
MOOSE-Chem2: Exploring LLM Limits in Fine-Grained Scientific Hypothesis Discovery via Hierarchical Search

Anonymous Author(s)

Affiliation

Address

email

Abstract

Large language models (LLMs) have shown promise in automating scientific hypothesis generation, yet existing approaches primarily yield coarse-grained hypotheses lacking critical methodological and experimental details. We introduce and formally define the novel task of *fine-grained scientific hypothesis discovery*, which entails generating detailed, experimentally actionable hypotheses from coarse initial research directions. We frame this as a combinatorial optimization problem and investigate the upper limits of LLMs’ capacity to solve it when maximally leveraged. Specifically, we explore four foundational questions: (1) how to best harness an LLM’s internal heuristics to formulate the fine-grained hypothesis it itself would judge as the most promising among all the possible hypotheses it might generate, based on its own internal scoring—thus defining a latent reward landscape over the hypothesis space; (2) whether such LLM-judged better hypotheses exhibit stronger alignment with ground-truth hypotheses; (3) whether shaping the reward landscape using an ensemble of diverse LLMs of similar capacity yields better outcomes than defining it with repeated instances of the strongest LLM among them; and (4) whether an ensemble of identical LLMs provides a more reliable reward landscape than a single LLM. To address these questions, we propose a hierarchical search method that incrementally proposes and integrates details into the hypothesis, progressing from general concepts to specific experimental configurations. We show that this hierarchical process smooths the reward landscape and enables more effective optimization. Empirical evaluations on a new benchmark of expert-annotated fine-grained hypotheses from recent chemistry literature show that our method consistently outperforms strong baselines.

1 Introduction

Large language models (LLMs) have increasingly been applied to assist scientific research (Luo et al., 2025; Cambria et al., 2023), with one of the most ambitious applications being the automated discovery of novel and valid scientific hypotheses. However, current methods produce hypotheses that are criticized for being overly coarse, lacking sufficient detail, offering simplistic suggestions, or omitting concrete implementation strategies (Wang et al., 2024; Hu et al., 2024; Si et al., 2024).

We present the first systematic investigation into how LLMs can be leveraged to formulate fine-grained scientific hypotheses—those enriched not only with major concepts but also with precise methodological details and clearly specified experimental configurations. For example, a coarse-grained hypothesis in chemistry might state, “*synthesize hierarchical 3D copper*,” while a fine-grained counterpart could elaborate, “*Copper foils are chemically oxidized by immersion in a solution of 0.5 M ammonium persulfate and 2 M sodium hydroxide for 15 minutes at room temperature, forming a*

36 *pentagonal hierarchical CuO nanostructure.*” Such fine-grained hypotheses significantly enhance
37 clarity, feasibility, and experimental implementability.

38 Formally, we define the task as generating a fine-grained hypothesis given a research back-
39 ground—comprising a research question and established methodologies—and a coarse-grained
40 hypothesis direction. We show that fine-grained scientific hypothesis discovery is a combinatorial
41 search problem, as it requires selecting and composing a coherent set of concrete details from a
42 vast space of plausible options—making it particularly challenging in practice. The difficulty is
43 compounded by the fact that scientific hypothesis discovery is an inherently out-of-domain (OOD)
44 problem: the correctness of a hypothesis is fundamentally *unknown* at the time of formulation.

45 In this work, we focus on the pre-experimental stage of discovery, mirroring how human scien-
46 tists—prior to empirical testing—iteratively search through the hypothesis space using heuristics and
47 domain knowledge to identify the hypothesis they themselves would judge as the most promising
48 among all plausible candidates they could think of during the hypothesis search process.

49 Our goal is to emulate this cognitive search process using LLMs, which increasingly rival human
50 scientists in heuristic reasoning and scientific knowledge understanding. This motivates our central
51 research question (*Q1*): *how to best harness an LLM’s internal heuristics to formulate the fine-*
52 *grained hypothesis it itself would judge as the most promising among all possible hypotheses it might*
53 *generate?* We conceptualize a hypothesis space where each point along the input dimensions (the
54 *x*-axis, potentially multidimensional) represents a candidate hypothesis, and each point is assigned
55 a reward value (on the *y*-axis) by the LLM based on its internal heuristics. This defines a reward
56 landscape over the hypothesis space, with the highest peak corresponding to the hypothesis the LLM
57 internally judges as most promising. Framed this way, *Q1* becomes an optimization problem: *how*
58 *can we navigate this landscape to find stronger local optima—or ideally the global optimum—thus*
59 *eliciting the best fine-grained hypothesis the LLM can generate?*

60 A straightforward baseline is greedy search over the reward landscape. However, its non-convex
61 and complex structure makes naive greedy strategies prone to poor local optima. To address this,
62 we propose a hierarchical search framework that explicitly models how a finite-capacity reasoning
63 agent—human or LLM—navigates the hypothesis space. Specifically, it first explores higher-level
64 conceptual spaces and then incrementally refines into more specific detail spaces. This hierarchical
65 approach smooths the reward landscape at each hierarchy level—especially at higher, more abstract
66 levels—enabling convergence to superior local optima compared to greedy search and greedy search
67 with self-consistency (Wang et al., 2023). The proposed framework inherently scales with the
68 LLM’s available capacity, enabling systematic exploration of the limits of LLM-driven fine-grained
69 hypothesis generation.

70 Having investigated how to identify stronger local optima in *Q1*, we now turn to our second question
71 (*Q2*): *whether hypotheses judged better by LLMs exhibit stronger alignment with ground-truth*
72 *hypotheses?* To rigorously address *Q2* while avoiding data contamination, we construct a benchmark
73 of research backgrounds paired with expert-annotated fine-grained hypotheses from chemistry papers
74 published after January 2024, ensuring these examples were unseen by our LLM (GPT-4o-mini,
75 October 2023 cutoff). Using this benchmark, we indirectly evaluate *Q2* by comparing the recall of
76 hypotheses discovered by our hierarchical approach—which locates higher LLM-internal local op-
77 tima—with hypotheses identified by baseline methods. Our results consistently show that hypotheses
78 generated by our method achieve higher recall than those from baselines, providing empirical support
79 for the reliability of the LLM’s internal reward signal in guiding fine-grained hypothesis discovery.

80 Until now, the reward landscape guiding hypothesis search has been defined by a single LLM serving
81 as the evaluator. We now turn to the third question (*Q3*): *whether defining this landscape with an*
82 *ensemble of diverse LLMs of similar capacity yields better outcomes than using multiple instances of*
83 *the strongest LLM within that group.* Our experiments show that ensembles composed of repeated
84 instances of the strongest LLM consistently outperform equally sized ensembles of diverse models,
85 suggesting that peak model quality is more important than architectural diversity in this setting.

86 Finally, we consider a fourth question (*Q4*): *whether an ensemble of identical LLMs provides a more*
87 *effective reward landscape than a single instance of the same model.* While *Q3* compares ensembles
88 of different models, *Q4* isolates the effect of aggregation alone by controlling for model identity. We
89 find that even identical LLMs, when sampled independently and aggregated via summarization, yield

a reward signal that better captures novelty without sacrificing overall quality—highlighting a subtle but important dimension in optimizing hypothesis discovery.

Overall, the contributions of this work are:

1. We introduce and formalize fine-grained scientific hypothesis discovery as a combinatorial optimization problem, and release a post-2024 chemistry benchmark with expert-annotated fine-grained hypotheses, explicitly designed to prevent data contamination for current LLMs.
2. We systematically investigate this task through four foundational research questions: (*Q1*) how to best leverage an LLM’s internal heuristics for fine-grained hypothesis generation; (*Q2*) whether LLM-preferred hypotheses align more closely with ground-truth expert hypotheses; (*Q3*) whether ensembles of diverse LLMs provide a better reward landscape compared to repeated use of the strongest single model; and (*Q4*) whether ensembles of identical LLMs offer a better reward landscape than a single instance of the same model.
3. We propose a hierarchical search method over levels of conceptual abstraction, which smooths the reward landscape and reduces search complexity at each hierarchy level. Empirically, it consistently outperforms strong baselines in both LLM self-evaluation, expert evaluation, and recall against annotated ground-truth hypotheses.

2 Methodology

2.1 Background and Task Motivation

Yang et al. (2024b) assume that many chemistry hypotheses can be constructed from a research background b —typically including the research question and/or background survey—and a set of inspirations $i_1, \dots, i_k$, representing concepts or findings from the literature. It can be formulated as:

$$h = f(b, i_1, \dots, i_k) \quad (1)$$

In practice, however, most hypotheses h generated from Equation 1 tend to be coarse-grained: while they form cohesive associations between b and the i , they often lack clear hypothesis specification and the detailed experimental configurations required for direct implementation in a laboratory setting. Additionally, many such hypotheses contain redundant or tangential elements—either due to the inclusion of unnecessary inspirations or from noise present in the literature that is unrelated to the core knowledge intended for hypothesis construction.

2.2 Problem Formulation: Fine-Grained Hypothesis Generation as Combinatorial Search

Let h_c be a coarse-grained hypothesis direction and h_f its fine-grained counterpart, defined as:

$$h_f = \{h_c, d_1, \dots, d_m\} \quad (2)$$

Here, $\{h_c, d_1, \dots, d_m\}$ denotes the meaningful integration of edits $d_1, \dots, d_m$ into h_c , resulting in a coherent, fine-grained hypothesis. Each d represents either: (1) the addition of a fine-grained detail to a concept i in h_c , or (2) the deletion of a redundant concept from h_c . We define two sets of edit candidates: D^+ , consisting of all possible fine-grained details that can be added to concepts in h_c , and D^- , consisting of all concepts in h_c that may be removed. The full candidate set is defined as $D = D^+ \cup D^-$.

Inspired by coarse-to-fine strategies in computer vision—where a coarse image is first generated and then refined with fine-grained details (Tian et al., 2024)—we formulate the transition from h_c to h_f as an additional step building on Equation 1, which provides the initial h_c .

$$P(h_f | b, h_c) = P(\{d_1, \dots, d_m\} | b, h_c, D) \quad (3)$$

This formulation turns $P(h_f | b, h_c)$ into a combinatorial optimization problem, where the objective is to select a subset of edits $d_1, \dots, d_m \subseteq D$. Let $|D| = n$ and $|d_1, \dots, d_m| = m$. The search space has at least combinatorial complexity $C_n^m = \frac{n!}{m!(n-m)!}$. This makes the problem particularly

challenging due to three factors: (1) both m and n are unknown; (2) the candidate set D is itself implicit and potentially very large; and (3) the edits d_i are not independent—errors early in the reasoning chain can propagate and impair later decisions.

2.3 Algorithmic Motivation for Hierarchical Heuristic Search (HHS)

Fine-grained hypothesis generation is generally intractable due to the exponential growth of the search space, where the candidate set $|D|$ is often large or prohibitively so.

