
MoAgent: A Hypothesis-Driven Multi-Agent Framework for Drug Mechanism of Action Discovery

Jun Hyeon Kim¹, Seokhyun Moon¹, Seonghwan Kim¹,
Junhyeok Jeon², Taein Kim¹, Jisu Seo¹, Songmi Kim¹, Woo Youn Kim^{1,2*}
¹KAIST ²HITS

Abstract

Determining the mechanism of action (MoA) of novel chemical compounds is a critical yet challenging task in drug discovery. We introduce MoAgent, a multi-agent framework that reframes MoA inference as a hypothesis-driven scientific discovery process. MoAgent integrates multi-modal data from chemical structure, gene expression, and biological pathways, deploying a committee of specialized agents to collaboratively generate and validate mechanistic hypotheses. The framework operates through an iterative cycle of evidence triangulation and hypothesis validation, where a bioinformatician agent assesses causal plausibility using a knowledge graph and a medicinal chemist agent verifies direct target engagement. Our experiments demonstrate that the integrated, hypothesis-driven strategy significantly enhances the accuracy and reliability of MoA inference, and maintains robust performance even in zero-shot scenarios. By emulating scientific reasoning, MoAgent offers a more effective paradigm for accelerating drug discovery.

1 Introduction

Understanding the mechanism of action (MoA) of drugs is a fundamental challenge in pharmacology and drug discovery. Knowledge of MoA helps distinguish intended multi-target effects from undesired off-target interactions [Davis, 2020, Watkins et al., 2010, De Savi et al., 2020], thereby balancing efficacy and safety in the early stages of development [Berger and Iyengar, 2011]. It also enables the identification of synthetic lethal relationships, which can guide the design of cell-specific therapeutic strategies, and supports drug repurposing or biomarker discovery for efficient candidate selection [Zhang et al., 2021, Previtali et al., 2024]. Though their elucidation is crucial, it remains challenging task for existing methods, as detailed in appendix A.

To address these challenges, we introduce MoAgent, an autonomous multi-agent framework that emulates the scientific process of MoA inference by simulating a collaborative scientific committee (see Figure 1). Building on the recognition that chemical similarity, transcriptomic signatures, and biological knowledge each provide only partial clues, MoAgent unifies them into a single iterative reasoning process. This process is centered on evidence triangulation, a principle wherein multiple, independent lines of evidence are systematically sought to corroborate a finding, thereby increasing the validity and reliability of the conclusion. A large language model (LLM) orchestrates specialized LLM agents to synthesize evidence, formulate hypotheses, and validate them through both mechanistic reasoning over biological resources and physical evidence from external databases. Through this orchestration, and by leveraging the prior biological knowledge embedded in LLMs, MoAgent transforms fragmented signals into coherent and testable mechanistic insights, offering a flexible and explainable framework for MoA discovery.

*Correspondence to: Woo Youn Kim <wooyoun@kaist.ac.kr>

2 Methods

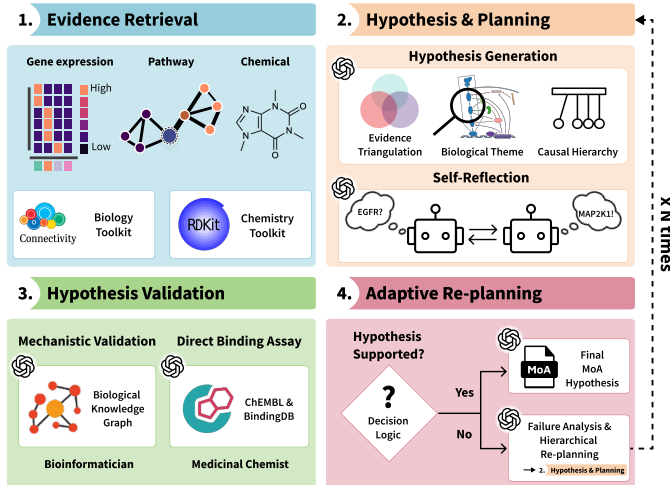


Figure 1: **The MoAgent framework for hypothesis-driven MoA discovery.** The framework emulates the scientific discovery process by first synthesizing multi-modal evidence to generate hypotheses. These hypotheses are then validated by a committee of specialist agents. A central controller evaluates these outcomes, generating a final MoA report upon successful validation or initiating an adaptive re-planning loop in case of refutation. This iterative, self-correcting architecture enables MoAgent to perform robust, evidence-based mechanistic inference.

2.1 Problem definition

The central challenge of MoA inference is to solve the inverse of this biological process: to infer the initial causal event from its downstream consequences. Formally, the problem is defined as follows: given a query compound’s molecular structure and its induced gene expression signature as inputs, the goal is to predict the set of protein targets that constitute its MoA. To simulate a realistic discovery scenario with no prior knowledge of a given compound, we anonymized it by withholding its common name and structure from the LLM in our MoAgent framework.

2.2 MoAgent: A multi-agent framework for drug MoA discovery

Implemented as a state graph using the LangGraph library, MoAgent employs a LLM to manipulate a collaborative, iterative workflow among specialized agents. As depicted in Figure 1, the modular architecture of MoAgent, which composed of four distinct modules corresponding to key stages of scientific inquiry, enables a flexible and robust reasoning process where subsequent actions are determined by the outcome of the previous stage. Details about each module and the overall workflow of MoAgent are introduced at Appendix B.1.

3 Results

To validate effectiveness of MoAgent, we designed a benchmark to test its ability to infer MoAs against established and generalist baselines. The evaluation focuses on the accuracy and biological completeness of MoA inferences for a set of well-characterized drugs. For detailed explanations about the dataset, baselines, and evaluation metrics, refer to Appendices C.1 to C.3.

3.1 Comparative performance on MoA inference

The performance of MoAgent compared to the baselines is summarized in Table 1. To first isolate and evaluate the underlying reasoning architecture of each framework, this initial experiment deliberately limits the input to a single data source: the compound’s gene expression signature. This establishes a

fair comparison, allowing for a direct assessment of how effectively each method infers MoA from transcriptomic evidence alone.

Table 1: MoA inference performance. MoAgent is compared against a signature similarity baseline (top-k union, k=1,3,5) and a general-purpose agent baseline (ReAct). Performance is evaluated using Precision, Recall, and F1-score (F1) across two aggregation strategies: Individual (*per-signature*), Union (*per-drug*). The best and second-best performances are indicated in bold and underlined, respectively.

Method	Individual			Union		
	Precision ($\uparrow$)	Recall ($\uparrow$)	F1 ($\uparrow$)	Precision ($\uparrow$)	Recall ($\uparrow$)	F1 ($\uparrow$)
Signature Sim. (k=1)	<u>0.0962</u>	0.1154	0.1026	0.0379	0.2222	0.0630
Signature Sim. (k=3)	0.0915	0.1667	<u>0.1067</u>	0.0412	0.2593	0.0613
Signature Sim. (k=5)	0.0599	<u>0.1806</u>	0.0818	0.0392	<u>0.2994</u>	0.0628
ReAct	0.0942	0.1346	0.0923	<u>0.1139</u>	0.2778	<u>0.1437</u>
MoAgent	0.4519	0.4385	0.4156	0.5278	0.4889	0.4598

MoAgent achieves a remarkable F1-score in the individual evaluation, an over four-fold improvement over the ReAct. This significant performance gap highlights the fundamental limitations of both conventional and generalist agent approaches. The signature similarity method, relying on a single data modality, is insufficient for this complex task. Moreover, the performance of the ReAct agent reveals that access to a powerful LLM and relevant tools is not sufficient to navigate the complexity of this scientific reasoning task. In contrast, MoAgent’s success is directly attributable to its architecture: its structured *hypothesis-driven workflow* and rigorous *verification process* provide the superior reasoning that the generalist agent lacks.

The traditional signature similarity method exhibits a classic trade-off. While combining results from replicates increases recall by capturing more potential true positives, it introduces a high number of false positives from noisy predictions. This leads to a drastic drop in precision that outweighs the gain in recall, ultimately resulting in a lower F1-score. In contrast, both agent-based approaches show a more favorable trend. While the ReAct agent also improves upon aggregation, the critical distinction lies in the absolute performance. MoAgent starts from a much higher baseline in the individual evaluation and further improves its scores, highlighting the effectiveness of considering biological context. By formulating and rigorously verifying hypotheses for each replicate, MoAgent generates high-confidence predictions that effectively filter out experimental noise. Consequently, the Union aggregation primarily serves to collect these high-quality predictions, capturing a more complete set of true MoAs. This empirically validates that our framework successfully emulates the scientific process of building a robust consensus from multiple pieces of evidence.

3.2 Reasoning beyond retrieval: MoA inference in a zero-shot setting

Having established the superiority of its architecture with a single data source, we next evaluated MoAgent in a more realistic, multi-modal setting. In this second experiment, we incorporated chemical structure as a key, independent line of evidence to assess MoAgent’s ability to synergize diverse information—a core aspect of its design. We also designed an experiment to contrast the performance of the model in two distinct scenarios: the first scenario (‘w/ Protein Target’), where the true protein target of the drug is present in the reference database, and a more challenging zero-shot scenario (‘w/o Protein Target’), where the true target is deliberately excluded from it. The latter simulates a true discovery task, compelling the agent to infer the MoA from a convergence of evidence without any direct target information to retrieve. This setup is crucial for assessing whether the agent can perform scientific discovery rather than simple fact-checking. To implement the zero-shot ‘w/o Protein Target’ scenario, we removed known interactions between a query drug and its ground-truth protein targets from the knowledge sources accessible to the agents (e.g., removing ‘EGFR’ for Gefitinib). For details, refer to section C.5.

As shown in Table 2, the results from the ‘w/ Protein Target’ scenario highlight that simpler, heuristic-based approaches can perform effectively when a direct, unambiguous signal is available. This is particularly evident with the structure similarity baseline; for the well-known drugs in our benchmark, structural information provides such a powerful signal that the task becomes more akin to information retrieving rather than complex reasoning. Consequently, all methods, including ReAct, achieve high F1-scores in this setting compared to Table 1.

