

Network-based drug repurposing for MYH9-related nephritis

Muhammad Ali ¹, Tommaso Gili,² and Guido Caldarelli ^{1,3,4}

¹*DSMN, Ca' Foscari University of Venice, Italy*

²*Networks Unit, Scuola IMT Alti Studi Lucca, Italy*

³*Institute of Complex Systems, CNR-ISC via dei Taurini 19, Rome, Italy*

⁴*London Institute for Mathematical Sciences (LIMS) Royal Institution, 21 Albemarle St., London W1S 4BS, England*



(Received 11 March 2026; accepted 17 June 2026; published 17 August 2026)

Using tools from network theory, we analyze the organization of a MYH9-oriented druglike library in chemical space using a multidescrptor framework. The dataset is drawn from ZINC, a publicly available database of commercially accessible compounds curated for virtual screening and drug discovery. Starting from 6004 molecules, preprocessing yields 5000 structurally valid and descriptor-complete compounds. Similarity is defined via Tanimoto distance on Morgan fingerprints and single-descriptor distances for xLogP, HBD, HBA, molecular weight, and rotatable bonds. For each representation we construct k -nearest-neighbor networks and identify communities using the Louvain–Leiden algorithm. All networks exhibit highly significant modularity ($Q = 0.91 - 0.99$) relative to degree-preserving null models, demonstrating pronounced nonrandom chemical organization across descriptors. Cross-descriptor robustness is quantified through a coclustering matrix over 1.25×10^7 molecular pairs, measuring how consistently compound pairs co-occur within the same community across descriptor-specific networks. Although most pairs show limited agreement, a sparse high-consensus core emerges, highlighting the complementarity of the descriptors. Minimum spanning trees derived from structural and consensus similarities reveal distinct backbone topologies: a scaffold-driven, sparse structure versus a compact, hub-rich consensus network. Betweenness centrality on these backbones identifies compounds that are both structurally central and descriptor-balanced. These results provide a statistically validated network representation of chemical space and a principled strategy to extract consensus-stable compounds for downstream screening.

DOI: [10.1103/v4k3-49dx](https://doi.org/10.1103/v4k3-49dx)

I. INTRODUCTION

Repurposing of drugs identifies new therapeutic uses for approved, failed, or investigational drugs [1,2]. From a complex systems perspective, it poses a challenging question: how to identify functional equivalence between distinct perturbations acting on a highly interconnected and heterogeneous network [3]. Compounds initially developed for different therapeutic indications may elicit similar physiological responses if they interact with the human organism through analogous structural and chemical mechanisms. Quantifying such similarity in a principled and scalable manner remains an open challenge.

Most existing approaches to drug repurposing are based on shared molecular targets, transcriptional signatures, or phenotypic correlations [4]. Although effective in specific domains, these strategies often neglect the physical constraints imposed by molecular geometry and chemical composition. However, from a system-level perspective, drug–body interactions are shaped by molecular shape, topology, and physicochemical properties, which determine binding affinities, off-target

effects, and the propagation of perturbations across biological interaction networks [5].

Machine-learning approaches have become increasingly prominent in drug repurposing over the past decade [6]. Supervised models, including support-vector machines, random forests, and gradient-boosted trees, classify compound–disease pairs from features such as molecular fingerprints, physicochemical descriptors, and gene-expression profiles. They have been successfully applied to quantitative structure–activity relationship (QSAR) modeling and to large-scale virtual screening when sufficient labeled training data are available. Graph neural networks (GNNs) have further extended this framework by operating directly on the bond topology of molecular graphs, encoding chemical structure without hand-crafted descriptors [7]. GNNs trained on multimodal networks that combine protein–protein interactions, drug–protein targets, and pharmacological side effects have demonstrated substantially higher predictive accuracy than feature-based baselines on drug–drug interaction prediction [8], and the same graph-learning architecture has been applied to drug repurposing by framing it as a link prediction problem on biological knowledge graphs [9]. A further class of methods relies on knowledge graphs that integrate heterogeneous biomedical resources, including genes, diseases, compounds, pathways, and anatomies, into a single unified network. Graph embeddings learned from these structures enable link prediction across drug–disease pairs at scale and

Published by the American Physical Society under the terms of the Creative Commons Attribution 4.0 International license. Further distribution of this work must maintain attribution to the author(s) and the published article's title, journal citation, and DOI.

have been applied to prioritize repurposing candidates across a broad range of indications [10,11].

Despite their predictive power, data-driven methods face structural limitations that are particularly acute for rare or poorly characterized conditions such as MYH9-related nephritis. Supervised models require large labeled datasets of confirmed active and inactive compounds; knowledge-graph approaches require dense prior annotation of drug–target and drug–disease associations. For rare nephropathies with no approved targeted therapy, such gold-standard labels are scarce, severely constraining both training and generalization. Deep learning models are furthermore typically opaque, which limits the chemical interpretability of their outputs and complicates the extraction of actionable hypotheses for medicinal chemistry. In addition, GNN and knowledge-graph methods operate primarily on the biological interaction layer, linking drugs through protein targets or phenotypes, and do not directly interrogate the intrinsic physicochemical and structural organization of the compound library itself. Compounds that lack target annotation may therefore be overlooked, even when they share the same structural scaffold or physicochemical profile as known actives.

Network theory [12,13] provides a natural framework for addressing this problem. Biological organization spans multiple interconnected networks, including protein–protein interaction networks [14], signaling pathways, and organ-level functional dependencies. In this context, drugs can be interpreted as external perturbations whose effects are mediated by their embedding within these networks. Network-based approaches have demonstrated that topological proximity and shared neighborhoods can predict functional similarity and drug efficacy [3,15,16].

To integrate structural and chemical information into this framework, it is necessary to move beyond single-layer representations. Multilayer network formalisms allow heterogeneous interaction domains to be encoded simultaneously within a unified mathematical structure [17,18]. In such representations, molecular shape similarity, physicochemical distance, and biological interaction profiles can be treated as coupled layers, enabling analysis of how local molecular features influence global network organization.

In this work we propose a network-theoretic approach to drug repurposing based on the similarity of drug–body interaction topologies. Drugs are represented as nodes in a multilayer network, with weighted edges that encode affinities derived from physical shape and chemical properties. Repurposing opportunities emerge as topological proximity, clustering, or community overlap between drugs associated with distinct therapeutic classes.

By reframing drug repurposing as a problem of similarity and distance in a complex network, this approach shifts the focus from isolated targets to system-level interaction patterns. The resulting framework is naturally aligned with the statistical physics of complex systems and provides a transparent, interpretable basis for identifying candidate drugs, complementary to existing data-driven methodologies [19]. The specific case study we present here is the exploration of a subset of an extensive database of approved chemical substances (with unknown efficacy), namely, the ZINC database [20], reduced

to all compounds that may be used for the treatment of MYH9 nephritis [21].

MYH9 encodes the heavy chain IIA of nonmuscle myosin, a structural component essential for cytoskeletal organization and podocyte stability in renal tissues. Mutations in MYH9 are directly associated with MYH9-related nephropathies, characterized by proteinuria, progressive renal decline, and an increased risk of end-stage kidney disease. Currently, there are no FDA-approved targeted therapies for MYH9-driven nephritis, and clinical management remains largely supportive. Drug-repurposing strategies that exploit existing drug-like chemical libraries offer a viable route to identify molecules capable of stabilizing MYH9-dependent pathways or modulating downstream dysfunction. Therefore, computational prioritization of structurally and physicochemically reliable compounds becomes critical to narrow the search space before biochemical evaluation.

Our approach allows us to explore a curated library of druglike compounds assembled for nephritis-related therapeutic exploration, with a particular focus on MYH9-associated nephritic disorders. The pipeline combines molecular fingerprints, descriptor-based similarity, community detection, null-model testing, and minimum spanning tree backbones within a unified network-science framework. This activity leverages the principles of complex systems to characterize connectivity patterns, quantify robustness, and identify structurally influential molecules that may serve as candidate compounds for downstream target-specific applications. By integrating complementary molecular representations and grounding them in rigorous statistical methods, we reveal the underlying organization of this MYH9-relevant chemical space and provide a reproducible, interpretable pathway for compound prioritization.

