

GRAIL: An Agentic AI Architecture for Interactive Grant Proposal Writing

Zhisheng Tang^{1,2}, Mayank Kejriwal^{1,2}

¹University of Southern California, Los Angeles, CA

²GRAIL, Los Angeles, CA

zhisheng@isi.edu, kejriwal@isi.edu

Abstract

Securing research funding remains fragmented and time-consuming: researchers must navigate separate databases across dozens of agencies while simultaneously drafting competitive proposals. We present GRAIL, a web-based platform that unifies grant discovery and proposal writing through conversational AI. Users describe their research interests in natural language to explore opportunities from a unified index of 11.8K U.S. federal and non-profit grant opportunities; within the document editor, integrated AI assistance supports real-time proposal revision and refinement. The system runs in any modern browser without installation. Conference attendees are invited to interact with the live system at the demo booth, explore the grant discovery and writing assistance workflows, and provide feedback on the user experience.

1 Introduction

Securing research funding is fundamental to academic innovation [Alberts *et al.*, 2014; Edwards and Roy, 2017], yet the grant proposal process remains fragmented and time-consuming [Gross and Bergstrom, 2019; Sandström and Van den Besselaar, 2018]. Researchers must identify opportunities across dozens of agencies and foundations while simultaneously navigating the demanding task of articulate proposal writing [Wessely, 1998; Langfeldt, 2001]. The stakes are high: while strong writing directly influences funding outcomes [Thelwall *et al.*, 2023], the process consumes valuable time that could be devoted to actual research. Recent advances in large language models (LLMs) have transformed scientific production by reducing barriers to knowledge exploration and interdisciplinary pivoting [Kusumegi *et al.*, 2025; Hill *et al.*, 2025], while quantifying both opportunities and challenges in this new research landscape [Gao and Wang, 2024; Liang *et al.*, 2025].

Grant proposal writing presents two critical bottlenecks that researchers struggle to overcome. First, **funding discovery** is inefficient. Researchers must navigate separate portals for the National Science Foundation (NSF), National Institutes of Health (NIH), DARPA, Department of Energy,

and others, each with different search interfaces, terminology, and application requirements. A researcher might spend hours searching multiple databases only to discover marginal matches, while missing relevant opportunities that use different terminology. No unified index exists; instead, fragmented databases force researchers to become database administrators rather than focusing on research ideas. While recent advances in AI-assisted tools have improved writing support [Seckel *et al.*, 2024; Godwin *et al.*, 2024], the fundamental problem of discovering relevant opportunities across fragmented systems remains largely unsolved.

Second, **proposal writing** is a labor-intensive, isolated task. Effective grant writing requires substantial revision and feedback, as shown by guidance from funding agencies themselves [Sohn, 2020; Botham *et al.*, 2017; Yuan *et al.*, 2016]. However, writing assistance tools remain disconnected from collaborative editing, and researchers typically draft proposals in isolation or through cumbersome back-and-forth email workflows. The absence of real-time, integrated assistance means researchers must rely on delayed feedback from collaborators, missing opportunities for iterative refinement and quality improvement [Gero *et al.*, 2023; Dhillon *et al.*, 2024; Weng *et al.*, 2025].

Recent advances in LLMs have demonstrated promise for supporting scientific writing, with studies showing improvements in writing quality, productivity, and creative exploration [Conroy, 2023; Huang and Tan, 2023; Li *et al.*, 2024; Wan *et al.*, 2024; Liang *et al.*, 2024]. However, research shows that AI remains most effective as a supplement to human judgment rather than a replacement [Quitadamo and Kurtz, 2007]. Trust calibration is essential for appropriate reliance on AI systems [Lee and See, 2004; He *et al.*, 2023; Ma *et al.*, 2024; Chong *et al.*, 2022], as users must develop nuanced understanding of both AI capabilities and limitations. Effective human-AI collaboration in writing contexts requires careful system design that supports appropriate task delegation and scaffolding [Lee *et al.*, 2024; Lubars and Tan, 2019; Dhillon *et al.*, 2024]. Meanwhile, funding agencies have established frameworks for AI use in proposals: the NSF and American Heart Association allow AI-assisted drafting with transparency [National Science Foundation, 2023; Association, 2024], while others warn against specific risks like confidentiality breaches [Lauer *et al.*, 2023]. Beyond policy compliance, concerns about au-

thorship and authenticity arise when integrating AI into creative and scholarly work [Hwang *et al.*, 2025; Guo *et al.*, 2025]. Multi-agent systems have shown promise for complex collaborative tasks [Yao *et al.*, 2023; Wu *et al.*, 2023; Chen *et al.*, 2024], suggesting that coordinated AI agents could support both discovery and drafting workflows.