Table 2: Comparison of MoA inference performance with and without the ground-truth protein target. MoAgent is compared against baseline methods under each condition. The best and second-best performances are indicated in bold and underlined, respectively.

Method	w/ Protein Target						w/o Protein Target					
	Individual			Union			Individual			Union		
	P($\uparrow$)	R($\uparrow$)	F1($\uparrow$)	P($\uparrow$)	R($\uparrow$)	F1($\uparrow$)	P($\uparrow$)	R($\uparrow$)	F1($\uparrow$)	P($\uparrow$)	R($\uparrow$)	F1($\uparrow$)
Structure Sim. (k=1)	0.673	0.420	0.403	0.648	0.404	0.388	0.000	0.000	0.000	0.000	0.000	0.000
Structure Sim. (k=3)	0.430	0.471	0.375	0.414	0.454	0.361	0.000	0.000	0.000	0.000	0.000	0.000
Structure Sim. (k=5)	0.309	<u>0.555</u>	0.357	0.298	0.534	0.344	0.000	0.000	0.000	0.000	0.000	0.000
ReAct	0.657	0.510	<u>0.517</u>	0.537	<u>0.555</u>	<u>0.480</u>	0.000	0.000	0.000	0.000	0.000	0.000
MoAgent	<u>0.663</u>	0.566	0.550	<u>0.622</u>	0.638	0.583	0.528	0.500	0.482	0.554	0.559	0.537

However, the evaluation for scientific inference comes from the ‘w/o Protein Target’ setting. Here, the performance of all baseline methods collapses, with F1-scores dropping to zero. The failure of the ReAct agent is particularly revealing: despite having access to the same tools as MoAgent, its unstructured process, driven by a single general-purpose agent, gets lost in the complex data landscape once the dominant signal is removed. In contrast, the performance of MoAgent remains remarkably robust, achieving a Union F1-score of 0.537. This success validates its core design for genuine scientific discovery. The strength of MoAgent lies in its *evidence triangulation*, a process powered by its specialized agents. Each agent examines the problem from a distinct biochemical perspective, allowing MoAgent to construct a unified biological context. This unified approach is particularly effective at amplifying weak signals that would be missed by relying on a single line of evidence. Furthermore, this framework is inherently robust; if an initial reasoning path fails to yield a coherent hypothesis, its adaptive re-planning mechanism allows it to reassess the evidence and forge a new, more promising strategy.

Strikingly, F1-score of MoAgent was higher in this challenging zero-shot setting than in the standard experiment where direct target information was available (Table 1). This counterintuitive result validates that *multi-modal synergy* is fundamental to the design of MoAgent. While there is no direct information about the ground-truth target, the integration of multi-modal evidences provides complementary but indirect signals. MoAgent is able to triangulate these weaker, distributed signals into a coherent mechanistic hypothesis. Consequently, it can sometimes perform better in scenarios where only indirect but diverse evidence is available. This result highlights that core strength of MoAgent lies not in simple retrieval of clear answers, but in synergizing heterogeneous evidence to uncover robust mechanistic insights.

4 Conclusion

In this work, we introduced MoAgent, a multi-agent framework that reframes mechanism of action (MoA) inference from a process reliant on correlation-based similarity to a hypothesis-driven process of scientific discovery. MoAgent integrates multi-modal data—spanning chemical structure, gene expression, and biological pathways—through a committee of specialized agents. It overcomes the distinct limitations of single-modality analysis through an iterative cycle of evidence triangulation and computational validation.

The integrated approach of MoAgent significantly improved prediction performance compared to conventional similarity-based methods and generalist single-agent systems. Crucially, its true strength was revealed in a multi-source, zero-shot scenario; while all baseline methods which collapsed completely, MoAgent successfully synergized distinct clues to maintain robust performance. This transition from a single-source to a multi-source context empirically validates that MoAgent excels at genuine scientific reasoning, constructing a coherent mechanistic narrative rather than simply retrieving facts from a dominant signal.

While evidence triangulation of the MoAgent is designed to be robust against gaps in any single knowledge source, its overall performance is naturally bounded by the collective completeness of its external resources. Promising directions for future work include extending the model by incorporating additional data modalities. Furthermore, to enhance generalizability of the model and practical applicability, future validation should encompass a broader range of experimental conditions, including diverse cell lines, dosages, and time points.

Acknowledgments and Disclosure of Funding

This work was supported by the Basic Science Research Programs through the National Research Foundation of Korea funded by the Ministry of Science and ICT (Grant No. RS-2023-00257479 and Grant No. NRF-2023R1A2C2004376).

References

- Ronald L Davis. Mechanism of action and target identification: a matter of timing in drug discovery. *Isience*, 23(9), 2020.
- John Watkins, Andrew Marsh, Paul C Taylor, and Donald RJ Singer. Personalized medicine: the impact on chemistry. *Therapeutic Delivery*, 1(5):651–665, 2010.
- Chris De Savi, David L Hughes, and Lisbet Kvaerno. Quest for a covid-19 cure by repurposing small-molecule drugs: mechanism of action, clinical development, synthesis at scale, and outlook for supply. *Organic Process Research & Development*, 24(6):940–976, 2020.
- Seth I Berger and Ravi Iyengar. Role of systems pharmacology in understanding drug adverse events. *Wiley interdisciplinary reviews: systems biology and medicine*, 3(2):129–135, 2011.
- Biyu Zhang, Chen Tang, Yanli Yao, Xiaohan Chen, Chi Zhou, Zhiting Wei, Feiyang Xing, Lan Chen, Xiang Cai, Zhiyuan Zhang, et al. The tumor therapy landscape of synthetic lethality. *Nature communications*, 12(1): 1275, 2021.
- Viola Previtali, Greta Bagnolini, Andrea Ciamarone, Giovanni Ferrandi, Francesco Rinaldi, Samuel Harry Myers, Marinella Roberti, and Andrea Cavalli. New horizons of synthetic lethality in cancer: current development and future perspectives. *Journal of Medicinal Chemistry*, 67(14):11488–11521, 2024.
- Aliyu Musa, Laleh Soltan Ghorai, Shu-Dong Zhang, Galina Glazko, Olli Yli-Harja, Matthias Dehmer, Benjamin Haibe-Kains, and Frank Emmert-Streib. A review of connectivity map and computational approaches in pharmacogenomics. *Briefings in bioinformatics*, 19(3):506–523, 2018.
- Bence Szalai and Dániel V Veres. Application of perturbation gene expression profiles in drug discovery—from mechanism of action to quantitative modelling. *Frontiers in Systems Biology*, 3:1126044, 2023.
- Maria-Anna Trapotsi, Layla Hosseini-Gerami, and Andreas Bender. Computational analyses of mechanism of action (moa): data, methods and integration. *RSC chemical biology*, 3(2):170–200, 2022.
- Justin Lamb, Emily D Crawford, David Peck, Joshua W Modell, Irene C Blat, Matthew J Wrobel, Jim Lerner, Jean-Philippe Brunet, Aravind Subramanian, Kenneth N Ross, et al. The connectivity map: using gene-expression signatures to connect small molecules, genes, and disease. *science*, 313(5795):1929–1935, 2006.
- Aravind Subramanian, Rajiv Narayan, Steven M Corsello, David D Peck, Ted E Natoli, Xiaodong Lu, Joshua Gould, John F Davis, Andrew A Tubelli, Jacob K Asiedu, et al. A next generation connectivity map: L1000 platform and the first 1,000,000 profiles. *Cell*, 171(6):1437–1452, 2017.
- Mark R Segal, Hao Xiong, Henrik Bengtsson, Richard Bourgon, and Robert Gentleman. Querying genomic databases: refining the connectivity map. *Statistical Applications in Genetics and Molecular Biology*, 11(2), 2012.
- Savina Jaeger, Junxia Min, Florian Nigsch, Miguel Camargo, Janna Hutz, Allen Cornett, Stephen Cleaver, Alan Buckler, and Jeremy L Jenkins. Causal network models for predicting compound targets and driving pathways in cancer. *Journal of biomolecular screening*, 19(5):791–802, 2014.
- Glyn Bradley and Steven J Barrett. Causalr: extracting mechanistic sense from genome scale data. *Bioinformatics*, 33(22):3670–3672, 2017.
- Anika Liu, Panuwat Trairatphisan, Enio Gjerga, Athanasios Didangelos, Jonathan Barratt, and Julio Saez-Rodriguez. From expression footprints to causal pathways: contextualizing large signaling networks with carnival. *NPJ systems biology and applications*, 5(1):40, 2019.
- Belinda B Garana, James H Joly, Alireza Delfarah, Hyunjun Hong, and Nicholas A Graham. Drug mechanism enrichment analysis improves prioritization of therapeutics for repurposing. *BMC bioinformatics*, 24(1):215, 2023.