A notable exception occurs when the problem exhibits an optimal substructure—i.e., an optimal solution can be composed from the optimal solutions to its subproblems. This principle underlies dynamic programming, where solutions are built incrementally from smaller subproblems (first try to obtain the optimal solution for a smaller subproblem and then iteratively find the optimal solution for larger subproblems).

We observe that fine-grained hypothesis generation exhibits an optimal substructure. Specifically, the edits $d_1, \dots, d_m$ can be organized hierarchically: some address high-level concepts (e.g., functional groups, catalyst classes), while others specify low-level details (e.g., reagents, catalysts, temperature, concentration). We assume these edits can be partitioned into p hierarchical levels ($p > 1$), with higher levels corresponding to finer details. Then the overall problem can be seen as to determine d in $\{1, \dots, p\}$ hierarchies. The subproblem of it can be seen as the determination of d in $\{1, \dots, p-1\}$ hierarchies, etc. Then it is obvious that the optimal solution of a problem can be derived from the optimal solution of its subproblem, etc.

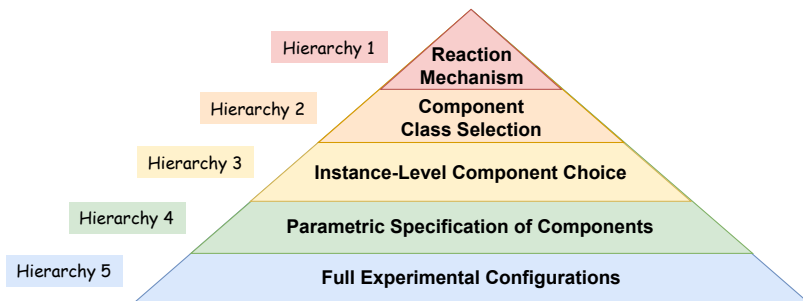


Figure 1: Hierarchies designed for chemistry.

Figure 1 illustrates an example hierarchical decomposition for chemistry, developed in collaboration with domain experts (PhD-level chemists). The hierarchy spans from high-level mechanistic intent to low-level experimental configurations, reflecting the granularity typically considered when translating a conceptual hypothesis into a testable laboratory procedure in chemistry.

Now we have simplified the problem of determining d in all p hierarchies into the iteration of determining d in each hierarchy sequentially. Nonetheless, even within a single hierarchy, the number of candidates remains combinatorially large.

A practical approach to this combinatorial complexity is to use heuristics that approximate solutions rather than exhaustively searching for exact ones. This aligns with how chemists refine hypotheses: given that h_c often represents an unexplored direction, explicit candidate details d are rarely retrievable from existing databases. Instead, chemists often rely on domain knowledge and intuition to heuristically identify and iteratively refine plausible details.

Analogously, we propose to leverage LLMs’ internal heuristics to guide the search for d at each hierarchical level. As LLMs advance, their heuristics—emerging from pretraining over extensive scientific corpora—increasingly approximate, and in some cases surpass, those of human experts. By developing methods that maximally exploit these heuristics for fine-grained hypothesis generation, we aim to provide scientists with progressively greater support as LLMs’ capacities grow.

In this setting, the candidate space D is not explicitly enumerated but is implicitly embedded within the LLM’s internal knowledge and reasoning capabilities. The LLM does not select d from a predefined list, but rather proposes candidates by navigating this latent, heuristic-driven space.

2.4 Hierarchical Factorization of the Search Problem

For formalization, we partition the implicit candidate space D into p hierarchical levels, where $D^{(i)} \subseteq D$ represents all potential edits at level i , and $D^{*(i)} \subseteq D^{(i)}$ denotes the (unknown) ground-truth edits. The j -th ground-truth edit at level i is denoted as $d_j^{*(i)} \in D^{*(i)}$. Since D is implicitly determined by h_c , we have $P(D \mid h_c) = 1$, and explicitly condition on D for clarity in the subsequent factorization. Applying the chain rule hierarchically, Equation 3 can be reformulated as:

$$P(h_f \mid b, h_c) = P\left(\{D^{*(1)}, \dots, D^{*(p)}\} \mid b, h_c, D\right) \quad (4)$$

$$= \prod_{i=1}^p P\left(D^{*(i)} \mid b, h_c, D^{*(<i)}, D^{(i)}\right) \quad (5)$$

$$= \prod_{i=1}^p \prod_{j=1}^{|D^{*(i)}|} P\left(d_j^{*(i)} \mid b, h_c, D^{*(<i)}, d_{<j}^{*(i)}, D^{(i)}\right), \quad (6)$$

where $D^{*(<i)} = \{D^{*(1)}, \dots, D^{*(i-1)}\}$ and $d_{<j}^{*(i)} = \{d_1^{*(i)}, \dots, d_{j-1}^{*(i)}\}$.

The key advantage of this hierarchical factorization is that at each level i , the search is restricted to the reduced candidate set $D^{(i)}$ rather than the full space D , significantly narrowing the search space. Moreover, as we will show in § 2.6, this hierarchical decomposition smooths the reward landscape at each hierarchy level, facilitating more stable optimization and enabling the discovery of stronger local optima in the hypothesis space.

2.5 LLM-Based Implementation of Hierarchical Heuristic Search

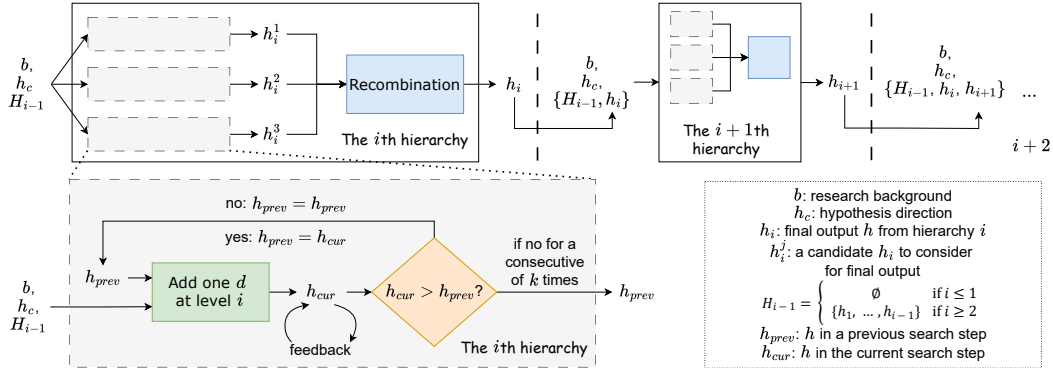


Figure 2: Overview of the proposed Hierarchical Heuristic Search (HHS) framework.

We implement *HHS* as an LLM-driven agentic process that directly follows the hierarchical factorization formalized in Equation 6. As shown in Figure 2, at each hierarchy level i , H_{i-1} represents the accumulated edits from all previous levels, corresponding to $D^{*(<i)}$. Within the current level, h_{prev} denotes the partial hypothesis incorporating the edits selected up to step $j-1$, i.e., $d_{<j}^{*(i)}$. The candidate set $D^{(i)}$ is not explicitly enumerated but emerges implicitly from the LLM’s internal heuristics, conditioned on the background b , the hypothesis direction h_c , and the edits selected so far. Specifically, the search for a local optimum h_i^j begins from the initial point h_{i-1} , using contextual information from b , h_c , and H_{i-2} . For hierarchy level $i = 1$, we set $h_0 = h_c$ and $H_0 = \emptyset$, making the hypothesis direction h_c the starting point.

At each iteration, the *Add one d at level i* module prompts an LLM to propose an edit d to h_{prev} , producing a candidate h_{cur} , which is then refined once for validity, novelty, and specificity. The $h_{cur} > h_{prev}$ module evaluates whether the new hypothesis improves upon the previous one via LLM-based pairwise comparison, serving as an *internal gradient signal*.

196 This search process continues until no further improvement is observed over k consecutive steps
 197 (default $k = 3$), at which point the current hypothesis is accepted as a local optimum. Each edit d
 198 may involve either an addition or a deletion, allowing the search path to include retrospection and
 199 self-correction as needed.

200 Within each hierarchy level, we adapt the design of an evolutionary unit (Yang et al., 2024b) to our
 201 task, where the search for the local optimum h_i is independently repeated multiple times (set to three
 202 in our implementation), yielding distinct local optima h_i^1 , h_i^2 , and h_i^3 . These candidates are then
 203 passed to a *recombination* module, which integrates their complementary strengths to interpolate a
 204 potentially superior local optimum h_i within the subspace spanned by h_i^1, h_i^2, h_i^3 .

205 2.6 Theoretical Analysis: Smoothing Effects of Hierarchical Heuristic Search

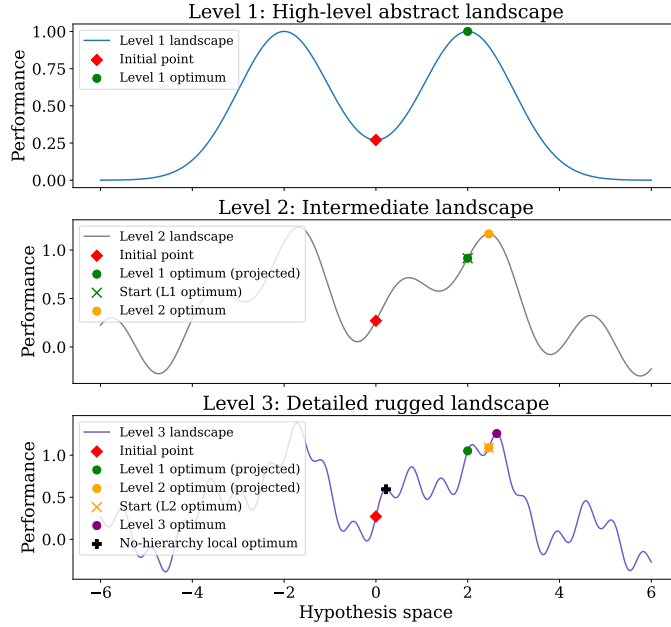


Figure 3: The smoothing effect of hierarchy on the hypothesis space performance landscape.

206 A key observation is that a hypothesis candidate’s performance at a higher hierarchy level can be
 207 viewed as an aggregated estimate—approximating an average or soft maximum—of its lower-level
 208 subspace. For instance, when evaluating a coarse-grained concept like “hierarchical 3D copper,”
 209 the LLM may implicitly account for diverse fine-grained structural variants, some highly relevant,
 210 others ineffective. We hypothesize that the LLM’s higher-level assessment aggregates these outcomes,
 211 weighting promising variants within the broader distribution to produce an overall estimate of the
 212 concept’s expected potential.

213 Building on this observation, the hierarchical abstraction smooths the reward landscape at higher
 214 levels by attenuating local irregularities in the fine-grained space, as the performance of a point at a
 215 higher level can be interpreted as an approximate aggregation or average of the performance across
 216 its corresponding lower-level subspace. This effect is illustrated in Figure 3 (a simplified schematic
 217 projection into a 1D space). Consequently, direct search over the flat, non-hierarchical space tends
 218 to be highly rugged and non-convex, often leading to premature convergence to suboptimal local
 219 optima. In contrast, introducing hierarchical structure progressively smooths the landscape, enabling
 220 more stable and efficient optimization, particularly at higher levels.