GRAIL addresses these bottlenecks by unifying funding discovery and proposal writing through an integrated AI agent system. Users interact with a conversational agent in a global chat panel to describe their research interests, and the agent searches a unified index of U.S. federal and non-profit opportunities to identify promising matches. Within the document editor itself, users can request real-time writing assistance, and the agent revises text in place while maintaining collaborative awareness. Rather than a detached chatbot, the agent operates as an integrated collaborator that understands both the broader research context and the specific content being drafted.

The platform also incorporates grant management features including tagging and calendar management, enabling researchers to organize opportunities and track important deadlines alongside their active proposals.

2 System Architecture & AI Capabilities

GRAIL is a web-based platform organized around three core layers: (1) a collaborative editor interface for document creation and revision, (2) an AI agent backend that handles user requests and tool execution, and (3) a data layer managing projects, conversations, and a unified index of 11.8K grant opportunities. Currently serving more than 3,000 users with 1,000+ collaborative projects, the platform enables interaction through two primary interfaces: a **global chat panel** for conversational grant discovery and project management, and an **in-editor sidebar** for real-time writing assistance within the document itself.

2.1 Conversational Grant Discovery

GRAIL maintains a unified index of funding opportunities from U.S. federal agencies (NSF, NIH, DARPA, DOE, NASA, NEH, USDA) and nonprofit foundations. Users describe their research in natural language (e.g., “I work on machine learning for wildfire detection”), and the agent automatically extracts keywords, searches the index, and presents five to ten relevant opportunities with titles, descriptions, and direct links to funding agency pages.

The agent follows a structured workflow: (1) interpret the user’s research domain, (2) extract domain-specific search keywords, (3) query the unified grant index, (4) summarize each opportunity’s description, and (5) refine the search if results are insufficient. Users can iteratively refine by asking follow-up questions like “Show me NSF programs” or “Look for international funding.”

2.2 In-Document Writing Assistance

Within a project document, users can request real-time revisions through the in-editor sidebar. The agent can: rewrite sections while preserving formatting and voice, answer questions about content, search across all project files to maintain

consistency, and suggest structural improvements. Revised text appears in the conversation panel, while the agent intelligently searches and analyzes all relevant files and project information to ensure coherence and consistency.

2.3 Project Organization

Users can create and manage projects with tags, organize multiple documents, and track progress. The agent assists with project setup, file management, and provides overviews of project structure and current status.

3 Interface & Use Cases

3.1 Grant Discovery Example

A researcher asks: “I’m working on machine learning for wildfire detection using satellite imagery. What grants could I apply for?” The agent extracts keywords, queries the grant index, and presents opportunities like NSF’s “Prediction of and Resilience against Extreme Events,” NASA’s “Earth Science Applications: Wildfires,” and DOE programs on environmental monitoring. If results are incomplete, the user can refine by asking “Show me programs with longer deadlines” or the agent can automatically search related areas like “climate resilience.”

4 AI Techniques & Innovations

GRAIL employs several key AI techniques and innovations:

Agentic LLM reasoning. An LLM-powered agent reasons over user requests in an iterative loop: construct a prompt from system instructions and conversation history, stream to the LLM, execute any tool calls returned, append results, and re-invoke the LLM until a final response is produced. This enables flexible, multi-step reasoning beyond simple text generation.

Context-aware tool orchestration. The agent’s available tools change dynamically based on context—grant search and web search in the global chat panel, but document analysis and project information retrieval in editor mode. This reduces prompt complexity and focuses the agent on relevant capabilities.

Domain-specific keyword extraction. For grant discovery, the agent extracts up to ten domain-specific keywords from natural language research descriptions, translating researcher vocabulary into effective search terms. The agent may also refine keywords iteratively if initial results are insufficient.

Document-aware writing assistance. When editing documents, the agent reads document structure and previous sections to maintain stylistic consistency, preserve formatting, and avoid AI-characteristic vocabulary.

Real-time streaming architecture. All LLM inference is fully streamed—content tokens, intermediate reasoning, tool invocations, and results flow to the frontend in real time, enabling responsive UX and transparency into agent reasoning.