- Shengqiao Gao, Lu Han, Dan Luo, Gang Liu, Zhiyong Xiao, Guangcun Shan, Yongxiang Zhang, and Wenxia Zhou. Modeling drug mechanism of action with large scale gene-expression profiles using gpar, an artificial intelligence platform. *BMC bioinformatics*, 22(1):17, 2021.
- Wenhui Wang, Sen Yang, and JING Li. Drug target predictions based on heterogeneous graph inference. In *Pacific symposium on biocomputing*, volume 18, pages 53–64. World Scientific, 2013.
- Shobeir Fakhraei, Bert Huang, Louiqa Raschid, and Lise Getoor. Network-based drug-target interaction prediction with probabilistic soft logic. *IEEE/ACM Transactions on Computational Biology and Bioinformatics*, 11(5):775–787, 2014.
- Neel S Madhukar, Prashant K Khade, Linda Huang, Kaitlyn Gayvert, Giuseppe Galletti, Martin Stogniew, Joshua E Allen, Paraskevi Giannakakou, and Olivier Elemento. A bayesian machine learning approach for drug target identification using diverse data types. *Nature communications*, 10(1):5221, 2019.
- Layla Hosseini-Gerami, Ixavier Alonzo Higgins, David A Collier, Emma Laing, David Evans, Howard Broughton, and Andreas Bender. Benchmarking causal reasoning algorithms for gene expression-based compound mechanism of action analysis. *BMC bioinformatics*, 24(1):154, 2023.
- Sophia Krix, Lauren Nicole DeLong, Sumit Madan, Daniel Domingo-Fernández, Ashar Ahmad, Sheraz Gul, Andrea Zaliani, and Holger Fröhlich. Multigml: Multimodal graph machine learning for prediction of adverse drug events. *Heliyon*, 9(9), 2023.
- Vanessa Lopez, Lam Hoang, Marcos Martinez-Galindo, Raúl Fernández-Díaz, Marco Luca Sbodio, Rodrigo Ordonez-Hurtado, Mykhaylo Zayats, Natasha Mulligan, and Joao Bettencourt-Silva. Enhancing foundation models for scientific discovery via multimodal knowledge graph representations. *Journal of Web Semantics*, 84:100845, 2025.
- Taffee T Tanimoto. Ibm internal report. *Nov*, 17:1957, 1957.
- Prisca Lo Surdo, Marta Iannuccelli, Silvia Contino, Luisa Castagnoli, Luana Licata, Gianni Cesareni, and Livia Perfetto. Signor 3.0, the signaling network open resource 3.0: 2022 update. *Nucleic acids research*, 51(D1):D631–D637, 2023.
- Barbara Zdrazil, Eloy Felix, Fiona Hunter, Emma J Manners, James Blackshaw, Sybilla Corbett, Marleen de Veij, Harris Ioannidis, David Mendez Lopez, Juan F Mosquera, Maria Paula Magarinos, Nicolas Bosc, Ricardo Arcila, Tevfik Kizilören, Anna Gaulton, A Patrícia Bento, Melissa F Adasme, Peter Monecke, Gregory A Landrum, and Andrew R Leach. The chembl database in 2023: a drug discovery platform spanning multiple bioactivity data types and time periods. *Nucleic Acids Research*, 52(D1):D1180–D1192, November 2023. ISSN 1362-4962. doi: 10.1093/nar/gkad1004. URL <http://dx.doi.org/10.1093/nar/gkad1004>.
- Tiqing Liu, Linda Hwang, Stephen K Burley, Carmen I Nitsche, Christopher Southan, W Patrick Walters, and Michael K Gilson. Bindingdb in 2024: a fair knowledgebase of protein-small molecule binding data. *Nucleic Acids Research*, 53(D1):D1633–D1644, November 2024. ISSN 1362-4962. doi: 10.1093/nar/gkae1075. URL <http://dx.doi.org/10.1093/nar/gkae1075>.
- Seokhyun Moon, Sang-Yeon Hwang, Jaechang Lim, and Woo Youn Kim. Pignet2: a versatile deep learning-based protein–ligand interaction prediction model for binding affinity scoring and virtual screening. *Digital Discovery*, 3(2):287–299, 2024. ISSN 2635-098X. doi: 10.1039/d3dd00149k. URL <http://dx.doi.org/10.1039/d3dd00149k>.
- Saro Passaro, Gabriele Corso, Jeremy Wohlwend, Mateo Reveiz, Stephan Thaler, Vignesh Ram Somnath, Noah Getz, Tally Portnoi, Julien Roy, Hannes Stark, David Kwabi-Addo, Dominique Beaini, Tommi Jaakkola, and Regina Barzilay. Boltz-2: Towards accurate and efficient binding affinity prediction. *bioRxiv*, 2025. doi: 10.1101/2025.06.14.659707.
- Pau Badia-i Mompel, Jesús Vélez Santiago, Jana Braunger, Celina Geiss, Daniel Dimitrov, Sophia Müller-Dott, Petr Taus, Aurelien Dugourd, Christian H Holland, Ricardo O Ramirez Flores, and Julio Saez-Rodriguez. decoupler: ensemble of computational methods to infer biological activities from omics data. *Bioinformatics Advances*, 2(1), January 2022. ISSN 2635-0041. doi: 10.1093/bioadv/vbac016. URL <http://dx.doi.org/10.1093/bioadv/vbac016>.
- Steven M Corsello, Joshua A Bittker, Zihan Liu, Joshua Gould, Patrick McCarren, Jodi E Hirschman, Stephen E Johnston, Anita Vrcic, Bang Wong, Mariya Khan, Jacob Asiedu, Rajiv Narayan, Christopher C Mader, Aravind Subramanian, and Todd R Golub. The drug repurposing hub: a next-generation drug library and information resource. *Nature Medicine*, 23(4):405–408, April 2017. ISSN 1546-170X. doi: 10.1038/nm.4306. URL <http://dx.doi.org/10.1038/nm.4306>.

- Craig Knox, Mike Wilson, Christen M Klinger, Mark Franklin, Eponine Oler, Alex Wilson, Allison Pon, Jordan Cox, Na Eun (Lucy) Chin, Seth A Strawbridge, Marysol Garcia-Patino, Ray Kruger, Aadhavya Sivakumaran, Selena Sanford, Rahil Doshi, Nitya Khetarpal, Omolola Fatokun, Daphnee Doucet, Ashley Zubkowski, Dorsa Yahya Rayat, Hayley Jackson, Karxena Harford, Afia Anjum, Mahi Zakir, Fei Wang, Siyang Tian, Brian Lee, Jaanus Liigand, Harrison Peters, Ruo Qi (Rachel) Wang, Tue Nguyen, Denise So, Matthew Sharp, Rodolfo da Silva, Cyrella Gabriel, Joshua Scantlebury, Marissa Jasinski, David Ackerman, Timothy Jewison, Tanvir Sajed, Vasuk Gautam, and David S Wishart. Drugbank 6.0: the drugbank knowledgebase for 2024. *Nucleic Acids Research*, 52(D1):D1265–D1275, November 2023. ISSN 1362-4962. doi: 10.1093/nar/gkad976. URL <http://dx.doi.org/10.1093/nar/gkad976>.
- Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. React: Synergizing reasoning and acting in language models, 2022. URL <https://arxiv.org/abs/2210.03629>.
- Sophia Müller-Dott, Eirini Tsirvouli, Miguel Vazquez, Ricardo O Ramirez Flores, Pau Badia-i Mompel, Robin Fallegger, Dénes Türei, Astrid Lægreid, and Julio Saez-Rodriguez. Expanding the coverage of regulons from high-confidence prior knowledge for accurate estimation of transcription factor activities. *Nucleic acids research*, 51(20):10934–10949, 2023.
- Michael Schubert, Bertram Klinger, Martina Klünemann, Anja Sieber, Florian Uhlig, Sascha Sauer, Mathew J Garnett, Nils Blüthgen, and Julio Saez-Rodriguez. Perturbation-response genes reveal signaling footprints in cancer gene expression. *Nature communications*, 9(1):20, 2018.

A Related Works

A dominant paradigm in MoA inference is rooted in biochemical similarity [Musa et al., 2018, Szalai and Veres, 2023, Trapotsi et al., 2022]. On the chemical side, compounds with similar molecular structures are often assumed to engage similar targets and thereby elicit comparable MoA. On the biological side, transcriptomic profiling resources such as the Connectivity Map (CMap) leverage gene expression similarity, under the assumption that drugs producing comparable gene expression signatures act through related pathways [Lamb et al., 2006, Subramanian et al., 2017, Segal et al., 2012]. A second, complementary paradigm relies on causal reasoning over biological knowledge. Approaches based on knowledge graphs (KGs) integrate curated interactions and pathways to trace upstream regulators [Jaeger et al., 2014, Bradley and Barrett, 2017, Liu et al., 2019], while pathway enrichment analyses highlight coordinated pathway activities [Garana et al., 2023, Gao et al., 2021].

These parallel paradigms highlight a critical challenge in MoA inference. Similarity-based approaches offer scalability but often yield superficial correlations [Wang et al., 2013, Fakhraei et al., 2014, Madhukar et al., 2019]. Conversely, methods centered on causal reasoning using knowledge graphs provide mechanistic depth, but their reliability is vulnerable to the inherent gaps and biases of any single, curated knowledge source. An inference process that depends solely on such a resource can be prone to failure when a causal path is not yet documented [Hosseini-Gerami et al., 2023]. Moreover, most existing approaches rely on a single modality—either molecular structure, gene expression, or biological networks—while MoA inference is inherently a multi-modal reasoning task that requires integrating chemical, biological, and knowledge-based evidence [Trapotsi et al., 2022, Krix et al., 2023, Lopez et al., 2025].

B Method details

This appendix provides supplementary details regarding the implementation of the MoAgent framework, intended to ensure the transparency and reproducibility of our work.

B.1 Detail modules of MoAgent

1. Evidence retrieval module This module gathers initial, multi-modal evidence to form a foundational understanding of the compound’s effects. It collects three distinct types of evidence from gene expression data:

- **Chemical evidence:** Based on the assumption that structurally similar compounds often interact with similar biological targets, this provides a strong, albeit heuristic, starting point. Specifically, we utilized the Tanimoto similarity [Tanimoto, 1957] to retrieve analogous compound from databases. The MoAs of known analogs offer valuable clues to the potential target class of the query compound (see the first example in Table 3).
- **Transcriptomic evidence:** This provides a functional, systems-level view of the compound’s impact. By comparing gene expression signatures, it groups compounds based on their ultimate phenotypic outcome, offering a complementary perspective to chemical structure. Specifically, we retrieved similar compounds from databases based on the cosine similarity of their signatures (see the second example in Table 3).
- **Mechanistic evidence:** This contextualizes the broad transcriptomic changes within known biological pathways. By inferring the activity of specific signaling pathways and transcription factors, it translates the high-dimensional gene expression data into more interpretable, mechanistic insights (see the third and fourth examples Table 3).

The evidence collected from these three modalities is then synthesized to identify convergent biological themes (e.g. EGFR signaling disruption). This consolidated output, which translates raw data into actionable insights, is then delivered to the Hypothesis planning module. For more details, refer to section C.4.