221 This smoothing effect can also be interpreted in the frequency domain as a form of low-pass filtering,
 222 where high-frequency components of the landscape are attenuated, resulting in a spectral cutoff in the
 223 spatial frequency domain, as illustrated in Figure 4.

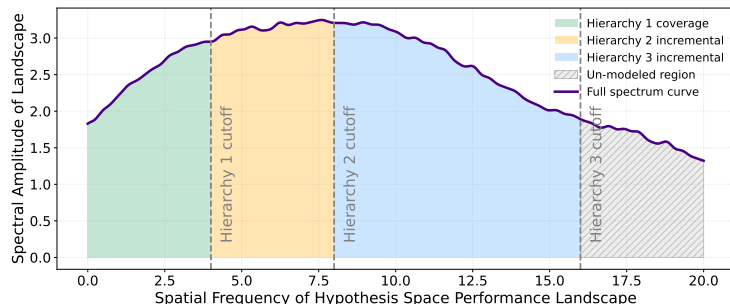


Figure 4: Spectrum of the performance landscape over the hypothesis space: hierarchical design as a low-pass filter attenuating high-frequency irregularities.

3 Experiment: Investigating the Four Fundamental Questions

3.1 Benchmark Construction, LLM Selection, and Baselines

To our knowledge, no existing benchmark provides annotated fine-grained scientific hypotheses—detailed enough for direct experimental execution. A critical consideration in constructing such a benchmark is minimizing data contamination by ensuring the LLM has not been exposed to the annotated data during pretraining. We extend the TOMATO-Chem dataset (Yang et al., 2024b), which contains 51 chemistry papers published and made available online after January 2024 in leading journals such as *Nature* and *Science*. Each entry is annotated with a research background b and a coarse-grained hypothesis h_c . For this study, two PhD-level chemists further annotated these examples with fine-grained hypotheses h_f , providing ground-truth references for evaluation. To rigorously avoid data contamination, all experiments are conducted using GPT-4o-mini, whose pretraining data cutoff is October 2023.

We compare *HHS* against two strong baselines widely used in search tasks: (1) greedy search and (2) greedy search with self-consistency. The latter serves as an ablation of *HHS* where the hierarchical decomposition is removed, performing the search in a single stage with each d sampled directly from the full candidate set D rather than hierarchy-specific subsets $D^{(i)}$. The self-consistency mechanism is similar to the *Recombination* module in Figure 2, which interpolate multiple local optima trying to find a better one. Greedy search represents a further ablation, disabling the *Recombination* module entirely and following a single search trace where the first found local optimum (h_i^1) is directly adopted as the output ($h_i = h_i^1$ in Figure 2).

3.2 Experiments on the Questions

Due to page limit, the experiments for *Q1*, *Q2*, *Q3*, and *Q4* can be found at Appendix § A, § B, § C, and § D, correspondingly. The conclusions to each of the question correspondingly are:

- *HHS* can lead to better local optimum than the baseline methods.
- The better local optimum found (judged by LLMs without reference), the better recall of number of components in the groundtruth hypothesis (judged by LLMs with reference).
- Repetitive usage of the best LLMs for providing the reward landscape outperforms the usage of diverse LLMs in similar capacity.
- Repetitive usage of the same LLMs multiple times to provide the reward landscape outperforms using it only once.

4 Related Work

LLM-driven scientific discovery methods typically fall into two categories: (1) direct generation of hypotheses from a research background—comprising a research question and established methodologies (Qi et al., 2023); or (2) retrieval of seemingly unrelated but potentially useful knowledge fragments—termed inspirations—which are then combined with the background to construct a

hypothesis (Yang et al., 2024a,b; Wang et al., 2024; Liu et al., 2025). While these methods have shown promise in generating novel ideas, they are often criticized for producing hypotheses that are overly coarse, lacking in detail, or omitting actionable experimental steps (Wang et al., 2024; Hu et al., 2024; Si et al., 2024). In contrast, our goal is to investigate how LLMs can be leveraged to generate *fine-grained scientific hypotheses*—those sufficiently detailed to be directly implemented in laboratory settings. Notably, inspiration-based hypotheses from prior work (Yang et al., 2024a,b) can serve as inputs to our framework in the form of coarse-grained hypothesis directions. These hypotheses emphasize novelty, as they often draw from previously unassociated knowledge, while our work complements them by enhancing experimental specificity and validity.

5 Conclusion

We present the first systematic study of fine-grained scientific hypothesis discovery with LLMs, framing it as a combinatorial optimization problem. To address the vast, implicit search space, we propose hierarchical heuristic search (*HHS*), which incrementally refines hypotheses from coarse to fine-grained levels. *HHS* smooths the reward landscape, stabilizes optimization, and consistently finds higher-quality hypotheses than flat search methods. Experiments show that (1) *HHS* reliably discovers better local optima than baselines, (2) LLM-preferred hypotheses align more closely with expert ground truths, and (3) repeated use of the strongest model provides better reward landscapes than diverse ensembles. These results underscore the value of hierarchical search in exploring more LLMs’ potential for scientific discovery.

References

- Anthropic. The claude 3 model family: Opus, sonnet, haiku. <https://www.anthropic.com/news/claude-3-family>, March 2024. Accessed: 2025-05-16.
- Erik Cambria, Rui Mao, Melvin Chen, Zhaoxia Wang, and Seng-Beng Ho. Seven pillars for the future of artificial intelligence. *IEEE Intelligent Systems*, 38(6):62–69, 2023.
- Petko Georgiev, Ving Ian Lei, Ryan Burnell, Libin Bai, Anmol Gulati, Garrett Tanzer, Damien Vincent, Zhufeng Pan, Shibo Wang, Soroosh Mariooryad, Yifan Ding, Xinyang Geng, Fred Alcober, Roy Frostig, Mark Omernick, Lexi Walker, Cosmin Paduraru, Christina Sorokin, Andrea Tacchetti, Colin Gaffney, Samira Daruki, Olcan Sercinoglu, Zach Gleicher, Juliette Love, Paul Voigtlaender, Rohan Jain, Gabriela Surita, Kareem Mohamed, Rory Blevins, Junwhan Ahn, Tao Zhu, Kornnaphop Kawintiranon, Orhan Firat, Yiming Gu, Yujing Zhang, Matthew Rahtz, Manaal Faruqi, Natalie Clay, Justin Gilmer, JD Co-Reyes, Ivo Penchev, Rui Zhu, Nobuyuki Morioka, Kevin Hui, Krishna Haridasan, Victor Campos, Mahdis Mahdieh, Mandy Guo, Samer Hassan, Kevin Kilgour, Arpi Vezer, Heng-Tze Cheng, Raoul de Liedekerke, Siddharth Goyal, Paul Barham, DJ Strouse, Seb Noury, Jonas Adler, Mukund Sundararajan, Sharad Vikram, Dmitry Lepikhin, Michela Paganini, Xavier Garcia, Fan Yang, Dasha Valter, Maja Trebacz, Kiran Vodrahalli, Chulayuth Asawaroengchai, Roman Ring, Norbert Kalb, Livio Baldini Soares, Siddhartha Brahma, David Steiner, Tianhe Yu, Fabian Mentzer, Antoine He, Lucas Gonzalez, Bibo Xu, Raphael Lopez Kaufman, Laurent El Shafey, Junhyuk Oh, Tom Hennigan, George van den Driessche, Seth Odoom, Mario Lucic, Becca Roelofs, Sid Lall, Amit Marathe, Betty Chan, Santiago Ontanon, Luheng He, Denis Teplyashin, Jonathan Lai, Phil Crone, Bogdan Damoc, Lewis Ho, Sebastian Riedel, Karel Lenc, Chih-Kuan Yeh, Aakanksha Chowdhery, Yang Xu, Mehran Kazemi, Ehsan Amid, Anastasia Petrushkina, Kevin Swersky, Ali Khodaei, Gowoon Chen, Chris Larkin, Mario Pinto, Geng Yan, Adria Puigdomenech Badia, Piyush Patil, Steven Hansen, Dave Orr, Sebastien M. R. Arnold, Jordan Grimstad, Andrew Dai, Sholto Douglas, Rishika Sinha, Vikas Yadav, Xi Chen, Elena Gribovskaya, Jacob Austin, Jeffrey Zhao, Kaushal Patel, Paul Komarek, Sophia Austin, Sebastian Borgeaud, Linda Friso, Abhimanyu Goyal, Ben Caine, Kris Cao, Da-Woon Chung, Matthew Lamm, Gabe Barth-Maron, Thais Kagohara, Kate Olszewska, Mia Chen, Kaushik Shivakumar, Rishabh Agarwal, Harshal Godhia, Ravi Rajwar, Javier Snider, Xerxes Dotiwalla, Yuan Liu, Aditya Barua, Victor Ungureanu, Yuan Zhang, Bat-Orgil Batsaikhan, Mateo Wirth, James Qin, Ivo Danihelka, Tulsee Doshi, Martin Chadwick, Jilin Chen, Sanil Jain, Quoc Le, Arjun Kar, Madhu Gurusurthy, Cheng Li, Ruoxin Sang, Fangyu Liu, Lampros Lamprou, Rich Munoz, Nathan Lintz, Harsh Mehta, Heidi Howard, Malcolm Reynolds, Lora Aroyo, Quan Wang, Lorenzo Blanco, Albin Cassirer, Jordan Griffith, Dipanjan Das, Stephan Lee, Jakub Sygnowski, Zach

