
BixBench: a Comprehensive Benchmark for LLM-based Agents in Computational Biology

Ludovico Mitchener^{* 1} Jon M Laurent^{* 1} Benjamin Tenmann² Siddharth Narayanan¹ Geemi P Wellawatte¹
Andrew White¹ Lorenzo Sani² Samuel G Rodriques¹

Abstract

Large Language Models (LLMs) and LLM-based agents show great promise in accelerating scientific research. Existing benchmarks for measuring this potential and guiding future development continue to evolve from pure recall and rote knowledge tasks, towards more practical work such as literature review and experimental planning. Bioinformatics is a domain where fully autonomous AI-driven discovery may be near, but no extensive benchmarks for measuring progress have been introduced to date. We therefore present the **Bioinformatics Benchmark** (BixBench), a dataset comprising over 50 real-world scenarios of practical biological data analysis with nearly 300 associated open-answer questions designed to measure the ability of LLM-based agents to explore biological datasets, perform long, multi-step analytical trajectories, and interpret the nuanced results of those analyses. We evaluate the performance of two frontier LLMs (GPT-4o and Claude 3.5 Sonnet) using a custom agent framework we open source. We find that even the latest frontier models only achieve 17% accuracy in the open-answer regime, and no better than random in a multiple-choice setting. By exposing the current limitations of frontier models, we hope BixBench can spur the development of agents capable of conducting rigorous bioinformatic analysis and accelerate scientific discovery.

1 Introduction

Accelerating scientific discovery across domains is a core aspiration for modern AI models and agents, driven by the ability of these systems to synthesize information and identify novel connections across disparate fields (Skarlinski et al., 2024). While significant progress has been made in developing systems for automating certain aspects of the scientific research process (Lu et al., 2024; Gao et al., 2024; Swanson et al., 2024; Schmidgall et al., 2025; Narayanan et al., 2024), the ultimate realization of fully automated scientific systems remains elusive.

The analysis of experimental data is fundamental to scientific research (Deshpande et al., 2024). Given the scale of data being produced by modern technologies and the complexity of analytical approaches, disciplines like bioinformatics and computational biology have risen to the forefront as fields in their own right. These specialties represent a microcosm of the scientific process where hypotheses and research questions drive analytical, statistical and computational approaches to uncover new knowledge from a variety of data sources, at times independently from traditional laboratory experiments. This disconnect from the physical laboratory presents an opportunity to develop AI systems for exploring data with the sophistication of expert bioinformaticians but at significantly larger scale and more rapid pace. Hence, bioinformatics stands as one of the first fields where fully autonomous scientific systems may become a reality.

As with any applied AI system, quantitatively assessing performance on benchmarks serves as a primary measure of progress (Chang et al., 2023; Huang et al., 2021; Wognum et al., 2024). Countless such benchmarks have been created for assessing capabilities such as image segmentation, natural language understanding, and protein-ligand binding (Zhou et al., 2017; Hendrycks et al., 2020; Quigley et al., 2024). More recently, groups have published benchmarking datasets focused on high complexity tasks with real-world biological research applicability (Laurent et al., 2024). These real-world benchmarks are already proving their value in driving development of scientific agents, including ones that are surpassing human expert performance on research activ-

^{*}Equal contribution ¹FutureHouse, San Francisco, USA ²ScienceMachine, London, UK. Correspondence to: Samuel G Rodriques <sam@futurehouse.org>, Lorenzo Sani <lorenzo@sciencemachine.ai>.

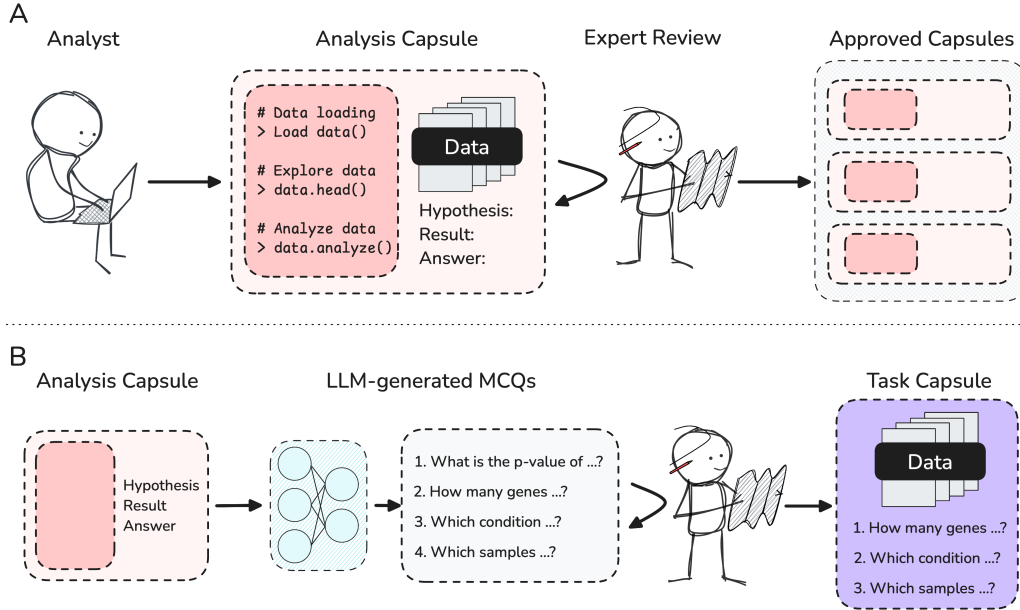


Figure 1: BixBench benchmark creation diagram. (A) To create the initial seed capsules, expert bioinformaticians assembled a code notebook and provided input data and metadata including a hypothesis, result, etc. Seed capsules were reviewed by other experts before being merged to final corpus. (B) To generate benchmark tasks, we first asked an LLM to propose candidate questions for each capsule. These questions were reviewed by multiple experts, yielding the final dataset.

ities like literature review and DNA construct engineering (Narayanan et al., 2024; Skarlinski et al., 2024). However, systems capable of truly open-ended scientific exploration and benchmarks for assessing them, remain elusive.

To help bridge this gap, we introduce the **Bioinformatics Benchmark** (BixBench). The benchmark is composed of 53 real-world analytical scenarios, each a pairing of guiding questions with heterogeneous input data files. These scenarios present a realistic environment in which an autonomous agent is expected to perform appropriate analyses given minimal prompting and context. Performing well on the benchmark requires being simultaneously skilled in understanding the subtleties of a research question, exploring data and understanding how it relates to the question, performing multi-step computational analyses, and finally interpreting the results in the context of the initial research question.

To assess the capabilities of current AI models for performing these types of analyses, we present results of two frontier models (GPT-4o from OpenAI (Hurst et al., 2024) and Claude 3.5 Sonnet from Anthropic (Anthropic, 2024)) integrated into a custom agent framework and evaluated on BixBench. We find that frontier models perform poorly overall ($\sim 17\%$ accuracy at best) in the primary open-answer evaluation format. When evaluated in a modified multiple-choice (MCQ) framework using majority voting, performance was no better than random guessing. We also provide extensive validation of our experimental setup, in-

vestigating potential data leakage, hyperparameter optimization, and frontier model choice.

1.1 Primary Contributions

The primary contributions of our work are as follows:

1. An expert-curated set of **53 realistic analytical scenarios** and **296 open-answer questions**, focused on the domain of bioinformatics and computational biology¹.
2. A framework for **assessing the performance of AI agents in biological data analysis**, including key metrics and calibrations².
3. An **open-source agent framework** for writing and executing Python, R and bash code in a Jupyter notebook environment³.

¹The benchmark can be found at <https://huggingface.co/datasets/futurehouse/BixBench>

²An evaluation harness can be found at <https://github.com/Future-House/BixBench>

³The data analysis agent can be found at <https://github.com/Future-House/data-analysis-crow>

2 Related Work

2.1 Data Science Benchmarks.

There has been an explosion in the number of benchmarks aimed at evaluating the ability of LLM-based agents to write and execute increasingly complex analytical workflows (Jing et al., 2024; Chan et al., 2024; Huang et al., 2024; Gu et al., 2024; Wijk et al., 2024). Tasks within these benchmarks range from simple data-frame operations to solving complex research engineering tasks. Benchmarks like RE-Bench (Wijk et al., 2024) require complex, iterative problem-solving, but rely on a simple reward function such as training loss to guide an agent’s performance. On the other end of the spectrum, benchmarks like BLADE (Gu et al., 2024) present broader analysis scenarios similar to those introduced here, but are artificially constrained to producing specific artifacts and making specific statistical choices.

2.2 Benchmarks for Scientific Agents.

Similarly, a number of benchmarking suites have been proposed for measuring AI agents’ ability to complete key scientific tasks, previously reserved for humans. These include ChemBench (Mirza et al., 2024), which tests models on a wide range of chemistry tasks, and LAB-Bench (Lau-
rent et al., 2024) which is composed of a large set of real-world biology research tasks ranging from literature review to DNA sequence manipulation and protocol troubleshooting. Closer to the bioinformatics space are benchmarks like BioCoder (Tang et al., 2024) that test code generation and function calling across a defined function set.