2. Hypothesis planning module This module performs an abductive reasoning task to bridge the gap from broad biological themes to specific, testable hypotheses, which are potential protein targets with respective rationales (e.g., Table 11). Its central challenge is to navigate a key complexity of

cellular signaling: distinct upstream perturbations, such as inhibiting different proteins within the same signaling pathway, can converge to produce similar downstream gene expression signatures. To disambiguate these scenarios and infer the most plausible causal origins, the reasoning of the module is guided by two core strategies, which are implemented in a structured, two-stage process involving distinct planning and refinement agents.

- **Stage 1: initial hypothesis generation by a planning agent:** The first agent applies the *evidence convergence strategy* to identify up to N consensus themes from the evidence. And then, for each theme, it compiles a comprehensive list of candidate targets by considering proteins (a) explicitly mentioned in the initial evidence, (b) known to be central to the theme, and (c) suggested as upstream regulators by the *causal hierarchy strategy*, which consider biological signaling cascade. Based on this expanded list, the agent formulates initial hypotheses.
- **Stage 2: hypothesis refinement by a reflector agent:** Each initial hypothesis is then passed to a second agent for refinement. This agent reviews the initial hypothesis and proposes additional, closely related target proteins, based on the biological themes. This step expands the hypotheses, ensuring a comprehensive search.

The final output provides a structured action plan containing prioritized target list of hypotheses ready for validation (see Table 11).

3. Hypothesis validation module This module rigorously tests the proposed hypotheses by requiring a consensus from two specialist agents, each assessing an *independent and orthogonal line of evidence*: causal plausibility and direct physical interaction.

- **The bioinformatician agent** investigates causal plausibility. Its primary tool is querying a biological knowledge graph (e.g., Signor [Lo Surdo et al., 2023]) to find established signaling pathways between a candidate target and the inferred active transcription factors. A short, known path provides strong evidence for mechanistic coherence.
- **The medicinal chemist agent** seeks direct evidence of molecular interaction. It is equipped with tools to query public bioactivity databases (e.g., ChEMBL [Zdrazil et al., 2023], BindingDB [Liu et al., 2024]) for experimental data (e.g., IC_{50} , K_i) confirming a high-affinity interaction between the compound and the target. When experimental data is unavailable, we discard the hypothesis. In the future, this can be addressed by employing computational methods such as molecular docking or AI-driven models (e.g., PIGNet [Moon et al., 2024], Boltz [Passaro et al., 2025]) to predict binding affinity.

4. Controller and iterative refinement loop The central LLM controller acts as the “principal investigator”. It receives analytical reports from the bioinformatician and medicinal chemist agents to determine the outcome. A hypothesis is considered successfully validated only if both agents independently classify it as *Supported*, confirming both causal plausibility and direct physical interaction. Upon finding the validated hypothesis, the framework initiates a *target expansion* phase to explore a broader mechanistic space around validated target, by leveraging the embedded biological knowledge of LLM, such as protein family relationships and drug polypharmacology (see Table 14 for example). These expanded targets form a new set of hypotheses, which are then fed back into the hypothesis validation module. This expansion-validation cycle repeats, allowing MoAgent to systematically map out potential multi-target effects. The number of expansion iterations is a configurable hyperparameter, which is set to one for this study.

If all initial hypotheses are refuted, controller initiates the adaptive replanning loop. In this *failure analysis* phase, it reviews the history of failed validation to discard invalidated assumptions and generates new hypotheses from alternative perspectives, returning to the planning module. To prevent infinite loops, this overall planning-validation cycle is constrained by a maximum number of rounds, which is set to three for this study. At the conclusion of the process, a final report is generated. If a hypothesis is successfully validated, the system generates a comprehensive *MoA hypothesis report* detailing the evidence-backed mechanism. Conversely, if all hypotheses are refuted within the maximum rounds, a *failure analysis report* is created. This report details the attempted hypotheses and the reasons for their refutation, providing critical insights for human scientists to inform the next steps of the investigation.

B.2 Implementation details

Frameworks and libraries The MoAgent system is built using Python 3.9. The core agentic workflow and state management are implemented using LangChain (v0.2.5) and LangGraph (v0.1.0), which provide a robust foundation for creating cyclical, stateful multi-agent systems. Key libraries for scientific computation and data handling include:

- **Biological analysis:** decoupler (v1.6.0) [Badia-i Mompel et al., 2022] for pathway and transcription factor activity analysis, and pandas (v2.2.2) for data manipulation.
- **Chemical informatics:** RDKit (v2023.09.6) for chemical structure processing and fingerprint generation, and cmapPy (v5.1.0) for parsing LINCS GCTX files.
- **Database and API interface:** py2neo (v2021.2.4) for interacting with the Neo4j graph database, and chembl-webresource-client (v0.10.9) for accessing the ChEMBL database.

Language models All reasoning, synthesis, and evaluation tasks within the MoAgent framework are performed by Large Language Models. For the experiments presented in this paper, we utilized OpenAI’s gpt-4o-mini model, accessed via their API.

B.3 Agent prompting strategy

Our prompting strategy is central to directing the behavior of the LLM-driven agents. All prompts are engineered to elicit structured reasoning and are designed to request JSON-formatted outputs to ensure reliable, deterministic parsing of agent responses.

MoAgent We employ a “Chain-of-Thought” approach, asking the agent to “think step-by-step” before providing a final answer. This encourages more logical and less error-prone reasoning. Key prompt templates are summarized below. All experiments were conducted using the same prompt.

Chemical structure analysis prompt

You are an expert medicinal chemist. Your task is to analyze a list of compounds that are structurally similar to a query compound and identify the dominant biological themes from their known targets.

Please analyze the following data:

— **Data from Structural Similarity Search:** {Structural Similarity Search Result} —

Analysis Instructions:

1. **Extract and Group Targets:** Examine the ‘target’ column for all listed compounds.
2. **Identify Biological Themes:** Group the identified protein targets into "Biological Themes." A theme represents a biological pathway (e.g., “PI3K/mTOR signaling”), a protein family (e.g., “Receptor Tyrosine Kinases”), or a key cellular process (e.g., “Cell Cycle Regulation”).
3. **Assess Theme Strength:** In the ‘comment’ field, explicitly state the strength of the theme based on the number of supporting compounds and the magnitude of their similarity scores (e.g., ‘Dominant theme supported by multiple high-similarity compounds’, ‘Secondary theme with moderate support’, ‘Weak signal from a single low-similarity compound’).
4. **Format Output:** Structure your entire response as a single JSON object. Do not include any text outside of the JSON structure.

JSON Output Format: “{“json ... ””

Gene expression analysis prompt

You are an expert bioinformatician. Your task is to analyze a list of compounds that produce a similar gene expression signature to a query compound and identify the dominant biological themes from their known targets.

Please analyze the following data:

— **Data from Gene Expression Similarity Search:** { Gene Expression Similarity Search Results } —

Analysis Instructions:

- Extract and Group Targets:** Examine the 'aggregated_targets' column for all listed compounds.
- Identify Biological Themes:** Group the identified protein targets into "Biological Themes." A theme represents a biological pathway (e.g., "PI3K/mTOR signaling"), a protein family (e.g., "Receptor Tyrosine Kinases"), or a key cellular process (e.g., "Cell Cycle Regulation").
- Assess Theme Strength:** In the 'comment' field, explicitly state the strength of the theme based on the number of supporting compounds and the magnitude of their similarity scores (e.g., 'Dominant theme supported by multiple high-similarity compounds', 'Secondary theme with moderate support', 'Weak signal from a single low-similarity compound').
- Format Output:** Structure your entire response as a single JSON object. Do not include any text outside of the JSON structure.

JSON Output Format: “json ... ”

Pathway activity analysis prompt

You are a systems biologist. Your task is to analyze predicted pathway activities (from PROGENy) and summarize them into high-level biological themes.

Please analyze the following data:

— **Data from Pathway Activity Prediction:** { Pathway Activity Prediction Results } —

Analysis Instructions:

- Focus on Significance:** Identify only the most significantly activated or inhibited pathways (e.g., absolute activity score > 10). Ignore pathways with minor activity.
- Quantify Theme Impact:** In the 'interpretation' field, explicitly describe the magnitude of the change using quantitative language (e.g., "Overwhelmingly strong inhibition with a score of -40.7", "Significant activation", "Moderate but notable inhibition"). This provides crucial context for the planning stage.
- Format Output:** Structure your entire response as a single JSON object. Do not include any text outside of the JSON structure.

JSON Output Format: “json ... ”

Hypothesis planning prompt

You are a lead scientist in a drug discovery team. Your task is to create a prioritized, multi-step verification plan to identify a compound's Mechanism of Action (MoA) based on several summarized pieces of evidence. Your reasoning must be sharp, critical, and adhere to scientific principles.

— Evidence Summaries —

{evidence_blocks}

— Core Principles for Scientific Reasoning —

Before formulating hypotheses, you must evaluate the evidence based on the following principles, in this order of importance:

****1. Principle of Convergence (Highest Priority):**** - A hypothesis is strongest when multiple, independent types of evidence (e.g., Pathway Activity, Expression Similarity) converge on the same theme. This is the most reliable starting point. - However, a single, exceptionally strong signal from one data type is still valuable and should be considered a valid hypothesis for testing, as it may represent a critical finding. - Use strong signals (Pathway |score| > 20, Expression |similarity| > 0.3, Structural similarity > 0.5) as the baseline for convergence, while exceptionally strong signals justify an isolated hypothesis.

****2. Principle of Specificity & Strength (Secondary Priority):**** - If no clear convergence exists, prioritize the most specific and overwhelmingly strong signal from a single source (e.g., a pathway activity score > 120). - A strong signal for a specific pathway (e.g., 'Estrogen Signaling') MUST be formulated into an equally specific hypothesis ('Inhibition of Estrogen Signaling'), not a broad one ('Hormonal Regulation').

