weak sources as independent evidence	External validation and sensitivity analysis	Controls target inclusion
Target selectivity	Is the candidate expected to favor relevant targets over undesirable alternatives?	Comparative target panel, activity evidence, molecular features, and safety targets	Estimate relative selectivity and prioritize profiling	Selectivity hypothesis	Inferring benefit from predicted preference	Broad biochemical and cellular profiling	Supports benefit–risk interpretation
Target engagement hypothesis	Could the candidate engage the proposed target in a relevant neural system?	Binding evidence, cellular expression, exposure plausibility, and compartment information	Integrate target, cell-context, and exposure evidence	Target-engagement hypothesis	Equating target prediction with intracellular engagement	Cellular target-engagement testing	Connects target prediction to validation

Off-target risk	Which unintended targets could create neural or systemic liabilities?	Broad target predictions, secondary pharmacology data, and known liability targets	Detect and rank off-target alerts	Off-target concern profile	Ignoring uncertain but high-severity liabilities	Secondary pharmacology and functional safety assays	Feeds safety-aware filtering
Pathway relevance	Are implicated pathways related to the specified phenotype?	Curated pathways, disease omics, perturbation studies, and target directionality	Enrichment and network-context analysis	Pathway-relevance hypothesis	Treating enrichment as causal evidence	Pathway perturbation and rescue studies	Supports mechanistic plausibility
Neuroinflammation relevance	Could the candidate affect inflammatory processes associated with neuronal injury?	Microglial, cytokine, innate-immune, and neural–immune evidence	Map targets to inflammatory modules	Neuroinflammation hypothesis	Assuming anti-inflammatory activity establishes neuroprotection	Microglial, neuronal, or co-culture validation	Defines one mechanistic domain
Oxidative stress relevance	Could the candidate modify redox imbalance or antioxidant responses?	Reactive oxygen species, redox enzymes, Nrf2-related evidence, and cellular stress data	Integrate redox-associated targets and pathways	Oxidative-stress modulation hypothesis	Treating generic antioxidant activity as therapeutic evidence	Neuron-relevant redox and survival assays	Defines one mechanistic domain
Mitochondrial dysfunction relevance	Could the candidate influence mitochondrial integrity or bioenergetics?	Respiration, membrane potential, mitophagy, calcium, and metabolic evidence	Map mitochondrial targets and phenotypes	Mitochondrial relevance hypothesis	Extrapolating general cytoprotection to mitochondrial rescue	Bioenergetic and mitochondrial-function assays	Defines one mechanistic domain
Protein aggregation relevance where supported	Could the candidate alter formation, clearance, or toxicity of disease-associated aggregates?	Aggregation assays, proteostasis targets, clearance mechanisms, and disease-specific evidence	Link candidate evidence to aggregation or proteostasis modules	Aggregation-related hypothesis	Assuming docking or pathway overlap prevents aggregation	Biophysical and disease-relevant aggregation studies	Conditional mechanistic domain
Synaptic or neuronal survival relevance where supported	Could the candidate preserve neuronal viability or synaptic function?	Neuronal survival, neurite, synaptic marker, electrophysiological, or functional evidence	Integrate functional and pathway signals	Neuronal-survival or synaptic hypothesis	Confusing generic cell survival with neuroprotection	Neuron-specific functional assays	Adds phenotype-level evidence
Omics evidence where supported	Do molecular profiles support the proposed target or pathway hypothesis?	Quality-controlled transcriptomic, proteomic, metabolomic, or genetic data	Signature matching, module detection, and multiomics integration	Omics-concordance profile	Batch effects, tissue mismatch, or reverse causality	Independent datasets and perturbation validation	Modifies target and pathway confidence
Text evidence where supported	What published evidence links compounds, targets, pathways, and phenotypes?	Peer-reviewed literature, normalized entities, and relation extraction	Retrieve and rank literature relationships	Traceable literature evidence graph	Publication bias, repetition, and extraction error	Manual verification of pivotal relations	Supports evidence discovery
Mechanistic plausibility	Is there a coherent evidence chain from compound to	Identity, target, network, exposure, and	Evidence synthesis with explicit uncertainty	Mechanistic-plausibility tier	Presenting coherence as causal proof	Orthogonal perturbation, rescue, and replication	Integrates mechanism-related evidence

	target, pathway, and phenotype?	functional evidence					
No mechanism- proof boundary	Which conclusion remains prohibited at the computational stage?	Claim taxonomy and predefined evidence requirements	Block unsupported causal terminology	Mechanistic hypothesis requiring validation	Rebranding association as validated mechanism	Target engagement, perturbation, rescue, and reproducibili ty	Mandatory claim-control boundary

Blood–brain barrier prediction

Blood–brain barrier prediction should begin with correctly represented chemistry. A model cannot compensate for uncertain constituent identity, unresolved mixtures, incorrect salt or protonation states, missing stereochemistry, or inconsistent structure standardization. Once the molecular record has been curated, descriptors related to size, polarity, charge, hydrogen bonding, lipophilicity, flexibility, and ionization may be used to estimate BBB-relevant behavior. Models developed specifically for natural products show that computational BBB classification can support early prioritization and can be paired with in vitro validation, but their outputs remain conditional on dataset composition, endpoint definition, and chemical-domain coverage [16]. Model development must also address class imbalance and resampling because apparently strong performance can arise from data distributions that do not reflect the intended candidate space [17].

CNS drug-likeness and BBB classification should be treated as related but non-equivalent concepts. CNS-oriented molecular profiles can indicate whether a compound resembles known brain-directed chemical space, whereas a BBB classifier estimates a defined permeability or penetration endpoint under the assumptions of its training data. Efficient machine-learning models can assist large-scale screening, but a positive output does not establish passive diffusion, transporter-mediated uptake, absence of efflux, adequate systemic exposure, sufficient unbound brain concentration, or persistence at the relevant site of action [18]. Uncertainty-aware BBB models provide a more defensible basis for prioritization because they can identify compounds for which the model should abstain, request additional evidence, or avoid binary interpretation [19].

A robust BBB assessment should compare multiple model classes and representations while examining whether agreement reflects independent evidence or shared bias.

Dynamic consensus approaches may reduce reliance on one algorithm, yet consensus cannot correct mislabeled data, narrow chemical coverage, endpoint inconsistency, or correlated training sets [20]. Emerging molecular language representations may identify structural patterns not captured by conventional descriptors, but they introduce additional requirements for transparent data splitting, leakage assessment, interpretability, calibration, and external evaluation [21]. Applicability-domain analysis and out-of-distribution detection are therefore required for natural products, whose scaffolds, stereochemical complexity, glycosylation, ionization patterns, and molecular sizes may differ substantially from compounds dominating conventional BBB datasets.