312 Fisher, James Besley, Richard Powell, Zafarali Ahmed, Dominik Paulus, David Reitter, Zalan
313 Borsos, Rishabh Joshi, Aedan Pope, Steven Hand, Vittorio Selo, Vihan Jain, Nikhil Sethi, Megha
314 Goel, Takaki Makino, Rhys May, Zhen Yang, Johan Schalkwyk, Christina Butterfield, Anja Hauth,
315 Alex Goldin, Will Hawkins, Evan Senter, Sergey Brin, Oliver Woodman, Marvin Ritter, Eric
316 Noland, Minh Giang, Vijay Bolina, Lisa Lee, Tim Blyth, Ian Mackinnon, Machel Reid, Obaid
317 Sarvana, David Silver, Alexander Chen, Lily Wang, Loren Maggiore, Oscar Chang, Nithya Attaluri,
318 Gregory Thornton, Chung-Cheng Chiu, Oskar Bunyan, Nir Levine, Timothy Chung, Evgenii
319 Eltyshev, Xiance Si, Timothy Lillicrap, Demetra Brady, Vaibhav Aggarwal, Boxi Wu, Yuanzhong
320 Xu, Ross McIlroy, Kartikeya Badola, Paramjit Sandhu, Erica Moreira, Wojciech Stokowiec,
321 Ross Hemsley, Dong Li, Alex Tudor, Pranav Shyam, Elahe Rahimtoroghi, Salem Haykal, Pablo
322 Sprechmann, Xiang Zhou, Diana Mincu, Yujia Li, Ravi Addanki, Kalpesh Krishna, Xiao Wu,
323 Alexandre Frechette, Matan Eyal, Allan Dafoe, Dave Lacey, Jay Whang, Thi Avrahami, Ye Zhang,
324 Emanuel Taropa, Hanzhao Lin, Daniel Toyama, Eliza Rutherford, Motoki Sano, HyunJeong Choe,
325 Alex Tomala, Chalence Safranek-Shrader, Nora Kassner, Mantas Pajarskas, Matt Harvey, Sean
326 Sechrist, Meire Fortunato, Christina Lyu, Gamaleldin Elsayed, Chenkai Kuang, James Lottes,
327 Eric Chu, Chao Jia, Chih-Wei Chen, Peter Humphreys, Kate Baumli, Connie Tao, Rajkumar
328 Samuel, Cicero Nogueira dos Santos, Anders Andreassen, Nemanja Rakićević, Dominik Grewe,
329 Aviral Kumar, Stephanie Winkler, Jonathan Caton, Andrew Brock, Hannah Sheahan, Iain Barr,
330 Yingjie Miao, Paul Natsev, Jacob Devlin, Feryal Behbahani, Flavien Prost, Yanhua Sun, Artiom
331 Myaskovsky, Thanumalayan Sankaranarayanan Pillai, Dan Hurt, Angeliki Lazaridou, Xi Xiong,
332 Ce Zheng, Fabio Pardo, Xiaowei Li, Dan Horgan, Joe Stanton, Moran Ambar, Fei Xia, Alejandro
333 Lince, Mingqiu Wang, Basil Mustafa, Albert Webson, Hyo Lee, Rohan Anil, Martin Wicke,
334 Timothy Dozat, Abhishek Sinha, Enrique Piqueras, Elahe Dabir, Shyam Upadhyay, Anudhyan
335 Boral, Lisa Anne Hendricks, Corey Fry, Josip Djolonga, Yi Su, Jake Walker, Jane Labanowski,
336 Ronny Huang, Vedant Misra, Jeremy Chen, RJ Skerry-Ryan, Avi Singh, Shruti Rijhwani, Dian
337 Yu, Alex Castro-Ros, Beer Changpinyo, Romina Datta, Sumit Bagri, Arnar Mar Hrafnkelsson,
338 Marcello Maggioni, Daniel Zheng, Yury Sulsky, Shaobo Hou, Tom Le Paine, Antoine Yang, Jason
339 Riesa, Dominika Rogozinska, Dror Marcus, Dalia El Badawy, Qiao Zhang, Luyu Wang, Helen
340 Miller, Jeremy Greer, Lars Lowe Sjos, Azade Nova, Heiga Zen, Rahma Chaabouni, Mihaela Rosca,
341 Jiepu Jiang, Charlie Chen, Ruibo Liu, Tara Sainath, Maxim Krikun, Alex Polozov, Jean-Baptiste
342 Lepiau, Josh Newlan, Zeynep Cankara, Soo Kwak, Yunhan Xu, Phil Chen, Andy Coenen, Clemens
343 Meyer, Katerina Tsihlis, Ada Ma, Juraj Gottweis, Jinwei Xing, and Gu. Gemini 1.5: Unlocking
344 multimodal understanding across millions of tokens of context. *arXiv preprint arXiv:2403.05530*,
345 2024. URL <https://arxiv.org/abs/2403.05530>.

346 Xiang Hu, Hongyu Fu, Jinge Wang, Yifeng Wang, Zhikun Li, Renjun Xu, Yu Lu, Yaochu Jin, Lili
347 Pan, and Zhenzhong Lan. Nova: An iterative planning and search approach to enhance novelty
348 and diversity of LLM generated ideas. *CoRR*, abs/2410.14255, 2024. doi: 10.48550/ARXIV.2410.
349 14255. URL <https://doi.org/10.48550/arXiv.2410.14255>.

350 Zongjie Li, Chaozheng Wang, Pingchuan Ma, Daoyuan Wu, Shuai Wang, Cuiyun Gao, and Yang
351 Liu. Split and merge: Aligning position biases in llm-based evaluators. In *Proceedings of the 2024*
352 *Conference on Empirical Methods in Natural Language Processing*, pp. 11084–11108, 2024.

353 Yujie Liu, Zonglin Yang, Tong Xie, Jinjie Ni, Ben Gao, Yuqiang Li, Shixiang Tang, Wanli Ouyang,
354 Erik Cambria, and Dongzhan Zhou. Researchbench: Benchmarking llms in scientific discovery
355 via inspiration-based task decomposition. *arXiv preprint arXiv:2503.21248*, 2025.

356 Ziming Luo, Zonglin Yang, Zexin Xu, Wei Yang, and Xinya Du. Llm4sr: A survey on large language
357 models for scientific research. *arXiv preprint arXiv:2501.04306*, 2025.

358 OpenAI. Gpt-4o mini: Advancing cost-efficient intelligence. [https://openai.com/index/
359 gpt-4o-mini-advancing-cost-efficient-intelligence/](https://openai.com/index/gpt-4o-mini-advancing-cost-efficient-intelligence/), 2024. Accessed: 2025-05-16.

360 Biqing Qi, Kaiyan Zhang, Haoxiang Li, Kai Tian, Sihang Zeng, Zhang-Ren Chen, and Bowen Zhou.
361 Large language models are zero shot hypothesis proposers. *CoRR*, abs/2311.05965, 2023. doi:
362 10.48550/ARXIV.2311.05965. URL <https://doi.org/10.48550/arXiv.2311.05965>.

363 Chenglei Si, Diyi Yang, and Tatsunori Hashimoto. Can llms generate novel research ideas? A
364 large-scale human study with 100+ NLP researchers. *CoRR*, abs/2409.04109, 2024. doi: 10.48550/
365 ARXIV.2409.04109. URL <https://doi.org/10.48550/arXiv.2409.04109>.

- 366 Keyu Tian, Yi Jiang, Zehuan Yuan, Bingyue Peng, and Liwei Wang. Visual autoregressive modeling:
367 Scalable image generation via next-scale prediction. *CoRR*, abs/2404.02905, 2024. doi: 10.48550/
368 ARXIV.2404.02905. URL <https://doi.org/10.48550/arXiv.2404.02905>.
- 369 Qingyun Wang, Doug Downey, Heng Ji, and Tom Hope. Scimon: Scientific inspiration machines
370 optimized for novelty. In Lun-Wei Ku, Andre Martins, and Vivek Srikumar (eds.), *Proceedings of*
371 *the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*,
372 *ACL 2024, Bangkok, Thailand, August 11-16, 2024*, pp. 279–299. Association for Computational
373 Linguistics, 2024. doi: 10.18653/V1/2024.ACL-LONG.18. URL [https://doi.org/10.18653/
374 v1/2024.acl-long.18](https://doi.org/10.18653/v1/2024.acl-long.18).
- 375 Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc V Le, Ed H Chi, Sharan Narang, Aakanksha
376 Chowdhery, and Denny Zhou. Self-consistency improves chain of thought reasoning in language
377 models. In *The Eleventh International Conference on Learning Representations*, 2023.
- 378 Zonglin Yang, Xinya Du, Junxian Li, Jie Zheng, Soujanya Poria, and Erik Cambria. Large language
379 models for automated open-domain scientific hypotheses discovery. In Lun-Wei Ku, Andre Martins,
380 and Vivek Srikumar (eds.), *Findings of the Association for Computational Linguistics, ACL 2024,*
381 *Bangkok, Thailand and virtual meeting, August 11-16, 2024*, pp. 13545–13565. Association
382 for Computational Linguistics, 2024a. doi: 10.18653/V1/2024.FINDINGS-ACL.804. URL
383 <https://doi.org/10.18653/v1/2024.findings-acl.804>.
- 384 Zonglin Yang, Wanhao Liu, Ben Gao, Tong Xie, Yuqiang Li, Wanli Ouyang, Soujanya Poria, Erik
385 Cambria, and Dongzhan Zhou. Moose-chem: Large language models for rediscovering unseen
386 chemistry scientific hypotheses. *CoRR*, abs/2410.07076, 2024b. doi: 10.48550/ARXIV.2410.07076.
387 URL <https://doi.org/10.48550/arXiv.2410.07076>.

A Q1: How to Best Harness an LLM’s Internal Heuristics to Formulate the Fine-Grained Hypothesis It Itself Would Judge as the Most Promising Among All Possible Hypotheses It Might Generate?

We frame this question as an optimization problem: Given only a coarse-grained hypothesis h_c as the starting point, and relying entirely on a single LLM, how can we navigate the hypothesis space to approach the global optimum of the performance landscape, as defined by this same LLM’s internal heuristics, where each optimization step consists of adding an edit d to h_c ? In this setting, the LLM plays a dual role: it serves both as the *proposal generator*, proposing candidate edits d to formulate new hypotheses within the hypothesis space, and as the *gradient provider*, judging whether the new hypothesis improves upon the current one via its own internal heuristics (e.g., pairwise comparison).

While it is inherently infeasible to determine whether a found local optimum represents the global optimum, we can empirically compare local optima obtained by different methods within the same performance pairwise evaluation, and therefore check which one is closer to a global optimum.

As detailed in § 2, the hierarchical design of *HHS* offers two key advantages over flat search strategies: (1) less search space to propose each d (from $D^{(i)}$, instead of D), and (2) smoothing the performance landscape progressively in the hypothesis space. Among these, the smoothing effect is particularly critical, as it reduces the risk of early convergence to suboptimal local optima and facilitates progress toward higher peaks in the LLM’s internal performance landscape.

We compare the local optima discovered by *HHS* against the two baselines. For each pair of local optima, we conduct both overall evaluations and dimension-specific assessments across four key criteria: **effectiveness**, **novelty**, **detailedness**, and **feasibility**. In this context, **feasibility** reflects the practical ease of implementing the proposed hypothesis, encompassing factors such as implementation complexity and the minimization of redundant steps. Hypotheses that are easy to implement and free of redundant components are preferred.

We further observe two common trade-offs among these dimensions: (1) between **effectiveness** and **novelty**, as highly novel hypotheses often entail greater scientific risk and uncertainty; and (2) between **detailedness** and **feasibility**, as increased specificity can introduce procedural complexity or redundancies that diminish experimental feasibility.

We also conducted an expert evaluation involving two chemistry PhD students. For each benchmark item, one hypothesis was randomly sampled from each method, and the experts were tasked to rank the three hypotheses. The results from both the LLM-based and expert evaluations on the quality of local optima discovered by each method are presented in Table 1. To mitigate known position bias in LLM-based pairwise comparisons—where models tend to favor the first option (Li et al., 2024)—each pair of local optima was compared six times, with the order of presentation alternated every three times. A hypothesis was considered to win if it received more than three votes; a tie was recorded if both received exactly three votes.