These compounds should ideally reflect robust similarity relationships independent of descriptor choice and therefore represent ideal candidates for subsequent docking, virtual screening, optimization, and repurposing workflows targeting MYH9-associated nephritis.

II. RESULTS

With a curated set of 5000 structurally valid and complete descriptor molecules, the first analytical step was to examine how each descriptor organizes chemical space. Because such descriptors capture different chemical principles—structural motifs, hydrophobicity, polarity, functional groups, size, and conformational freedom—the networks constructed from them are expected to differ substantially. Such differences determine the types of communities that emerge, the degree to which the descriptors agree, and the need for a consensus analysis. Therefore, each descriptor-specific network offers a distinct view of how the MYH9-focused compound library is organized.

A. Topology

To observe these descriptor-specific geometries, we constructed six k -nearest-neighbor (kNN) similarity networks (Fig. 1), one for each descriptor: Simplified molecular input line entry system (SMILES) fingerprints, xLogP, HBD, HBA, MW, and ROTB (see Sec. V D for

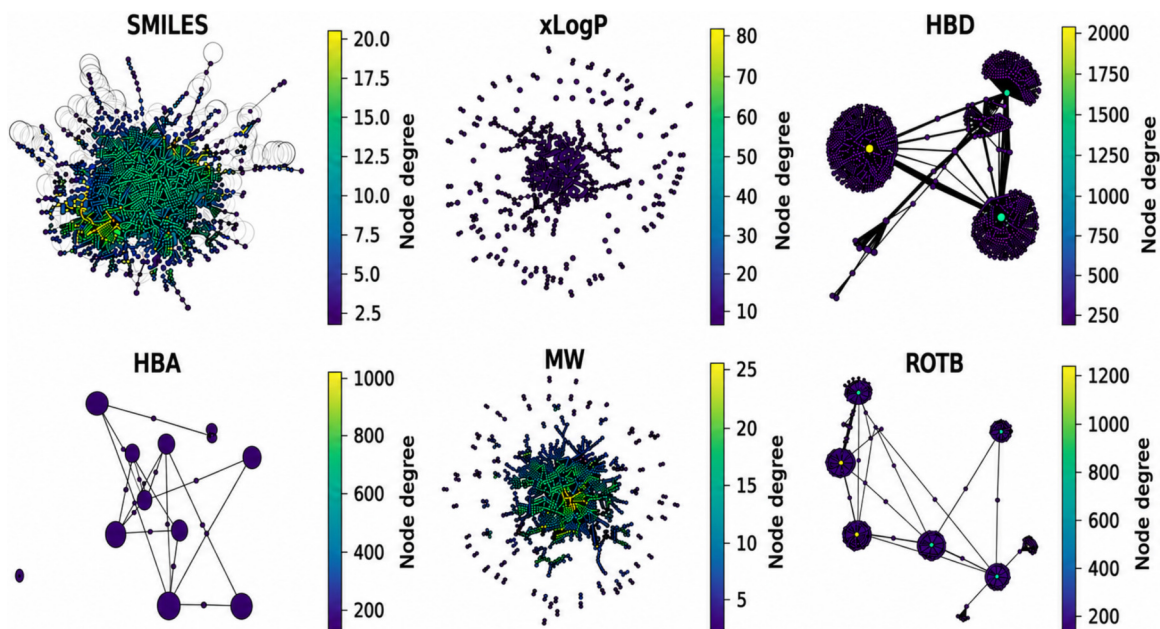


FIG. 1. Descriptor-specific kNN networks for SMILES, xLogP, HBD, HBA, MW, and ROTB.

definitions of each). Each network retains at least $k = 5$ strongest neighbors per node. Despite identical node sets, the graphs differ in edge density and mesoscale organization, reflecting principles of descriptor-specific organization. Each network contains the same 5000 nodes, but their connectivity patterns differ because each descriptor encodes chemical similarity differently. For the SMILES network, the Tanimoto similarity S_{ij}^{Tan} from Eq. (4) was used in Morgan fingerprints (radius 2, 2048 bits) to quantify the substructural relatedness among molecules. A force-directed layout shows a dense core with peripheral lobes, reflecting structurally related molecular families.

For descriptor-based networks, we use z -score normalization, followed by the similarity function of Eq. (5). The edge counts were $\text{xLogP} = 20\,466$; $\text{HBD} = 26\,961$; $\text{HBA} = 22\,896$; $\text{MW} = 15\,865$; $\text{ROTB} = 24\,760$ (Table II). These differences anticipate the observed variation in the community structure and consensus stability between descriptors. The SMILES network had $N = 5000$ nodes and $E = 17\,112$ edges, with the largest connected component containing 4962 nodes, average degree 6.844, average path length 7.21, and density 0.00137. Descriptor-based networks showed comparable sparsity but different mesoscale structures.

The SMILES network forms a dense, scaffold-driven structure. The xLogP network forms a compact hydrophobicity-driven core. HBD yields discrete donor-count clusters; HBA shows a functional-group-driven substructure; MW forms continuous gradient windows; ROTB reveals rigid vs flexible modules.

Because each descriptor produces a distinct topology, descriptor-specific community detection becomes essential.

B. Communities

We then proceeded by applying the Louvain–Leiden community detection method [22] to each descriptor-specific network. When the vertices are reordered according to their community membership, the adjacency matrices reorganize into a block-diagonal form. Dense diagonal blocks correspond to communities with strong internal connectivity, while off-diagonal regions are sparse. This rearrangement does not alter edge weights or network structure but only changes the ordering of vertices. The resulting pattern visually confirms the presence of well-defined communities in the network. These are shown in Fig. 2. SMILES communities reflect scaffold-level diversity, grouping molecules around recurring aromatic frameworks, heterocycles, or fragment motifs (Fig. 2), while

TABLE I. Communities per descriptor.

Descriptor	Communities
SMILES	98
xLogP	149
HBD	5
HBA	11
MW	154
ROTB	9

TABLE II. Original and randomized edge counts for degree-preserving null models.

Descriptor	Original edges	Randomized edges
SMILES	17 112	17 110
xLogP	20 466	20 466
HBD	26 961	26 961
HBA	22 896	22 896
MW	15 865	15 865
ROTB	24 760	24 760

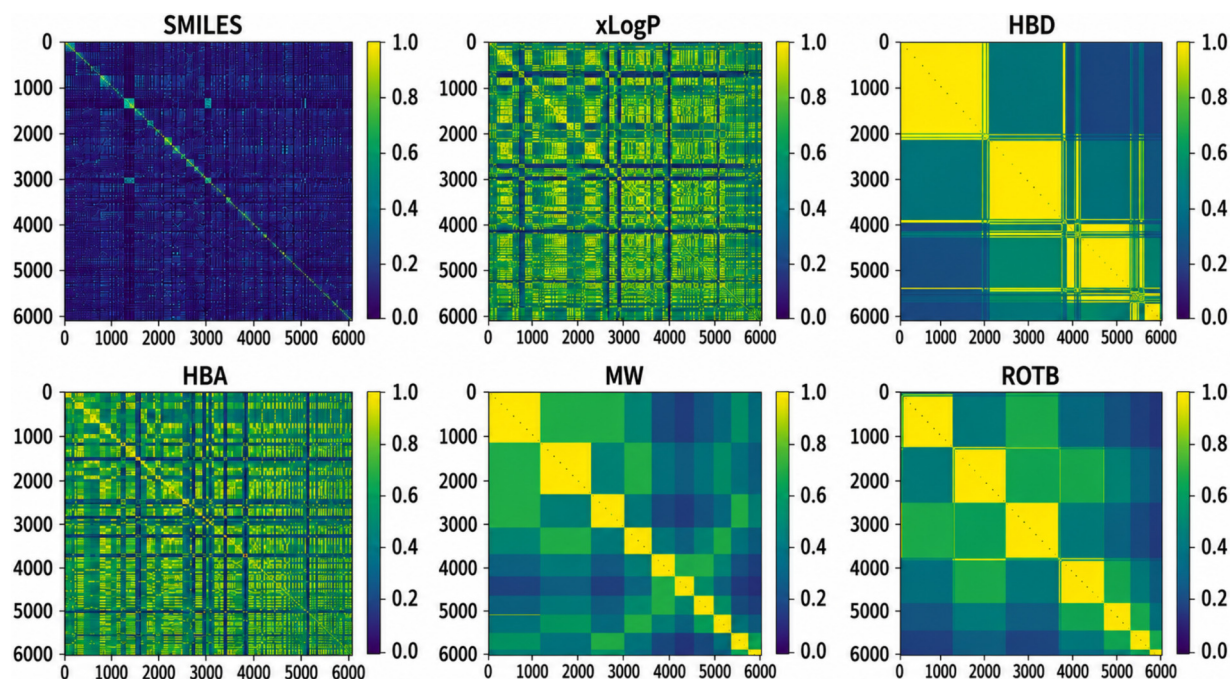


FIG. 2. Clustered similarity matrices (SMILES, xLogP, HBD, HBA, MW, ROTB).