****3. Principle of Causal Hierarchy (Upstream First):**** - After identifying a theme, apply this principle. Biological pathways are hierarchical. When you see effects across a known signaling cascade (e.g., A -> B -> C), your primary hypothesis must target the ****most upstream plausible component****. - ****Mental Model:**** Ask yourself, "Could inhibiting Target X **cause** the other observed effects?" - ****Mandatory Test Strategy:**** When a major signaling pathway is implicated (e.g., MAPK, PI3K), your 'verification_targets' ****must**** include a representative set of key nodes to ensure full coverage: a plausible ****upstream receptor**** (e.g., RTK), a ****critical intermediate kinase**** (e.g., MAP2K1, AKT), and a ****downstream effector**** (e.g., MAPK1, MTOR). This avoids missing intermediate targets.

— Handling Uncertainty and Contradictory Evidence —

1. ****Acknowledge Contradictions:**** If evidence sources are in strong conflict (e.g., Pathway data suggests activation, but Expression data suggests inhibition), state the contradiction in the 'rationale' and formulate hypotheses to test the conflicting possibilities. 2. ****Flag Low Confidence:**** When evidence is weak, reflect this using the 'confidence' field ("High", "Medium", "Low") and explain why in the 'rationale'.

— Instructions for Plan Formulation —

Based on your rigorous evaluation of the evidence, create a step-by-step verification plan. For each hypothesis, you must follow a three-step process for identifying targets:

1. ****List Evidence Targets:**** In the 'all_evidence_targets' field, list only the protein targets that appeared **directly** in the provided evidence summaries. 2. ****Expand with Biological Knowledge:**** For the 'all_related_targets' field, use your extensive biological knowledge to create a comprehensive list of all known, druggable targets associated with the 'hypothesis_theme'. 3. ****Prioritize for Verification:**** From your expanded 'all_related_targets' list, select a small number (up to 5) of the most critical and representative targets for initial testing and list them in 'verification_targets', following the ****Mandatory Test Strategy**** above.

You can propose up to 5 hypotheses in total.

****JSON Output Format:**** Your entire response must be a single JSON object. Targets should be represented by their HUGO gene symbols (e.g., "EGFR", "AKT1", "MTOR"), and should be in uppercase. Do not use synonyms or other identifiers (e.g., Use KDR not VEGFR2, MAP2K1 not MEK1, MAPK1 not ERK2).

“{json ...}”

Hypothesis re-planning prompt

You are a lead scientist adapting your team's research strategy based on new experimental results. Your initial MoA verification plan has failed. Your task is to analyze the failure, review ALL the original evidence with fresh eyes, and formulate a new, smarter verification plan.

— Original Evidence Summaries —

{evidence_block}

— Previous Verification Plan (Now Invalidated) — “{json {previous_hypothesis}}”

— Verification History (What We've Learned) — “{json {history_str}}”

— Instructions for Re-planning —

1. **Analyze the Failure:** Understand *why* the previous hypotheses were refuted. The ‘verification_history’ shows that direct binding to the selected targets was not observed. This suggests the compound's effects are likely *indirect* or that we tested the wrong targets.

2. **Re-evaluate ALL Evidence with a New Mindset:** Go back to the *original evidence summaries*. Your previous interpretation was wrong. Look for new connections. - **Focus on Upstream Causes:** The observed pathway activity (e.g., MAPK inhibition) is a real effect, but it's downstream. Your primary goal now is to find the upstream cause. - **Systematically Work Backwards:** If the strongest signal was "MAPK pathway inhibition," your new top hypothesis *must* be to test known upstream regulators. Formulate a hypothesis like "Inhibition of MAPK Pathway via Upstream Regulators" and for ‘verification_targets’, you *must* include key upstream kinases like *RAF*, *MEK*, and relevant RTKs (e.g., *EGFR*, *ALK*, *MET*, *ROS1*). - **Promote Secondary Themes:** If systematic backtracking doesn't yield a clear hypothesis, re-examine secondary themes from the original evidence (e.g., a theme with moderate support from expression data) that you previously dismissed. This could be the real MoA.

3. **Consider supported and unsupported targets:** Review the ‘supported_targets’ and ‘unsupported_targets’ from the binding affinity expert. Use this information to avoid previously tested targets and focus on new, plausible candidates. - **Supported Targets:** These targets had some evidence of binding. They might be part of the MoA but were not sufficient alone. Consider them as part of a broader pathway or complex. “{json {supported_str}}” - **Unsupported Targets:** These targets were directly tested and refuted. Avoid them in your new plan. “{json {unsupported_str}}”

Your Goal: Do not simply test the next hypothesis from the old plan. Create a *new* plan that demonstrates you have learned from the experimental failure and are now thinking about *indirect effects and upstream causes*.

JSON Output Format: Your entire response must be a single JSON object containing the new ‘verification_plan’.

Targets should be represented by their HUGO gene symbols (e.g., "EGFR", "AKT1", "MTOR"), and should be in uppercase. Do not use synonyms or other identifiers (e.g., Use KDR not VEGFR2, MAP2K1 not MEK1, MAPK1 not ERK2). “{json ...”

Hypothesis reflection prompt

You are a pharmacology expert specializing in experimental design for the following biological theme: `**{hypothesis_theme}**`. Your task is to review a single, draft MoA hypothesis and ensure its verification plan is optimal and comprehensive.

****Your Goal:**** Refine the `'verification_targets'` for this specific hypothesis to maximize the scientific insight from the first round of experiments. ****Crucially, you MUST NOT change the `'priority'`, `'hypothesis_theme'`, `'confidence'`, or `'rationale'`.** ****** Your focus is solely on improving the experimental design.

— Draft Hypothesis — `“json {hypothesis_str} ““`

— ****Incorporate Search Insights:**** Review the `'search_result'` for any new insights or targets that may have emerged from recent literature or databases. If relevant, integrate these findings into your refined hypothesis. `“json {search_result_str} ““`

— Expert Review and Refinement Instructions —

1. ****Analyze the Target Pool:**** As an expert in this specific theme, examine the `'all_related_targets'`. This list represents the complete set of potential targets for this biological theme. 2. ****Critique the Verification Targets:**** Look at the `'verification_targets'` selected for the experiment. Ask yourself: - ****Completeness:**** Does this list cover the key nodes of the pathway (e.g., upstream receptors, intermediate kinases, downstream effectors)? - ****Upstream Cause:**** Is the most likely upstream driver of this pathway included? - ****Crosstalk:**** Is there a critical, commonly-known crosstalk partner that should be measured to understand the system-level response? - ****Representativeness:**** Are these the best targets to confirm or refute the hypothesis? 3. ****Propose Refinements:**** Rewrite the hypothesis with your expert suggestions. Copy the `'priority'`, `'hypothesis_theme'`, `'confidence'`, and `'rationale'` exactly as they were. Modify only the `'verification_targets'` to be more comprehensive and scientifically robust based on your expert critique. If you believe the original targets are already optimal, you may return the original hypothesis unchanged. 4. ****Consider supported and unsupported targets:**** Review the `'supported_targets'` and `'unsupported_targets'` from the binding affinity expert. Use this information to avoid previously tested targets and focus on new, plausible candidates. If there are not either supported or unsupported targets, you can ignore this step. - ****Supported Targets:**** These targets had some evidence of binding. They might be part of the MoA but were not sufficient alone. Consider them as part of a broader pathway or complex. `“json {supported_targets_str} ““` - ****Unsupported Targets:**** These targets were directly tested and refuted. Avoid them in your new plan. `“json {unsupported_targets_str} ““`

****JSON Output Format:**** Your entire response must be a single JSON object representing the single, refined hypothesis. The structure must be identical to the input. Targets should be represented by their HUGO gene symbols (e.g., "EGFR", "AKT1", "MTOR"), and should be in uppercase. Do not use synonyms or other identifiers (e.g., Use KDR not VEGFR2, MAP2K1 not MEK1).

`“json ... ””`

Bioinformatician prompt

You are an expert in bioinformatics and knowledge graphs, specializing in gene regulatory networks.

Your task is to determine if the proposed hypothesis is supported by evidence of a connection between activated Transcription Factors (TFs) and the proposed protein targets.

— Input Data —

1. **Current Hypothesis:** `“{hypothesis_str}”`
2. **Predicted Transcription Factor Activity:** (Positive scores indicate activation, negative scores indicate inhibition) `“{tf_activity_str}”`
3. **KG-derived Paths between Targets and TFs:** (This data shows which proposed targets have a documented path to the most active TFs in the knowledge graph). `“{kg_path_str}”`

— Instructions for Review —

For each target in the hypothesis's 'verification_targets' list, perform the following analysis:

1. **Identify Active TFs:** Look at the 'Predicted Transcription Factor Activity' data and identify the top 5-10 TFs with the highest absolute scores. Note their activation/inhibition status.
2. **Check for Connectivity:** Refer to the 'KG-derived Paths' data to see if the target has a documented connection to these active TFs. The 'total_tf_coverage' field indicates how many of the top TFs are connected to the target.
3. **Assess Consistency:** The connection is strongest if the direction of the hypothesis (e.g., inhibition of a pathway) is consistent with the TF activity. For example, if a pathway is inhibited, you might expect TFs that activate downstream genes to also be inhibited.
4. **Make a Decision:** - **Supported:** Classify the target as "Supported" if there is strong evidence of connectivity ('total_tf_coverage' > 0) to relevantly activated/inhibited TFs. - **Refuted:** Classify the target as "Refuted" if there is no path ('total_tf_coverage' == 0) or if the connection contradicts the TF activity.
5. **Provide Rationale:** Briefly explain your decision. For example: "Supported due to a direct connection to the highly activated TF 'STAT3' (score: 25.5)" or "Refuted as no path exists in the KG to any of the top 10 active TFs."

JSON Output Format: Your entire response must be a single JSON object.

`“{json ...}”`

Medicinal chemist prompt

You are an expert pharmacologist specializing in interpreting protein-ligand binding affinity data from databases like BindingDB and ChEMBL.