4.1 Core Innovations

Beyond the technical implementation, GRAIL advances the state of the art through four core innovations:

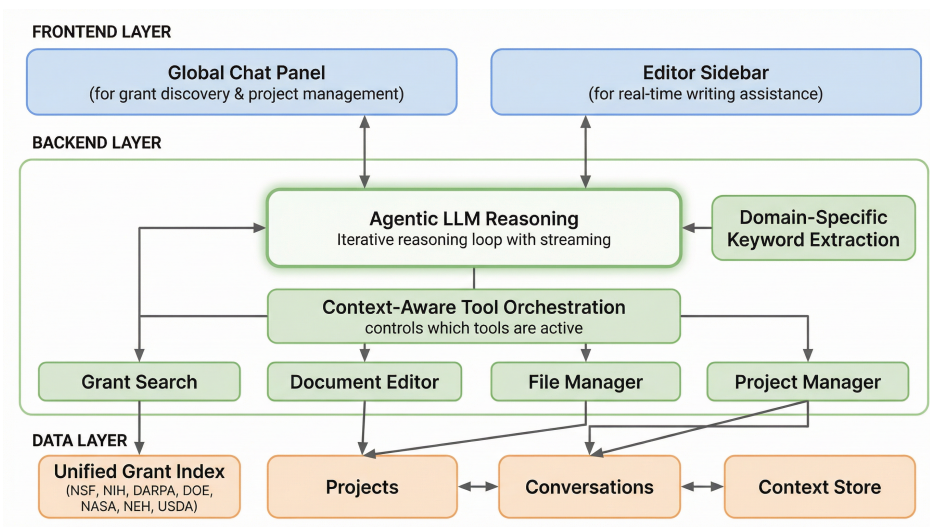


Figure 1: GRAIL system architecture: Frontend interface with chat panel and editor sidebar, agentic LLM reasoning with context-aware tool orchestration, and integrated data layer for grant indexing and project management.

1. Unified integration of discovery and drafting. Existing workflows force researchers to navigate separate tools: Grants.gov or proprietary databases for grant search, then Google Docs or Overleaf for writing, with separate AI writing assistants disconnected from both. GRAIL unifies these workflows within a single platform. The agent maintains coherent context across both activities—understanding which grants a researcher is exploring, what their research focus is, and applying that knowledge when revising proposals. This eliminates context-switching and ensures grant-specific relevance throughout the writing process.

2. Context-aware assistance across the entire project. Generic writing assistants (ChatGPT, Grammarly) operate on individual documents or paragraphs with no project awareness. GRAIL’s in-document agent can search across all project files to maintain internal consistency: when revising a specific section, it reads your research statement, methodology, related proposal sections, and project metadata to ensure coherence. Revisions are tailored not just to the immediate text but to your entire proposal’s narrative and voice.

3. Agentic tool-use orchestration. Rather than choosing between pure chat interfaces and point-and-click UIs, GRAIL employs an LLM-based agent that reason over user requests and autonomously selects which tools to invoke—grant search, document editing, file retrieval, project management. This single conversational interface mirrors recent advances in LLM agents (ReAct, AutoGPT), allowing users to naturally describe what they want (“Find grants for climate research and draft an abstract”) while the agent plans and executes a multi-step workflow.

4. Unified grant index eliminating database fragmentation. Federal funding is scattered across separate portals: Grants.gov, NSF FastLane, NIH RePORTER, DARPA, DOE, NASA, NEH, USDA—each with different search interfaces and terminology. GRAIL indexes opportunities from all major U.S. federal and nonprofit funders in a single, naturally

searchable database. Researchers no longer navigate multiple login portals and inconsistent search paradigms; instead, natural language queries retrieve relevant opportunities across all major funding bodies.

5 Interactive Demo Elements

Attendees at the demo booth can experience the full GRAIL workflow: live grant discovery by describing research in natural language and refining results through follow-up questions, in-document editing with real-time agent rewrites of proposal text, iterative search refinement across multiple passes, multi-user collaboration with synchronized editing and conflict resolution, and project organization with agent-assisted setup.

6 Conclusions

GRAIL is accessible at <https://grailai.io/> with no installation required; it works in any modern web browser. A demo video covering all major features is available at <https://drive.google.com/file/d/1UfdeggLs61fSupxoVMtHyEz5vuGpDYqC/view?usp=sharing>, demonstrating grant discovery workflow, in-editor writing assistance, iterative refinement, multi-user collaboration, and project management. By integrating an AI agent directly into the document editor, GRAIL eliminates the friction of navigating fragmented funding databases and enable researchers to focus on proposal quality rather than administrative overhead.

References

- [Alberts *et al.*, 2014] Bruce Alberts, Marc W. Kirschner, Shirley Tilghman, and Harold Varmus. Rescuing us biomedical research from its systemic flaws. *Proceedings of the National Academy of Sciences*, 111(16):5773–5777, 2014.