BBB interpretation must ultimately be integrated with passive-permeability hypotheses, transporter relevance, efflux risk, metabolic stability, protein binding, systemic pharmacokinetics, and other ADME evidence. A compound predicted to cross the barrier may nevertheless fail to achieve meaningful brain exposure because of rapid clearance, low absorption, extensive binding, active efflux, metabolism, or insufficient unbound concentration. Conversely, an uncertain passive-permeability estimate does not exclude transporter-mediated disposition. Reviews of machine-learning BBB models emphasize that heterogeneous labels, assays, datasets, and evaluation practices prevent a single computational result from functioning as proof of CNS exposure [22]. The framework therefore uses the terms BBB prediction, BBB-relevant evidence, CNS-relevance hypothesis, and brain-exposure hypothesis deliberately; demonstrated brain exposure requires appropriate experimental measurement and cannot be inferred from prediction alone.

Figure 1 illustrates how neuroprotective natural product evidence can be linked to network biology, target selectivity, BBB prediction, CNS relevance, and safety-aware interpretation without converting computational predictions into therapeutic claims.

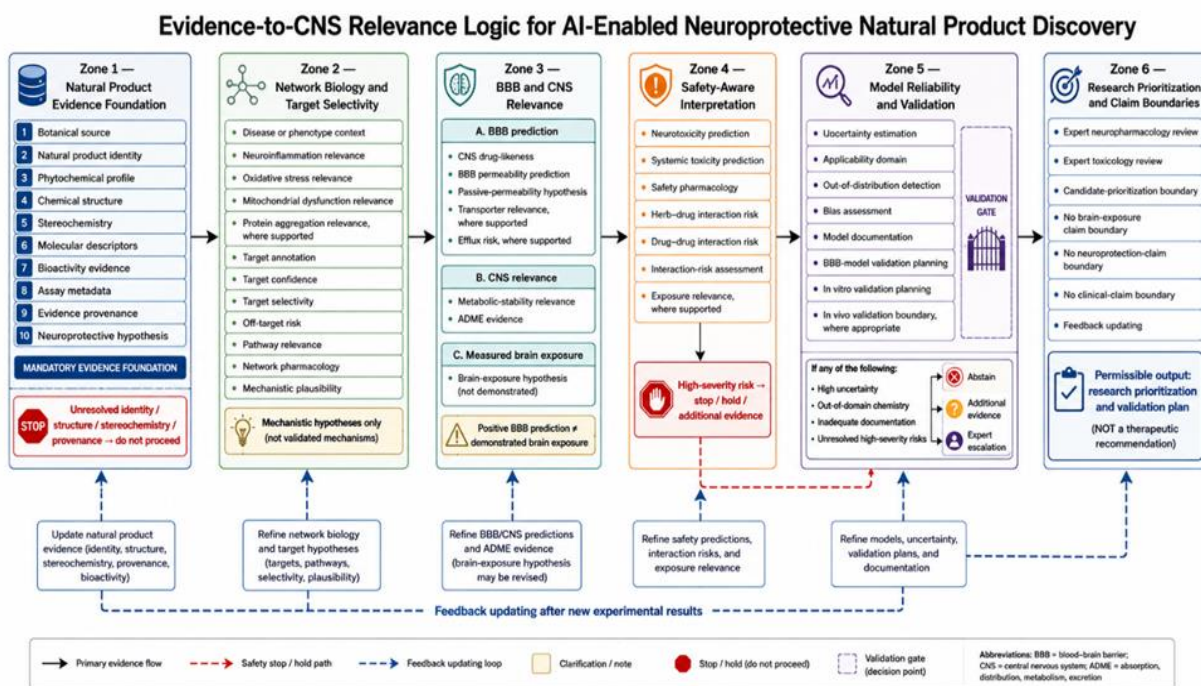


Figure 1. Evidence-to-CNS Relevance Logic for AI-Enabled Neuroprotective Natural Product Discovery

Alt-text description

A conceptual evidence-to-CNS relevance diagram showing natural product identity connected to neuroprotective hypotheses, network biology, target selectivity, BBB prediction, CNS relevance, safety evidence, uncertainty assessment, validation gates, and claim boundaries.

Safety-aware discovery logic

Safety-aware discovery should operate as a parallel decision stream rather than as a final filter applied after activity and BBB predictions have already favored a candidate. Multi-endpoint ADMET platforms can organize evidence related to absorption, distribution, metabolism, excretion, physicochemical behavior, organ toxicity, and other developability concerns, but their outputs remain endpoint-specific predictions whose reliability depends on training data, applicability domain, and model calibration [23]. Broad toxicity-prediction systems can additionally flag structural alerts, toxicity pathways, adverse molecular targets, and systemic hazards, yet the absence of a predicted alert cannot establish that a natural product or natural product-inspired candidate is safe [24]. The framework therefore records each endpoint separately and prevents a favorable aggregate profile from compensating for a severe but uncertain liability.

Neurotoxicity requires a dedicated evaluation layer because a compound proposed for neuroprotective research may also alter neuronal viability, excitability, differentiation, synaptic function, mitochondrial integrity, or neural-network behavior in undesirable ways.

Molecular machine-learning models can support early chemical neurotoxicity screening, but their interpretation is limited by label quality, endpoint heterogeneity, class imbalance, and incomplete coverage of complex natural-product chemistry [25]. Higher-content experimental systems, including human midbrain organoid-based approaches, may provide a more biologically informative validation step for selected neurotoxicity hypotheses by connecting computational alerts to cellular and phenotypic responses [26]. Such models remain context-specific, however, and cannot independently capture chronic exposure, metabolite effects, multicellular interactions, behavioral consequences, or organism-level toxicity. Interaction risk is particularly important when a candidate derives from a botanical preparation, co-occurs with other constituents, or could be investigated in settings involving concomitant medicines. Deep-learning approaches can predict drug–drug and drug–food interaction relationships and may help identify combinations requiring mechanistic review, but predicted interactions do not by themselves define direction, magnitude, dose dependence, exposure dependence, or clinical severity [27]. Herb–drug interactions may arise through enzyme inhibition or induction, transporter modulation, altered absorption, pharmacodynamic synergy or antagonism, and variability in botanical composition [28]. The proposed logic therefore distinguishes constituent-level interaction hypotheses from product-level interaction risk and requires metabolism, transporter, exposure, and composition evidence before any clinically oriented interpretation.