In science, most analytical tasks are ambiguous, open-ended, and do not benefit from having a clear optimization metric to verify performance. We developed BixBench to bridge this gap and provide a benchmark that evaluates challenging multi-step analytical capabilities in complex, ambiguous research tasks without being artificially constrained to specific environments, tools, or tasks.

3 Methods

3.1 Benchmark creation

3.1.1 ANALYST RECRUITMENT

To develop BixBench, we systematically assembled a diverse collection of analytical trajectories representing canonical tasks in bioinformatics research (Figure 2). We recruited expert analysts through multiple channels: our own professional networks, reaching out to authors of bioinformatics papers, and recruiting through affiliated research institutions. Our team of contracted analysts consists exclusively of PhD holders or candidates in bioinformatics and related fields,

each with extensive experience in biological data analysis as their primary research focus.

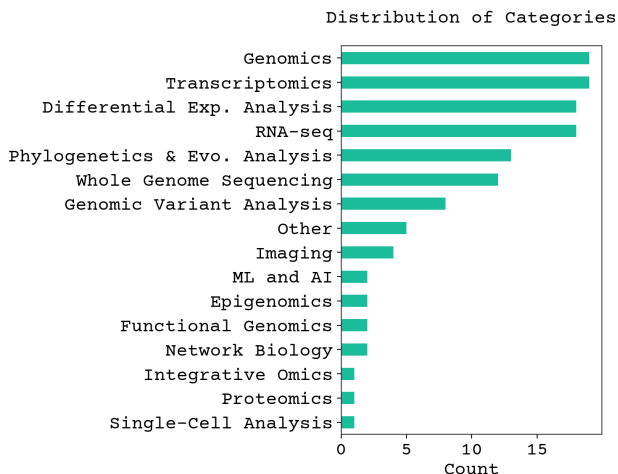


Figure 2: Histogram of counts for capsules in a broad set of self-selected categories. BixBench covers a wide array of analysis types.

3.1.2 CAPSULE CREATION

We refer to our gathered analytical trajectories and associated data as analysis “capsules”. An analysis capsule consists of three primary components: 1) A `hypothesis` or research question driving the analysis, 2) `input data` going into the analysis, and 3) `code` for carrying out the analysis. We also captured a `result` (a few sentences describing the primary result of the analysis as it pertains to the hypothesis) and `answer`, a true/false assessment of whether or not the hypothesis was supported by the analysis, as well as some additional metadata.

Contract analysts were asked to recapitulate published analyses or produce interesting *de novo* analytical trajectories. To retain flexibility while maintaining a standardized environment, we capture code in Jupyter notebooks using Google Colab. We provided a user interface through which analysts initiate a new capsule, triggering creation of a new empty Google Drive folder for the capsule and pre-populating it with a template Colab notebook containing instructions and with specific formatting and structure. The analysts were instructed to write their analysis code in the Colab notebook and upload data files to the notebook environment. Files uploaded to the Colab notebook environment itself are transient, and thus analysts are further instructed to upload their data files to the Drive folder for us to capture. Analysts were allowed to use Python, R, or bash commands inside the notebook environment, and install any packages necessary.

Capsules and their trajectories were thoroughly reviewed by the authors and in some cases by other analysts, to arrive at

Table 1: BixBench comparison to related benchmarks

Metric	Time (h)	Task #	Eval method	Multi-lang.	Science-focused	Avg lines
DA-Code	0.1	500	Task-specific verifier	✓	✗	85
DSBench	17	540	Task-specific verifier	✗	✗	75
MLE Bench	2.5	75	Reward metric	✗	✗	-
REBench	8	7	Reward metric	✓	✗	650
BLADE	1.5	188	MCQ	✗	✓	75
ScienceAgentBench	2.5	102	Task-specific verifier	✗	✓	58
BixBench (ours)	4.2	296	Open-ended	✓	✓	106

a final approved set of **53 analytical scenarios**.

3.1.3 TASK GENERATION

To generate the MCQ-based tasks comprising BixBench, we devised a multi-stage generation and review process. We employed an LLM, specifically Claude 3.5 Sonnet from Anthropic (20241022 release), to generate initial drafts of multiple-choice questions (MCQs) related to the analysis in the capsule. We provided a modified version of the code notebook, the hypothesis and result, and a specially crafted prompt to produce MCQs.

We generated question drafts for each capsule in two rounds of four questions each, for a total of eight MCQ drafts per capsule. The first question in each round was instructed in the prompt to be phrased around the hypothesis as a mechanism for capturing understanding of the hypothesis in a non-binary manner.

The second stage of task generation was human expert review. Our contract analysts were provided an interface through which they could view the model-generated questions and answer options. They were allowed full edit access for each, and were provided review instructions to Approve or Reject, with or without editing. They were also able to see and re-review previously reviewed questions to identify mistakes in the review process itself. To provide complete context for the questions, the reviewers were also provided access to the original capsules, including notebooks and data files.

Because we generated MCQs in two separate rounds, the final stage of task generation after expert review was duplicate filtering. The set of Approved questions for a capsule were provided to an LLM (Claude 3.5 Sonnet) which was prompted to flag any duplicated questions. This step was performed in triplicate and manually assessed for concordance across the three responses (we estimated approximately 95% concordance). Duplicates were manually verified and removed if appropriate, and the duplicate flag request to the LLM was repeated with the trimmed set of questions. This

was repeated until no questions were flagged as duplicates. The final task set consists of 296 questions covering 53 analytical capsules, each having three to eight associated questions (averaging 5.6 questions per capsule).

3.2 Benchmark Evaluation

3.2.1 OVERVIEW

BixBench was developed to serve as a benchmark for real-world bioinformatics capabilities. As such, we aimed to present realistic scenarios that a bioinformatician or computational biologist might face. Thus, at evaluation time, we present an agent with an empty Jupyter notebook, a set of input data files, and several relevant questions for the agent to answer. The agent thus has free reign to explore the data, make its own assumptions and plan its own analyses in order to answer the questions presented. This approach addresses key challenges in evaluating complex scientific tasks: the need for iterative analysis, access to specialized software environments, and open-answering. In practice, bioinformaticians do not have access to multiple-choice options and so open-answer is our preferred evaluation method. However, we are also able to perform multiple-choice evaluations with options to determine how that may influence performance in answering the questions.

3.2.2 AGENT INFRASTRUCTURE

We implemented our agent framework using Aviary, an extensible gymnasium for language agents (Narayanan et al., 2024). Aviary was selected for its ability to provide a controlled environment supporting tool use and multi-step reasoning. The framework enables reproducible evaluation by maintaining consistent software environments and tool access across different model evaluations. The agentic prompting approach we used is based on ReAct (Yao et al., 2023) due to its demonstrated effectiveness at tasks requiring balancing reasoning and execution.

To ensure reproducibility and isolate model evaluation from environment setup complexity, all analy-

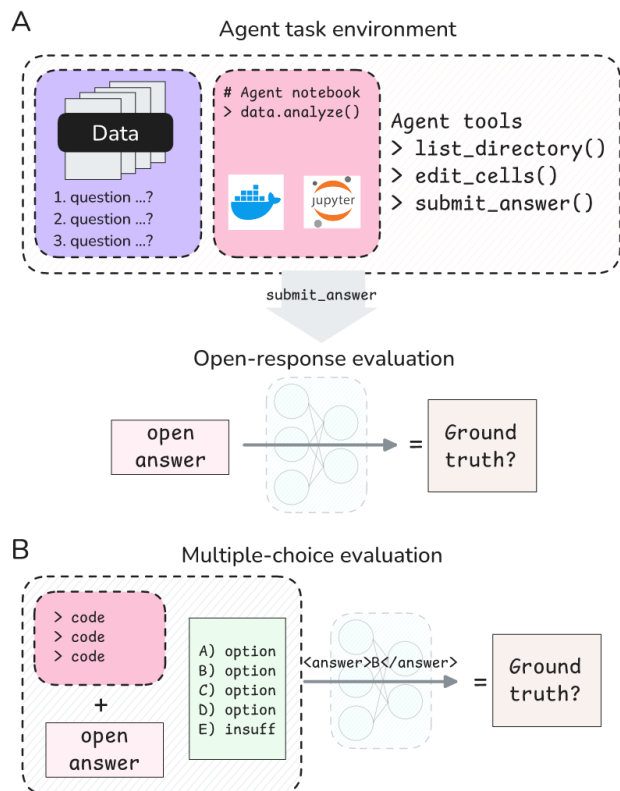


Figure 3: BixBench evaluation. An agent is provided a task capsule, consisting of the capsule data and the associated questions. The agent environment contains an empty code notebook in a docker environment pre-loaded with many popular bioinformatics packages and has three tools to manipulate the environment. In the open-reponse regime (A) the agent’s submitted answer is compared directly to the ground truth answer by a separate LLM call. In the MCQ regime (B) a second LLM is given the agent notebook and answer, and asked to choose from the available options.

ses are executed within a pre-built Docker container (BixBench-env:v1.0). This container includes extensive Python, R, and Bash packages commonly used in bioinformatics workflows. This ensures that evaluation remains focused on problem-solving capabilities rather than on software installation or dependency resolution—mimicking real-world bioinformatics workflows, where researchers often operate within standardized, pre-configured environments. Agents must still identify and load packages necessary for the analyses they want to carry out.