- [Association, 2024] American Heart Association. 2024 aha postdoctoral fellowship. <https://professional.heart.org/en/research-programs/aha-funding-opportunities/postdoctoral-fellowship>, 2024. Accessed: 2024-01-15.
- [Botham *et al.*, 2017] Crystal M. Botham, Joshua A. Arribere, Sky W. Brubaker, and Kevin T. Beier. Ten simple rules for writing a career development award proposal. *PLoS Computational Biology*, 13(12):e1005863, 2017.
- [Chen *et al.*, 2024] Weize Chen, Yusheng Su, Jingwei Zuo, Cheng Yang, and Chenfei Yuan. Agentverse: Facilitating multi-agent collaboration and exploring emergent behaviors. In *The Twelfth International Conference on Learning Representations*, 2024.
- [Chong *et al.*, 2022] Leah Chong, Guanglu Zhang, Kosa Goucher-Lambert, Kenneth Kotovsky, and Jonathan Cagan. Human confidence in artificial intelligence and in themselves: The evolution and impact of confidence on adoption of ai advice. *Computers in Human Behavior*, 2022.
- [Conroy, 2023] Gemma Conroy. Scientists used chatgpt to generate an entire paper from scratch—but is it any good? *Nature*, 619(7970):443–444, 2023.
- [Dhillon *et al.*, 2024] Paramveer S. Dhillon, Somayeh Molaei, Jiaqi Li, Maximilian Golub, Shaochun Zheng, and Lionel Peter Robert. Shaping human-ai collaboration: Varied scaffolding levels in co-writing with language models. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, 2024.
- [Edwards and Roy, 2017] Marc A. Edwards and Siddhartha Roy. Academic research in the 21st century: Maintaining scientific integrity in a climate of perverse incentives and hypercompetition. *Environmental Engineering Science*, 34(1):51–61, 2017.
- [Gao and Wang, 2024] Jian Gao and Dashun Wang. Quantifying the use and potential benefits of artificial intelligence in scientific research. *Nature human behaviour*, 8(12):2281–2292, 2024.
- [Gero *et al.*, 2023] Katy Ilonka Gero, Tao Long, and Lydia B Chilton. Social dynamics of ai support in creative writing. In *Proceedings of the 2023 CHI conference on human factors in computing systems*, pages 1–15, 2023.
- [Godwin *et al.*, 2024] R. C. Godwin, J. J. DeBerry, B. M. Wagener, D. E. Berkowitz, and R. L. Melvin. Grant drafting support with guided generative ai software. *SoftwareX*, 27:101784, 2024.
- [Gross and Bergstrom, 2019] Kevin Gross and Carl T. Bergstrom. Contest models highlight inherent inefficiencies of scientific funding competitions. *PLOS Biology*, 17(1):e3000065, 2019.
- [Guo *et al.*, 2025] Alicia Guo, Shreya Sathyanarayanan, Leijie Wang, Jeffrey Heer, and Amy X. Zhang. From pen to prompt: How creative writers integrate ai into their writing practice. In *Proceedings of the 2025 Conference on Creativity and Cognition*, 2025.
- [He *et al.*, 2023] Gaole He, Lucie Kuiper, and Ujwal Gadiraju. Knowing about knowing: An illusion of human competence can hinder appropriate reliance on ai systems. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*, 2023.
- [Hill *et al.*, 2025] Ryan Hill, Yian Yin, Carolyn Stein, Xizhao Wang, Dashun Wang, and Benjamin F Jones. The pivot penalty in research. *Nature*, 642(8069):999–1006, 2025.
- [Huang and Tan, 2023] J. Huang and M. Tan. The role of chatgpt in scientific communication: writing better scientific review articles. *American Journal of Cancer Research*, 13(4):1148, 2023.
- [Hwang *et al.*, 2025] Angel Hsing-Chi Hwang, Q. Vera Liao, Su Lin Blodgett, Alexandra Olteanu, and Adam Trischler. It was 80 percent me, 20 percent ai: Seeking authenticity in co-writing with large language models. *Proceedings of the ACM on Human-Computer Interaction*, 2025.
- [Kusumegi *et al.*, 2025] Keigo Kusumegi, Xinyu Yang, Paul Ginsparg, Mathijs de Vaan, Toby Stuart, and Yian Yin. Scientific production in the era of large language models. *Science*, 390(6779):1240–1243, 2025.
- [Langfeldt, 2001] Liv Langfeldt. The decision-making constraints and processes of grant peer review, and their effects on the review outcome. *Social Studies of Science*, 31(6):820–841, 2001.
- [Lauer *et al.*, 2023] Mike Lauer, Stephanie Constant, and Amy Wernimont. Using ai in peer review is a breach of confidentiality. *NIH Office of Extramural Research Blog*, 2023.
- [Lee and See, 2004] John D. Lee and Katrina A. See. Trust in automation: Designing for appropriate reliance. *Human Factors*, 2004.
- [Lee *et al.*, 2024] Mina Lee, Katy Ilonka Gero, John Joon Young Chung, Simon Buckingham Shum, and Vipul Raheja. A design space for intelligent and interactive writing assistants. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, pages 1–35, 2024.
- [Li *et al.*, 2024] Zhuoyan Li, Chen Liang, Jing Peng, and Ming Yin. The value, benefits, and concerns of generative ai-powered assistance in writing. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, 2024.
- [Liang *et al.*, 2024] Weixin Liang, Yuhui Zhang, Hancheng Cao, Binglu Wang, Daisy Yi Ding, Xinyu Yang, Kailas Vondrahalli, Siyu He, Daniel Scott Smith, Yian Yin, Daniel A. McFarland, and James Zou. Can large language models provide useful feedback on research papers? a large-scale empirical analysis. *NEJM AI*, 2024.
- [Liang *et al.*, 2025] Weixin Liang, Yaohui Zhang, Zhengxuan Wu, Haley Lepp, Wenlong Ji, Xuandong Zhao, Hancheng Cao, Sheng Liu, Siyu He, Zhi Huang, et al. Quantifying large language model usage in scientific papers. *Nature Human Behaviour*, 9(12):2599–2609, 2025.