Your task is to review the following experimental binding data for a compound against a list of protein targets. Based *only* on the provided data, determine if the data supports the hypothesis that the compound is a direct and potent binder for each target.

— Input Data —

- **Protein Targets to Evaluate:** `{targets_str}` - **Retrieved Binding Data:** `{data_str}`

— Instructions for Review —

For each target, perform the following analysis:

1. **Review the Data:** Look at the provided binding affinity values (e.g., IC50, Ki, Kd, EC50). Potent binding is typically indicated by low nanomolar (nM) or at worst, low micromolar (< 1 μ M) values.
2. **Make a Decision:** - If the data shows potent binding (e.g., IC50 < 100 nM), classify the target as "Supported". - If no data is available for a target, or if the reported affinity values are weak (e.g., > 10 μ M), classify the target as "Refuted".
3. **Provide Rationale:** Briefly explain your decision based on the data (or lack thereof). For example: "Supported due to a reported IC50 of 50 nM," or "Refuted as no binding data was found in the database."

JSON Output Format: Your entire response must be a single JSON object. The 'has_supported_targets' field should be 'true' if at least one target is marked as "Supported".

`“{json ...}”`

ReAct agent For the ReAct agent, we utilized the standard prompt ('hwchase17/react') from the huggingface hub. This prompt enables the agent to follow a "Chain-of-Thought" process, allowing it to think step-by-step before acting. To adapt this generalist agent to our specific scientific task, we augmented the base prompt with a clear set of instructions, guiding it to analyze the compound's MoA and structure its final output in a standardized JSON format. This approach ensures that while the agent has the freedom to use the provided tools (Table 3), its ultimate goal and the format of its conclusion are well-defined. The specific instructions provided to the agent are detailed below.

ReAct agent prompt
Analyze the mechanism of action (MoA) for the compound associated with signature ID {signature_id}. Your goal is to identify the primary target proteins and explain the MoA. Provide your conclusive answer as a JSON object with two keys: - "predicted_targets": A list of gene symbols for the most likely primary target(s). Gene symbols should be capitalized and be HUGO standard. - "summary": A text summary that states the primary target(s) and synthesizes all the evidence (structural, expressional, pathway, TF, binding, and KG) that supports your conclusion.

Table 3: Tools for MoA inference

Tool Name	Description
find_similar_compounds_by_structure	Finds compounds that are structurally similar to the compound associated with the given signature ID. This tool helps understand the potential MoA based on known MoAs of similar chemical structures. The input to this tool is a single signature ID string (e.g., "CPC014_A549_24H_BRD-K64052750-001-04-3_10"). Returns a markdown table of similar compounds, their known targets, and Tanimoto similarity scores.
find_similar_compound_signatures	Finds other compound signatures that are similar to the input signature based on gene expression profiles (L1000 data). This is useful for finding compounds that have a similar biological effect, even if their structures are different. The input to this tool is a single signature ID string (e.g., "CPC014_A549_24H_BRD-K64052750-001-04-3_10"). Returns a markdown table summarizing compounds with similar signatures and their aggregated targets.
predict_pathway_activity	Predicts the activity of biological pathways (e.g., MAPK, PI3K) based on the gene expression signature. This helps to understand the high-level biological processes affected by the compound. The input to this tool is a single signature ID string (e.g., "CPC014_A549_24H_BRD-K64052750-001-04-3_10"). Returns a markdown table of pathways and their predicted activity scores.
predict_transcription_activity	Predicts the activity of transcription factors (TFs) based on the gene expression signature. This can reveal the key regulatory proteins that mediate the compound's effects. The input to this tool is a single signature ID string (e.g., "CPC014_A549_24H_BRD-K64052750-001-04-3_10"). Returns a markdown table of the top activated/inactivated transcription factors and their scores.
query_binding_affinity	Queries public databases (ChEMBL, BindingDB) for measured binding affinities between the compound (identified by sig_id) and a list of potential target proteins (gene symbols). This is a crucial step for verifying direct interaction. The input must be a single string containing the signature ID, followed by a comma, and then a Python-style list of gene symbols. For example: "CPC006_HCC515_24H, ['TP53', 'EGFR']". Returns a string summarizing the findings for each target.
analyze_tf_target_connections	Analyzes the connections in a biological knowledge graph between a list of potential target genes and a list of activated transcription factors (TFs). This helps to check if a proposed target can plausibly explain the observed TF activity. The input must be a single string containing a Python-style list of target gene symbols, followed by a comma, and then a Python-style list of TF gene symbols. For example: "['TP53', 'EGFR'], ['MYC', 'JUN']".

B.4 Hyperparameter configuration

The behavior of the MoAgent is governed by several key hyperparameters, which were fixed for all experiments to ensure consistency. The configuration is detailed in Table 4.

Table 4: Hyperparameter settings for the MoAgent framework.

Hyperparameter	Value
LLM Model Name	gpt-4o-mini
LLM Temperature	0.0
Max Evidence Compounds	20
Max Re-planning Attempts	3
Plan Reflection Step	True
Max Target Expansion Attempts	1

C Experimental details

C.1 Dataset

We constructed a benchmark dataset from the Broad Institute’s LINCS L1000 library of drug-induced gene expression signatures [Subramanian et al., 2017]. To mitigate confounding variables and ensure a rigorous comparison, we exclusively selected profiles generated under consistent experimental conditions: a single cell line (MCF7), a fixed dosage (10 μ M), and a 24-hour time point. This controlled setting isolates the mature transcriptomic response to a standardized perturbation, leading to a total of 26 individual signatures, consisting of randomly sampled three biological replicates for eight drugs and two for one. Our final benchmark comprises 9 targeted anti-cancer drugs, chosen to represent a diverse range of well-established MoAs (section C.5). The ground truth MoAs were rigorously compiled and cross-referenced from leading expert-annotated databases, including clue.io, the Drug Repurposing Hub, and DrugBank [Subramanian et al., 2017, Corsello et al., 2017, Knox et al., 2023]. For our similarity baseline, we utilized the gene expression signatures from the LINCS L1000 dataset, constrained to a single cell line (MCF7), a fixed dosage (10 μ M), and time points of 6 and 24 hours.

C.2 Baselines

We compare MoAgent against two representative baselines to contextualize its performance. The first, signature similarity, represents the conventional, non-agentic approach to this problem, while the second, a ReAct Agent [Yao et al., 2022], serves as an architectural ablation to evaluate our collaborative multi-agent framework against a single-agent system.

- **Signature similarity:** This method represents the conventional transcriptomic-based approach to MoA inference, as popularized by the CMap [Lamb et al., 2006]. It operates on the assumption that similar gene expression signatures imply similar MoAs. For a given query signature, we retrieve the most similar signatures from the reference library based on their cosine similarity to the query. The MoAs of the top- k most similar cases are aggregated (union) to make a final prediction. This baseline test evaluates whether MoAgent, a multi-agent-based system with tools beyond similarity measurement, provides a benefit over simple pattern matching.
- **ReAct agent:** We implemented a generalist agent using the ReAct framework, providing it with the exact same set of tools as MoAgent. However, it operates as a single agent following a simpler *think-act-observe* loop, lacking the multi-agent collaborative strategy and hypothesis-driven process of MoAgent. This comparison evaluates whether MoAgent’s performance gains stem from its structured, collaborative workflow rather than just the availability of tools. More detailed explanation can be found in Appendix B.3.

C.3 Evaluation metrics

MoA inference is a multi-label classification task, as a single drug can have multiple MoAs. We evaluate performance using standard metrics: precision (P), recall (R), and F1-score (F1). We report results under two distinct aggregation strategies:

- **Individual:** Metrics are calculated for each of the 26 individual signatures and then averaged. This evaluates the *per-signature* accuracy in MoA inference.
- **Union:** For each drug, the MoA inferences from its replicate signatures are combined into a single set using a union operation. This approach is important from a biological standpoint, as scientists typically form robust hypotheses from a consensus of multiple experiments rather than a single result. This aggregation is also crucial for evaluation integrity because it mitigates the impact of any single noisy replicate. Such outliers, often caused by factors like batch effects, could otherwise lead to a misleading assessment of model performance. Therefore, this metric assesses the per-drug accuracy of the model in MoA inference.

C.4 Knowledge sources and pre-processing

Chemical structural similarity To quantify structural similarity, compounds were represented as 2048-bit Morgan fingerprints with a radius of 2. The Tanimoto similarity between these fingerprints was then calculated using the RDKit library.

Gene expression data We utilized the LINCS L1000 dataset from the clue.io repository. Specifically, we used the level 5 data, which consists of moderated Z-scores representing the differential expression of 12,328 genes upon perturbation, normalized relative to controls. For details, see Subramanian et al. [2017].

Biological knowledge graph Our knowledge graph was constructed in Neo4j. Protein interactions involved in signal transduction were obtained from SIGNOR 3.0 [Lo Surdo et al., 2023], a database that provides manually curated causal interactions annotated with direction and signed information. Transcription factor–target interactions were collected from CollecTRI [Müller-Dott et al., 2023], which integrates multiple public datasets through manual curation and includes 43,175 interactions across 1,186 transcription factors. Pathway–gene relationships were derived from PROGENy [Schubert et al., 2018], a method that identifies pathway-responsive genes by analyzing gene expression patterns from numerous pathway perturbation experiments. These relationships also account for the effects of post-translational regulation. The graph primarily consists of ‘protein’ nodes and directed edges representing regulatory relationships, such as ‘UPREGULATE’ and ‘DOWNREGULATE’.

External bioactivity databases The Chemical Verification Agent performs real-time queries to the ChEMBL (v34) and BindingDB databases to fetch experimental bioactivity data. No pre-processing was performed on these sources.