xLogP communities form broader hydrophobicity-based clusters reflecting lipophilicity regimes. HBD and HBA induce functional-group-driven clusters based on hydrogen-bonding capacity. MW produces numerous microclusters across narrow mass ranges. ROTB separates rigid molecules from highly flexible ones. Number of communities per descriptor are shown in Table I. This differentiation confirms that chemical similarity depends strongly on the selected descrip-

tor; no single representation adequately captures the chemical space, and descriptor-specific networks uncover complementary molecular relationships. These findings justify the use of multiview frameworks rather than relying solely on structural or physicochemical models.

The numbers of communities for the different cases is reported in Table I. To determine whether the descriptor-specific communities identified by Louvain clustering reflect

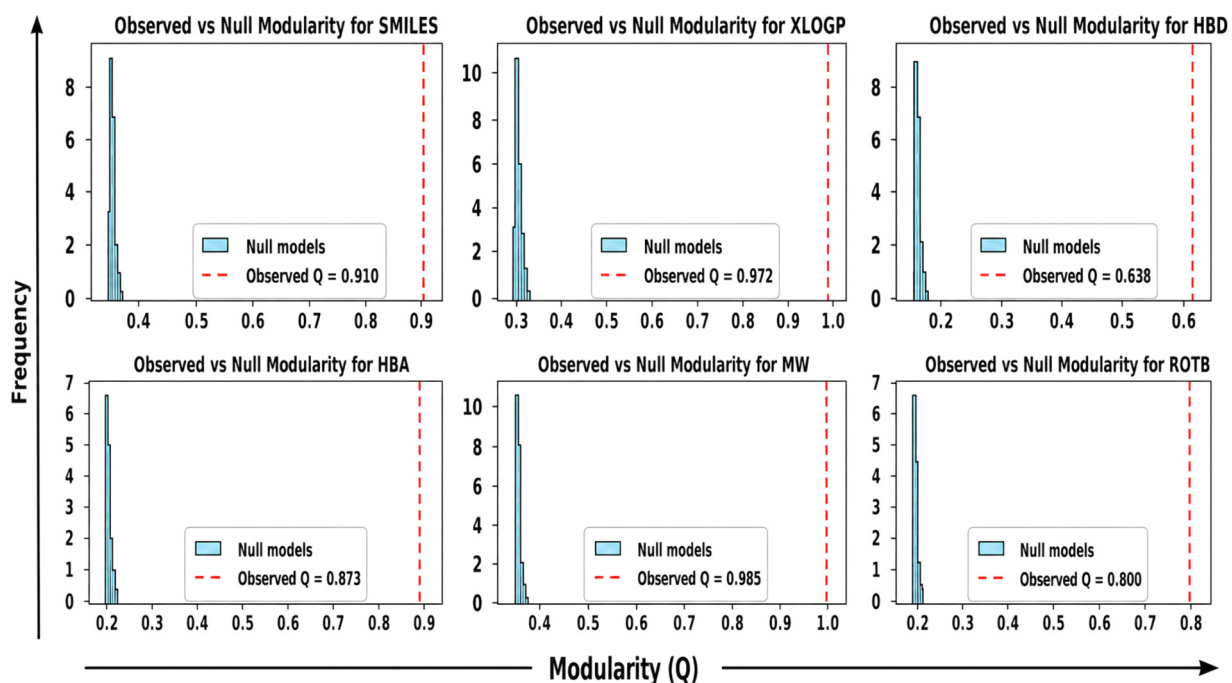


FIG. 3. Observed vs null modularity distributions for the six descriptor-specific kNN networks (SMILES, xLogP, HBD, HBA, MW, and ROTB).

TABLE III. Observed modularity values (Q) compared with degree-preserving null models.

Network	Q_{obs}	Null mean $\pm$ SD
SMILES	0.910	0.350 ± 0.001
xLogP	0.972	0.306 ± 0.002
HBD	0.638	0.152 ± 0.000
HBA	0.873	0.228 ± 0.001
MW	0.985	0.366 ± 0.002
ROTB	0.800	0.190 ± 0.001

genuine chemical structure rather than random graph artifacts, we compared the observed modularity values with those from 50-degree-preserving randomized networks for each descriptor. Original and randomized edge counts are shown in Table II.

The distribution of modularity values for the null models was consistently lower than that of the empirical networks. The observed modularity values (Q_{obs}) were significantly higher ($p < 0.001$), confirming that the community structures detected were statistically robust and not artifacts of degree correlation. Modularity is defined as Eq. (6).

Figure 3 shows the observed vs null modularity distribution for the six descriptor-specific kNN networks (SMILES, xLogP, HBD, HBA, MW, and ROTB). For each network the blue histogram shows the modularity distribution of 50-degree-preserving randomized null models, while the red dashed line marks the observed modularity Q_{obs} . In all cases, Q_{obs} lies far outside the null distribution, indicating a substantial and statistically significant modular organization driven by a descriptor-specific similarity structure in Table III. For each descriptor, none of the 50-degree-preserving null networks achieved a modularity value equal to or greater than the observed value, producing an empirical p value of 0. This confirms that the detected communities are statistically significant and reflects a meaningful chemical organization rather than random partitioning. Because descriptor-level community structures are robust, it is appropriate to examine agreement across descriptors through coclustering analysis.

C. Cross-descriptor consensus via coclustering

To quantify how consistently the six chemical descriptors agree on molecular similarity, we constructed the coclustering matrix C based on their corresponding community assignments using Eq. (8). For each descriptor v , molecules were clustered independently, and for each molecular pair (i, j) , C_{ij} counts the number of descriptors for which molecules i and j appear in the same community. With 5000 molecules, the resulting 5000×5000 matrix captures all 12 497 500 unique molecular pairs, allowing a systematic assessment of the agreement, complementarity, and consensus structure of the descriptor. This analysis was performed over community structures derived from the six kNN-based similarity networks built on SMILES fingerprints, xLogP, HBD, HBA, MW, and ROTB.

TABLE IV. Number of molecular pairs at each coclustering level across the six descriptors.

Coclustering level	Descriptor agreement	Number of pairs
0	0 descriptors	9 332 928
1	1 descriptor	6 608 897
2	2 descriptors	1 803 216
3	3 descriptors	248 689
4	4 descriptors	21 487
5	5 descriptors	1466
6	All 6 descriptors	4323

1. Global distribution of coclustering values

The distribution of C_{ij} across all 12.5 million molecular pairs is strongly skewed toward low cross-descriptor agreement. Figure 4 illustrates the full distribution of coclustering frequencies; Table IV summarizes the exact frequencies. Figure 4 illustrates the full distribution of coclustering frequencies across the 12.5 million molecular pairs. The steep decay shows that most pairs share zero or one descriptor-based community, confirming strong descriptor complementarity. At the same time, the histogram has a long right tail that represents a very small number of high-consensus molecular relationships (shared in 5–6 descriptors), which form robust, descriptor-invariant connections.