Uncertainty estimation, applicability-domain analysis, and bias assessment must accompany every safety endpoint. A structurally unfamiliar candidate, poorly represented toxicity class, uncertain stereochemical assignment, or variable mixture should trigger abstention, additional testing, or expert escalation rather than a reassuring prediction. Expert toxicology review should examine the full evidence dossier, including high-severity low-confidence alerts, possible neuroactive liabilities, metabolic concerns, predicted interactions, exposure

assumptions, assay limitations, and missing endpoints. The permissible output is a research-oriented safety-risk hypothesis and a prioritized validation plan. It is not a declaration that the candidate is safe, non-toxic, suitable for administration, compatible with medicines, or appropriate for treatment.

Table 2 maps BBB prediction, CNS relevance, and safety-aware discovery requirements for neuroprotective natural product prioritization.

Table 2. BBB Prediction, CNS Relevance, and Safety-Aware Discovery Requirements for Neuroprotective Natural Product Prioritization: CNS Drug-Likeness, BBB Permeability Prediction, Transporter and Efflux Relevance, Brain Exposure Hypotheses, ADMET, Neurotoxicity, Systemic Toxicity, Interaction Risk, Applicability Domain, Uncertainty, and Claim Boundaries

Evaluation component	Core question	Required input or evidence	Prediction or safety-aware function	Potential output	Failure or uncertainty risk	Validation requirement	Decision boundary
Chemical structure	Is the candidate represented correctly for CNS and safety modeling?	Standardized structure, stereochemistry, protonation, tautomer, salt state, purity, and constituent identity	Structure normalization and quality control	Model-ready molecular record	Misassigned identity, unresolved mixtures, or stereochemical ambiguity	Analytical verification and expert curation	No prediction from unresolved molecular identity
Molecular descriptors	Which physicochemical features influence modeled CNS disposition?	Size, polarity, charge, hydrogen bonding, lipophilicity, flexibility, and ionization	Compute descriptor profile	BBB- and CNS-relevance descriptor set	Calculation-method and protonation assumptions	Comparison with measured properties where available	Descriptors support interpretation but do not prove exposure
CNS drug-likeness	Does the candidate resemble molecular profiles associated with CNS development?	Curated molecular representation and physicochemical evidence	Apply CNS-oriented heuristic or multivariate assessment	CNS-likeness hypothesis	Bias toward conventional synthetic chemical space	Evaluation against chemically relevant reference compounds	CNS likeness is not BBB penetration or efficacy
BBB permeability prediction	Does a validated model predict a defined BBB endpoint?	Curated structure–BBB dataset, endpoint definition, and model documentation	Classification or regression	Predicted BBB class, probability, or interval	Label noise, leakage, class imbalance, and endpoint inconsistency	External testing and BBB-relevant assays	Prediction alone does not establish BBB passage
Passive permeability hypothesis	Could passive diffusion contribute to barrier crossing?	Ionization, polarity, lipophilicity, size, and membrane-assay evidence	Estimate passive-permeability potential	Passive-diffusion hypothesis	Active transport, metabolism, and binding omitted	PAMPA-BBB, transwell, or related testing	Passive potential is not total brain exposure

Transporter relevance where supported	Could uptake or efflux transporters influence CNS disposition?	Transporter substrate or inhibitor evidence, structure, expression, and assay context	Predict or annotate transporter relationships	Transporter -relevance profile	Sparse labels, species differences, and concentration dependence	Transporter-specific functional assays	Transporter prediction requires confirmation
Efflux risk where supported	Could active efflux limit CNS exposure?	Efflux-model outputs, bidirectional-permeability evidence, and transporter information	Flag possible P-gp, BCRP, or related liability	Efflux-risk hypothesis	Shared model bias and variable assay systems	Bidirectional assays with relevant controls	Absence of predicted efflux does not prove exposure
Metabolic stability relevance	Could metabolism prevent sufficient systemic or CNS availability?	Microsomal, hepatocyte, enzyme, and metabolite evidence	Predict clearance and metabolic liability	Metabolic-stability concern	Species, enzyme, concentration, and metabolite uncertainty	In vitro metabolism and pharmacokinetic studies	Positive BBB prediction cannot override unstable exposure
ADME evidence	Are disposition properties compatible with further testing?	Predicted and measured absorption, distribution, metabolism, excretion, binding, and clearance evidence	Integrate endpoint-specific ADME information	ADME-readiness profile	Heterogeneous assays and discordant endpoints	Endpoint-specific experimental confirmation	ADME prediction remains screening evidence
Brain exposure hypothesis	Is there a plausible path to measurable unbound brain exposure?	BBB, systemic PK, binding, metabolism, transporter, and distribution evidence	Integrate exposure determinants	Brain-exposure hypothesis	Compounding uncertainty across multiple models	Appropriate in vivo or translational exposure studies	No statement of demonstrated brain exposure without measurement
ADMET prediction	Which developability liabilities require investigation?	Curated structure and validated endpoint models	Multi-endpoint ADMET screening	Endpoint-specific liability profile	Composite scores may conceal severe individual risks	Confirmation of pivotal endpoints	No overall developability or safety claim
Neurotoxicity prediction	Could the candidate adversely affect neural cells or functions?	Chemical structure, neurotoxicity labels, phenotypic evidence, and model domain	Classify or rank neurotoxicity concern	Neurotoxicity-risk hypothesis	Narrow endpoints, weak labels, and OOD chemistry	Neural-cell, organoid, or other relevant testing	Negative prediction does not establish neurological safety
Systemic toxicity prediction	Could the candidate produce non-neural organ or general toxicity?	Acute, chronic, organ, genetic, developmental, and structural evidence	Multi-endpoint toxicity triage	Systemic-toxicity alert profile	Endpoint imbalance and missing chronic effects	Appropriate toxicological assays	No declaration of systemic safety