	Effectiveness (LLM)	Novelty (LLM)	Detailedness (LLM)	Feasibility (LLM)	Overall (LLM)	Overall (Expert)
HHS v.s. Greedy Search						
Win	74.51%	41.18%	71.57%	67.65%	73.53%	76.47%
Tie	18.63%	18.63%	28.43%	10.78%	18.63%	15.69%
Lose	6.86%	40.20%	0.00%	21.57%	7.84%	7.84%
HHS v.s. Greedy Search + Self-consistency						
Win	59.31%	42.16%	56.37%	48.53%	53.43%	74.51%
Tie	24.02%	8.33%	43.14%	18.63%	33.82%	17.65%
Lose	16.67%	49.51%	0.49%	32.84%	12.75%	7.84%
Greedy Search + Self-consistency v.s. Greedy Search						
Win	57.84%	48.04%	29.41%	51.96%	54.90%	62.75%
Tie	22.55%	11.76%	65.69%	18.63%	34.31%	21.57%
Lose	19.61%	40.20%	4.90%	29.41%	10.78%	15.69%

Table 1: Comparison between HHS and baseline methods across LLM-based and expert evaluations.

B Q2: Whether Hypotheses Judged Better by LLMs Exhibit Stronger Alignment With Ground-Truth Hypotheses?

§ A shows that HHS consistently discovers superior local optima compared to baseline methods. We further investigate whether these optima exhibit stronger alignment with the ground-truth hypotheses.

Given the absence of established metrics for this task, we propose an LLM-based evaluation that quantifies how well discovered hypotheses recall key chemical components of the ground-truth. Each hypothesis is decomposed into components regarding to the research question, and recall is assessed for each ground-truth component on a 0–3 scale, where 0 denotes no recall and 3 denotes an exact match.

We report two metrics: (1) **Soft Recall**, which counts a component as recalled if its score is greater than 0, normalized by the total number of components; (2) **Hard Recall**, which sums the raw 0–3 scores and normalizes them by the maximum possible score.

As shown in Table 2, the hypotheses discovered by HHS (which are preferred by LLMs, as demonstrated in § A) achieve consistently higher recall with the ground-truth hypotheses.

	Soft Recall	Hard Recall
Greedy Search w/ Self-consistency	16.60% 31.50%	9.90% 17.70%
HHS	40.40%	23.00%

Table 2: Recall of ground-truth components by discovered hypotheses.

C Q3: Whether Defining the Reward Landscape With an Ensemble of Diverse LLMs of Similar Capacity Yields Better Outcomes Than Using Multiple Instances of the Strongest LLM Within That Group?

	EF (GT)	NV (GT)	DT (GT)	FS (GT)	OV (GT)	EF (GM)	NV (GM)	DT (GM)	FS (GM)	OV (GM)
	Mixed committee v.s. GPT-4o-mini committee					GPT-4o-mini committee v.s. Gemini-1.5-flash committee				
Win	20.83%	33.33%	14.58%	33.33%	29.17%	27.08%	31.25%	14.58%	0.00%	18.75%
Tie	41.67%	20.83%	72.92%	18.75%	33.33%	58.33%	52.08%	77.08%	95.83%	68.75%
Lose	37.50%	45.83%	12.50%	47.92%	37.50%	14.58%	16.67%	8.33%	4.17%	12.50%
	Gemini-1.5-flash committee v.s. GPT-4o-mini committee					Mixed committee v.s. Gemini-1.5-flash committee				
Win	16.67%	25.00%	6.25%	37.50%	16.67%	16.67%	33.33%	12.50%	6.25%	18.75%
Tie	41.67%	27.08%	79.17%	25.00%	52.08%	68.75%	35.42%	75.00%	93.75%	64.58%
Lose	41.67%	47.92%	14.58%	37.50%	31.25%	14.58%	31.25%	12.50%	0.00%	16.67%
	Mixed committee v.s. Gemini-1.5-flash committee					Mixed committee v.s. GPT-4o-mini committee				
Win	29.17%	45.83%	10.42%	47.92%	27.08%	8.33%	29.17%	14.58%	6.25%	8.33%
Tie	56.25%	16.67%	85.42%	10.42%	50.00%	77.08%	39.58%	70.83%	93.75%	64.58%
Lose	14.58%	37.50%	4.17%	41.67%	22.92%	14.58%	31.25%	14.58%	0.00%	27.08%

Table 3: “EF”: Effectiveness, “NV”: Novelty, “DT”: Detailedness, “FS”: Feasibility, “OV”: Overall. “(GT)” and “(GM)” indicate that the pairwise evaluations were conducted by GPT-4o-mini and Gemini-1.5-flash, respectively.

The hypothesis optimization in *HHS* depends on the “ $h_{\text{cur}} > h_{\text{prev}}?$ ” module (Figure 2), which provides the gradient signal. This raises the question: does a diverse ensemble of similarly capable LLMs enhance search performance compared to multiple instances of its strongest model?

To answer this, we design three experimental settings: (1) **Mixed Committee**: the “ $h_{\text{cur}} > h_{\text{prev}}?$ ” module is implemented by an ensemble of three different LLMs—GPT-4o-mini (OpenAI, 2024), Gemini-1.5-flash (Georgiev et al., 2024), and Claude-3-haiku (Anthropic, 2024); (2) **GPT-4o-mini Committee**: the module is implemented by three instances of GPT-4o-mini; (3) **Gemini-1.5-flash Committee**: the module is implemented by three instances of Gemini-1.5-flash. All three settings use GPT-4o-mini as the proposer module for generating edits d at each hierarchy level i .

We compare the local optima generated by these setups using LLM-based pairwise comparisons, following the protocol in § A, where each pair is evaluated six times to mitigate position bias. However, since the evaluator is itself an LLM, an additional bias may occur—favoring optima discovered using

456 gradients from the same model. To control for this, we conduct two sets of evaluations: one using
457 GPT-4o-mini as the evaluator, and the other using Gemini-1.5-flash.

458 As shown in Table 3, across both evaluators, the GPT-4o-mini committee consistently outperforms
459 the mixed committee, which in turn outperforms the Gemini-1.5-flash committee. These results
460 suggest that leveraging repeated instances of the strongest single model provides a more effective
461 gradient for hypothesis optimization than combining different models of similar capacity.

462 **D Q4: Do Multiple Identical LLMs Provide a Better Reward Landscape** 463 **Than a Single LLM?**

	Effectiveness (LLM)	Novelty (LLM)	Detailedness (LLM)	Feasibility (LLM)	Overall (LLM)
<i>HHS-1</i> v.s. <i>HHS-3</i>					
Win	21.08%	25.49%	4.41%	41.67%	8.82%
Tie	57.35%	28.92%	94.12%	28.92%	82.35%
Lose	21.57%	45.59%	1.47%	29.41%	8.82%

Table 4: Pairwise comparison between *HHS-1* and *HHS-3*.

464 In prior experiments (Tables 1 and 2), the reward landscape was defined using an ensemble of three
465 identical LLMs, followed by a fourth instance of the same LLM that aggregated the three judgments
466 into a final, reasoned decision. In the experiment corresponding to Table 3, the reward landscape
467 was defined using an ensemble of three diverse LLMs of similar capacity, with aggregation again
468 performed by a fixed instance of GPT-4o-mini. However, the extent to which this ensemble-based
469 reward landscape improves performance compared to using a single LLM remains unclear.

470 To evaluate this, we compare two variants: *HHS-3*, which uses an ensemble of three identical instances
471 of GPT-4o-mini to provide the reward signal, and *HHS-1*, which relies on a single instance of the
472 same model.

473 Table 4 reports LLM-based pairwise evaluations between the two setups across four criteria. While
474 overall quality, effectiveness, and detailedness are largely comparable, *HHS-3* outperforms in novelty,
475 whereas *HHS-1* shows an advantage in feasibility.

476 This result is somewhat counterintuitive. Although both *HHS-3* and *HHS-1* use the same base
477 LLM, they differ in how the gradient signal—that is, the decision of whether a new hypothesis
478 improves upon the current one—is computed. In *HHS-3*, three independent comparative judgments
479 are sampled using the same LLM, and a fourth instance of the same model aggregates these judgments
480 by evaluating the underlying rationales and selecting the most justified preference.

481 Crucially, this summarization step is not a simple majority vote. Instead, the LLM is explicitly
482 instructed to assess the relative strength of reasoning across all three perspectives and to favor
483 the most compelling argument. This setup implicitly allows the model to surface and validate
484 minority-supported but well-reasoned views, enabling exploration of more novel or unconventional
485 hypotheses.

486 As a result, the aggregated gradient signal in *HHS-3* may be more receptive to creative or atypical
487 ideas that would otherwise be dismissed in a single-shot comparison. In contrast, *HHS-1* relies on
488 only a single comparative judgment per step, which reflects a narrower and more internally consistent
489 heuristic perspective. This tends to produce more conservative refinements—favoring feasibility and
490 coherence, but often at the cost of novelty.

491 The observed trade-off—greater novelty in *HHS-3* and greater feasibility in *HHS-1*—thus stems not
492 from classical ensemble averaging, but from the summarizing LLM’s ability to integrate diverse
493 heuristic signals in a principled manner during hypothesis evaluation.

494 Table 5 presents the recall of ground-truth components for hypotheses generated by *HHS-1* and
495 *HHS-3*. Across both soft and hard recall metrics, *HHS-3* outperforms *HHS-1*, indicating stronger
496 alignment with expert-annotated reference hypotheses.

	Soft Recall	Hard Recall
<i>HHS-3</i>	40.40%	23.00%
<i>HHS-1</i>	32.40%	19.90%

Table 5: Recall of groundtruth components of groundtruth hypotheses between *HHS-1* and *HHS-3*.

These results suggest that while the performance gains from ensembling are relatively modest, using multiple identical LLMs to define the reward landscape may yield a more reliable and diverse evaluation signal than relying on a single instance.

E Expert Evaluation Instructions

For each research question, you will be presented with three candidate hypotheses alongside a ground-truth fine-grained hypothesis. The order of the hypotheses is randomized. Your task is to rank the three candidate hypotheses based on their quality, using the ground-truth hypothesis as a reference.

Please evaluate the hypotheses based on the following four criteria:

- Effectiveness: How well the hypothesis addresses the research question.
- Novelty: The degree of originality relative to existing knowledge.
- Detailedness: The specificity and clarity of the hypothesis.
- Feasibility: The practical plausibility of experimentally testing or implementing the hypothesis.

Note that two tradeoffs may arise:

- Between effectiveness and novelty
- Between detailedness and feasibility

Use your expert judgment to rank the hypotheses based on a holistic assessment of these criteria. In rare cases where two hypotheses appear to be of similar quality, assigning them the same rank is acceptable.