The heatmap in Fig. 5 provides a visual understanding of how often selected molecules cocluster across descriptors. Cells closer to yellow indicate stronger agreement with the descriptor.

Most molecular pairs fall into the low-agreement range (0–2 descriptors), reflecting significant descriptor diversity. To facilitate interpretation, the levels were grouped into three agreement categories (Table V).

Low agreement dominates the dataset, indicating extremely strong complementarity among the descriptors.

2. Descriptor complementarity overview

The predominance of low coclustering frequencies reflects strong complementarity among the six descriptors. Table VI shows the percentage of pairs at each exact agreement level. Table VI reports the proportion of molecular pairs at each exact coclustering level.

The extremely high proportion of level-0 pairs shows the rarity of strong descriptor overlap. These results confirm that different descriptors emphasize distinct structural or physicochemical dimensions, leading to divergent community assignments.

TABLE V. Aggregated coclustering categories.

Category	Descriptor agreement	Number of pairs	% of all pairs
Low agreement	0–2 descriptors	17 745 041	$\approx 99.0\%$
Moderate agreement	3–4 descriptors	270 176	$\approx 2.16\%$
High agreement	5–6 descriptors	5789	$\approx 0.046\%$

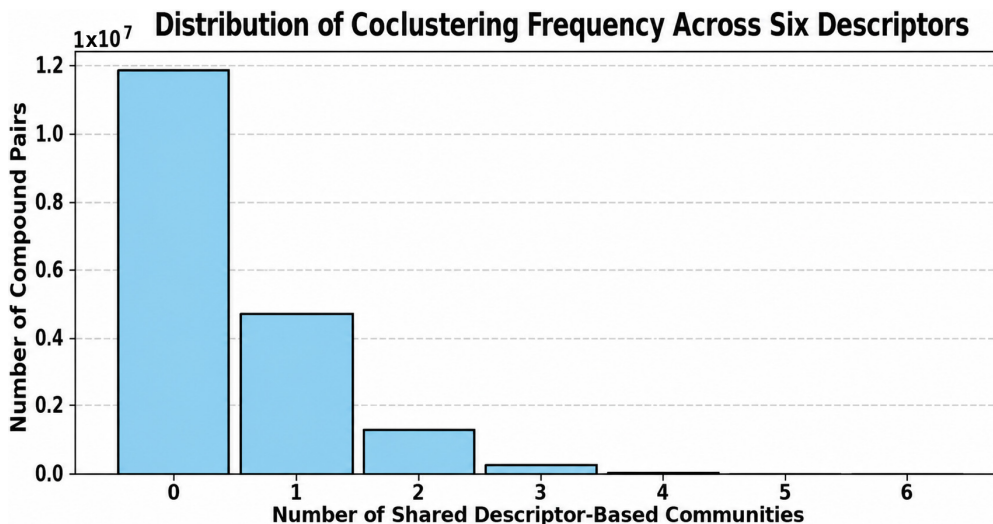


FIG. 4. Distribution of coclustering frequency across six descriptors.

3. Per-compound descriptor stability

To examine molecular behavior at the compound level, we computed three metrics for each of the 5000 molecules.

Unique cluster membership. First, the number of distinct cluster labels a molecule i receives across descriptors is

$$U(i) = |\{c_i^{(v)} : v = 1, \dots, V\}|. \quad (1)$$

Number of molecules sharing ≥ 1 cluster. Second, the number of molecules sharing at least one cluster with i is

$$N_{\geq 1}(i) = |\{j : C_{ij} \geq 1\}|. \quad (2)$$

Total coclustering strength. Third, the global descriptor agreement for molecule i is given by the total coclustering strength

$$S_{\text{tot}}(i) = \sum_j C_{ij}. \quad (3)$$

Illustrative examples for five molecules are reported in Tables VII–IX.

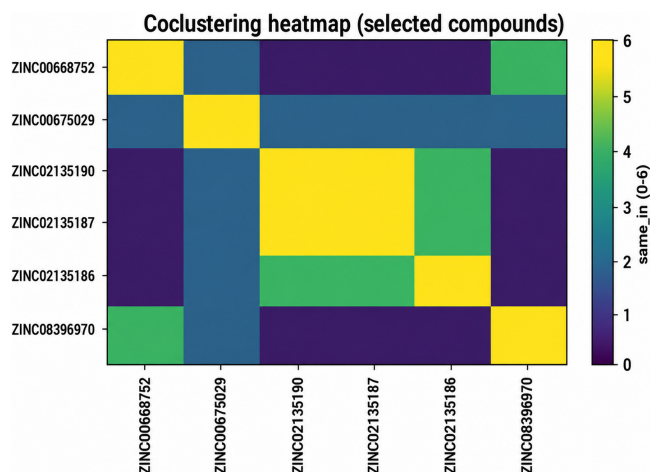


FIG. 5. Coclustering heatmap for a representative subset of molecules.

Unique cluster counts indicate how consistently each molecule is placed across descriptors. Most molecules fall between three and five clusters, indicating moderate descriptor coherence: three clusters reflect strong agreement and stable neighborhoods, while five to six clusters indicate increasing disagreement among descriptors and multimodal behavior.

Molecules typically share a cluster with 3000 others, reflecting a densely connected chemical. Each molecule shares a cluster with 2900–3100 others, demonstrating that the chemical space is highly interconnected even before enforcing consensus space.

Total coclustering strength measures global consensus for each molecule across all descriptors. Higher values correspond to molecules with strong, consistent similarity patterns across the descriptors. Lower values indicate descriptor-specific or inconsistent clustering behavior.

4. Consolidated interpretation

Table X consolidates the main descriptor-based coclustering insights.

The combined analysis of descriptor-based coclustering reveals a chemically diverse, multimodal landscape in which different descriptors capture distinct structural and physico-chemical relationships. The overwhelmingly high proportion of low-agreement molecular pairs demonstrates that these descriptors encode largely complementary information. Despite

TABLE VI. Proportion of molecular pairs at each exact coclustering level.

Level	Percent of all pairs
0	74.7%
1	52.9%
2	14.4%
3	1.99%
4	0.17%
5	0.011%
6	0.034%

TABLE VII. Unique cluster membership across descriptors (example 5 compounds).

ZINC ID	Unique cluster count
ZINC00668752	4
ZINC00675029	3
ZINC02135190	5
ZINC02135187	5
ZINC02135186	4

this heterogeneity, a small but meaningful subset of molecules consistently aligns across multiple descriptors, forming a stable consensus core that reflects shared underlying chemical principles.

These results highlight the importance of using multi-descriptor frameworks for chemical-space exploration and drug repurposing. By integrating neighborhoods derived from structurally distinct descriptors, one can obtain a richer, more informative representation of molecular similarity. This enhanced view helps identify robust chemical families, reveals stable connectivity patterns, and improves the reliability of downstream tasks such as target prediction, clustering, and screening. Overall, the coclustering analysis demonstrates both the diversity and the hidden coherence within the chemical landscape, reinforcing the value of multiperspective molecular characterization.

The coclustering results show that only 0.046% of molecular pairs achieve agreement across five or six descriptors, forming a small but highly stable consensus core. These high-consensus relationships are statistically rare yet chemically robust, reflecting descriptor-independent similarity. Because the consensus minimum spanning tree (MST) aims to preserve only the most reliable cross-descriptor relationships, these high-agreement molecular pairs naturally serve as structural anchors. Their consistent behavior across multiple chemical perspectives makes them ideal backbone edges for constructing a unified, consensus-based topology.

D. Structural and consensus backbones via minimum spanning trees

To derive global, interpretable backbones of chemical space under both structural and consensus similarity, we constructed two MSTs: one based on SMILES-derived structural similarity and one based on consensus similarity aggregated from all six descriptors (Fig. 6). Distances for both MSTs were computed as in Eq. (9), where S_{ij} denotes the pairwise similarity. For the consensus MST, similarities were

TABLE VIII. Number of molecules sharing at least one community with each example compound.

