

RAGent: A Self-Learning RAG Agent for Adaptive Data Science Education

Mariia Vetluzhskikh ✉ 

Science Academy, University of Maryland, College Park, MD, USA

Fardina Fathmiul Alam¹ ✉ 

Department of Computer Science, University of Maryland, College Park, MD, USA

Abstract

Undergraduate data science education faces a scalability challenge: addressing a high volume of diverse student questions stemming from varying levels of prior knowledge, technical skills, and learning styles – while ensuring timely and accurate responses. Traditional solutions like manual replies or generic chatbots often fall short in terms of contextual relevance, speed, and efficiency. To tackle this, we introduce RAGent, a Retrieval-Augmented Generation (RAG) agent tailored for a university-level data science course at the University of Maryland. RAGent integrates course-specific materials – lecture notes, assignments, and syllabi – to deliver fast, context-aware answers while maintaining low computational overhead. A central innovation of RAGent is its query classification system, which categorizes student questions into: (i) directly answerable, (ii) relevant but unresolved (requiring instructor input), and (iii) irrelevant or out-of-scope. This system uses semantic similarity, keyword relevance, and dynamic thresholds to drive a targeted prompting strategy, enhancing response accuracy. Another key feature is RAGent’s self-learning loop, which continuously improves performance by integrating resolved queries into its knowledge base and flagging unresolved ones for review and retraining. This dual mechanism ensures both immediate adaptability and long-term scalability. We evaluate RAGent using standard NLP metrics (accuracy, precision, recall, F1-score) and report strong performance in filtering and answering student queries. In a user study with 125 students, over 94% expressed a desire to keep RAGent in the course, citing improved clarity and helpfulness. These results suggest that RAGent significantly enhances support in data science education by providing accurate, contextual responses and reducing instructor workload – offering a scalable, adaptive alternative to conventional support methods. Future work will explore deployment across additional courses and institutions to further validate the RAGent’s adaptability.

2012 ACM Subject Classification Applied computing → Education; Information systems → Information retrieval; Computing methodologies → Artificial intelligence; Computing methodologies → Machine learning

Keywords and phrases RAG, Agent, Chatbot, Data Science, Education, Query Classification, Information Retrieval, LLM

Digital Object Identifier 10.4230/OASICS.ICPEC.2025.8

1 Introduction

Data-science courses attract students from computer science, mathematics, biology, and other disciplines [11]. This disciplinary breadth brings uneven technical preparation, posing instructional challenges. Heterogeneous backgrounds and fast-moving content generate a steady flow of conceptual, technical, and administrative questions. When instructors or TAs cannot reply promptly and precisely, learning process suffers. Static resources – discussion forums and FAQs – are often either overwhelming or too generic to capture question-specific nuance. Scalable, course-aware support is needed so staff can focus on

¹ corresponding author



© Mariia Vetluzhskikh and Fardina Fathmiul Alam;
licensed under Creative Commons License CC-BY 4.0

6th International Computer Programming Education Conference (ICPEC 2025).

Editors: Ricardo Queirós, Mário Pinto, Filipe Portela, and Alberto Simões; Article No. 8; pp. 8:1–8:10

OpenAccess Series in Informatics



OASICS Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl Publishing, Germany

higher-order interactions. Advances in large language models (LLMs) [12] and Retrieval-Augmented Generation (RAG) [9] offer such scalability. Although prior work validates LLMs in education [17, 13], generic models lack course grounding, and student queries are frequently ambiguous. We introduce RAGent, a specialized RAG assistant for University of Maryland undergraduate data-science courses. RAGent couples GPT-4o with a knowledge base built from lecture notes, the syllabus, and supplementary materials. Initially piloted in graduate courses, it is evaluated here with undergraduates to test cross-population adaptability. This work contributes: (1) an implementation of a university-level RAG agent built with LangChain, FAISS, GPT-4o, and Streamlit; (2) a multi-dimensional query classifier labeling questions as relevant-known, relevant-unknown, or irrelevant; (3) a self-learning loop capturing unanswered queries for knowledge base expansion; (4) a rigorous evaluation framework with categorized question sets and metrics; and (5) an analysis of feedback from 125 undergraduate students on usability and experience.