3.2.3 NOTEBOOK ANALYSIS WITH AGENTIC EXECUTION

Given that bioinformatics tasks involve complex computations requiring specialized software environments, we implement an **agentic scaffold** within a controlled execution

environment. This scaffold allows the LLM to iteratively refine its analysis in a structured manner. We leverage the Aviary framework to provide the agent with three distinct tools:

- `edit_cell`: Allows the agent to select, modify and execute a cell within a Jupyter Notebook.
- `list_workdir`: Enables the model to inspect its own workspace directory recursively to ascertain the structure and format of the dataset.
- `submit_answer`: Used by the model to finalize and submit an open-answer.

At each timestep, the model must decide which tool to invoke. The agent then receives observations from the tool’s execution. Each code modification triggers a full rerun of the notebook, allowing the model to view results in tabular or plot formats as well as debug tracebacks. Prompt engineering was explored extensively as a means to improve performance. An example of the final prompt used to initiate an agent trajectory is shared in the [Appendix A](#).

3.2.4 EVALUATION

Open-answer To more fully mimic evaluation of a real-world scenario of an autonomous analysis agent, we are primarily interested in open-answer responses to the prompted questions rather than the more popular MCQ paradigm. During the agent analysis trajectory, it is able to call the `submit_answer` tool when it is ready to answer the questions, which triggers the end of the trajectory. The final submitted answer is then automatically evaluated by a judge LLM (Claude 3.5 Sonnet) by comparing the agent-generated response against a ground-truth solution, with correctness assigned as a binary score (1 if correct, 0 otherwise). To account for stochastic trajectories, we run each analysis in parallel 10 times to calculate overall performance.

Multiple-choice Due to the challenging nature of the open-answer benchmark, we also provide the option to evaluate agents’ performance via MCQ. We believe this will serve as a useful proxy evaluation on the journey to building a truly autonomous bioinformatician. Specifically, after the analyses are complete, we can pass the complete analysis notebooks to an LLM, along with the initial question now provided with the multiple-choice answer options, and the open-answer response from the initial agent run. This second LLM is prompted to choose the most appropriate answer option. Importantly, we also provide an “opt-out” for the model in answering the MCQ by providing an “Insufficient information” refusal option. To evaluate performance on the MCQs, we further employ **majority voting** over the 10 runs to derive a majority-based consensus across the iterations.

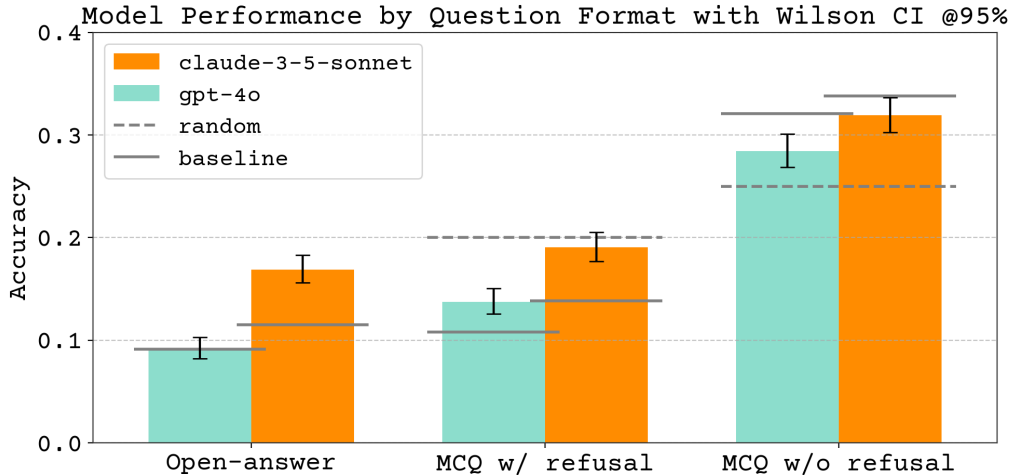


Figure 4: Overall model performance. Models perform poorly in the purely open-answer regime (left bars) and display increasing performance in the MCQ with refusal regime and further increased performance with MCQs without a refusal option. Performance does not surpass a baseline assessed as performance on the questions given to the model without access to any analysis notebook to base answers in (i.e. pure model recall.)

3.2.5 MODEL SELECTION AND EVALUATION

We have restricted our current evaluation to GPT-4o and Claude due to their demonstrated capabilities in structured output generation and handling of extended contexts, prerequisites for complex analytical workflows. During our preliminary tests we found that reasoning models such as o1 and DeepSeek R1 struggled to perform such tasks due to the long contexts and the structured outputs required for agentic tool use.

4 Results

4.1 Benchmark assembly and description

The BixBench dataset comprises **53 real-world analytical scenarios** or “capsules” with **296 associated open-answer questions**. The capsules are each composed of input data files and a set of guiding open-answer questions only answerable upon completion of a valid analysis (Figure 3). In Table 1 we compare BixBench to other similar benchmarks. BixBench differentiates itself as the first benchmark focused on open-ended scientific data analysis, evaluated on long and complex multi-step trajectories.

Capsules contain data files in various formats, type (e.g. csv, rds) and directory structure requiring effective workspace navigation and appropriate use of libraries to load and process heterogeneous data. While frontier models have demonstrated proficiency in code writing given user-provided prompts, they do not typically come pre-built with configurable code interpreters. Thus, to assess the capabilities

of current frontier models to perform open-ended analytical tasks in the BixBench benchmark, we open-source a configurable agentic notebook environment based on Aviary, a public agent scaffold (Narayanan et al., 2024). Using our framework, we evaluate SOTA performance on BixBench with two frontier AI models, GPT-4o from OpenAI and Claude 3.5 Sonnet from Anthropic.

We also assess the impact of scaling inference-time compute, by performing 10 parallel iterations of each capsule resulting in the collection of 530 trajectories across two different models (4o and Claude) and two different modalities (with and without images) resulting in 2,120 total trajectories. Performance on a question is primarily calculated as accuracy, the fraction of correct answers given across all parallel runs over all questions provided. In the MCQ case, we also conduct majority voting across the 10 iterations and report precision (questions answered correctly over all questions answered, i.e. those that were not answered with the “Refusal” option.) The evaluation scheme is diagrammed in Figure 3. Finally, we conduct ablations on the presence of the refusal option and the use of images/plots in the notebooks.

Results over all questions and trajectories are presented in Figure 4 (open-answer, left-most bars, MCQ, right two bar pairs.) Overall, BixBench proved very challenging for existing frontier models, even when provided with a custom agent scaffold and associated tools. Claude 3.5 Sonnet bested GPT-4o with 17% accuracy vs. 9% in the open-answer evaluation regime. As a further calibration, we measure the pure recall performance of both models by

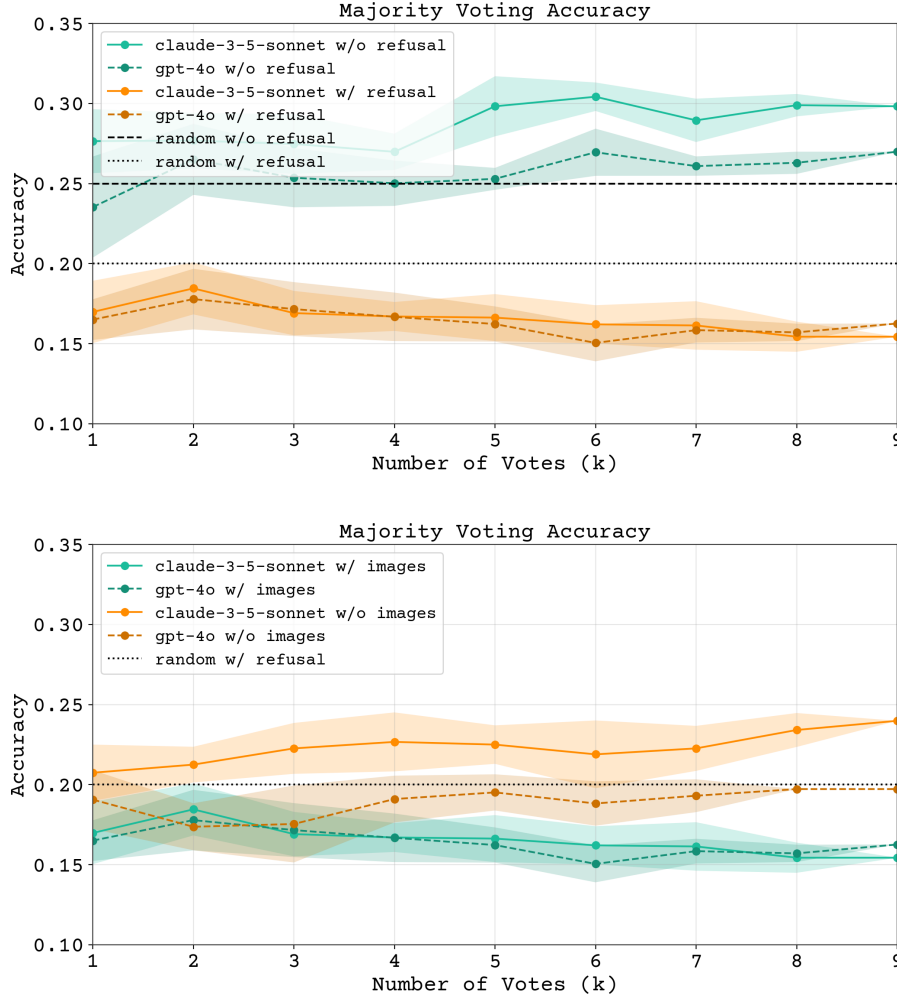


Figure 5: Performance of models in the MCQ regime relative to number of votes with ablations of refusal (top plot) and image generation (bottom.) Note that in the vision ablation the refusal option is present and the agents are prompt engineered to not produce images/plots. In the refusal ablation agents are free to use images/plots.

asking the BixBench questions without any notebook or other context, as indicated by the solid gray line in Figure 4. In the MCQ regime, overall performance between models was closer than open-answer (Figure 4,) though Claude 3.5 Sonnet remained the clear higher performer between the two models.