ZINC ID	Molecules sharing ≥ 1 cluster
ZINC00668752	2917
ZINC00675029	3073
ZINC02135190	3013
ZINC02135187	3013
ZINC02135186	2968

TABLE IX. Total coclustering strength for example molecules.

ZINC ID	Total coclustering strength
ZINC00668752	4133
ZINC00675029	4093
ZINC02135190	4028
ZINC02135187	4028
ZINC02135186	3969

defined by Eq. (10), with C_{ij} representing the number of times molecules i and j coclustered across the six descriptor-specific networks.

The SMILES-based MST consisted of 4972 nodes and 4972 edges, with 2847 leaf nodes (57.2%). The maximum degree was 9, the diameter reached 43, and the average path length was 17.4. This topology reflects a scaffold-dominated organization, characterized by elongated chains of structural variants connected by relatively sparse bridging compounds.

In contrast, the consensus MST contained 4995 nodes and 4995 edges, including 2601 leaf nodes (52.1%). It exhibited a higher maximum degree (12), a reduced diameter (31), and a shorter average path length (12.8). These properties indicate a more compact and integrated backbone, where molecules consistently grouped across multiple descriptors form a tighter, more cohesive chemical neighborhood. Taken together, the SMILES MST and consensus MST reveal two distinct connectivity regimes: structural adjacency and multi-descriptor chemical coherence. These form the foundation for identifying central molecular hubs.

E. Identification of backbone key compounds

To determine which molecules play the most influential bridging roles within the structural and consensus MSTs, we computed betweenness centrality for every node using Eq. (11). Betweenness centrality quantifies how often a molecule lies along the shortest paths between pairs of other molecules. High-centrality molecules serve as critical connectors between chemical families, and their importance differs depending on whether structural or consensus similarity is used. Table XI shows the global properties of SMILES-based and consensus-based MSTs.

TABLE X. Summary of descriptor-based coclustering insights.

Aspect	Key observation
Descriptor complementarity	$\approx 99\%$ of molecular pairs agree in only 0–2 descriptors
Moderate consensus	2.16% of pairs align in 3–4 descriptors
High-consensus core	Only 0.046% of pairs agree in 5–6 descriptors
Unique cluster stability	Most molecules appear in 3–5 clusters across descriptors
Descriptor interconnectedness	Each molecule coclusters with ~ 3000 others in at least one descriptor.
Global coclustering strength	Values typically range from 3900–4200 across molecules.

TABLE XI. Global properties of SMILES-based and consensus minimum spanning trees (MSTs).

Property	SMILES-based MST	Consensus MST
Nodes	4972	4995
Edges	4972	4995
Leaf nodes	2847 (57.2%)	2601 (52.1%)
Maximum degree	9	12
Diameter	43	31
Average path length	17.4	12.8

Table XII summarizes betweenness-centrality statistics for the top 50 nodes in each MST.

The top ten high-centrality compounds in the SMILES MST are reported in Table XIII.

These molecules form key structural connectors in scaffold space, bridging otherwise distant structural motifs.

The top ten high-centrality compounds in the consensus MST are summarized in Table XIV.

In contrast to the SMILES MST hubs, these consensus hubs represent molecules that exhibit multidescrptor coherence, aligning across structural, functional, hydrophobic, and conformational dimensions. Together, the two hub sets provide a dual view of chemical importance within the MYH9-oriented library: one reflecting pure structural transitions in scaffold space, the other reflecting balanced multiparameter chemical similarity. The latter, in particular, highlights consensus-stable compounds that are well suited as prioritized candidates for subsequent MYH9-focused docking, virtual screening, or repurposing studies.

III. DISCUSSION

This work established a comprehensive multidescrptor analysis framework for characterizing the structural organization of a MYH9-oriented, druglike chemical space, integrating

TABLE XII. Summary statistics of betweenness centrality for top 50 nodes in each MST.

Network	Count	Mean	Median	SD	Min.	Max.
Consensus MST	50	0.3743	0.4092	0.0821	0.2352	0.6557
SMILES MST	50	0.3738	0.3516	0.0791	0.2540	0.5279

similarity modeling, network representations, community detection, null-model validation, coclustering, and backbone extraction. By systematically comparing descriptor-specific similarity networks and identifying consensus-stable molecular relationships, we demonstrated how a chemically meaningful structure emerges across multiple representational perspectives within a therapeutically relevant library. The resulting findings provide a reproducible foundation for interpreting molecular organization and identifying compounds that play a central role in the similarity landscape of a MYH9-focused repurposing set.

We observed that structural diversity manifests itself differently in the six descriptor-specific networks. The SMILES-based network forms a dense scaffold-driven space characterized by extensive structural transitions among aromatic frameworks, heterocycles, and hybrid scaffolds. In contrast, descriptor-derived networks—particularly HBD and ROTB—produce discrete partitions that reflect categorical chemical counts rather than continuous gradients. xLogP, MW, and HBA exhibit intermediate behavior in which molecules form cohesive neighborhoods based on hydrophobicity or sizes. These observations confirm that chemical similarity is strongly representation-dependent, that no single descriptor captures the entire landscape, and that descriptor-specific networks uncover complementary molecular relationships.

Community detection applied to each descriptor-specific network reveals clear mesoscale organization, with distinct patterns emerging across descriptors. Structural fingerprints (SMILES), MW, and xLogP produce highly granular

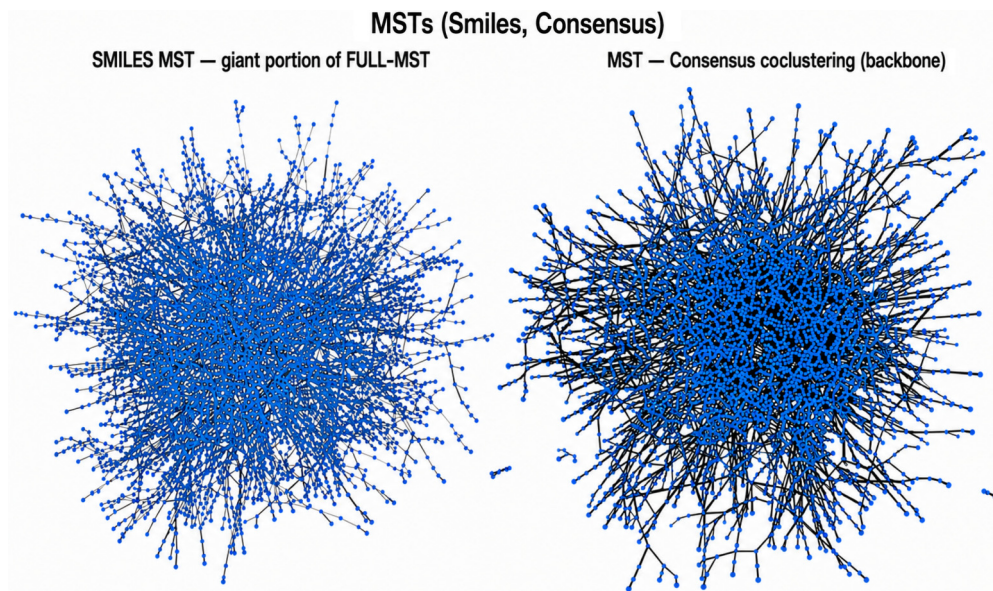


FIG. 6. SMILES MST vs consensus MST.

TABLE XIII. Top ten high-centrality compounds in the SMILES MST.

Rank	Compound ID	Betweenness
1	ZINC00670787	0.527908
2	ZINC01895029	0.517964
3	ZINC08399011	0.503845
4	ZINC00674567	0.486375
5	ZINC02134860	0.469088
6	ZINC00670791	0.463146
7	ZINC02134862	0.459930
8	ZINC02134855	0.458765
9	ZINC08399014	0.458300
10	ZINC00675029	0.454215

partitions that reflect subtle and continuous gradients in molecular variation. Conversely, discrete descriptors yield fewer chemically interpretable communities tied to donor/acceptor counts or conformational flexibility. Degree-preserving null models consistently yield substantially lower modularity distributions than observed values, demonstrating that the identified communities are not artifacts of network topology or density. This statistical validation confirms that the detected clusters correspond to reproducible, chemically meaningful similarity relationships rather than to algorithmic bias.