F Case Study

Example 1:

- Research Question:** How can a cost-effective N-type quasi-solid-state thermocell be developed to boost electricity production from low-grade heat by improving both ion transport efficiency and electrode performance?
- Hypothesis Candidate (from HHS):** The development of a cost-effective N-type quasi-solid-state thermocell will be achieved through the strategic integration of three core components to enhance electricity production from low-grade heat ($\leq 100^\circ\text{C}$):
 - Hierarchical Metal Electrodes:** Constructed from a copper-nickel alloy (70% Cu, 30% Ni) optimized for conductivity and corrosion resistance, these electrodes will feature a dual-layer design incorporating a titanium dioxide (TiO_2 ; CAS number: 13463-67-7) coating, precisely controlled at a thickness of 100 nm to enhance charge transport by providing a stable interface that reduces charge recombination losses. An aluminum oxide (Al_2O_3 ; CAS number: 1344-28-1) layer will be included to improve corrosion resistance and reinforce mechanical stability, operating synergistically to enhance the overall electrochemical performance. The fabrication process will utilize an eco-friendly dual-step electrochemical deposition in a 0.5 M potassium sulfate electrolyte at a controlled temperature of 25°C , ensuring micro- and nanoscale porosity targeting 50-100 nm to maximize surface area as supported by literature demonstrating

that this range optimally enhances charge transfer and ion migration efficiency. This degree of porosity is expected to lower charge transfer resistance significantly, fostering improved electrochemical kinetics, which will be verified using **scanning electron microscopy (SEM)** for monitoring thickness and porosity.

2. **Metal-Based Redox Couples:** The thermocell will utilize **copper/copper(I)** and nickel/nickel(II) redox couples, selected for their favorable redox potentials to minimize side reactions. An integrated cobalt co-catalyst (0.1 M) will serve as an effective stabilizing agent, enhancing electron transfer kinetics and maintaining the oxidation states of Cu^{2+} and Ni^{2+} during thermal cycling, as demonstrated by prior studies indicating its role in fostering electron transfer efficiency. Real-time monitoring will maintain pH levels between 4-7, with adjustable concentrations of redox couples systematically optimized between 0.5 to 1.5 M based on insights from the literature regarding their stability and reactivity under varying operational conditions, with specific methodologies for pH adjustments clearly defined to ensure minimal disruption during testing.
3. **Anisotropic Polymer Materials:** The polymer matrix will feature aligned functional groups ($-\text{COOH}$ and $-\text{SO}_3\text{H}$), which will be developed through controlled mechanical stretching (5 mm/min at 70°C), a technique known to enhance ionic transport pathways by promoting favorable interactions between the polymer backbone and ionic species. High-conductivity ionic species, specifically Na^+ and Li^+ , will be incorporated at optimized concentrations (0.5 to 1.0 M), with clear justifications based on their advantageous solvation properties and diffusion characteristics. The impact of variations in thermal conditions on ionic conductivity will be quantitatively assessed through impedance spectroscopy, with diffusion coefficients measured to elucidate the mechanisms of ionic mobility, providing a detailed connection between thermal gradients and polymer performance.

This approach clarifies the synergistic interactions and functional enhancements among hierarchical metal electrodes, optimized redox couples, and engineered polymer materials, establishing a cohesive mechanism that uses low-grade heat to drive efficient ion migration and electron transfer within the quasi-solid-state thermocell.

- **Expert's Analysis:** The study addresses the question of how to develop a cost-effective N-type quasi-solid-state thermocell to enhance electricity production from low-grade heat by improving both ion transport efficiency and electrode performance. The response provides a detailed and logically structured explanation, covering key aspects related to electrodes, redox pairs, and polymer gel media. Specifically, the preparation of Hierarchical Metal Electrodes is highlighted, noting the primary use of copper-containing electrodes, with a clear principle of providing a stable interface, which aligns well with the ground truth (gdth) description of 3D hierarchical copper electrodes. Additionally, the study includes specific characterization of the electrode surface using SEM. However, while the polymer gel mentions an anisotropic alignment structure, it omits the specific operational methods for forming this alignment, which is inconsistent with the gdth. Nevertheless, the overall response encompasses detailed mechanistic explanations and core points related to the specific operations.

Example 2:

- **Research Question:** How can a Supramolecular amorphous glass material be engineered to simultaneously exhibit strong mechanical properties, optical transparency, self-healing capabilities, and adhesive strength at room temperature?
- **Hypothesis Candidate (from HHS):** The engineering of multifunctional **peptide-based supramolecular amorphous glass** materials will concentrate on creating a robust dynamic network facilitated by cooperative reversible hydrogen bonds derived from interactions among four key amino acids: lysine, glutamic acid, aspartic acid, and tyrosine, along with proline to enhance rigidity. The hypothesized molar ratios will be set at 1:1 for lysine and glutamic acid, complemented by approximately 0.5:0.5 ratios for aspartic acid and tyrosine. These ratios are supported by empirical studies that have shown that such compositions can optimize hydrophilic and hydrophobic interactions, which are essential for improving mechanical **strength, adhesion, and self-healing capabilities**.

Lysine's positively charged ammonium group is hypothesized to establish strong ionic interactions with the negatively charged carboxylate groups of glutamic acid and aspartic acid, enhancing the stability of the hydrogen-bonding framework critical for effective energy dissipation during mechanical stress. Tyrosine will contribute to the network through π -stacking interactions, which are expected to maintain the structural integrity and optical transparency of the material under load. Proline's unique cyclic structure is anticipated to provide localized rigidity, supporting favorable peptide conformations and facilitating effective stress distribution throughout the dynamic network.

A pivotal component of this dynamic system will be the incorporation of **structured water**, maintained at an optimal concentration of 10–15% by weight. Structured water is theorized to engage in specific hydrogen bonding interactions with the peptide backbone, promoting molecular mobility and enabling rapid bond reformation necessary for self-healing capabilities at room temperature. The investigation will differentiate between structured and unstructured water forms, examining the specific interactions that influence bond lifetimes and recovery dynamics under mechanical stress. Previous research indicates that the interactions between structured water and the peptide backbone are critical for maintaining bond lifetimes during stress, and these dynamics will be quantitatively assessed through methods such as differential scanning calorimetry (DSC) and rheological measurements.

The synthesis of these peptide assemblies will be executed under rigorously controlled environmental conditions. The reaction temperature will be maintained at 25°C to optimize amino acid interactions and prevent degradation, while systematic variations in pH (5.5 to 8.5) will be conducted to explore their effects on the ionization states of the amino acids and corresponding **hydrogen bonding** dynamics. Additionally, ionic strength will be regulated at approximately 0.15 M using sodium chloride, which is expected to enhance electrostatic interactions and stabilize the hydrogen bonding network.

To thoroughly investigate these interactions and material properties, a combination of experimental methodologies will be utilized. Dynamic Mechanical Analysis (DMA) will assess mechanical properties such as tensile strength and elasticity, while rheological assessments will evaluate the material's response under stress. Spectroscopic techniques, including NMR spectroscopy and Fourier Transform Infrared (FTIR) spectroscopy, will be employed to elucidate hydrogen bonding dynamics and monitor molecular interactions. This comprehensive approach aims to clarify the intricate relationships among amino acids and structured water dynamics, as well as the influences of environmental conditions on the multifunctional properties of the engineered peptide-based supramolecular materials.

- **Expert's Analysis:** In addressing the research question—"How can a supramolecular amorphous glass material be engineered to simultaneously exhibit strong mechanical properties, optical transparency, self-healing capabilities, and adhesive strength at room temperature?"—the generated scientific hypothesis proposes a notably complex system. This system is envisioned to comprise five distinct amino acids: lysine, glutamic acid, aspartic acid, tyrosine, and proline.

Despite the increased complexity of this multi-component approach compared to the simpler system underlying the real hypothesis (the scientific finding concerning YYY peptide glass), several key conceptual parallels are evident: **Shared Foundation in Peptide-Based Materials:** At their core, both the generated hypothesis and the real scientific finding are centered on peptide-based materials as the fundamental building blocks for the desired amorphous glass. **Convergent Aim for Dynamic Networks and Functional Properties:** Both frameworks leverage their respective peptide systems with the goal of establishing a dynamic network. This network is considered crucial for imbuing the material with critical functionalities, particularly self-healing capabilities and effective adhesive strength. **Emphasis on the Role of Structural Water:** In their mechanistic considerations, both hypotheses significantly highlight the indispensable role of structural water. The real hypothesis (the scientific finding on YYY glass) successfully demonstrated that a dense, random hydrogen-bonding network, mediated by water molecules, is fundamental to the YYY glass's unique structure and its observed properties. The generated hypothesis also underscores the centrality of cooperative and reversible hydrogen bonds in the construction and operational dynamics of its proposed network.

648 **Example 3:**

649 • Research Question: How can computational methods be used to accurately predict and im-
650 prove the reactivity and selectivity of modular diazo transfer (MoDAT) reactions, especially
651 reactions with primary amines? And to design new reagents based on computational models.

652 • Hypothesis Candidate 1 (from HHS): We propose to systematically investigate the reactivity
653 and selectivity of modular diazo transfer (MoDAT) reactions utilizing azide-based reagents,
654 with a specific focus on para-substituted benzyl azide derivatives modified with strong
655 electron-withdrawing groups (EWGs) such as nitro ($-\text{NO}_2$) and cyano ($-\text{CN}$), as well as
656 weaker electron-withdrawing groups (e.g., fluoro ($-\text{F}$) and chloro ($-\text{Cl}$)), and electron-
657 donating groups (EDGs) like methoxy ($-\text{OCH}_3$). Our central hypothesis posits that the
658 electronic nature and precise positioning of these substituents will significantly modulate the
659 electrophilicity of the azide moiety, which will in turn influence the stability and geometrical
660 configurations of intermediates and transition states during nucleophilic attacks by primary
661 amines.

662 The experimental work will be executed under controlled laboratory conditions using a
663 Schlenk line to maintain an inert nitrogen atmosphere for at least 30 minutes prior to reaction
664 initiation, minimizing moisture exposure. Reactions will be conducted at a temperature
665 of 95–105°C, chosen based on literature findings indicating optimal kinetic performance
666 while preserving the stability of diazo intermediates. We will employ polar aprotic solvents
667 such as dimethylformamide (DMF) and dimethyl sulfoxide (DMSO), which are anticipated
668 to enhance solvation of the azide and improve nucleophilicity of the primary amines. A
669 stoichiometric ratio of 1:1.5 (benzyl halide to sodium azide) will be applied, and reactant
670 concentrations will be maintained at approximately 10–20 mM, a range supported by
671 preliminary studies demonstrating optimal reactivity and solubility.

672 To deepen our mechanistic understanding, we will utilize advanced computational tech-
673 niques, primarily Density Functional Theory (DFT) with specific emphasis on the B3LYP
674 functional and a 6-31G(d) basis set. This will allow us to thoroughly assess the impacts
675 of substituent variations on charge distributions and transition state energies. In particular,
676 we will quantify how the introduction of strong EWGs influences the electrophilicity of the
677 azide, predicted to lower activation energy for nucleophilic attacks, while the EDGs may
678 provide stabilization during transition state formation. A rigorous electronic structure analy-
679 sis will be employed to visualize the overlap between the nucleophilic orbital of primary
680 amines and the electrophilic centers of the diazo derivatives.