The questions contained in BixBench are intended by design not to be answerable by model recall, i.e. they should require completion of some analysis to answer. During the MCQ answering process we prompt engineered the agent to answer only using the analysis notebook. To investigate the “ethics” of the models in this context (i.e. their willingness to answer even if they don’t like the options), we performed the MCQ evaluation in two scenarios: *with refusal*, where the agent is provided an “insufficient information” answer as

one of the options, and *without refusal*, where no such option is provided. Overall, when provided the refusal option, both models performed worse than random (Figure 4, middle bars) indicating their tendency to opt-out of answering when given the option in the context of a complex analysis. When only given real answer options without a refusal option, performance is higher but only marginally above random (Figure 4.) When further evaluating this effect in a majority voting scenario across the 10 trajectories for each question, we can see a slight decrease in performance as vote count goes up (Figure 5, top, orange trajectories) suggesting an accumulation of refusal choices. This contrasts with the no-refusal scenario, which sees a slight uptick in performance as vote count increases (Figure 5, top, green trajectories).

Interestingly, we noted during question generation and other

testing that the models seemed to do a poor job of interpreting plots in the notebooks generated both by humans and the models themselves. We reasoned that even the multimodal models used here may be able to reason more effectively over the underlying text or data rather than plots themselves, and thus performed an ablation where the agent was prompted not to generate images during its analysis. Consistent with our earlier observation, forcing the agent to avoid generating plots results in higher performance for both underlying models, though this effect is more pronounced for Claude 3.5 Sonnet (Figure 5, bottom).

5 Discussion

Biological data analysis and associated domains offer a clear opportunity to build the first open-ended agents for scientific discovery, and thus assessing the state of progress is critical. We’ve introduced here BixBench, a benchmark and evaluation framework for measuring agent performance on real-world bioinformatics tasks. BixBench presents 53 real-world analytical scenarios in bioinformatics, along with a set of 296 associated questions, supporting both open-answer or multiple-choice evaluation regimes.

In measuring the performance of two current frontier models on our benchmark, we find that there is significant development necessary to achieve the goal of an autonomous bioinformatician. When assessed in an open-answer environment, the models perform at just 17% overall accuracy in answering questions related to the analyses. When using multiple-choice versions of the same questions to interpret answers from the analyses, overall performance increases, though not substantially and only marginally above random (Figure 3). Of note, when removing the option to abstain from answering the performance again increases, which we speculate is due in large part to the model relying on answering via recall rather than the information contained in the analysis.

Ultimately, BixBench highlights a fundamental shortcoming in current LLM capabilities for conducting rigorous bioinformatics research, and will thus be a valuable resource to guide development of these capabilities and accelerate discovery across biology.

5.1 Future Work

5.1.1 FURTHER SAMPLING OF THE FIELD

While we consider BixBench to be a highly representative benchmark, building a truly comprehensive benchmark of tasks for AI models in bioinformatics would be extraordinarily difficult, and we acknowledge that there are many important workflows, pipelines, statistical approaches, data types, and other parameters that are missing from BixBench. These will be important additional improvements as related

benchmarks and agents continue to be developed.

5.1.2 HUMAN BASELINE COMPARISON

An important extension of the current work would be the additional of a human baseline comparison, where human bioinformatics experts are tasked with performing the same analyses and answering the same questions as the agents here. As the seed capsules used to generate BixBench tasks were generated de novo by such human experts, we anticipate that additional human experts would perform significantly higher than the agent performance reported here, and thus did not prioritize gathering this data. We are currently exploring options for doing so.

5.1.3 EVALUATING ADDITIONAL MODELS

Reasoning models, introduced recently in (OpenAI et al., 2024; DeepSeek-AI et al., 2025) and others have rapidly topped many benchmarks, including MLE-bench. It is still early to incorporate them into environments, with few details or results on tool calling. They hold promise for future work due to their ability to improve performance from binary reward signals, like the capsules we introduce here (Team, 2025; Lambert et al., 2025; DeepSeek-AI et al., 2025).

Impact Statement

We believe that the primary effect of automation within bioinformatics and computational biology will be a faster pace in biological discovery, ultimately to the great benefit of patients in need for treatments. The worldwide effort to develop vaccines and treatments for the COVID-19 pandemic highlighted the importance of rapid, targeted discovery using only small amounts of data.

As much as bioinformatics automation could be used to discover good things, it could also be used to discover bad things (e.g. chemical/biological weapons). However, as substantial resources and infrastructure are required to produce these, we do not expect this benchmark by itself to play a role in enabling new parties to begin their development.

Finally, it’s worth highlighting the impact of increased automation on the workforce and employment rates within certain professions. While recent AI advancements have been impressive on tasks previously thought to be the monopoly of humans, bioinformatics and computational roles are still among the most sought-after skills in the industry. As with other technological leaps in the past, we believe that AI will create new professions that were not previously existent and lift human work onto more rewarding, interesting duties.