Safety pharmacology	Could the candidate affect vital physiological systems?	Safety-target panels, functional evidence, and exposure assumptions	Identify cardiovascular, respiratory, neurological, or other liabilities	Safety-pharmacology concern map	Missing targets and insufficient concentration context	Functional secondary-pharmacology testing	Prediction only prioritizes testing
Off-target safety risk	Which unintended targets could create significant harm?	Broad target predictions, liability-target knowledge, and pharmacology data	Rank high-severity off-target hypotheses	Off-target safety profile	Uncertain targets may be over- or undervalued	Broad biochemical and cellular profiling	Off-target prediction is not a confirmed adverse effect
Herb–drug interaction risk	Could constituents alter the action or disposition of medicines?	Product composition, enzymes, transporters, targets, and exposure evidence	Map plausible pharmacokinetic and pharmacodynamic mechanisms	Herb–drug interaction hypothesis	Variable preparations, unidentified constituents, and dose uncertainty	Mechanistic interaction and exposure studies	No medication-use recommendation
Drug–drug interaction risk	Could an isolated or inspired candidate interact with concomitant drugs?	Metabolism, transporters, targets, co-exposure, and pharmacokinetic information	Predict potential interaction relationships	Prioritized DDI concern	Incomplete labels and missing dose context	Mechanistic and exposure-based evaluation	No clinical severity inference from prediction alone
Metabolism-related concern	Could inhibition, induction, or reactive metabolism increase risk?	CYP, UGT, transporter, metabolite, and structural-alert evidence	Identify metabolic interaction and toxicity mechanisms	Metabolism-related risk profile	Unidentified metabolites and species differences	Microsomal, hepatocyte, and metabolite studies	No interaction or toxicity conclusion without confirmation
Dose or exposure relevance where supported	Are predicted liabilities plausible at achievable exposure?	Concentration–response evidence, PK, binding, and toxicokinetic information	Contextualize hazards by exposure	Exposure-adjusted research concern	Exposure predictions may be highly uncertain	Pharmacokinetic and toxicokinetic testing	No safe dose or therapeutic-window inference
Neuroactive liability concern	Could intended or unintended CNS pharmacology impair function?	Receptor, transporter, ion-channel, behavioral, and functional evidence	Identify sedative, excitatory, motor, cognitive, or dependence-related signals	Neuroactive-liability hypothesis	Dose, metabolite, and context dependence	Functional CNS safety evaluation	No claim of benign CNS activity
Applicability domain	Does the model adequately represent the candidate's chemistry?	Training-set features and candidate representation	Similarity, leverage, density, or coverage analysis	In-domain, borderline, or out-of-domain status	Domain metric may not capture all novelty	Evaluation with challenge compounds	OOD outputs cannot independently guide advancement
Out-of-distribution detection	Is the candidate unfamiliar to the model in feature or latent space?	Candidate and reference representations	Detect distributional shift	OOD warning	Detector sensitivity and shared representation bias	Prospective testing on novel scaffolds	OOD status triggers abstention or added evidence

Uncertainty estimation	How reliable is each BBB, ADMET, toxicity, or interaction prediction?	Ensembles, calibrated probabilities, conformal intervals, or other UQ measures	Quantify predictive reliability	Confidence interval, uncertainty tier, or abstention status	Miscalibration and correlated model error	Independent calibration and external evaluation	High uncertainty blocks reassuring interpretation
Bias assessment	Are relevant compounds, endpoints, or evidence sources underrepresented?	Dataset composition, labels, assay origins, and chemical-space coverage	Audit imbalance, leakage, and subgroup performance	Bias and coverage report	Hidden historical, publication, or assay bias	Stratified and external testing	Biased outputs cannot independently exclude or advance candidates
Expert toxicology review	Are predicted liabilities biologically credible and adequately investigated?	Complete safety dossier, uncertainty, exposure assumptions, and validation plan	Human adjudication and prioritization	Reviewed safety-risk hypothesis	Automation bias and incomplete evidence	Independent expert assessment	Review does not certify safety
No brain-exposure claim boundary	What BBB-related conclusion is prohibited?	Claim taxonomy and evidence threshold	Prevent prediction-to-exposure substitution	BBB-relevant or brain-exposure hypothesis language	Positive model output may be overstated	Direct exposure evidence	No claim that the candidate reaches the brain
No safety-claim boundary	What safety conclusion is prohibited?	Claim-control rules and endpoint requirements	Block “safe,” “non-toxic,” or equivalent language	Explicitly qualified safety-risk statement	Absence of alerts misread as absence of risk	Comprehensive staged testing	No computational safety assurance
No treatment-recommendation boundary	Can these outputs guide treatment, supplementation, or monitoring?	Clinical efficacy, dose, interaction, safety, and patient evidence	Enforce research-only interpretation	Non-clinical research statement	Candidate prioritization converted into self-treatment advice	Appropriate clinical development and professional evaluation	No treatment, dose, substitution, or monitoring recommendation

Proposed computational framework

The proposed framework is designed as a staged evidence-and-decision architecture rather than as a single predictive model. Each stage produces a bounded research output, an uncertainty statement, and a defined requirement for progression. Uncertainty quantification is treated as a decision-control mechanism because molecular-property uncertainty methods differ in calibration, ranking behavior, and reliability across datasets and chemical domains [29]. Explainability is incorporated to help experts interrogate molecular features, target associations, pathway links, BBB predictions, and toxicity alerts, but an explanation generated by a model is not accepted as causal or mechanistic evidence [30]. The framework therefore combines prediction, explanation, applicability-domain

analysis, and explicit abstention rather than interpreting confident output as validated knowledge.

The first stage defines the research question, intended phenotype, disease context, evidence requirements, and prohibited claims. The second establishes the natural product evidence foundation through botanical provenance, phytochemical identity, structural standardization, stereochemistry, bioactivity evidence, assay metadata, and traceable sources. The third classifies the neuroprotective hypothesis according to phenotype and mechanistic domain. Network biology mapping and target-selectivity assessment then organize disease modules, predicted or measured targets, target confidence, off-target concerns, pathway relevance, omics concordance, and mechanistic plausibility. BBB prediction and CNS-relevance assessment follow as separate stages integrating

permeability models, passive-diffusion hypotheses, transporter and efflux evidence, metabolic stability, ADME, and brain-exposure uncertainty. Evaluation metrics and validation criteria must be chosen according to the endpoint and scientific decision rather than selected solely to maximize apparent model performance [31].

Safety-aware filtering proceeds in parallel with mechanistic and CNS assessment. ADMET, neurotoxicity, systemic toxicity, safety pharmacology, metabolism-related concerns, and interaction risks are recorded as separate evidence dimensions with endpoint-specific uncertainty and veto conditions. Validation gates are then sequenced according to the strongest unresolved claim and the most consequential risk. An active-learning strategy may direct experiments toward candidates or endpoints expected to provide the greatest information gain, but experimental selection remains governed by scientific relevance, safety, feasibility, and human oversight [32]. Model evaluation, data splitting, baselines, external testing, reproducibility, and documentation must follow fit-for-purpose chemical machine-learning practices throughout the workflow [33].

Expert neuropharmacology review evaluates whether the disease context, neural phenotype, target network, pathway directionality, CNS relevance, and proposed validation model are biologically coherent. Expert toxicology review evaluates neurotoxicity, systemic liabilities, off-target

pharmacology, metabolism, interactions, exposure assumptions, and the adequacy of planned safety studies. Candidate prioritization occurs only after these reviews and should use transparent multi-criteria logic with non-compensatory vetoes. For example, high pathway relevance should not cancel an unresolved severe toxicity alert, and favorable BBB prediction should not compensate for uncertain molecular identity. The output is a research-priority category, evidence-gap map, and validation plan rather than a validated candidate ranking or therapeutic recommendation.