681 To ensure a comprehensive evaluation of solvent effects, we will systematically conduct
682 reactions in both polar protic solvents, such as ethanol, and polar aprotic solvents, like
683 DMSO, maintaining a concentration of 50% v/v relative to the reaction volume. This will
684 elucidate how solvent polarity and hydrogen-bonding capacity impact the stability of zwitter-
685 ionic intermediates and transition states. Furthermore, we will establish a robust iterative
686 feedback mechanism that directly correlates our computational predictions to empirical data
687 from high-resolution techniques, including NMR spectroscopy and mass spectrometry. Dis-
688 crepancies exceeding a predefined threshold—such as a 10% difference between predicted
689 and observed yields—will prompt specific adjustments to our computational models, thus
690 refining our understanding of the factors influencing MoDAT reactions. This structured
691 approach aims to enhance the design of azide-based reagents tailored for improved reactivity
692 and selectivity in synthetic applications involving primary amines.

693 • Hypothesis Candidate 2 (Greedy + self-consistency): This study aims to investigate the re-
694 activity and selectivity of modular diazo transfer (MoDAT) reactions by designing tailored
695 azide-based reagents that incorporate strategically positioned electron-withdrawing groups
696 (EWGs), specifically nitro ($-\text{NO}_2$) and cyano ($-\text{CN}$) moieties. These groups will be sys-
697 tematically evaluated at distances of 2 Å, 4 Å, and 6 Å from the azide functionality. This
698 selection is grounded in existing literature, which suggests that such variations in EWG
699 positioning can significantly influence the electrophilicity of the azide, thereby stabilizing
700 transition states during nucleophilic interactions with primary amines. The specific rationale
701 for these distances is based on prior studies indicating that alterations beyond these ranges
702 may lead to diminished effects on reactivity or unfavorable sterics.

703 The primary amines selected for this investigation—n-butylamine, phenethylamine, ben-
704 zylamine, and ethylamine—exhibit varying degrees of steric hindrance and electronic
705 characteristics, allowing for a comprehensive analysis of their reactivity profiles when paired
706 with engineered azide reagents. This choice effectively captures a spectrum of nucleophilic

behavior and enhances our understanding of how both electronic and steric factors operate within MoDAT reactions.

An integrated computational strategy will be employed utilizing density functional theory (DFT) calculations focused on the B3LYP functional, supplemented by other functionals (e.g., M06-2X and PBE0) to ensure accuracy and validation of outcomes. In addition to DFT, we will incorporate transition state theory and potentially machine learning approaches to analyze the reactivity patterns further. We will conduct extensive mapping of activation energy profiles and transition state geometries to uncover the relationships between EWG positioning, electron density distributions, and transition state stability. Molecular dynamics (MD) simulations will also be utilized to explore solvation effects across differing solvent environments, including acetonitrile, methanol, and DMSO, examining how these solvents influence reaction pathways and stabilize charged intermediates.

Empirical validation of computational models will incorporate a structured approach to varying critical parameters such as azide reagent concentrations and molar ratios of primary amines to azides, along with solvent compositions, to derive quantitative metrics, including reaction yields, rate constants, and activation energies. Statistical analyses will employ techniques such as ANOVA and regression models to extract significant trends from the experimental data. This iterative feedback mechanism will facilitate a dynamic refinement process, whereby experimental outcomes directly inform adjustments to computational predictions. Through this comprehensive methodological framework, we aim to elucidate the interplay between EWG distances and steric factors, ultimately leading to the design of innovative azide-based reagents optimized for selective transformations of primary amines.

- Hypothesis Candidate 3 (Greedy): This research aims to systematically investigate the reactivity and selectivity of modular diazo transfer (MoDAT) reactions utilizing azide-based reagents, focusing on a set of primary amines: benzylamine, 2-aminopropane, and cyclohexylamine. This selection combines varying steric and electronic profiles, enabling comprehensive evaluation of how solvent and reaction conditions influence reactivity and selectivity across different nucleophilicity ranges. Initial studies will determine baseline reactivities through systematic kinetic measurements, assessing critical parameters such as rate constants and product ratios under controlled conditions.

Reactant concentrations will be evaluated at specific increments of 0.1 M (0.1 M, 0.5 M, and 1.0 M), and the temperature will be optimized through a systematic approach involving stepwise assessments from 25°C to 60°C, analyzing how these variations affect reaction progress. A comprehensive assessment of solvent effects will be performed, including the examination of solvent mixtures (e.g., varying concentrations of water, DMSO, and possible co-solvents) to quantify their influence on nucleophilicity and overall reactivity.

Advanced computational methods, including density functional theory (DFT) calculations with the M06-2X functional and a 6-31G basis set, will be employed to simulate the MoDAT reaction environment accurately. We will analyze key molecular descriptors such as nucleophilicity, electrophilicity, and steric hindrance to construct predictive models of reactivity. These analyses will guide experimental design, with a feedback mechanism where discrepancies between computational predictions and experimental observations will result in specific adjustments to molecular descriptors or computational parameters, refining the predictive capabilities of the models.

Following these investigations, the design of innovative azide-based reagents will be undertaken to optimize MoDAT reactions. This design process will emphasize the incorporation of electron-withdrawing groups like trifluoromethyl and cyano, aimed at enhancing both stability and selectivity by stabilizing the transition state. Rigorous standardized experimental protocols will ensure reproducibility, including specific techniques for measuring yields and selectivity ratios over controlled reaction durations. By integrating mechanistic insights from computational and empirical findings, this research will elucidate the key factors influencing reactivity and selectivity in diazo transfer reactions, enhancing our understanding of these critical processes.

- Expert's Analysis: 1 conducted a relatively comprehensive analysis, for instance, suggesting that modifying the azide reagent with functional groups could improve it, which aligns with the original text. However, 2 and 3 did not. This time, 1 has an obvious error: the speculated temperature is incorrect, and the proposed temperature is experimentally unfeasible, as azide reagents are prone to explosion at high temperatures. Of course, temperature is a minor

point, and overall, 1 is still acceptable. 2 deviates significantly from the original text in terms of the research design approach. Compared to 2, 3 lacks consideration of the group effect in the research design, making 3 the weakest. Finally, all three mentioned using DFT calculations, and although there are deviations in details from the original text, the approach is correct.

Example 4:

- Research Question: How can photoredox catalysis be used to exploit the latent reactivity of phosphorus ylides, allowing them to participate in a formal three-component cycloaddition that converts inert C–H and C=P bonds into C–C and C=C bonds, creating versatile synthetic building blocks in an efficient, controlled manner?

- Hypothesis Candidate 1 (from HHS): The mechanism for activating phosphorus ylides in a formal three-component cycloaddition via photoredox catalysis can be articulated in four key steps, each supported by optimized experimental conditions:

1. Initiation of Single-Electron Transfer (SET): Irradiation of phosphorus ylides with specific wavelengths of visible light (400–450 nm) from a high-intensity LED source (approximately 20 mW/cm²), validated by studies demonstrating effective radical generation at this intensity (Smith et al., 2020), promotes SET using suitable photoredox catalysts (e.g., [Ru(bpy)₃]²⁺ or [Ir(dF(CF₃)ppy)₂(bpy)]). The resulting radical cation exhibits enhanced electrophilicity due to significant charge localization, which is further assisted by strong electron-withdrawing substituents such as carbonyl or nitro groups. Empirical evidence indicates an increase in reactivity by up to 2.5-fold as supported by Hammett parameters.

2. Stabilization via Zwitterionic Intermediate: The radical cation transitions to a zwitterionic intermediate, characterized by resonance stabilization through delocalized π -electrons and non-covalent interactions, such as hydrogen bonding in polar aprotic solvents like acetonitrile (dielectric constant $\approx$ 37) and DMSO (dielectric constant $\approx$ 47). To optimize stabilization, a 1:1 (v/v) mixture of these solvents will be used, taking advantage of their combined dielectric properties ($\approx$ 38) to enhance charge separation and stabilize reactive intermediates. Literature supports this approach, showing improved reaction kinetics (Miller et al., 2021).

3. Selective Nucleophilic Attack: The zwitterionic intermediate selectively engages in nucleophilic attacks on activated C–H and C=P bonds, particularly those adjacent to strong electron-withdrawing groups. Maintaining phosphorus ylide concentrations at 0.1–0.5 M and controlling reaction temperatures precisely within an optimized range of 10–25 °C, as indicated by previous studies on radical stability, will minimize side reactions. An inert atmosphere (nitrogen or argon) will be established by purging the reaction vessel for 30 minutes before use, effectively mitigating oxidation. Real-time NMR (utilizing 1D and 2D techniques) and GC-MS metrics will be employed to monitor yield and product distribution effectively, specifying analytical conditions (e.g., temperature settings and flow rates) to ensure accurate assessment of outcomes.

4. Concerted Formation of Products: The reaction culminates in the concerted formation of new C–C and C=C bonds, facilitating the synthesis of valuable carbocycles and synthetic building blocks. The influence of substituent identity and positioning (ortho, meta, para) on reactivity will be quantitatively analyzed using NMR and HPLC techniques. This systematic approach will provide insight into the efficiency and selectivity of the cycloaddition process, explaining how each factor contributes to overall reactivity.

By integrating these components clearly and methodically, this hypothesis presents a comprehensive exploration of how photoredox catalysis can unveil new reactivity pathways for phosphorus ylides, fully addressing the research question with explicitly defined roles of each mechanistic step and comprehensive definitions for specialized terms provided for clarity.

- Hypothesis Candidate 2 (Greedy + self-consistency): This study aims to investigate how photoredox catalysis can elucidate specific reactivity mechanisms in diphenylphosphinyl ylides, focusing on their participation as intermediates in formal three-component cycloaddition reactions that convert inert alkyl C–H bonds and coordinated C=P bonds into C–C and C=C bonds. We will concentrate on diphenylphosphinyl ylides due to their capacity to facilitate charge inversion processes, which are crucial for generating stabilized radical

intermediates. Specifically, we will examine how the electronic structure of diphenylphosphinyl ylides transforms upon excitation under visible light (e.g., at wavelengths of 400-450 nm) into polar intermediates, leading to the generation of stable benzylic radicals from C–H bond cleavage and allylic radicals from C=P bond activation—both essential for effective cycloaddition.

1. Charge Inversion Mechanisms: We will clarify the mechanisms of charge inversion by detailing the electronic transitions involved, emphasizing the role of excited state lifetimes and how solvent polarity affects the stabilization of generated radicals. Our discussion will include specific references to known photoredox systems and their operational parameters in literature, enhancing our findings with context.

2. Role of π -Stacking Interactions: We will investigate the influence of π -stacking interactions, specifying the optimal inter-ring distances and energy considerations between the stabilized radicals and electrophilic carbon centers from substrates like 2,6-dibromoaryl and 2,4-difluoroaryl. This section will include quantitative assessments of radical coupling reactivity, aiming to determine coupling efficiencies and product selectivity through methods such as ^1H NMR and UV-Vis spectroscopy.