- [Lubars and Tan, 2019] Brian Lubars and Chenhao Tan. Ask not what ai can do, but what ai should do: Towards a framework of task delegability. In *Advances in Neural Information Processing Systems*, 2019.
- [Ma *et al.*, 2024] Shuai Ma, Xinru Wang, Ying Lei, Chuhan Shi, Ming Yin, and Xiaojuan Ma. Are you really sure? understanding the effects of human self-confidence calibration in ai-assisted decision making. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, 2024.
- [National Science Foundation, 2023] National Science Foundation. Notice to research community: Use of generative artificial intelligence technology in the nsf merit review process, 2023.
- [Quitadamo and Kurtz, 2007] Ian J. Quitadamo and Martha J. Kurtz. Learning to improve: using writing to increase critical thinking performance in general education biology. *CBE—Life Sciences Education*, 6(2):140–154, 2007.
- [Sandström and Van den Besselaar, 2018] Ulf Sandström and Peter Van den Besselaar. Funding, evaluation, and the performance of national research systems. *Journal of Informetrics*, 12(1):365–384, 2018.
- [Seckel *et al.*, 2024] E. Seckel, B. Y. Stephens, and F. Rodriguez. Ten simple rules to leverage large language models for getting grants. *PLoS Computational Biology*, 20(3):e1011863, 2024.
- [Sohn, 2020] Emily Sohn. Secrets to writing a winning grant. *Nature*, 577(7788):133–135, 2020.
- [Thelwall *et al.*, 2023] Mike Thelwall, Kayvan Kousha, Mahieddine Abdoli, Emma Stuart, Meiko Makita, Cristina Font-Julián, Paul Wilson, and Jonathan Levitt. Is research funding always beneficial? a cross-disciplinary analysis of uk research 2014-20. *Quantitative Science Studies*, 4(2):501–534, 2023.
- [Wan *et al.*, 2024] Qian Wan, Siying Hu, Yu Zhang, Piao-hong Wang, Bo Wen, and Zhicong Lu. It felt like having a second mind: Investigating human-ai co-creativity in prewriting with large language models. *Proceedings of the ACM on Human-Computer Interaction*, 2024.
- [Weng *et al.*, 2025] Yixuan Weng, Minjun Zhu, Guangsheng Bao, Hongbo Zhang, Jindong Wang, Yue Zhang, and Linyi Yang. Cycleresearcher: Improving automated research via automated review. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [Wessely, 1998] Simon Wessely. Peer review of grant applications: what do we know? *The Lancet*, 352(9124):301–305, 1998.
- [Wu *et al.*, 2023] Qingyun Wu, Gagan Bansal, Jieyu Zhang, Yiran Wu, Beibin Li, and Erkang Zhu. Autogen: Enabling next-gen llm applications via multi-agent conversation, 2023. arXiv preprint arXiv:2308.08155.
- [Yao *et al.*, 2023] Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. React: Synergizing reasoning and acting in language models. In *International Conference on Learning Representations*, 2023.
- [Yuan *et al.*, 2016] Ke Yuan, Lei Cai, S. P. Ngok, L. Ma, and Crystal M. Botham. Ten simple rules for writing a postdoctoral fellowship. *PLoS Computational Biology*, 12(7):e1004934, 2016.