Feedback updating closes the framework. New analytical, biochemical, cellular, BBB-relevant, pharmacokinetic, toxicological, omics, or phenotypic results are returned to the evidence record, and affected models, confidence tiers, target networks, risk profiles, and validation priorities are revised under version control. Negative or inconclusive results are retained to reduce selective learning and prevent repeated investigation of unsupported hypotheses. The framework’s original contribution is this coordinated logic linking natural product evidence, network biology, selectivity, CNS relevance, safety, uncertainty, validation, and expert review while preserving explicit boundaries against claims of demonstrated neuroprotection, brain exposure, safety, clinical utility, or treatment readiness.

Table 3 presents the proposed computational framework for AI-enabled neuroprotective natural product discovery.

Table 3. Proposed Computational Framework for AI-Enabled Neuroprotective Natural Product Discovery: Natural Product Evidence Foundation, Network Biology, Target Selectivity, BBB Prediction, CNS Relevance, Safety-Aware Filtering, Interaction-Risk Assessment, Uncertainty Review, Validation Gates, Expert Oversight, Candidate Prioritization, Feedback, and Claim Boundaries

Framework stage	Core purpose	Required input	Computational function	Potential output	Validation need	Failure risk	Feedback mechanism	Decision boundary
Research question definition	Define the phenotype, disease context, intended use, and permissible claims	Research objective, biological context, endpoint definitions, stakeholders, and claim policy	Encode scope, terminology, and evidence thresholds	Versioned research specification	Neuropharmacology and methodological review	Vague endpoints, post hoc interpretation, or scope drift	Revise specification after evidence review	No analysis without a defined and bounded question
Natural product evidence foundation	Establish chemical identity, provenance, and evidence traceability	Botanical source, taxonomy, phytochemical profile, structure, stereochemistry, purity, assays, and provenance	Standardize entities and connect evidence	Curated natural-product candidate record	Analytical verification and metadata quality control	Misidentified constituent, variable mixture, or unsupported database compound	Update with new analytical or sourced evidence	Unresolved identity blocks computational prioritization
Neuroprotective hypothesis classification	Convert broad neuroprotective language	Disease or phenotype context,	Classify phenotype and	Context-qualified	Expert review and phenotype-relevant testing	Treating “neuroprotective” as an	Reclassify after	Hypothesis language only

	into a testable research hypothesis	biological system, endpoint, and mechanistic evidence	mechanism domain	neuroprotective hypothesis		established property	functional evidence	
Network biology mapping	Determine how candidate-associated targets relate to disease systems	Curated targets, PPI networks, pathways, disease genes, omics, and text evidence	Module detection, proximity, topology, and enrichment analysis	Network-relevance profile	Independent resources and perturbation studies	Annotation bias, database circularity, and false connectivity	Update network after target validation	Network association is not causal mechanism
Target annotation and confidence	Identify and grade plausible molecular targets	Measured bioactivity, prediction models, target databases, model domain, and provenance	Rank and evidence-weight targets	Confidence-tiered target set	Binding and functional confirmation	Predicted target treated as measured target	Revise after orthogonal assays	Low-confidence targets cannot drive mechanism claims
Target selectivity assessment	Compare intended activity with alternative and safety-relevant targets	Comparative target evidence, molecular features, activity data, and liability targets	Estimate relative selectivity and prioritize profiling	Selectivity and off-target hypothesis	Broad biochemical and cellular profiling	Single-target bias or unrecognize d promiscuity	Update after secondary pharmacology	Predicted selectivity is not therapeutic benefit
Target engagement hypothesis	Assess whether the candidate could engage a target in a relevant neural system	Binding, expression, cellular localization, exposure plausibility, and functional evidence	Integrate target and context evidence	Target-engagement hypothesis	Cellular engagement and perturbation studies	Target prediction equated with in-cell engagement	Revise after engagement testing	No target-engagement claim without direct evidence
BBB prediction	Estimate a defined BBB permeability or penetration endpoint	Standardized chemistry, descriptors, validated datasets, and model documentation	Classification, regression, consensus modeling, and UQ	BBB-prediction profile	External datasets and BBB-relevant assays	Leakage, imbalance, inconsistent labels, or OOD chemistry	Update with measured permeability	No brain-exposure claim
CNS relevance assessment	Integrate BBB prediction with other determinants of CNS disposition	Passive permeability, transporters, efflux, metabolism, binding, PK, and CNS target context	Multi-evidence CNS suitability assessment	CNS-relevance and brain-exposure hypothesis	Transport, PK, and exposure studies	Overweighting one favorable BBB result	Revise after ADME and exposure evidence	CNS relevance does not establish neuroprotection
ADME and metabolic assessment	Identify disposition constraints and metabolite-related concerns	Absorption, binding, clearance, enzyme, transporter, and metabolite evidence	Multi-endpoint ADME integration	ADME concern and testing profile	Endpoint-specific experiments	Predicted exposure treated as measured exposure	Update after PK and metabolism studies	No exposure or dose conclusion