References

Anthropic. Introducing the next generation of claude, March

2024. URL <https://www.anthropic.com/news/claude-3-5-sonnet>. Accessed: 2024-07-11.
- Chan, J. S., Chowdhury, N., Jaffe, O., Aung, J., Sherburn, D., Mays, E., Starace, G., Liu, K., Maksin, L., Patwardhan, T., et al. Mle-bench: Evaluating machine learning agents on machine learning engineering. *arXiv preprint arXiv:2410.07095*, 2024.
- Chang, Y., Wang, X., Wang, J., Wu, Y., Yang, L., Zhu, K., Chen, H., Yi, X., Wang, C., Wang, Y., Ye, W., Zhang, Y., Chang, Y., Yu, P. S., Yang, Q., and Xie, X. A survey on evaluation of large language models, 2023. URL <https://arxiv.org/abs/2307.03109>.
- DeepSeek-AI, Guo, D., Yang, D., Zhang, H., Song, J., Zhang, R., Xu, R., Zhu, Q., Ma, S., Wang, P., Bi, X., Zhang, X., Yu, X., Wu, Y., Wu, Z. F., Gou, Z., Shao, Z., Li, Z., Gao, Z., Liu, A., Xue, B., Wang, B., Wu, B., Feng, B., Lu, C., Zhao, C., Deng, C., Zhang, C., Ruan, C., Dai, D., Chen, D., Ji, D., Li, E., Lin, F., Dai, F., Luo, F., Hao, G., Chen, G., Li, G., Zhang, H., Bao, H., Xu, H., Wang, H., Ding, H., Xin, H., Gao, H., Qu, H., Li, H., Guo, J., Li, J., Wang, J., Chen, J., Yuan, J., Qiu, J., Li, J., Cai, J. L., Ni, J., Liang, J., Chen, J., Dong, K., Hu, K., Gao, K., Guan, K., Huang, K., Yu, K., Wang, L., Zhang, L., Zhao, L., Wang, L., Zhang, L., Xu, L., Xia, L., Zhang, M., Zhang, M., Tang, M., Li, M., Wang, M., Li, M., Tian, N., Huang, P., Zhang, P., Wang, Q., Chen, Q., Du, Q., Ge, R., Zhang, R., Pan, R., Wang, R., Chen, R. J., Jin, R. L., Chen, R., Lu, S., Zhou, S., Chen, S., Ye, S., Wang, S., Yu, S., Zhou, S., Pan, S., Li, S. S., Zhou, S., Wu, S., Ye, S., Yun, T., Pei, T., Sun, T., Wang, T., Zeng, W., Zhao, W., Liu, W., Liang, W., Gao, W., Yu, W., Zhang, W., Xiao, W. L., An, W., Liu, X., Wang, X., Chen, X., Nie, X., Cheng, X., Liu, X., Xie, X., Liu, X., Yang, X., Li, X., Su, X., Lin, X., Li, X. Q., Jin, X., Shen, X., Chen, X., Sun, X., Wang, X., Song, X., Zhou, X., Wang, X., Shan, X., Li, Y. K., Wang, Y. Q., Wei, Y. X., Zhang, Y., Xu, Y., Li, Y., Zhao, Y., Sun, Y., Wang, Y., Yu, Y., Zhang, Y., Shi, Y., Xiong, Y., He, Y., Piao, Y., Wang, Y., Tan, Y., Ma, Y., Liu, Y., Guo, Y., Ou, Y., Wang, Y., Gong, Y., Zou, Y., He, Y., Xiong, Y., Luo, Y., You, Y., Liu, Y., Zhou, Y., Zhu, Y. X., Xu, Y., Huang, Y., Li, Y., Zheng, Y., Zhu, Y., Ma, Y., Tang, Y., Zha, Y., Yan, Y., Ren, Z. Z., Ren, Z., Sha, Z., Fu, Z., Xu, Z., Xie, Z., Zhang, Z., Hao, Z., Ma, Z., Yan, Z., Wu, Z., Gu, Z., Zhu, Z., Liu, Z., Li, Z., Xie, Z., Song, Z., Pan, Z., Huang, Z., Xu, Z., Zhang, Z., and Zhang, Z. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning, 2025. URL <https://arxiv.org/abs/2501.12948>.
- Deshpande, D., Chhugani, K., Ramesh, T., Pellegrini, M., Shiffman, S., Abedalthagafi, M. S., Alqahtani, S., Ye, J., Liu, X. S., Leek, J. T., et al. The evolution of computational research in a data-centric world. *Cell*, 187(17): 4449–4457, 2024.
- Gao, S., Fang, A., Huang, Y., Giunchiglia, V., Noori, A., Schwarz, J. R., Ektefaie, Y., Kondic, J., and Zitnik, M. Empowering biomedical discovery with ai agents, 2024. URL <https://arxiv.org/abs/2404.02831>.
- Gu, K., Shang, R., Jiang, R., Kuang, K., Lin, R.-J., Lyu, D., Mao, Y., Pan, Y., Wu, T., Yu, J., et al. Blade: Benchmarking language model agents for data-driven science. *arXiv preprint arXiv:2408.09667*, 2024.
- Hendrycks, D., Burns, C., Basart, S., Zou, A., Mazeika, M., Song, D., and Steinhardt, J. Measuring massive multitask language understanding. *arXiv preprint arXiv:2009.03300*, 2020.
- Huang, K., Fu, T., Gao, W., Zhao, Y., Roohani, Y., Leskovec, J., Coley, C. W., Xiao, C., Sun, J., and Zitnik, M. Therapeutics data commons: Machine learning datasets and tasks for drug discovery and development. *arXiv preprint arXiv:2102.09548*, 2021.
- Huang, Y., Luo, J., Yu, Y., Zhang, Y., Lei, F., Wei, Y., He, S., Huang, L., Liu, X., Zhao, J., et al. Da-code: Agent data science code generation benchmark for large language models. *arXiv preprint arXiv:2410.07331*, 2024.
- Hurst, A., Lerer, A., Goucher, A. P., Perelman, A., Ramesh, A., Clark, A., Ostrow, A., Welihinda, A., Hayes, A., Radford, A., et al. Gpt-4o system card. *arXiv preprint arXiv:2410.21276*, 2024.
- Jing, L., Huang, Z., Wang, X., Yao, W., Yu, W., Ma, K., Zhang, H., Du, X., and Yu, D. Dsbench: How far are data science agents to becoming data science experts? *arXiv preprint arXiv:2409.07703*, 2024.
- Lambert, N., Morrison, J., Pyatkin, V., Huang, S., Ivison, H., Brahman, F., Miranda, L. J. V., Liu, A., Dziri, N., Lyu, S., Gu, Y., Malik, S., Graf, V., Hwang, J. D., Yang, J., Bras, R. L., Tafford, O., Wilhelm, C., Soldaini, L., Smith, N. A., Wang, Y., Dasigi, P., and Hajishirzi, H. Tulu 3: Pushing frontiers in open language model post-training, 2025. URL <https://arxiv.org/abs/2411.15124>.
- Laurent, J. M., Janizek, J. D., Ruza, M., Hinks, M. M., Hammerling, M. J., Narayanan, S., Ponnampati, M., White, A. D., and Rodrigues, S. G. Lab-bench: Measuring capabilities of language models for biology research. *arXiv preprint arXiv:2407.10362*, 2024.
- Lu, C., Lu, C., Lange, R. T., Foerster, J., Clune, J., and Ha, D. The ai scientist: Towards fully automated open-ended scientific discovery. *arXiv preprint arXiv:2408.06292*, 2024.

- Mirza, A., Alampara, N., Kunchapu, S., Ríos-García, M., Emoekabu, B., Krishnan, A., Gupta, T., Schilling-Wilhelmi, M., Okereke, M., Aneesh, A., Elahi, A. M., Asgari, M., Eberhardt, J., Elbeheiry, H. M., Gil, M. V., Greiner, M., Holick, C. T., Glaubitz, C., Hoffmann, T., Ibrahim, A., Klepsch, L. C., Köster, Y., Kreth, F. A., Meyer, J., Miret, S., Peschel, J. M., Ringleb, M., Roesner, N., Schreiber, J., Schubert, U. S., Stafast, L. M., Wonnanke, D., Pieler, M., Schwaller, P., and Jablonka, K. M. Are large language models superhuman chemists?, 2024. URL <https://arxiv.org/abs/2404.01475>.
- Narayanan, S., Braza, J. D., Griffiths, R.-R., Ponnampati, M., Bou, A., Laurent, J., Kabeli, O., Wellawatte, G., Cox, S., Rodrigues, S. G., et al. Aviary: training language agents on challenging scientific tasks. *arXiv preprint arXiv:2412.21154*, 2024.
- OpenAI, :, Jaech, A., Kalai, A., Lerer, A., Richardson, A., El-Kishky, A., Low, A., Helyar, A., Madry, A., Beutel, A., Carney, A., Iftimie, A., Karpenko, A., Passos, A. T., Neitz, A., Prokofiev, A., Wei, A., Tam, A., Bennett, A., Kumar, A., Saraiva, A., Vallone, A., Duberstein, A., Kondrich, A., Mishchenko, A., Applebaum, A., Jiang, A., Nair, A., Zoph, B., Ghorbani, B., Rossen, B., Sokolowsky, B., Barak, B., McGrew, B., Minaiev, B., Hao, B., Baker, B., Houghton, B., McKinzie, B., Eastman, B., Lugaresi, C., Bassin, C., Hudson, C., Li, C. M., de Bourcy, C., Voss, C., Shen, C., Zhang, C., Koch, C., Orsinger, C., Hesse, C., Fischer, C., Chan, C., Roberts, D., Kappler, D., Levy, D., Selsam, D., Dohan, D., Farhi, D., Mely, D., Robinson, D., Tsipras, D., Li, D., Oprica, D., Freeman, E., Zhang, E., Wong, E., Proehl, E., Cheung, E., Mitchell, E., Wallace, E., Ritter, E., Mays, E., Wang, F., Such, F. P., Raso, F., Leoni, F., Tsimpouras, F., Song, F., von Lohmann, F., Sulit, F., Salmon, G., Parascandolo, G., Chabot, G., Zhao, G., Brockman, G., Leclerc, G., Salman, H., Bao, H., Sheng, H., Andrin, H., Bagherinezhad, H., Ren, H., Lightman, H., Chung, H. W., Kivlichan, I., O’Connell, I., Osband, I., Gilaberte, I. C., Akkaya, I., Kostrikov, I., Sutskever, I., Kofman, I., Pachocki, J., Lennon, J., Wei, J., Harb, J., Twore, J., Feng, J., Yu, J., Weng, J., Tang, J., Yu, J., Candela, J. Q., Palermo, J., Parish, J., Heidecke, J., Hallman, J., Rizzo, J., Gordon, J., Uesato, J., Ward, J., Huizinga, J., Wang, J., Chen, K., Xiao, K., Singhal, K., Nguyen, K., Cobbe, K., Shi, K., Wood, K., Rimbach, K., Gu-Lemberg, K., Liu, K., Lu, K., Stone, K., Yu, K., Ahmad, L., Yang, L., Liu, L., Maksin, L., Ho, L., Fedus, L., Weng, L., Li, L., McCallum, L., Held, L., Kuhn, L., Kondraciuk, L., Kaiser, L., Metz, L., Boyd, M., Trebacz, M., Joglekar, M., Chen, M., Tintor, M., Meyer, M., Jones, M., Kaufer, M., Schwarzer, M., Shah, M., Yatbaz, M., Guan, M. Y., Xu, M., Yan, M., Glaese, M., Chen, M., Lampe, M., Malek, M., Wang, M., Fradin, M., McClay, M., Pavlov, M., Wang, M., Wang, M., Murati, M., Bavarian, M., Rohaninejad, M., McAleese, N., Chowdhury, N., Chowdhury, N., Ryder, N., Tezak, N., Brown, N., Nachum, O., Boiko, O., Murk, O., Watkins, O., Chao, P., Ashbourne, P., Izmailov, P., Zhokhov, P., Dias, R., Arora, R., Lin, R., Lopes, R. G., Gaon, R., Miyara, R., Leike, R., Hwang, R., Garg, R., Brown, R., James, R., Shu, R., Cheu, R., Greene, R., Jain, S., Altman, S., Toizer, S., Toyer, S., Miserendino, S., Agarwal, S., Hernandez, S., Baker, S., McKinney, S., Yan, S., Zhao, S., Hu, S., Santurkar, S., Chaudhuri, S. R., Zhang, S., Fu, S., Papay, S., Lin, S., Balaji, S., Sanjeev, S., Sidor, S., Broda, T., Clark, A., Wang, T., Gordon, T., Sanders, T., Patwardhan, T., Sottiaux, T., Degry, T., Dimson, T., Zheng, T., Garipov, T., Stasi, T., Bansal, T., Creech, T., Peterson, T., Eloundou, T., Qi, V., Kosaraju, V., Monaco, V., Pong, V., Fomenko, V., Zheng, W., Zhou, W., McCabe, W., Zaremba, W., Dubois, Y., Lu, Y., Chen, Y., Cha, Y., Bai, Y., He, Y., Zhang, Y., Wang, Y., Shao, Z., and Li, Z. Openai o1 system card, 2024. URL <https://arxiv.org/abs/2412.16720>.
- Quigley, I. K., Blevins, A., Halverson, B. J., and Wilkinson, N. Belka: The big encoded library for chemical assessment. In *NeurIPS 2024 Competition Track*, 2024.
- Schmidgall, S., Su, Y., Wang, Z., Sun, X., Wu, J., Yu, X., Liu, J., Liu, Z., and Barsoum, E. Agent laboratory: Using llm agents as research assistants. *arXiv preprint arXiv:2501.04227*, 2025.
- Skarlinski, M. D., Cox, S., Laurent, J. M., Braza, J. D., Hinks, M., Hammerling, M. J., Ponnampati, M., Rodrigues, S. G., and White, A. D. Language agents achieve superhuman synthesis of scientific knowledge. *arXiv preprint arXiv:2409.13740*, 2024.
- Swanson, K., Wu, W., Bulaong, N. L., Pak, J. E., and Zou, J. The virtual lab: Ai agents design new sars-cov-2 nanobodies with experimental validation. *bioRxiv*, pp. 2024–11, 2024.
- Tang, X., Qian, B., Gao, R., Chen, J., Chen, X., and Gerstein, M. Biocoder: A benchmark for bioinformatics code generation with large language models, 2024. URL <https://arxiv.org/abs/2308.16458>.
- Team, N. Sky-t1: Train your own o1 preview model within \$450. <https://novasky-ai.github.io/posts/sky-t1>, 2025. Accessed: 2025-01-09.
- Wijk, H., Lin, T., Becker, J., Jawhar, S., Parikh, N., Broadley, T., Chan, L., Chen, M., Clymer, J., Dhyani, J., et al. Rebench: Evaluating frontier ai r&d capabilities of language model agents against human experts. *arXiv preprint arXiv:2411.15114*, 2024.
- Wognum, C., Ash, J. R., Aldeghi, M., Rodríguez-Pérez, R., Fang, C., Cheng, A. C., Price, D. J., Clevert, D.-A.,