C.5 Test set and experimental setup for zero-shot scenario

Table 5: Test drugs and their target proteins

Drug Name	Target Protein(s)
gefitinib	EGFR
selumetinib	MAP2K1, MAP2K2
everolimus	MTOR
alectinib	ALK, MEK
vemurafenib	BRAF, RAF1
bexarotene	RXRA, RXRB, RXRG
nilotinib	BCR, ABL1, KIT, PDGFRA, PDGFRB, CSF3R, DDR1, DDR2
ruxolitinib	JAK1, JAK2, JAK3, TYK2, PLAUR
tivozanib	FLT1, KDR, FLT4, KIT, PDGFRA, PDGFRB, FLT3, PTK6, MET

For our evaluation in the “w/o Protein Target” setting, a specific adjustment was made for gefitinib. We additionally removed the related proteins ERBB2, ERBB3, and ERBB4 from the candidate targets, given that EGFR inhibitors often share high target similarity with other members of the ERBB family. This adjustment, as clearly demonstrated in Tables 6 and 7, results in the disappearance of EGFR and ERBB family inhibitors from the list of most structurally similar compounds.

Table 6: Retrieved evidence by chemical structure (before target protein removal)

Name	Target Protein(s)	Similarity
canertinib	ERBB2, EGFR, ERBB4	0.697368
ZM-447439	AURKB, AURKA	0.653333
ZM-306416	KDR, FLT1	0.485294
foretinib	KDR, MET	0.483146
WHI-P154	JAK3, EGFR	0.442857
tandutinib	PDGFRA, KIT, FLT3, PDGFRB	0.411111
BIBX-1382	EGFR	0.407407
KIN001-055	JAK3	0.400000
AV-412	ERBB2, EGFR	0.397849
BIBU-1361	EGFR	0.383721

Table 7: Retrieved evidence by chemical structure (after target protein removal)

Name	Target Protein(s)	Similarity
ZM-447439	AURKB, AURKA	0.653333
ZM-306416	FLT1, KDR	0.485294
foretinib	MET, KDR	0.483146
WHI-P154	JAK3	0.442857
tandutinib	FLT3, KIT, PDGFRB, PDGFRA	0.411111
KIN001-055	JAK3	0.400000
LY-456236	GRM1	0.347826
iloperidone	HTR7, DRD2, HTR6, HTR2A, HTR1A	0.315217
copanlisib	PIK3CD, PIK3CA	0.312500
eprobemide	MAOA	0.293333

D Case study: a walkthrough for gefitinib

To provide a concrete illustration of MoAgent’s reasoning process, we present a condensed walk-through for the well-known drug Gefitinib, a small molecule inhibitor.

1. Evidence retrieval The process begins by analyzing three independent lines of evidence to identify convergent biological themes. (a) The analysis of structurally similar compounds (Table 8) proposed several kinase inhibitor activities, with Aurora Kinase and VEGFR signaling being the most prominent. (b) Concurrently, the analysis of similar gene expression signatures (Table 9) suggested broader cancer-related mechanisms such as Topoisomerase Inhibition and Cell Cycle Regulation. (c) Finally, the direct pathway activity analysis (Table 10) provided a decisive report on a potent and specific inhibitory signal targeting the EGFR pathway.

Table 8: Potential MoA themes for gefitinib based on structural similarity.

theme	targets	compounds	similarity	comment
Aurora Kinase Inhibition	AURKB, RKA	AU- ZM-447439	0.653333	Dominant theme supported by a single high-similarity compound.
Vascular Endothelial Growth Factor Receptor Signaling	FLT1, KDR	ZM-306416, foretinib, ZM-323881	0.485294, 0.483146, 0.261364	Dominant theme supported by multiple compounds with moderate structural similarity.
Janus Kinase Inhibition	JAK3	WHI-P154, KIN001-055	0.442857, 0.4	Secondary theme with moderate support.
Receptor Tyrosine Kinase Inhibition	FLT3, PDGFRB, PDGFRA	KIT, tandutinib	0.411111	Secondary theme with moderate support.
Serotonin Receptor Modulation	HTR7, HTR6, HTR1A	DRD2, HTR2A, iloperidone	0.315217	Weak signal from a single low-similarity compound.
Phosphoinositide 3-Kinase Inhibition	PIK3CD, PIK3CA	copanlisib	0.3125	Weak signal from a single low-similarity compound.
Monoamine Oxidase Inhibition	MAOA	eprobemide, moclobemide	0.293333, 0.266667	Weak signal from two low-similarity compounds.
Histone Methyltransferase Inhibition	EHMT2	UNC-0321, BIX-01294	0.28125, 0.27957	Weak signal from two low-similarity compounds.
Serotonin Receptor Modulation	HTR4	RS-23597-190, cisapride	0.280488, 0.278351	Weak signal from two low-similarity compounds.
Cannabinoid Receptor Modulation	CNR2	GW-405833	0.268817	Weak signal from a single low-similarity compound.
Mitogen-Activated Protein Kinase Inhibition	MAPK12, MAPK11, MAPK14	doramapimod	0.262626	Weak signal from a single low-similarity compound.
Sigma Receptor Modulation	SIGMAR1	API-001, cutamesine	0.261905, 0.25974	Weak signal from two low-similarity compounds.

2. Hypothesis generation Based on the convergent themes identified from the retrieved evidence, the hypothesis generation module formulates and prioritizes a set of actionable scientific hypotheses. As detailed in Table 11, the module assigns the highest priority to the “Inhibition of Aurora Kinase Signaling” hypothesis, citing the strong, convergent evidence from both structural and gene expression analyses. The “Inhibition of EGFR Pathway” is identified as the second-highest priority hypothesis, supported with high confidence due to the potent and specific signal from the pathway analysis. Other plausible MoAs, such as “Topoisomerase and PI3K inhibition,” are also identified but ranked lower due to their reliance on evidence from a single source. This ranked list of hypotheses provides a clear roadmap for the subsequent verification stage.

Table 9: Potential MoA themes for gefitinib based on gene expression similarity.

theme	targets	compounds	similarity	comment
Topoisomerase Inhibition	TOP2A	Doxorubicin	0.502557	Dominant theme supported by a high-similarity compound.
Cell Cycle Regulation	CCND1, CDK4	Arcyriaflavin-a	0.498831	Strong theme supported by a high-similarity compound.
p53 Pathway	TP53, HSPA1A	Pifithrin-mu	0.49714	Strong theme supported by a high-similarity compound.
Retinoic Acid Signaling	RXRG, RARA, RORB, RARB, RARG, RXRA, RXRB	Retinol	0.491672	Moderate theme supported by a single compound.
DNA Damage Response	CHEK1	PF-477736	0.490488	Moderate theme supported by a single compound.
Dopaminergic Signaling	ADRA2C, DRD4, DRD2, DRD3, ADRA2B, HTR2B, HTR2C	Bromocriptine	0.489856	Moderate theme supported by a single compound.
Histamine Receptor Signaling	HRH1	Ketotifen, Terfenadine	0.48569, 0.470357	Moderate theme supported by multiple compounds.
PI3K Signaling	PIK3CD, PRKDC, PIK3CA, PLK1, PIK3CG, PIK3CB	LY-294002	0.472557	Dominant theme supported by a high-similarity compound.
Serotonin Receptor Signaling	HTR4, KCNH2	Cisapride	0.470699	Moderate theme supported by a single compound.
PARP Inhibition	PARP1, PARP2	Veliparib	0.463792	Moderate theme supported by a single compound.
Receptor Tyrosine Kinases	MERTK, MST1R, AXL, TYRO3, MET	BMS-777607	0.463594	Moderate theme supported by a single compound.
Aurora Kinase Inhibition	AURKB	Barasertib-HQPA	0.459572	Moderate theme supported by a single compound.

Table 10: Potential MoA themes for gefitinib based on pathway activity.

theme	pathways	activity scores	comment
Inhibition of EGFR Pathway	EGFR	-10.251666	Significant inhibition of the EGFR pathway, indicating a strong impact on signaling related to cell proliferation and survival.

3. Hypothesis validation The final stage involves a validation module that systematically tests the prioritized hypotheses by dispatching tasks to a team of validation specialists: a bioinformatician and a medicinal chemist. The module first investigates the top-ranked hypothesis, “Inhibition of Aurora Kinase Signaling,” by evaluating its proposed targets (Table 12). Although the bioinformatician confirms plausible biological pathways for several targets (e.g., AURKA, AURKB), the medicinal chemist consistently refutes these findings, citing weak or non-existent binding affinity (e.g., $K_d > 10,000$ nM). Due to this critical lack of biochemical support, the Aurora Kinase hypothesis is ultimately rejected.

Proceeding to the second hypothesis, the module examines the “Inhibition of EGFR Pathway” (Table 13). In this case, a strong consensus is achieved for the primary target, EGFR. The bioinformatician validates its connection to active transcription factors, while the medicinal chemist confirms potent binding with a reported IC₅₀ of 0.1 nM. This successful, concordant verification, contrasted with the failure of the higher-ranked but unsupported hypothesis, allows MoAgent to confidently confirm that the primary Mechanism of Action for Gefitinib is the direct inhibition of EGFR.

With EGFR established as the primary target, the investigation then broadens to assess the compound’s effects on closely related proteins, such as other ErbB family members and key downstream signaling molecules, to fully delineate its pharmacological scope (Table 14).