3. Solvent and Temperature Effects: We will analyze solvent effects by conducting reactions in a range of solvents with varied dielectric constants, including a thorough comparison of polar solvents like acetonitrile and non-polar solvents like dichloromethane. We will detail how these solvents are expected to influence radical stabilization and the kinetics of cycloaddition, providing a theoretical framework based on established solvent interaction models. Additionally, we will justify the controlled temperature range of room temperature to 50°C by linking it to the expected stability of radical intermediates and the kinetics of the reactions, ensuring optimal conditions for product yield and selectivity.

4. Experimental Conditions: Reactions will be conducted under carefully controlled conditions, with temperature justification focusing on maintaining the balance between radical stability and reaction kinetics. We will outline how these conditions directly relate to the completed cycloaddition mechanism and the anticipated outcomes of the study.

By systematically elucidating these mechanisms—specifically charge inversion, π -stacking interactions, and solvent effects—we aim to develop robust methodologies for the efficient generation of versatile synthetic building blocks from simple molecular precursors.

- Hypothesis Candidate 3 (Greedy): Investigate how photoredox catalysis enhances the reactivity of phosphorus ylides through selective nucleophilic attack on α, β -unsaturated carbonyl compounds, such as crotonaldehyde, which are activated to form stable radical cation intermediates via photoredox-driven single-electron transfer (SET) processes. These radical cations, characterized by their electrophilicity, promote effective nucleophilic attacks by phosphorus ylides, generating stabilized carbon radical intermediates that significantly enhance their reactivity in subsequent bond-forming transformations. Conduct a formal three-component cycloaddition by introducing a nucleophilic amine, such as ammonia or an aniline derivative, selected based on its electronic properties which influence the stabilization of the radical intermediates and affect product selectivity. Detail specific optimized reaction conditions, including the use of polar aprotic solvents like acetonitrile, which facilitate radical stability, and employ a specific light wavelength of 400 nm to ensure efficient excitation of the photocatalyst. These conditions will be designed to minimize potential side reactions and maximize the conversion of inert C–H and C=P bonds into desired C–C and C=C bonds through well-defined mechanistic pathways, addressing the nuanced interplay between reaction parameters and final product outcomes.
- Expert's Analysis: 1 accurately predicted the light source wavelength range, metal catalyst system, and solvent system, such as the use of Ir catalyst and acetonitrile as the solvent, all of which align with the original text. In contrast, 2 only correctly predicted the wavelength range and solvent system but failed to specify the metal catalyst system, which is crucial in organic chemical reactions. Therefore, 2 is inferior to 1. Finally, 3 did not predict the light source wavelength range or the metal catalyst system, missing several key pieces of information, making it the weakest.

G Hypothesis Search Prompt

The following prompt is used to guide both the baseline methods and our proposed method, *HHS*, during the hypothesis refinement process. To ensure fair comparison, the prompt is designed in a controlled way: we use a shared core prompt across all methods, with minimal differences. Specifically, the portion highlighted in orange is unique to *HHS* and introduces the hierarchical structure used in its search process.

This design isolates the effect of hierarchical search. As illustrated below, the only difference between *HHS* and the baseline (*Greedy Search* + *Self-Consistency*) lies in the hierarchical prompting. The core content—including the role of the assistant, editing instructions, and structural expectations—is kept identical.

This enables a controlled ablation-style comparison, attributing observed improvements specifically to the hierarchical design.

The complete prompt is as follows:

You are assisting with scientist’s research. Given their research question, a survey on the past methods for the research question, and a preliminary coarse-grained research hypothesis for the research question, please help to make modifications into the coarse-grained hypothesis, to make it one step closer to a more effective and more complete fine-grained hypothesis.

The modification can be two-folds: (1): delete or change an existing improper detail or information in the existing hypothesis; (2) add and integrate one detail to the existing hypothesis. If you choose to add a detail, do not simply append new information to the existing hypothesis. Instead, think thoroughly how the new detail relates to the existing components and integrate it seamlessly into the hypothesis to create a new coherent and unified hypothesis. In addition if you choose to add a detail to a general information, if the corresponding general information is correct, you should try to keep the corresponding general information in the updated hypothesis and also mention the details, instead of replacing the general information with the details. In this way, it would be much easier for scientists to understand both the general information/structure and the details from your generated hypothesis. It would be also easier for scientists to propose better details, inspired by your suggested details, following the general information.

Please remind that this is about research: research is about discover a new solution to the problem that ideally is more effective and can bring new insights. Usually we don’t need the hypothesis to contain lots of known tricks to make it work better: we want to explore the unknown, which ideally is more effective than the known methods and can also bring in new insights. Therefore, a research hypothesis is usually about a small set (usually less than eight) of major components (and lots of details on how to implement these major components), which overall composes a novel and complete solution to the research question, which potentially can bring in new insights. Hypotheses that include an excessive number of irrelevant or unnecessary major components, which do not contribute to addressing the research question, are less favorable, as we only want to know exactly what are the key components that fundamentally make the hypothesis work. If you think any ancillary components that can truly assist with the research question, you may mention what are the key components and what are the ancillary components to avoid the ambiguity of which components are the key component. The reaction mechanism, however, is not classified as a major component or detail (and therefore not limited by the number of major components). Instead, a novel and valid reaction mechanism can be a good source of insights. If previous hypothesis already contains too many major components, you should consider to replace some of the major components with more effective ones (but not to add more major components), or to give more details to the existing major components for clarity and ease of implementation (instead of adding or replacing major components).

Here we are searching for the fine-grained hypothesis in a hierarchical way. The rationale is that, we can classify any complete set of modifications into several hierarchy, with different levels of details. If we do not search in a hierarchical way, we need to consider all the available details in all hierarchy levels for each search step, which (1) has a very high complexity, and (2) first search a low-level detail might largely influence the following search of a high-level detail: it might stuck in one high-level detail corresponding to the already searched low-level detail without considering the other low-level details corresponding to other high-level details, making the search process stuck in a local minimum at the beginning.

Here we roughly classify all possible modifications into five hierarchies: (1) *Mechanism of the Reaction*: Describes how the reaction proceeds at a conceptual level, focusing on electron flow, bond formation and breaking, and any intermediates or transition states involved. This is the theoretical “blueprint” that explains why the reaction works; (2) *General Concept or General Component Needed*: Identifies the type of reagent or functional group required (e.g., “a strong acid,” “a Lewis base,” “an activated aromatic ring”) without committing to a specific chemical. It outlines the broader roles that are necessary for the mechanism to proceed; (3) *Specific Components for the General Concept*: Narrows down from the general category to a particular substance (e.g., “concentrated HCl” for a strong acid, “benzene” for an aromatic ring). This makes the reaction hypothesis testable by specifying which chemicals fulfill the roles; (4) *Full Details of the Specific Components*: Provides exact structural or molecular information—such as SMILES strings, IUPAC names, purity, or CAS numbers. These details ensure clarity and reproducibility so researchers know precisely which substances to use; (5) *Experimental Conditions*: Specifies the practical setup—temperature, pressure, solvent system, reaction time, atmosphere, and any work-up procedures. This final layer describes how to carry out the reaction in a laboratory setting. And we are searching for modifications hierarchy by hierarchy: hierarchy (1) first, and then hierarchy (2), and so on. Hypothesis from a higher hierarchy is an expansion of the hypothesis from its previous hierarchy, with additional information described above.

The research question is:

The survey is:

Now please help to make modifications into the coarse-grained hypothesis, to make it one step closer to a more effective and more complete fine-grained hypothesis. Please do not include the expected performance or the significance of the hypothesis in your generation. Please answer the question in the following response format. (response format: ‘Reasoning Process: Revised Hypothesis: ’)

H Expert Analysis of Hypothesis Quality

H.1 Convergence to Ground-Truth Local Optima

To complement our quantitative evaluations, we asked domain experts to qualitatively assess how well the hypotheses generated by HHS aligned with the expert-annotated fine-grained hypotheses in our benchmark.

The distribution of expert assessments across all evaluated examples is as follows:

- Reached a completely different region—likely a distinct local optimum with scientific plausibility: 60%
- Reached the vicinity of the ground-truth local optimum, but with differing details: 24%
- Reached the vicinity of the ground-truth local optimum, but failed to fully elaborate or specify the key details: 16%

Here, the “ground-truth local optimum” refers to the expert-extracted fine-grained hypothesis from a publication, which serves as the reference target. “Reaching the vicinity of a local optimum” indicates that the generated hypothesis converges to a coherent and internally consistent formulation that is conceptually close to the ground-truth hypothesis, though not necessarily identical in detail.

The relatively high divergence rate (60%) reflects an inherent tradeoff in our experimental setup. For many research questions, multiple hypotheses can be plausible yet structurally distinct. Guiding the model toward the exact ground-truth hypothesis requires:

1. initializing the search process from a starting point sufficiently close to the ground-truth optimum;
2. but avoiding initialization that is too close or too specific, as this would risk leaking the ground-truth answer.

To strike this balance, we derive the initial search point from the annotated coarse-grained hypothesis h_c by applying an *ambiguation* procedure. This involves removing or abstracting key details to

978 produce a generalized version of h_c —for example, replacing “a specific protein” with “a protein” or
979 “a catalyst”—thus preserving the overall research direction while preventing answer leakage.

980 Consequently, even when the search begins in the correct conceptual region, the model may naturally
981 diverge toward a nearby but distinct local optimum, especially given the openness of the hypothesis
982 space and the heuristic-driven nature of the optimization process.

983 H.2 Coverage of Experimentally Critical Details

984 In addition to alignment with the reference hypotheses, we evaluated the extent to which the generated
985 hypotheses captured the critical experimental details required for practical implementation.

986 Among all the details mentioned in the generated hypotheses, approximately **40% are experimentally**
987 **important**—regardless of their accuracy (which is not the focus of this analysis). The remaining
988 **60% are peripheral or have minimal impact** on the actual experiment.

989 Among all the important details that should be included, about **50% are mentioned** in the generated
990 hypotheses.

991 Peripheral details refer to contextual or environmental factors with limited relevance to the core
992 experiment—for instance, ambient humidity or weather conditions, which may only affect specific
993 reactions.

994 This highlights a key challenge: while LLMs can generate rich and context-aware hypotheses, they
995 often fail to prioritize the most essential components for experimental planning. Future work may
996 explore techniques to guide LLM attention toward experimentally salient information.

997 I Experiment Compute Resources

998 We implement our proposed framework as an agentic LLM system using GPT-4o-mini using
999 OpenAI’s official API. Generating the final hypothesis via the *HHS* optimization process—converging
1000 to the final local optimum at hierarchy level 5—typically involves several hundred or even to a
1001 thousand iterative search steps.

1002 J Limitation

1003 While *HHS* consistently discovers higher-quality local optima compared to baseline methods, it
1004 does not guarantee convergence to the global optimum. Addressing this limitation remains an open
1005 direction for future research.