Engkvist, O., and Walters, W. P. A call for an industry-led initiative to critically assess machine learning for real-world drug discovery. *Nature Machine Intelligence*, pp. 1–2, 2024.

Yao, S., Zhao, J., Yu, D., Du, N., Shafran, I., Narasimhan, K., and Cao, Y. React: Synergizing reasoning and acting in language models, 2023. URL <https://arxiv.org/abs/2210.03629>.

Zhou, B., Zhao, H., Puig, X., Fidler, S., Barriuso, A., and Torralba, A. Scene parsing through ade20k dataset. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 633–641, 2017.

A Appendix

In the appendix we include a number of additional figures and plot displaying other metrics and analyses performance on the benchmark and evaluation results.

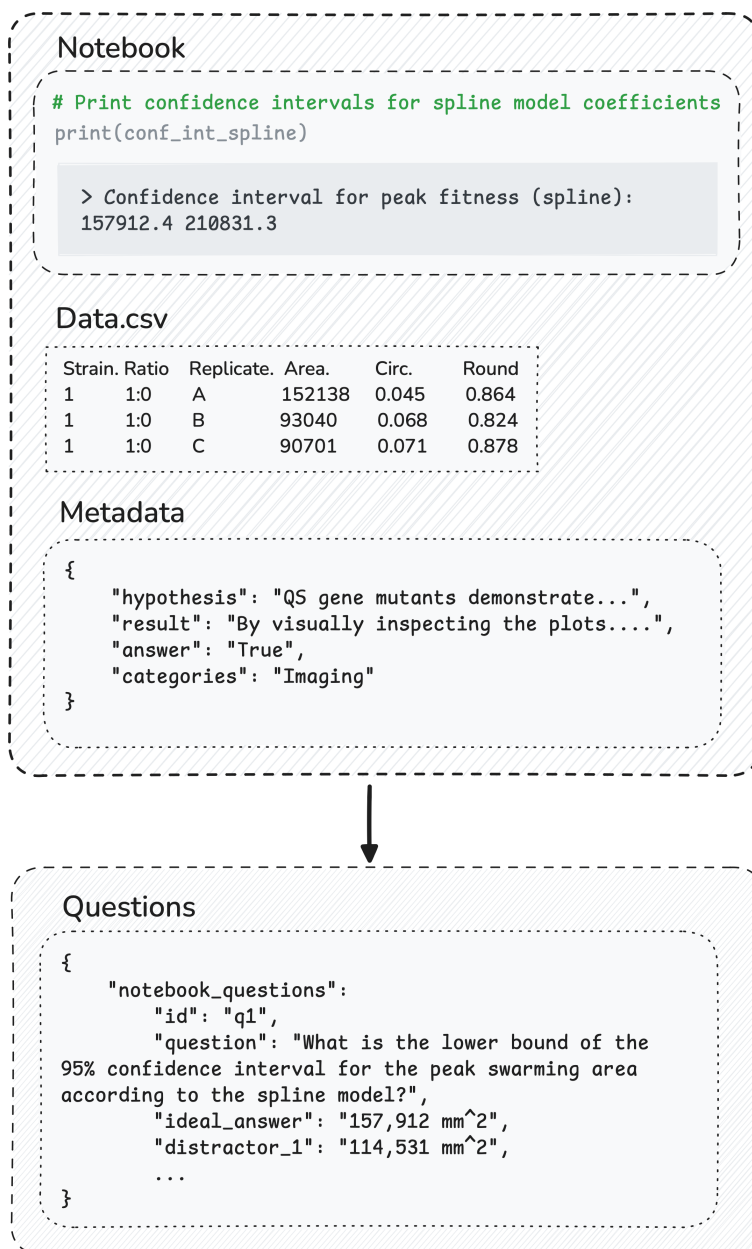


Figure 6: An example of an actual analysis capsule within BixBench and associated question tasks. The primary capsule consists of a notebook of analysis code generated by an expert analyst, any input data files for the analysis, and some metadata including a hypothesis, result, and answer. The Questions are initially generated by an LLM, and reviewed and edited by expert humans. See Methods for a complete description.

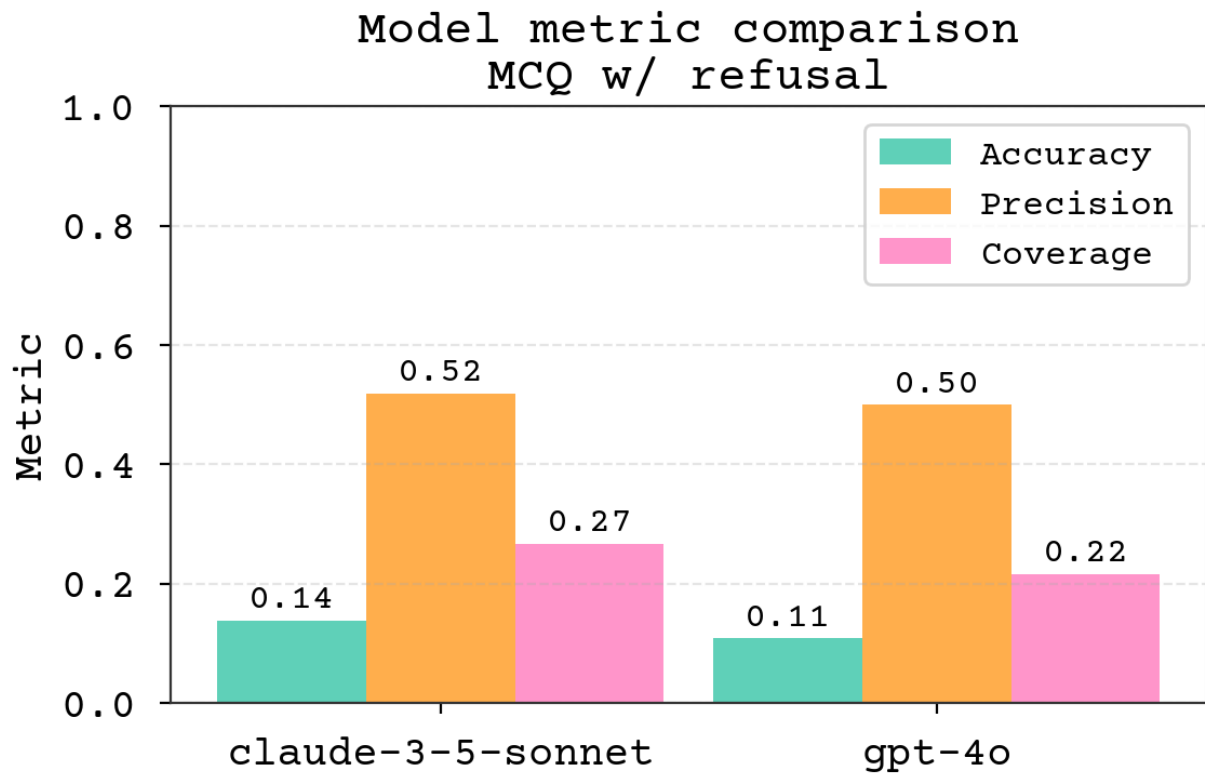


Figure 7: Accuracy, precision and recall across baseline model runs.