The coclustering analysis reveals a striking observation: most molecular pairs agree on only a few descriptors. Approximately 99% pairs match in zero, one, or two descriptors, highlighting strong representational complementarity across molecular features. Only $\sim 0.046\%$ pairs exhibit agreement on five or six descriptors, forming a small but extremely stable consensus core. These high-agreement relationships represent descriptor-independent similarities, indicating molecular pairs that consistently appear close in structural, hydrophobic, bonding, and flexibility scales. This subset serves as the foundation for a consensus-based network backbone.

Minimum spanning trees expose backbone connectivity under both structural (SMILES-based) and consensus similarity. The SMILES-based MST reflects a scaffold-driven structure with long branches, limited hubs, and large diameter, illustrating gradual structural transitions. The consensus MST

TABLE XIV. Top ten high-centrality compounds in the consensus MST.

Rank	Compound ID	Betweenness
1	ZINC00662981	0.655728
2	ZINC02101627	0.486345
3	ZINC00995741	0.462674
4	ZINC08399817	0.461278
5	ZINC08399809	0.452311
6	ZINC00655084	0.442167
7	ZINC00675029	0.438614
8	ZINC08399990	0.436192
9	ZINC00668752	0.434819
10	ZINC08396970	0.431860

is more compact, with higher-degree hubs and shorter path lengths, because multi-descriptor-consistent molecules act as central connectors. This duality is critical for drug discovery: scaffold hopping emerges naturally from the structural backbone, while multiparameter optimization is reflected in the consensus backbone. Ideally, druglike molecules lie near consensus hubs.

Betweenness centrality further highlights the distinct classes of central molecules. Structural-MST central compounds bridge scaffold clusters and link remote regions of chemical space. The consensus-MST central compounds are multiproperty balanced molecules capable of connecting descriptor-stable regions. Such molecules are ideal candidates for prioritizing balanced chemical profiles or representative scaffolds and may serve as lead structures for subsequent screening or analog development.

Finally, it is important to emphasize that this study focuses on structural organization and consensus-driven prioritization rather than direct prediction of biological interactions. Docking-based interaction assessment and biological validation (e.g., podocyte-specific assays, MYH9-dependent pathway readouts) constitute natural downstream steps once a statistically justified subset of candidates has been identified. By constructing an interpretable and robust chemical-space backbone, the present framework provides a rational prefilter for subsequent target-specific screening pipelines.

IV. CONCLUSION

This study presented an integrated chemical-space analysis framework tailored to a collection of MYH9-relevant, drug-like compounds. Using multiple molecular representations and network-science principles, we systematically identified descriptor-stable relationships, statistically meaningful community structures, and backbone connectivity patterns that reveal how molecules organize within this therapeutic search space. The comparison between structural and consensus-based networks demonstrated that only a very small subset of molecules consistently align across multiple descriptors, identifying a chemically robust core that forms reliable connectivity backbones. These consensus hubs, together with their high centrality in the network, provide interpretable indicators of balanced chemical similarity and structural representativeness, thus offering rational entry points for downstream drug-repurposing workflows related to MYH9-associated nephritis.

Importantly, the prioritization provided here serves as a necessary intermediate computational stage before targeted-specific screening or biological evaluation. Rather than performing unbiased docking across thousands of compounds, our pipeline reduces chemical redundancy, highlights meaningful scaffold transitions, and isolates compounds that exhibit reproducible multiparameter coherence. These results establish an informed foundation for target-oriented modeling, including structure-based screening, pharmacophore refinement, network pharmacology, and pathway-level validation of MYH9-modulating candidates. In general, this work demonstrates that structurally grounded, statistically validated backbone extraction can serve as a powerful filtering and prioritization layer in repurposing pipelines targeting rare or

underexplored nephropathic mechanisms, including MYH9-associated pathological phenotypes.

V. METHODS

A. Data preprocessing and descriptor resolution

To ensure uniform downstream analysis across all similarity spaces, descriptor headers were normalized through a robust synonym-matching routine. When any of the five descriptors were missing, but a valid SMILES was available, we computed the values directly from the structure using RDKit (Crippen logP for xLogP; Lipinski donor/acceptor counts; and exact molecular weight). Molecules with invalid or unparsable SMILES strings and rows lacking essential descriptors were removed from the working set. This quality-control step led to a consistent index of 5 000 chemically valid entries across all network analyses. This reduction from 6 004 to 5 000 compounds is expected in cheminformatics workflows when fingerprints or descriptor computations fail for a subset of structures; keeping a standard index enables one-to-one comparison between descriptor-specific networks and prevents missing-node artifacts in integrated analyses.

B. Network definition

Starting from the above data each **Node** in the graph corresponds to a single compound in the dataset, identified by its ZINC ID, and the node set is identical across all networks considered. **Edges** vary in their definition but always represent similarity relationships between nodes defined through a k -nearest-neighbor (kNN) graph construction. We made two different methods for two different questions. In one case we considered the structural properties, and in the second, the physicochemical properties.

For the second case, the edges encode similarity along a single physicochemical dimension (defined below) rather than structural overlap. Two compounds are connected if they are among each other's nearest neighbors in a z-score-normalized descriptor space, with edge weights reflecting numerical proximity in that specific property.

C. Networks from molecular structure

For the first case, starting from the SMILES coding of a compound, we considered its *Morgan fingerprints* (see below) and created an edge between nodes i, j when the Tanimoto/Jaccard similarity S_{ij} , defined as

$$S_{ij}^{\text{Tan}} = \frac{|A_i \cap A_j|}{|A_i \cup A_j|} = \frac{|A_i \cap A_j|}{|A_i| + |A_j| - |A_i \cap A_j|}, \quad (4)$$

is above a certain threshold; the edge weights equal the corresponding Tanimoto coefficient [23,24] on circular Morgan fingerprints (radius 2, 2048 bits) to quantify substructural relatedness between molecules. From the full similarity matrix, a k -nearest-neighbor (kNN) graph with $K = 5$ was constructed to retain the most informative edges while ensuring sparse, interpretable topology. Edge weights encode similarity in $[0, 1]$. The resulting network contains 5000 nodes and 17 112 weighted edges. A force-directed layout reveals a dense, cohesive core and several peripheral lobes, consistent

with families of structurally related compounds linked by strong intrafamily edges and bridged by weaker interfamily ties.

Morgan fingerprints are a widely used molecular representation in cheminformatics designed to encode the topological structure of a molecule into a fixed-length numerical vector, enabling efficient comparison, similarity search, and machine-learning analysis. They belong to the family of circular fingerprints and are most commonly known for their implementation as extended connectivity fingerprints (ECFPs) [25]. The core idea behind Morgan fingerprints is to represent a molecule by systematically characterizing the local atomic environments around each atom. Initially, each atom is assigned an identifier based on basic atomic properties such as atomic number, valence, charge, and aromaticity. These identifiers are then iteratively updated by aggregating information from neighboring atoms, effectively expanding the atomic environment to increasing radii. At each iteration, the resulting substructures are encoded and mapped, typically via hashing, into a high-dimensional binary or count-based vector. The final fingerprint captures the presence of specific substructural patterns within the molecule, reflecting its connectivity and chemical functionality.

D. Networks from descriptor-based networks

To complement structural similarity, we constructed five descriptor-specific kNN networks using z-score normalization followed by a one-dimensional similarity function

$$S_{ij} = \frac{1}{(1 + |z_i - z_j|)}, \quad (5)$$

where z_i and z_j are z-score normalized descriptor values for compounds i and j . This transformation scales molecular similarity between 0 and 1, allowing uniform comparison across descriptors.

(1) **xlogP** that captures lipophilicity; the network shows multiple moderately dense regions with hubs indicative of lipophilic series. Formally it is defined as $x\log P = \log_{10}(c_o/c_w)$, where c_o and c_w denote the equilibrium concentrations in octanol and water, respectively.