Safety-aware filtering	Identify neural and systemic hazards before advancement	Neurotoxicity, systemic toxicity, safety pharmacology, target, and structural evidence	Multi-endpoint hazard triage with veto rules	Safety-risk hypothesis and test priorities	Relevant toxicology and functional studies	Composite score conceals severe endpoint	Revise after safety testing	Filtering does not prove safety
Interaction-risk assessment	Identify plausible herb–drug and drug–drug interaction mechanisms	Composition, enzymes, transporters, targets, pharmacodynamics, and exposure	Mechanistic and network interaction prediction	Interaction concern profile	Enzyme, transporter, exposure, and combination studies	Missing constituents, dose, or co-exposure context	Update after interaction studies	No clinical interaction advice
Uncertainty estimation	Quantify reliability of each computational output	Calibrated models, ensembles, prediction intervals, and endpoint data	Estimate confidence and support abstention	Endpoint-specific uncertainty tier	Independent calibration and external testing	Numerically confident but unreliable output	Recalibrate after new data	High uncertainty prevents positive inference
Applicability-domain review	Determine whether each model adequately covers the candidate	Candidate and training-set representations	Similarity, leverage, density, and coverage analysis	In-domain, borderline, or OOD classification	Prospective evaluation on novel chemistry	Hidden domain mismatch	Expand training data or trigger experiments	OOD outputs cannot independently advance a candidate
Bias assessment	Identify systematic limitations in evidence or models	Dataset composition, assay source, labels, subgroup coverage, and publication patterns	Audit imbalance, leakage, and differential performance	Bias and coverage report	Stratified and independent evaluation	Historical or assay bias reproduced by the model	Correct data and re-evaluate	Biased predictions require conservative use
Model documentation	Preserve transparency, reproducibility, and intended-use limits	Data sources, preprocessing, features, algorithms, splits, metrics, versions, and limitations	Produce auditable model records	Model card or equivalent documentation	Independent review and reproducibility checks	Undocumented preprocessing or untraceable output	Update after every material change	Undocumented models cannot support decisions
Validation gate sequencing	Match each hypothesis and liability to the next decisive experiment	Evidence gaps, uncertainty, risk severity, feasibility, and claim requirements	Prioritize orthogonal testing	Staged validation plan	Predefined criteria and appropriate controls	Convenient rather than informative experiments	Results return to all relevant stages	Advancement requires gate-specific evidence
BBB model validation planning	Select experiments matching the predicted transport mechanism	BBB output, passive/active transport hypothesis, and uncertainty	Choose barrier, transporter, and integrity assays	BBB-validation plan	Fit-for-purpose barrier models and controls	Model endpoint does not match assay endpoint	Update BBB and CNS profiles	No exposure inference before validation
In vitro validation planning	Test target, pathway, phenotype, and safety hypotheses	Mechanistic chain, cell context, exposure range, and controls	Select orthogonal biochemical and cellular studies	In vitro evidence plan	Replication, concentration relevance, and appropriate comparators	Assay interference or irrelevant cell system	Feed results into confidence tiers	In vitro activity is not clinical efficacy

In vivo validation boundary where appropriate	Define when organism-level studies may be justified	Strong identity, mechanistic, BBB, exposure, and safety rationale	Integrate prior-gate evidence	Justified or unjustified escalation decision	Ethical and scientifically appropriate study design	Premature escalation from weak computational evidence	Return PK, efficacy-related, and toxicity evidence	In vivo evidence does not automatically establish human relevance
Expert neuropharmacology review	Evaluate disease, pathway, neural, and CNS plausibility	Complete mechanism and exposure dossier	Human evidence adjudication	Reviewed neuropharmacology decision	Independent documented assessment	Automation or confirmation bias	Request analysis, evidence, or revised hypothesis	Review cannot create efficacy evidence
Expert toxicology review	Evaluate safety, interaction, and exposure concerns	Complete safety dossier and uncertainty record	Human risk adjudication	Reviewed toxicology decision	Independent documented assessment	False reassurance from negative predictions	Request safety studies or redesign	Review does not certify safety
Candidate prioritization	Allocate research resources transparently	Evidence tiers, validation status, uncertainty, feasibility, and risk	Multi-criteria analysis with non-compensatory vetoes	Research-priority category	Sensitivity analysis and committee review	Arbitrary weighting or compensation across fatal risks	Re-rank after each evidence update	Priority is not validated candidacy
Feedback updating	Learn from positive, negative, and inconclusive results	Analytical and experimental outcomes, errors, and failure analysis	Update records, models, networks, and priorities	Versioned framework state	Audit trail and change review	Selective learning from favorable outcomes	Continuous design–test–learn cycle	Prior outputs remain traceable
No neuroprotection-claim boundary	Prevent computational evidence from being stated as validated neuroprotection	Claim taxonomy and evidence thresholds	Automated and editorial claim checking	Qualified hypothesis terminology	Final scientific and editorial review	Association presented as therapeutic effect	Correct language after every stage	No validated neuroprotection claim without appropriate evidence
No clinical-claim boundary	Prevent research outputs from becoming clinical recommendations	Intended-use policy and clinical evidence requirements	Block treatment-oriented output	Research-use-only conclusion	Governance and manuscript review	Candidate priority interpreted as treatment guidance	Reassess only after appropriate development	No clinical utility, treatment, dose, or patient-selection claim

Figure 2 presents the proposed computational framework for AI-enabled neuroprotective natural product discovery, integrating network biology, target selectivity, BBB

prediction, safety-aware filtering, validation gates, expert oversight, and feedback updating.

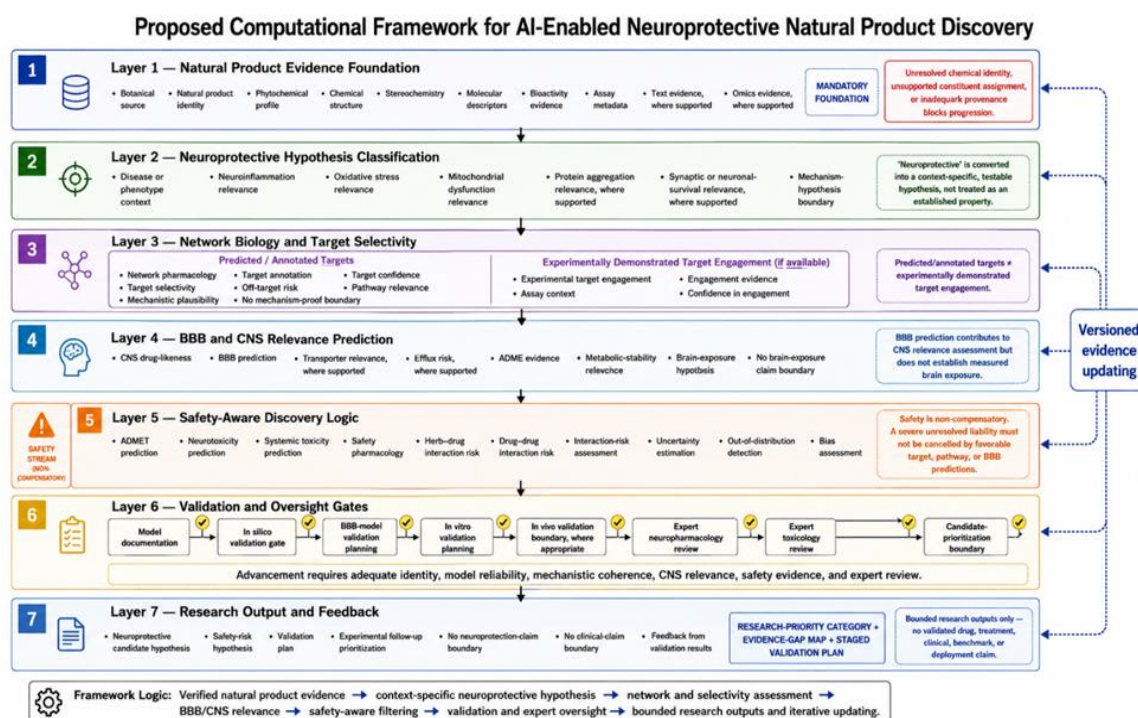


Figure 2. Proposed Computational Framework for AI-Enabled Neuroprotective Natural Product Discovery

Alt-text description

A conceptual computational framework showing natural product evidence flowing through neuroprotective hypothesis classification, network biology, target selectivity, BBB prediction, CNS relevance assessment, safety-aware filtering, interaction-risk review, validation gates, expert oversight, candidate prioritization, and feedback loops.