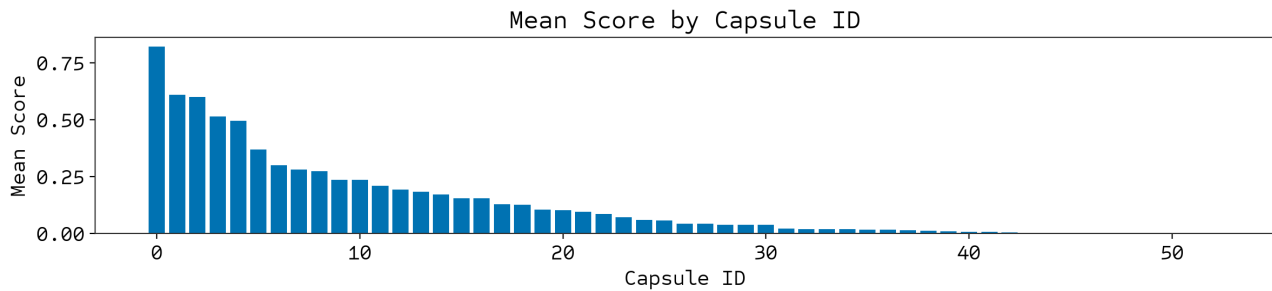


Figure 8: Average rate of correct answer (accuracy) per capsule across replicates.

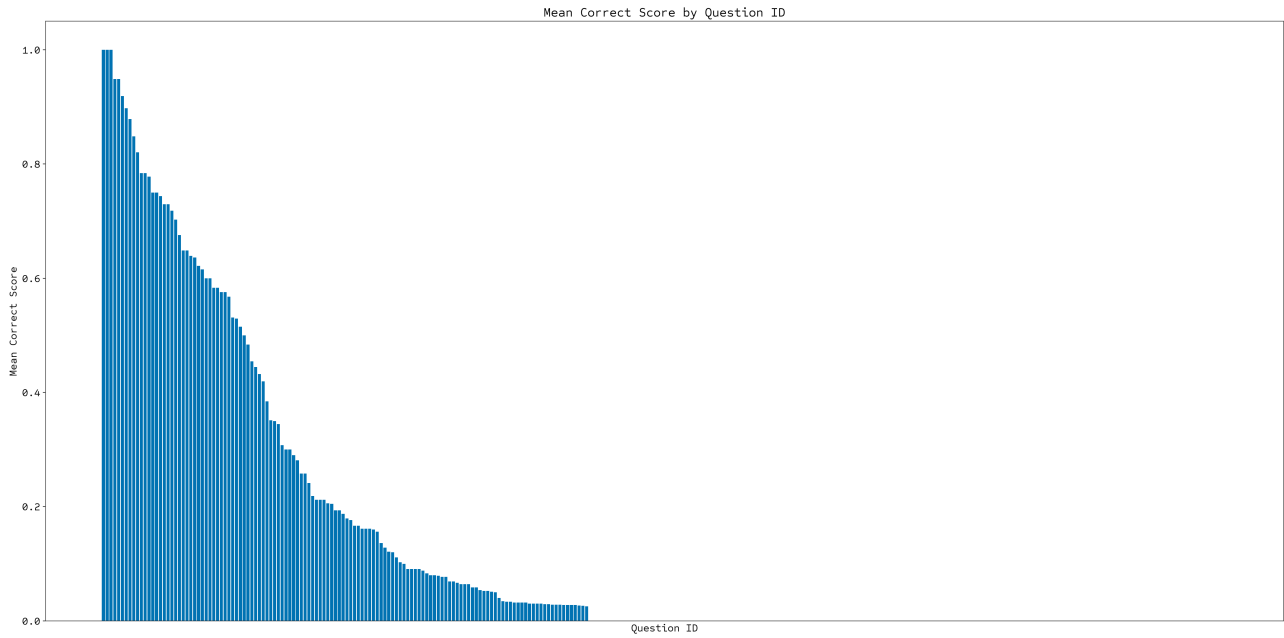


Figure 9: Average rate of correct answer (accuracy) per individual question across replicates.

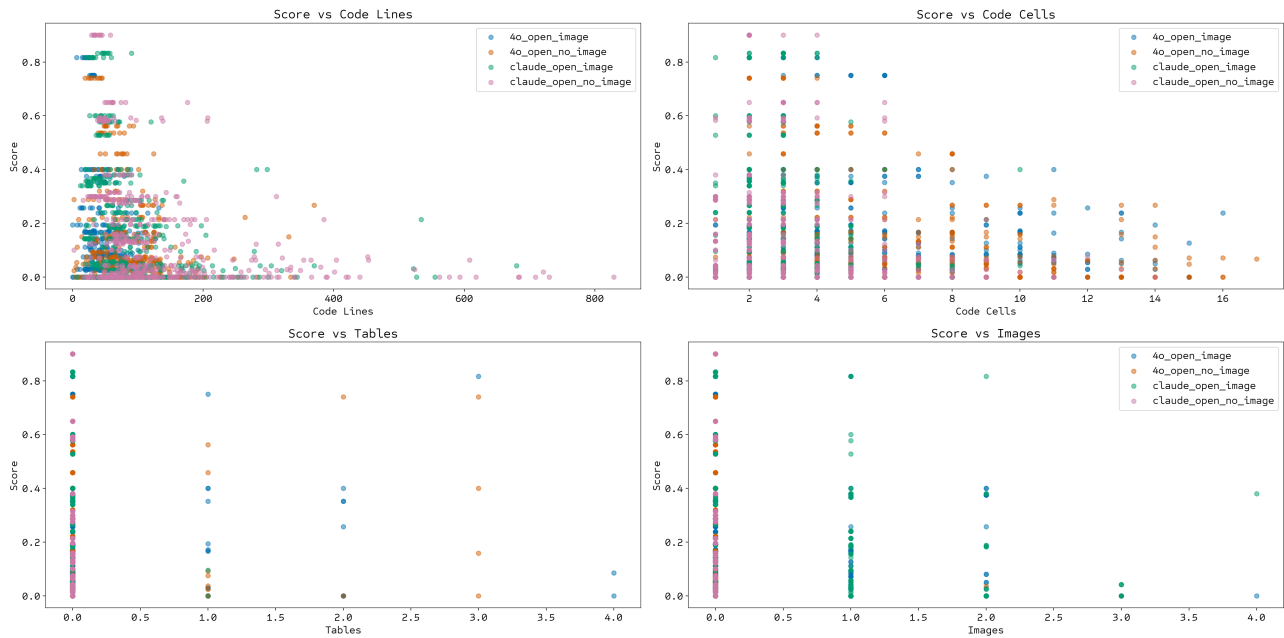


Figure 10: Distributions of performance vs. (top left) number of code lines, (top right) number of code cells, (bottom left) number of tables produced, and (bottom right) number of images produced for each model split according to being allowed to produce images or not.

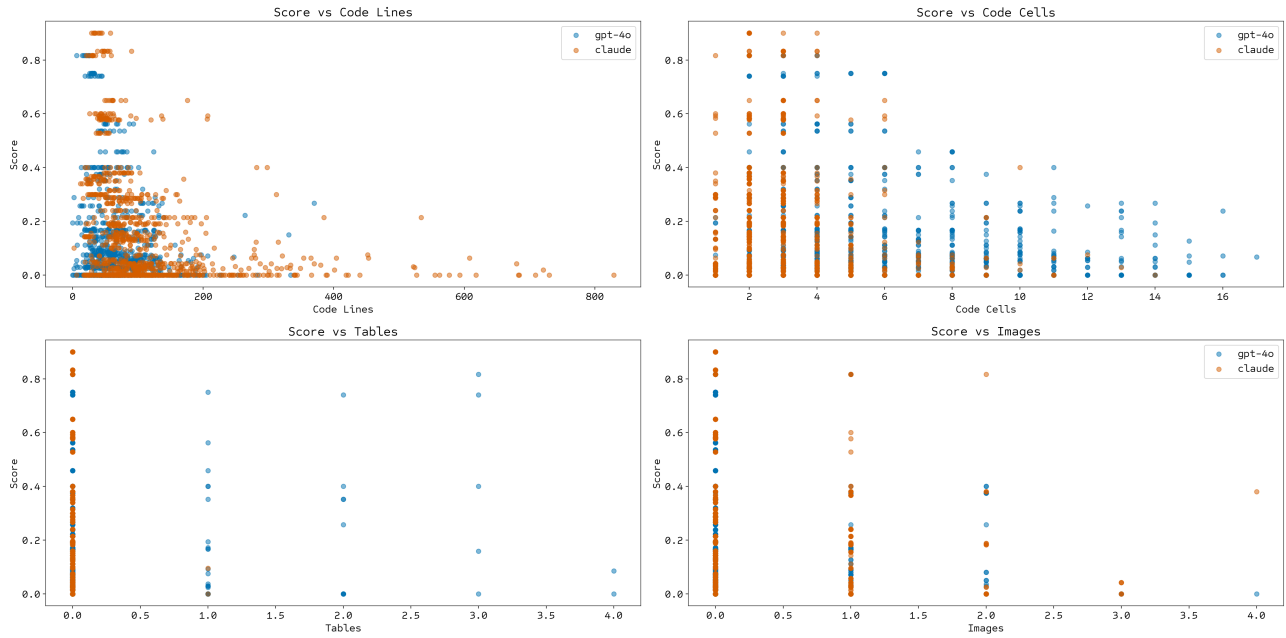


Figure 11: Distributions of performance vs. (top left) number of code lines, (top right) number of code cells, (bottom left) number of tables produced, and (bottom right) number of images produced for each model.