(2) **HBD**, i.e., hydrogen bond donors. This is a molecular descriptor that counts the number of functional groups in a molecule capable of donating a hydrogen atom in a hydrogen bond. Typical HBD groups include: OH (alcohols, phenols); -NH, -NH₂ (amines, amides); -SH (thiols, less commonly). The quantity is integer valued. The network exhibits a small number of large, cohesive communities driven by donor-count similarity.

(3) **HBA**, i.e., hydrogen bond acceptor. This is a molecular descriptor that counts the number of atoms in a molecule capable of accepting a hydrogen bond, typically atoms with lone pairs. Common HBA atoms/groups include: oxygen atoms (e.g., carbonyl O, ether O, hydroxyl O); nitrogen atoms (e.g., amines, imines, nitriles); sulfur atoms (e.g., thioethers, sulfones) that form a moderately fragmented mesoscale structure, consistent with acceptor-count diversity across the library.

(4) **MW**, i.e., molecular weight. MW represents the size and mass of a molecule, and it serves as a proxy for how

easily the compound can cross biological membranes. We can assume that $MW \leq 500\text{Da}$ is an empirical upper bound for compounds that tend to show acceptable oral absorption. Indeed, large values of MW are connected with larger polar surface area, increased steric bulk, greater flexibility (often more rotatable bonds, see below) that produce lower passive membrane permeability, poorer aqueous solubility, higher conformational entropy penalties upon binding. These factors collectively reduce the likelihood that a molecule can diffuse through lipid bilayers.

(5) **ROTB**, i.e., rotatable-bond counts. Introduced as refinement of the above four parameters [26], this quantity takes into account the conformational flexibility and allows one to distinguish between otherwise similar molecules. Indeed, two molecules with identical MW and logP can behave very differently if one is rigid and the other highly flexible.

The five descriptors described above are taken from the (generally accepted) Lipinski's rule of five [27], that is "A compound is more likely to be orally druglike if it does not violate more than one of the following criteria": (i) hydrogen bond donors (HBD) ≤ 5 ; (ii) hydrogen bond acceptors (HBA) ≤ 10 ; (iii) molecular weight (MW) $\leq 500\text{ Da}$; (iv) octanol-water partition coefficient (xlogP) ≤ 5 .

Each network retains $K = 5$ strongest neighbors per node. Despite identical node sets, the graphs differ in edge density and mesoscale organization, reflecting descriptor-specific organization principles. Edge counts were xLogP = 20, 466; HBD = 26, 961; HBA = 22, 896; MW = 15, 865; ROTB = 24, 760. These differences anticipate the observed variation in community structure and consensus stability across descriptors.

E. Modularity and Louvain method of computing

The modularity (Q) of a network, which measures the density of links within communities compared to random expectation, is defined as

$$Q = (1/2m) \sum_{i,j} [A_{ij} - (k_i k_j / 2m)] \delta(c_i, c_j), \quad (6)$$

where A_{ij} is the edge weight between nodes i and j , k_i and k_j are their degrees, and $\delta(c_i, c_j)$ equals 1 if both belong to the same community. Higher modularity values indicate stronger community structures.

The Louvain algorithm [22] (one of the most widely used methods for community detection in large networks, owing to its computational efficiency and its ability to uncover meaningful mesoscopic structure) is based on recursive modularity optimization. The algorithm proceeds iteratively in two main phases: first, each node is assigned to its own community, and nodes are then greedily reassigned to neighboring communities if this results in an increase in modularity; second, the identified communities are aggregated into supernodes, yielding a reduced network on which the process is repeated. This hierarchical aggregation continues until no further modularity improvement is possible. As a result, the Louvain algorithm naturally produces a multilevel, hierarchical partition of the network and scales efficiently to networks with millions of nodes, making it particularly suitable for the analysis of

large, complex systems where community structure plays a central role.

F. Null models and modularity validation

Degree-preserving randomized networks were generated via double-edge swaps [28]. For each null, modularity was re-computed, yielding a distribution $\{Q_r\}$. Statistical significance was estimated as

$$Z_Q = \frac{Q_{\text{obs}} - \mu_{Q_r}}{\sigma_{Q_r}}, \quad (7)$$

where Q_{obs} is the observed modularity, and μ_{Q_r} and σ_{Q_r} are the mean and standard deviation of modularity over the null ensemble [29]. Degree-preserving null models isolated the contribution of the empirical topology from that of the degree distribution, providing a rigorous baseline for modularity significance and following principles of statistical network comparison widely used in complex network theory [12,29]

G. Cross-view consensus via coclustering

A coclustering matrix was constructed across multiple network views. For each descriptor v and each molecular pair (i, j) , we computed

$$C_{ij} = \sum_{v=1}^V \mathbf{1}(c_i^{(v)} = c_j^{(v)}), \quad (8)$$

where $c_i^{(v)}$ is the community of compound i in view v , $V = 6$ is the number of descriptors, and $\mathbf{1}(\cdot)$ is the indicator function. This matrix quantified how consistently molecule pairs cooccurred in the same community across different similarity definitions [30–32].

H. Minimal spanning trees

The Minimum Spanning Tree (MST) captures the backbone of each network by connecting all nodes with minimal total distance. Indeed, given a connected graph with weighted edges, an MST is defined as the acyclic subgraph that connects all nodes while minimizing the total edge weight. By construction, the MST contains exactly $N - 1$ edges for a network of N nodes, preserving global connectivity while discarding redundant or weaker links. This property makes MSTs particularly valuable for analyzing complex systems in which the full network is dense, noisy, or difficult to interpret, as the MST highlights the most significant relationships encoded in the weight structure while removing unnecessary links [12,13]. The structure is obtained by ranking the edges according to their value of similarities and connecting nodes from the most similar couple of edges. Whenever an edge creation would form a loop, this move is canceled and we pass to consider the next couple of nodes. The procedure stops when all the nodes are connected. For backbone analysis, similarities were converted into distances,

$$d_{ij} = 1 - S_{ij}, \quad (9)$$

and to get the MSTs we used Kruskal's algorithm [33,34]. For the consensus MST, similarities were defined as

$$S_{ij}^{\text{cons}} = \frac{C_{ij}}{V}, \quad d_{ij}^{\text{cons}} = 1 - S_{ij}^{\text{cons}}, \quad (10)$$

with C_{ij} from Eq. (8) representing the number of times molecules i and j coclustered across the six descriptor-specific networks.

I. Network diagnostics and candidate prioritization

Network diagnostics included degree distributions, path lengths, diameters, edge-weight statistics, and leaf ratios, aligning with network-analysis frameworks in chemical-space mapping [35]. Cross-view agreement was measured using edge-set Jaccard overlap and degree correlations.

Candidate prioritization utilized betweenness centrality [36],

$$BC(v) = \sum_{s \neq v \neq t} \frac{\sigma_{st}(v)}{\sigma_{st}}, \quad (11)$$

where $\sigma_{st}(v)$ is the count of shortest paths between nodes s and t that pass through node v . Nodes with high betweenness typically occupy bridging positions between modules

and therefore act as mediators of structural or functional diversity [32,37].

J. Consensus network integration and robustness analysis

Similarity network fusion (SNF) [38] was applied to integrate networks across descriptors. SNF iteratively diffused local similarity across network layers, producing a consensus representation that preserved shared structure while integrating complementary views. Its adaptation to molecular similarity networks followed the principles of data integration in complex systems [32]. Robustness was evaluated under varying kNN values and repeated community detection [35].

K. Statistical reporting and reproducibility

All analyses were performed with fixed random seeds. Observed values were reported alongside null distributions with means, standard deviations, and empirical p values. Core tools included RDKit [39], SCIKIT-LEARN [40], and NETWORKX [41].

ACKNOWLEDGMENT

G.C. thanks the COGNIA visitor program at LIMS and the PRIN 2022 SNO-MINK.

DATA AVAILABILITY

The data that support the findings of this article are openly available [20], embargo periods may apply.