Research implications

For natural product informatics and computational neuropharmacology, the framework implies that data infrastructure is as important as algorithm choice. Botanical provenance, chemical identity, stereochemistry, mixture composition, assay metadata, negative findings, and experimental context must be represented in forms that permit traceable integration across chemical, target, network, BBB, ADME, toxicity, and interaction datasets. Machine learning is most useful when embedded within multidisciplinary discovery workflows that connect predictions to explicit experimental decisions rather than operating as an isolated candidate-ranking layer [34]. This requires collaboration among natural product chemists, analytical scientists, neuropharmacologists, computational biologists, pharmacokineticists, toxicologists, and data-governance specialists. For BBB modeling and safety pharmacology, future research should prioritize endpoint harmonization, external-domain testing, prospective validation, and models capable of abstaining when evidence is inadequate.

Natural-product chemistry should be evaluated explicitly rather than assumed to fall within conventional drug-like training distributions. Neurotoxicity research would benefit from links between structure-based models, molecular initiating events, neural-cell phenotypes, organoid systems, and organism-level evidence. At the same time, the field should resist claims that improvements in benchmark performance necessarily translate into more reliable mechanisms, safer candidates, successful development, or clinical impact [35]. The most meaningful performance criterion is whether a model improves the quality, efficiency, transparency, or safety of subsequent research decisions under realistic conditions. Experimental validation design should be aligned with the specific uncertainty produced by each framework stage. Identity uncertainty calls for analytical clarification; target uncertainty calls for binding or functional testing; network uncertainty calls for perturbation and rescue studies; BBB uncertainty calls for barrier, transporter, pharmacokinetic, or exposure experiments; and safety uncertainty calls for endpoint-specific toxicology. Standardized assay metadata, model documentation, versioned evidence records, negative-result retention, and explicit decision thresholds would make computational findings easier to reproduce and revise. The framework should therefore be interpreted as infrastructure for disciplined hypothesis prioritization and evidence updating, not as an autonomous system for discovering treatments or determining clinical use.

CONCLUSION

This article proposes a computational framework for AI-enabled neuroprotective natural product discovery that integrates verified natural product evidence, neuroprotective hypothesis classification, network biology, target annotation and selectivity, blood–brain barrier prediction, CNS relevance, ADME and safety-aware filtering, neurotoxicity and systemic toxicity assessment, interaction-risk evaluation, uncertainty and applicability-domain review, validation gates, expert neuropharmacology and toxicology oversight, candidate prioritization, and feedback updating. Its defining contribution is not a new predictive model or a claim of validated discovery performance, but a structured logic for connecting heterogeneous evidence while preventing unsupported inference. Computationally prioritized natural products and natural product-inspired candidates remain research hypotheses until their identity, target engagement, pathway effects, neuroprotective function, brain exposure, safety, and therapeutic relevance have been established through appropriate experimental and translational evidence.

Acknowledgments: None

Conflict of interest: None

Financial support: None

Ethics statement: None

REFERENCES

- [1] Saldívar-González FI, Aldas-Bulos VD, Medina-Franco JL, Plisson F. Natural product drug discovery in the artificial intelligence era. *Chem Sci*. 2022;13(6):1526-46. doi:10.1039/D1SC04471K
- [2] Othman ZK, Ahmed MM, Kasimieh O, Musa SS, Branda F, Cue EG, et al. Artificial intelligence for natural product drug discovery and development: current landscape, applications, and future directions. *Intell Based Med*. 2025;12:100316. doi:10.1016/j.ibmed.2025.100316
- [3] Vatansever S, Schlessinger A, Wacker D, Kaniskan HÜ, Jin J, Zhou MM, et al. Artificial intelligence and machine learning-aided drug discovery in central nervous system diseases: state-of-the-arts and future directions. *Med Res Rev*. 2021;41(3):1427-73. doi:10.1002/med.21764
- [4] Fontanella F, D'Alessandro T, Nardone E, De Stefano C, Vicidomini C, Roviello GN. Artificial intelligence for natural products drug discovery in neurodegenerative therapies: A review. *Biomolecules*. 2026;16(1):129. doi:10.3390/biom16010129
- [5] Lim DW, Lee JE, Lee C, Kim YT. Natural products and their neuroprotective effects in degenerative brain diseases: A comprehensive review. *Int J Mol Sci*. 2024;25(20):11223. doi:10.3390/ijms252011223
- [6] Rahman MH, Bajgai J, Fadriqela A, Sharma S, Trinh TT, Akter R, et al. Therapeutic potential of natural products in treating neurodegenerative disorders and their future prospects and challenges. *Molecules*. 2021;26(17):5327. doi:10.3390/molecules26175327
- [7] Moratilla-Rivera I, Sánchez M, Valdés-González JA, Gómez-Serranillos MP. Natural products as modulators of Nrf2 signaling pathway in neuroprotection. *Int J Mol Sci*. 2023;24(4):3748. doi:10.3390/ijms24043748
- [8] Piccirillo S, Preziuso A, Serfilippi T, Cerqueni G, Terenzi V, Lariccia V, et al. Unveiling the potential neuroprotective effect of bioactive compounds from plants with sedative and mood-modulating properties: innovative approaches for the prevention of Alzheimer's and Parkinson's diseases. *Curr Neuropharmacol*. 2025;23(10):1169-83. doi:10.2174/011570159X345397241210103538
- [9] Luo Z, Yu G, Chen X, Liu Y, Zhou Y, Wang G, et al. Integrated phytochemical analysis based on UHPLC-LTQ-Orbitrap and network pharmacology approaches to explore the potential mechanism of *Lycium ruthenicum* Murr. for ameliorating Alzheimer's disease. *Food Funct*. 2020;11(2):1362-72. doi:10.1039/C9FO02840D
- [10] Gan X, Shu Z, Wang X, Yan D, Li J, Ofaim S, et al. Network medicine framework reveals generic herb-symptom effectiveness of traditional Chinese medicine. *Sci Adv*. 2023;9(43):eadh0215. doi:10.1126/sciadv.adh0215
- [11] Qiang B, Lai J, Jin H, Zhang L, Liu Z. Target prediction model for natural products using transfer learning. *Int J Mol Sci*. 2021;22(9):4632. doi:10.3390/ijms22094632
- [12] Cockroft NT, Cheng X, Fuchs JR. StarFish: A stacked ensemble target fishing approach and its application to natural products. *J Chem Inf Model*. 2019;59(11):4906-20. doi:10.1021/acs.jcim.9b00489
- [13] Yu X, Lai S, Chen H, Chen M. Protein-protein interaction network with machine learning models and multiomics data reveal potential neurodegenerative disease-related proteins. *Hum Mol Genet*. 2020;29(8):1378-87. doi:10.1093/hmg/ddaa065
- [14] Hajjo R, Sabbah DA, Abusara OH, Al Bawab AQ. A review of the recent advances in Alzheimer's disease research and the utilization of network biology approaches for prioritizing diagnostics and