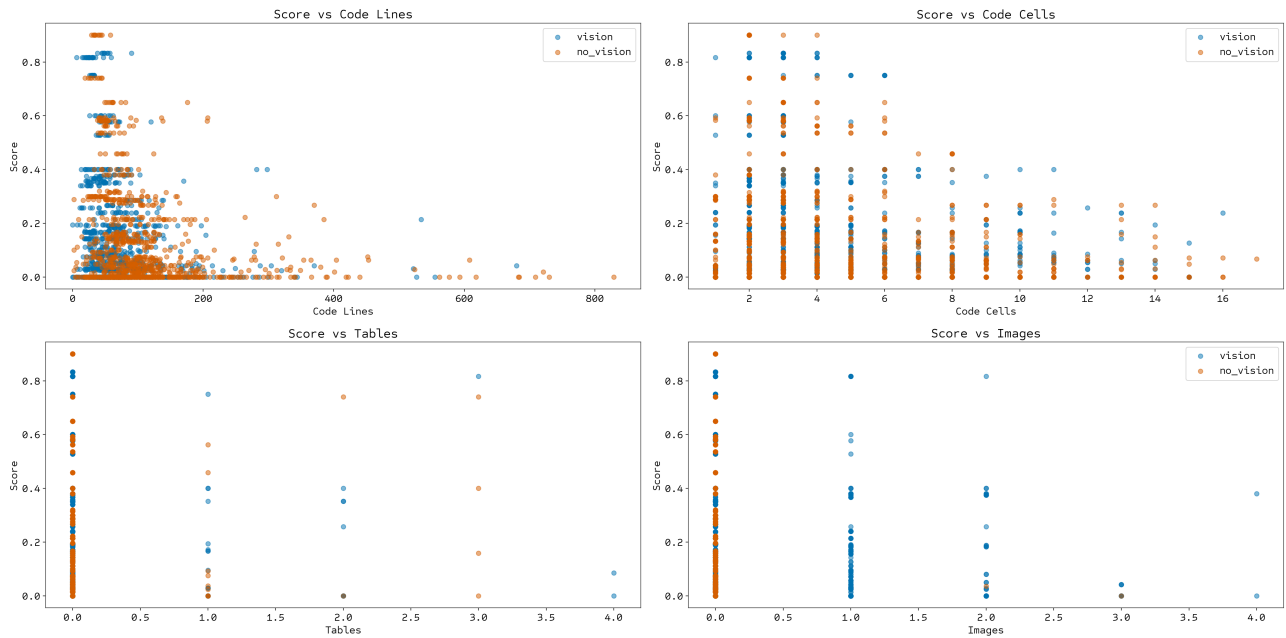


Figure 12: Distributions of performance vs. (top left) number of code lines, (top right) number of code cells, (bottom left) number of tables produced, and (bottom right) number of images produced for combined models being allowed to produce images or not.

Example prompt used to initialise an agent trajectory

You are an expert bioinformatician and seasoned biological data scientist. Your task is to create a comprehensive Jupyter notebook named 'notebook.ipynb' that analyzes data to answer a series of questions.

The notebook should contain all necessary artifacts (plots, tables, print outputs) to fully answer these questions, structured in a way that another person could use to derive the answers.

Here are the questions you need to address:

<questions>

q1: What percentage of genes differentially expressed in strain 97 are also differentially expressed in strain 99?

q2: How many genes are uniquely differentially expressed in strain 98 but not in either of strains 97 or 99?

q3: How many genes have a dispersion value below $1e-05$ after DEseq analysis?

q4: Which strain and media condition do not cluster together based on pairwise correlation after regularized log transformation of the differential expression analysis results?

</questions>

Example prompt used to initialise an agent trajectory continued 1

Follow these steps to create your notebook, using chain-of-thought reasoning at each stage:

1. List Directory Contents:

<analysis_planning>

- Consider how to use the list_workdir tool to recursively list the directory contents.
- Think about how to organize and present this information clearly in the notebook.
- List potential challenges in interpreting the directory structure.
- Consider how the directory structure might inform your approach to the analysis.

</analysis_planning>

Place the output of the list_workdir tool inside <directory_contents> tags.

2. Load Data and Perform Descriptive Statistics:

<analysis_planning>

- Identify which data files are most relevant to answering the questions. List these files.
- Plan how to load these files efficiently in R or Python.
- List the specific descriptive statistics you plan to use (e.g., summary(), str(), head()).
- Consider potential issues like missing data or unexpected formats. How will you handle each?
- Plan how to present this information clearly in the notebook.
- Write down key statistics you expect to see and how you'll interpret them.
- Consider potential data quality issues and how you'll address them.

</analysis_planning>

Execute your plan to load data and perform descriptive statistics.

3. Develop Analysis Plan:

<analysis_planning>

- Break down each question into testable components. List these components.
- For each component, list appropriate statistical tests or visualizations.
- Consider alternative approaches for each component and justify your choices.
- Identify potential confounding factors and how to address them.
- Plan the sequence of your analysis steps, explaining the rationale for each.
- Consider how this analysis plan will be documented in the notebook.
- List potential statistical assumptions for your chosen methods and how you'll test them.
- Think about how your analysis plan addresses each of the original questions.

</analysis_planning>

Write out your analysis plan as comments in the notebook.

4. Execute Analysis Plan:

<analysis_planning>

- For each step in your analysis plan, list the R, Python or bash functions and libraries you'll use.
- Think about how to structure your code for readability and efficiency.
- Plan how to document your code with clear comments.
- Consider how to present results clearly, using tables or visualizations where appropriate.
- Ensure that all outputs are clearly labeled and explained in the context of the questions.
- Plan how you'll interpret each result in relation to the original questions.
- Consider potential unexpected results and how you'll handle them.

</analysis_planning>

Execute your analysis plan, creating new cells as needed.

5. Conclude and Submit Answer:

<thought_process>

- Reflect on how your results relate to each question.
- List any limitations or uncertainties in your analysis.
- Plan a concise summary of your findings for each question.
- Think about how to phrase your conclusions as clear statements.
- Ensure that the notebook contains all necessary information for another model to derive these answers.
- Consider any additional insights or patterns you've noticed during the analysis.
- Think about potential follow-up questions or areas for further investigation.

</thought_process>

Use the submit_answer tool to submit your final answer as a dictionary with keys as the question number and a short answer¹⁷ to the question.

If the question asks for a number, be precise to 2 decimal places.

Example prompt used to initialise an agent trajectory continued 2

General Guidelines:

- Write small to medium-sized cells for easier debugging.
 - Edit existing cells by their index number when fixing bugs, rather than creating new ones.
 - Check dataframe shapes before printing. Use head() for large dataframes.
 - Ensure each cell executes successfully before moving to the next.
 - Assume you already have the packages you need installed and only install new ones if you receive errors.
 - If you need to install packages, use mamba or conda.
- IMPORTANT: R vs Python vs bash
- You can use either Python, R or bash cells to complete the analysis.
 - All cells are by default Python cells. However, you can use both bash and R cells by adding %%bash or %%R to the first line of the cell.
 - The first cell has already been loaded with %load_ext rpy2.ipynb so you can use %%R cells from the second cell onwards
 - AVOID USING PLOTS. USE TABLES AND PRINT OUTPUTS INSTEAD AS MUCH AS POSSIBLE.

R-Specific Guidelines:

1. Load packages using this format to minimize verbose output:

```
```r
if (!requireNamespace("package_name", quietly = TRUE)) {
 install.packages("package_name")
}
suppressPackageStartupMessages(library(package_name))
```
```

2. For data operations, suppress messages about column name repairs:

```
```r
variable_name <- read_excel("<fpath>.csv", col_names = FALSE, .name_repair = "
minimal")
```
```

3. When printing dataframes, always wrap them in print() statements:

```
```r
print(head(dataframe))
```
```

The final answers to the questions must be submitted using the submit_answer tool. You must use the submit_answer tool to end the episode. You must use JSON format as follows:

Example output:

```
```
submit_answer({
 "q1": "Short answer to question 1",
 "q2": "Short answer to question 2",
 "q3": "Short answer to question 3",
 "q4": "Short answer to question 4"
})
```
```

YOU MUST ANSWER ALL FOUR QUESTIONS in the json format above. Remember to provide clear, concise answers that directly address each question based on your analysis.