-
- [1] T. T. Ashburn and K. B. Thor, Drug repositioning: Identifying and developing new uses for existing drugs, *Nat. Rev. Drug Discov.* **3**, 673 (2004).
 - [2] J. Lamb, E. D. Crawford, D. Peck, J. W. Modell, I. C. Blat, M. J. Wrobel, J. Lerner, J.-P. Brunet, A. Subramanian, K. N. Ross, *et al.*, The connectivity map: Using gene-expression signatures to connect small molecules, genes, and disease, *Science* **313**, 1929 (2006).
 - [3] A.-L. Barabási, N. Gulbahce, and J. Loscalzo, Network medicine: A network-based approach to human disease, *Nat. Rev. Genet.* **12**, 56 (2011).
 - [4] Y. Zhou, Y. Hou, J. Shen, Y. Huang, W. Martin, and F. Cheng, Network-based drug repurposing for novel coronavirus 2019-nCoV/SARS-CoV-2, *Cell Discov.* **6**, 14 (2020).
 - [5] A. Baldi, Computational approaches for drug design and discovery: An overview, *Syst. Rev. Pharm.* **1**, 99 (2010).
 - [6] J. Vamathevan, D. Clark, P. Czodrowski, I. Dunham, E. Ferran, G. Lee, B. Li, A. Madabhushi, P. Shah, M. Spitzer, and S. Zhao, Applications of machine learning in drug discovery and development, *Nat. Rev. Drug Discov.* **18**, 463 (2019).
 - [7] H. Chen, O. Engkvist, Y. Wang, M. Olivecrona, and T. Blaschke, The rise of deep learning in drug discovery, *Drug Discov. Today* **23**, 1241 (2018).
 - [8] M. Zitnik, M. Agrawal, and J. Leskovec, Modeling polypharmacy side effects with graph convolutional networks, *Bioinformatics* **34**, i457 (2018).
 - [9] V. N. Ioannidis, D. Zheng, and G. Karypis, Few-shot link prediction via graph neural networks for Covid-19 drug-repurposing, *arXiv:2007.10261*.
 - [10] D. S. Himmelstein, A. Lizée, C. Hessler, L. Brueggeman, S. L. Chen, D. Hadley, A. Green, P. Khankhanian, and S. E. Baranzini, Systematic integration of biomedical knowledge prioritizes drugs for repurposing, *eLife* **6**, e26726 (2017).
 - [11] P. Perdomo-Quinteiro and A. Belmonte-Hernández, Knowledge graphs for drug repurposing: A review of databases and methods, *Brief. Bioinform.* **25**, bbae461 (2024).
 - [12] G. Caldarelli, *Scale-Free Networks: Complex Webs in Nature and Technology* (Oxford University Press, Oxford, 2007).
 - [13] A.-L. Barabási and M. Pósfai, *Network Science* (Cambridge University Press, Cambridge, 2016).
 - [14] D. Szklarczyk, A. Franceschini, S. Wyder, K. Forslund, D. Heller, J. Huerta-Cepas, M. Simonovic, A. Roth, A. Santos, K. P. Tsafou, *et al.*, STRING v10: Protein–protein interaction networks, integrated over the tree of life, *Nucleic Acids Res.* **43**, D447 (2015).
 - [15] E. Guney, J. Menche, M. Vidal, and A.-L. Barabási, Network-based in silico drug efficacy screening, *Nat. Commun.* **7**, 10331 (2016).
 - [16] A. L. Hopkins, Network pharmacology: The next paradigm in drug discovery, *Nat. Chem. Biol.* **4**, 682 (2008).
 - [17] M. De Domenico, A. Solé-Ribalta, E. Cozzo, M. Kivelä, Y. Moreno, M. A. Porter, S. Gómez, and A. Arenas, Mathematical

- formulation of multilayer networks, *Phys. Rev. X* **3**, 041022 (2013).
- [18] S. Battiston, G. Caldarelli, and A. Garas, *Multiplex and Multi-level Networks* (Oxford University Press, Oxford, 2018).
- [19] O. Sporns, Structure and function of complex brain networks, *Dialogues Clin. Neurosci.* **15**, 247 (2013).
- [20] <https://zinc.docking.org>.
- [21] N. Tabibzadeh, D. Fleury, D. Labatut, F. Bridoux, A. Lionet, N. Jourde-Chiche, F. Vrtovsniak, N. Schlegel, and P. Vanhille, MYH9-related disorders display heterogeneous kidney involvement and outcome, *Clin. Kidney J.* **12**, 494 (2019).
- [22] V. D. Blondel, J.-L. Guillaume, R. Lambiotte, and E. Lefebvre, Fast unfolding of communities in large networks, *J. Stat. Mech.: Theory Exp.* (2008) P10008.
- [23] P. Jaccard, The distribution of the flora in the Alpine zone, *New Phytol.* **11**, 37 (1912).
- [24] T. T. Tanimoto, *An Elementary Mathematical Theory of Classification and Prediction* (International Business Machines Corporation, New York, NY, 1958).
- [25] D. Rogers and M. Hahn, Extended-connectivity fingerprints, *J. Chem. Inf. Model.* **50**, 742 (2010).
- [26] D. F. Veber, S. R. Johnson, H.-Y. Cheng, B. R. Smith, K. W. Ward, and K. D. Kopple, Molecular properties that influence the oral bioavailability of drug candidates, *J. Med. Chem.* **45**, 2615 (2002).
- [27] C. A. Lipinski, F. Lombardo, B. W. Dominy, and P. J. Feeney, Experimental and computational approaches to estimate solubility and permeability in drug discovery and development settings, *Adv. Drug Deliv. Rev.* **64**, 4 (2012).
- [28] S. Maslov and K. Sneppen, Specificity and stability in topology of protein networks, *Science* **296**, 910 (2002).
- [29] G. Cimini, T. Squartini, F. Saracco, D. Garlaschelli, A. Gabrielli, and G. Caldarelli, The statistical physics of real-world networks, *Nat. Rev. Phys.* **1**, 58 (2019).
- [30] S. Monti, P. Tamayo, J. Mesirov, and T. Golub, Consensus clustering: A resampling-based method for class discovery and visualization of gene expression microarray data, *Mach. Learn.* **52**, 91 (2003).
- [31] A. Lancichinetti and S. Fortunato, Consensus clustering in complex networks, *Sci. Rep.* **2**, 336 (2012).
- [32] G. Caldarelli and M. Catanzaro, *Networks: A Very Short Introduction* (Oxford University Press, Oxford, 2012).
- [33] J. B. Kruskal, On the shortest spanning subtree of a graph and the traveling salesman problem, *Proc. Am. Math. Soc.* **7**, 48 (1956).
- [34] R. L. Graham and P. Hell, On the history of the minimum spanning tree problem, *IEEE Ann. Hist. Comput.* **7**, 43 (1985).
- [35] V. F. Scalfani, V. D. Patel, and A. M. Fernandez, Visualizing chemical space networks with RDKit and NetworkX, *J. Cheminf.* **14**, 87 (2022).
- [36] L. C. Freeman, A set of measures of centrality based on betweenness, *Sociometry* **40**, 35 (1977).
- [37] A.-L. Barabási, Scale-free networks: A decade and beyond, *Science* **325**, 412 (2009).
- [38] B. Wang, A. M. Mezlini, F. Demir, M. Fiume, Z. Tu, M. Brudno, B. Haibe-Kains, and A. Goldenberg, Similarity network fusion for aggregating data types on a genomic scale, *Nat. Methods* **11**, 333 (2014).
- [39] G. Landrum, RDKit: Open-source cheminformatics software, 2016, <https://www.rdkit.org>.
- [40] F. Pedregosa *et al.*, Scikit-learn: Machine learning in Python, *J. Mach. Learn. Res.* **12**, 2825 (2011).
- [41] A. A. Hagberg, D. A. Schult, and P. J. Swart, Exploring network structure, dynamics, and function using NetworkX, in *Proceedings of the 7th Python in Science Conference*, edited by G. Varoquaux, T. Vaught, and J. Millman (SciPy, 2008), pp. 11–15.