Table 11: Hypothesis prioritization and supporting evidence

priority	hypothesis_theme	confidence	rationale	all_evidence_targets	all_related_targets	verification_targets
1	Inhibition of Aurora Kinase Signaling	High	This hypothesis is supported by strong convergence from both structural similarity evidence (high similarity with ZM-447439) and gene expression evidence (high similarity with Barasertib-HQPA), indicating a significant role of Aurora Kinases in the compounds mechanism of action.	AURKB, AURKA	AURKA, AURKB, CDK1, CDK2, PLK1, CCNB1, TP53, BCL2, MAPK1, CDKN1A	AURKA, AURKB, CDK1, CCNB1, TP53, CDKN1A
2	Inhibition of EGFR Pathway	High	The pathway activity evidence shows significant inhibition of the EGFR pathway, which is a strong signal indicating a direct impact on cell proliferation and survival mechanisms.	EGFR	EGFR, ERBB2, ERBB3, ERBB4, PIK3CA, KRAS, AKT1, MTOR, RAPIA, SOS1	EGFR, KRAS, MTOR, SOS1, PIK3CA, RAPIA
3	Inhibition of Topoisomerase Activity	Medium	The gene expression evidence indicates a strong similarity with Doxorubicin, a known topoisomerase inhibitor, suggesting that the compound may also inhibit topoisomerase activity.	TOP2A	TOP1, TOP2B, CHK1, TP53, PARP1	TOP2A, ATM, BRCA1, TP53, MRE11, PARP1
4	Inhibition of PI3K Signaling	Medium	The gene expression evidence shows a moderate similarity with LY-294002, a PI3K inhibitor, indicating potential involvement of the PI3K pathway in the compounds mechanism.	PIK3CD, PIK3CA	PIK3CA, PIK3CB, PIK3CD, AKT1, MTOR, PIK3R1, PTEN, RPS6KB1, GSK3B	PIK3CA, PIK3R1, RPS6KB1, GSK3B, AKT1, PTEN
5	Inhibition of Cell Cycle Regulation	Low	The evidence suggests a moderate similarity with Arcyriaflavin-a, which is known to affect cell cycle regulation, but the signal strength is not robust enough to prioritize this hypothesis highly.	CCND1, CDK4	CCND1, CDK4, CDK6, RB1, E2F1, CDKN1A, CCNE1, TP53, MDM2	CCND1, CDK4, CDKN1A, CCNE1, TP53, MDM2, RB1, E2F1

Table 12: Hypothesis verification and rationale 1

target_name	is_supported	rationale
AURKA	False	[Bioinformatician] Supported. AURKA has a total TF coverage of 3, showing connections to active TFs, supporting the hypothesis of inhibiting Aurora Kinase signaling. [Medicinal chemist] Refuted as the reported Kd is > 10,000 nM, indicating weak binding.
AURKB	False	[Bioinformatician] Supported. AURKB has a total TF coverage of 3, indicating connections to active TFs, which aligns with the hypothesis. [Medicinal chemist] Refuted as the reported Kd is > 10,000 nM, indicating weak binding.
CDK1	False	[Bioinformatician] Supported. CDK1 also has a total TF coverage of 3, showing connections to active TFs, supporting the hypothesis regarding Aurora Kinase signaling inhibition. [Medicinal chemist] Refuted as no binding evidence was found in the database.
CCNB1	False	[Bioinformatician] Refuted. CCNB1 has no documented path to any of the top active TFs, indicating a lack of support for the hypothesis. [Medicinal chemist] Refuted as no binding evidence was found in the database.
TP53	False	[Bioinformatician] Supported. TP53 has a total TF coverage of 3, indicating connections to multiple active TFs, which aligns with the hypothesis of inhibiting Aurora Kinase signaling. [Medicinal chemist] Refuted as no binding evidence was found in the database.
CDKN1A	False	[Bioinformatician] Supported. CDKN1A has a total TF coverage of 3, indicating it is connected to active TFs, which is consistent with the hypothesis. [Medicinal chemist] Refuted as no binding evidence was found in the database.

4. Final conclusion and report generation The successful validation of the “EGFR Inhibition” hypothesis, following the refutation of the top-ranked Aurora Kinase hypothesis, serves as the stopping condition for the agent’s iterative process of hypothesis generation and verification. With a supported MoA now confirmed, the Controller & Reasoning Module proceeds to its final task: synthesizing all analyzed evidence into a comprehensive scientific report. Table 15 presents this final output, which constructs a complete narrative around the primary MoA by detailing the causal chain, discussing network crosstalk, and critically evaluating both supporting and contradictory evidence.

Table 13: Hypothesis verification and rationale 2

target_name	is_supported	rationale
EGFR	True	[Bioinformatician] Supported. EGFR has a documented connection to 3 active TFs, and the inhibition of the EGFR pathway aligns with the activation of TFs like SREBF1 and SREBF2. [Medicinal chemist] Supported due to a reported IC50 of 0.1 nM, indicating potent binding.
PIK3CA	False	[Bioinformatician] Supported. PIK3CA is connected to 3 active TFs, and the inhibition of the EGFR pathway is consistent with the activation of TFs that may regulate downstream effects. [Medicinal chemist] Refuted as the only available data shows a weak Kd of >10,000 nM.
KRAS	False	[Bioinformatician] Supported. KRAS has a documented connection to 3 active TFs, and the inhibition of the EGFR pathway supports the activation of these TFs. [Medicinal chemist] Refuted as no binding evidence was found in the database.
AKT1	False	[Bioinformatician] Supported. AKT1 is connected to 3 active TFs, and the inhibition of the EGFR pathway is consistent with the activation of these TFs. [Medicinal chemist] Refuted as the only available data shows a weak Kd of >10,000 nM.
MTOR	False	[Bioinformatician] Supported. MTOR has a documented connection to 3 active TFs, and the inhibition of the EGFR pathway aligns with the activation of these TFs. [Medicinal chemist] Refuted as the only available data shows a weak Kd of >10,000 nM.
RAP1A	False	[Bioinformatician] Supported. RAP1A is connected to 2 active TFs, and the inhibition of the EGFR pathway is consistent with the activation of these TFs. [Medicinal chemist] Refuted as no binding evidence was found in the database.
SOS1	False	[Bioinformatician] Supported. SOS1 has a documented connection to 2 active TFs, and the inhibition of the EGFR pathway aligns with the activation of these TFs. [Medicinal chemist] Refuted as no binding evidence was found in the database.

Table 14: Target expansion

target	rationale
ERBB2	ERBB2 (HER2) is a member of the ErbB family and often forms heterodimers with EGFR. Inhibitors of EGFR have been shown to also affect HER2 signaling, making it a relevant target for potential cross-reactivity.
ERBB3	ERBB3 (HER3) is another member of the ErbB family that can form heterodimers with EGFR. Given its role in signaling pathways activated by EGFR, it is important to assess its potential inhibition.
MAPK1	MAPK1 (ERK2) is part of the Ras-Raf-MEK-ERK pathway activated by EGFR. Since the compound targets EGFR, it is plausible that it may also affect downstream signaling components like MAPK1.
PIK3CA	PIK3CA is a key component of the PI3K-Akt pathway, which is activated by EGFR. Given the interconnectedness of these pathways, testing PIK3CA could reveal additional effects of the compound.
AKT1	AKT1 is a downstream effector in the PI3K-Akt pathway, which is activated by EGFR signaling. Since compounds targeting EGFR may also influence AKT1 activity, it is a relevant target for verification.

Table 15: Final report on the mechanism of action of gefitinib

Final Report on the Mechanism of Action of the Compound

1. Primary Mechanism of Action Hypothesis

a. State the Primary MoA

The primary mechanism of action for this compound is the direct inhibition of the Epidermal Growth Factor Receptor (EGFR). This conclusion is supported by experimental data indicating a potent binding affinity to EGFR, with an IC50 of 0.1 nM, confirming its role as a significant target.

b. Build the Causal Chain

Direct binding to EGFR leads to profound inhibition of the EGFR pathway, as evidenced by a significant activity score of -10.251666. This inhibition is critical as it impacts downstream signaling pathways associated with cell proliferation and survival. The inhibition of EGFR signaling is further supported by the activation of transcription factors SREBF1 and SREBF2, which are known to be positively regulated by EGFR activity. The predicted transcription factor activities of SREBF1 (4.596963) and SREBF2 (5.759338) indicate a compensatory response to the inhibition of EGFR, suggesting a shift in cellular signaling dynamics.

c. Explain Network Effects (Crosstalk)

In addition to the direct effects on the EGFR pathway, the compound's action appears to have secondary effects on related pathways. For instance, the data indicates a significant inhibition of the PI3K signaling pathway, which is a known crosstalk partner of the EGFR pathway, with an activity score of -12.5. This suggests that the primary inhibition of EGFR signaling not only disrupts its direct downstream effects but also leads to a broader suppression of associated signaling networks, further emphasizing the compound's impact on cellular signaling.

2. Corroborating Evidence from Similar Compounds

The hypothesis of EGFR inhibition is strongly corroborated by expression similarity data. The compound shows a high gene expression similarity (score: 0.472557) to **LY-294002**, a well-documented PI3K inhibitor, which aligns with the observed inhibition of the PI3K pathway as a consequence of EGFR inhibition. This phenocopying of cellular responses to a known drug with overlapping effects provides robust evidence for the proposed mechanism.

Additionally, the compound exhibits structural similarity (score: 0.653333) to **ZM-447439**, a known Aurora Kinase inhibitor, which, while not directly targeting EGFR, supports the notion of the compound's potential multi-target profile. However, the primary focus remains on the EGFR inhibition as the core mechanism.

3. Analysis of Unexplained & Contradictory Evidence

a. Propose Secondary / Compensatory Mechanisms

The analysis also reveals an unexpected activation of the p53 pathway, indicated by the involvement of targets such as TP53 and HSPA1A, which are known to be activated in response to cellular stress and DNA damage. This activation may represent a compensatory mechanism in response to the inhibition of EGFR signaling, as cells often engage stress response pathways when proliferative signals are disrupted.

b. Identify Potential Off-Target Effects

One unexplained observation is the moderate inhibition of the Topoisomerase pathway (related target: TOP2A), which is not a canonical downstream effect of EGFR signaling. This suggests that the compound may possess additional, currently unidentified off-target activities. Further investigation, such as broad kinase screening, is recommended to identify the protein(s) responsible for this effect.

4. Conclusive Summary

In conclusion, the compound primarily acts through the inhibition of EGFR, leading to significant downstream effects on the EGFR and PI3K pathways, as well as the activation of transcription factors SREBF1 and SREBF2. The observed network effects highlight the compound's ability to disrupt multiple signaling pathways, reinforcing its potential as a therapeutic agent. However, the presence of unexplained off-target activities, particularly the inhibition of the Topoisomerase pathway, warrants further investigation to fully elucidate the compound's pharmacological profile and potential side effects.