- therapeutics. *Diagnostics* (Basel). 2022;12(12):2975. doi:10.3390/diagnostics12122975
- [15] Friedrich L, Cingolani G, Ko YH, Iaselli M, Miciaccia M, Perrone MG, et al. Learning from nature: from a marine natural product to synthetic cyclooxygenase-1 inhibitors by automated de novo design. *Adv Sci* (Weinh). 2021;8(16):e2100832. doi:10.1002/advs.202100832
- [16] Zhang X, Liu T, Fan X, Ai N. In silico modeling on ADME properties of natural products: classification models for blood-brain barrier permeability, its application to traditional Chinese medicine and in vitro experimental validation. *J Mol Graph Model*. 2017;75:347-54. doi:10.1016/j.jmgm.2017.05.021
- [17] Wang Z, Yang H, Wu Z, Wang T, Li W, Tang Y, et al. In silico prediction of blood-brain barrier permeability of compounds by machine learning and resampling methods. *ChemMedChem*. 2018;13(20):2189-2201. doi:10.1002/cmdc.201800533
- [18] Shaker B, Yu MS, Song JS, Ahn S, Ryu JY, Oh KS, et al. LightBBB: computational prediction model of blood-brain-barrier penetration based on LightGBM. *Bioinformatics*. 2021;37(8):1135-9. doi:10.1093/bioinformatics/btaa918
- [19] Tong X, Wang D, Ding X, Tan X, Ren Q, Chen G, et al. Blood-brain barrier penetration prediction enhanced by uncertainty estimation. *J Cheminform*. 2022;14(1):44. doi:10.1186/s13321-022-00619-2
- [20] Mazumdar B, Deva Sarma PK, Mahanta HJ, Sastry GN. Machine learning based dynamic consensus model for predicting blood-brain barrier permeability. *Comput Biol Med*. 2023;160:106984. doi:10.1016/j.combiomed.2023.106984
- [21] Huang ETC, Yang JS, Liao KYK, Tseng WCW, Lee CK, Gill M, et al. Predicting blood-brain barrier permeability of molecules with a large language model and machine learning. *Sci Rep*. 2024;14(1):15844. doi:10.1038/s41598-024-66897-y
- [22] Nabi AE, Pouladvand P, Liu L, Hua N, Ayubcha C. Machine learning in drug development for neurological diseases: A review of blood brain barrier permeability prediction models. *Mol Inform*. 2025;44(3):e202400325. doi:10.1002/minf.202400325
- [23] Fu L, Shi S, Yi J, Wang N, He Y, Wu Z, et al. ADMETlab 3.0: An updated comprehensive online ADMET prediction platform enhanced with broader coverage, improved performance, API functionality and decision support. *Nucleic Acids Res*. 2024;52(W1):W422-31. doi:10.1093/nar/gkae236
- [24] Banerjee P, Kemmler E, Dunkel M, Preissner R. ProTox 3.0: A webserver for the prediction of toxicity of chemicals. *Nucleic Acids Res*. 2024;52(W1):W513-20. doi:10.1093/nar/gkae303
- [25] Jiang C, Zhao P, Li W, Tang Y, Liu G. In silico prediction of chemical neurotoxicity using machine learning. *Toxicol Res (Camb)*. 2020;9(3):164-72. doi:10.1093/toxres/tfaa016
- [26] Monzel AS, Hemmer K, Kaoma T, Smits LM, Bolognin S, Lucarelli P, et al. Machine learning-assisted neurotoxicity prediction in human midbrain organoids. *Parkinsonism Relat Disord*. 2020;75:105-9. doi:10.1016/j.parkreldis.2020.05.011
- [27] Ryu JY, Kim HU, Lee SY. Deep learning improves prediction of drug-drug and drug-food interactions. *Proc Natl Acad Sci U S A*. 2018;115(18):E4304-11. doi:10.1073/pnas.1803294115
- [28] Sharma AK, Kapoor VK, Kaur G. Herb-drug interactions: A mechanistic approach. *Drug Chem Toxicol*. 2022;45(2):594-603. doi:10.1080/01480545.2020.1738454
- [29] Hirschfeld L, Swanson K, Yang K, Barzilay R, Coley CW. Uncertainty quantification using neural networks for molecular property prediction. *J Chem Inf Model*. 2020;60(8):3770-80. doi:10.1021/acs.jcim.0c00502
- [30] Jiménez-Luna J, Grisoni F, Schneider G. Drug discovery with explainable artificial intelligence. *Nat Mach Intell*. 2020;2(10):573-84. doi:10.1038/s42256-020-00236-4
- [31] Vishwakarma G, Sonpal A, Hachmann J. Metrics for benchmarking and uncertainty quantification: quality, applicability, and best practices for machine learning in chemistry. *Trends Chem*. 2021;3(2):146-56. doi:10.1016/j.trechm.2020.12.004
- [32] Reker D. Practical considerations for active machine learning in drug discovery. *Drug Discov Today Technol*. 2019;32-33:73-9. doi:10.1016/j.ddtec.2020.06.001
- [33] Bender A, Schneider N, Segler M, Walters WP, Engkvist O, Rodrigues T. Evaluation guidelines for machine learning tools in the chemical sciences. *Nat Rev Chem*. 2022;6(6):428-42. doi:10.1038/s41570-022-00391-9
- [34] Vamathevan J, Clark D, Czodrowski P, Dunham I, Ferran E, Lee G, et al. Applications of machine learning in drug discovery and development. *Nat Rev Drug Discov*. 2019;18(6):463-77. doi:10.1038/s41573-019-0024-5
- [35] Bender A, Cortés-Ciriano I. Artificial intelligence in drug discovery: what is realistic, what are illusions? Part 1: ways to make an impact, and why we are not there yet. *Drug Discov Today*. 2021;26(2):511-24. doi:10.1016/j.drudis.2020.12.009