2 Background and Related Work

2.1 AI Chatbots in STEM Education

Recent advances in NLP and LLMs have accelerated the integration of virtual assistants in higher education, particularly benefiting STEM disciplines with their rapidly evolving content [3, 2]. Notable implementations include Georgia Tech’s Jill Watson, which manages administrative queries with 76.7% accuracy [14] but lacks sophisticated query classification for technical support. Another noteworthy example is OwlMentor from Saarland University, which employs RAG for scientific literature comprehension [15] without self-learning capabilities. Commercial systems have also made significant inroads in educational AI. McGraw Hill’s ALEKS, which creates personalized learning paths in mathematics and chemistry [10] but requires domain-specific algorithm development for expansion to new subjects. Virtual Agent developed by IT department of University of Maryland reliably answers routine questions and gracefully flags out-of-scope queries [16] functioning as a multimodal tutoring system that supports both text and visual interactions, yet it does not include the continuous self-learning mechanism incorporated in RAGent. Recent research contributions include Kumar et al.’s KatzBot, demonstrating enhanced accuracy through domain-specific training [7] but without query classification mechanisms and Aleedy’s comprehensive survey classifying educational chatbots by functionality and domain [1]. While pioneering various approaches, these solutions generally lack comprehensive query classification and systematic self-learning mechanisms comparable to our implementation.

2.2 Multi-layered Query Classification and Self-learning

Educational virtual assistants typically implement basic binary classification, determining whether questions can be answered based on available knowledge. Research indicates classification accuracy in educational chatbots ranges from 70-90% [8], heavily depending on domain specificity. RAGent transcends these limitations by integrating keyword matching with semantic similarity scoring and confidence thresholds, resulting in higher accuracy rates. This hybrid classification addresses ethical concerns regarding content accuracy and system transparency highlighted by Kooli [6], who specifically emphasizes privacy risks and data security challenges in educational chatbot deployments, explicitly identifying when queries exceed the agent’s knowledge domain rather than providing potentially misleading information. RAGent’s classification novelty lies in its comprehensive methodology integrating keyword

matching against course-specific terminology, semantic similarity measurements using vector embeddings, dynamic confidence threshold assessment, and contextual relevance evaluation – providing nuanced understanding of student queries while maintaining alignment with educational objectives. A significant limitation of many educational AI solutions is their static nature. Commercial systems like EUDE’s Virtual Co-Tutor capture unanswered questions but lack formalized integration processes [5]. By contrast, RAGent implements a structured self-learning workflow bridging question collection and knowledge enhancement, ensuring increasing capability while providing instructors with insights into student misconceptions for targeted instructional interventions.

3 Methodology and System Overview

RAGent employs RAG to ground responses in verified course-specific knowledge sources [9], addressing limitations in traditional LLMs such as hallucination tendencies. The system implements a hybrid architecture combining dense retrieval with generative modeling through five stages, as illustrated in Figure 1, with detailed component interactions shown in Figure 2.

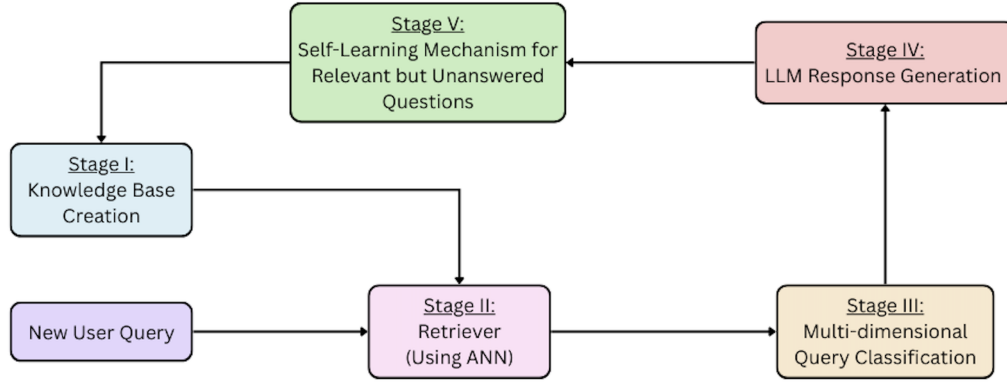
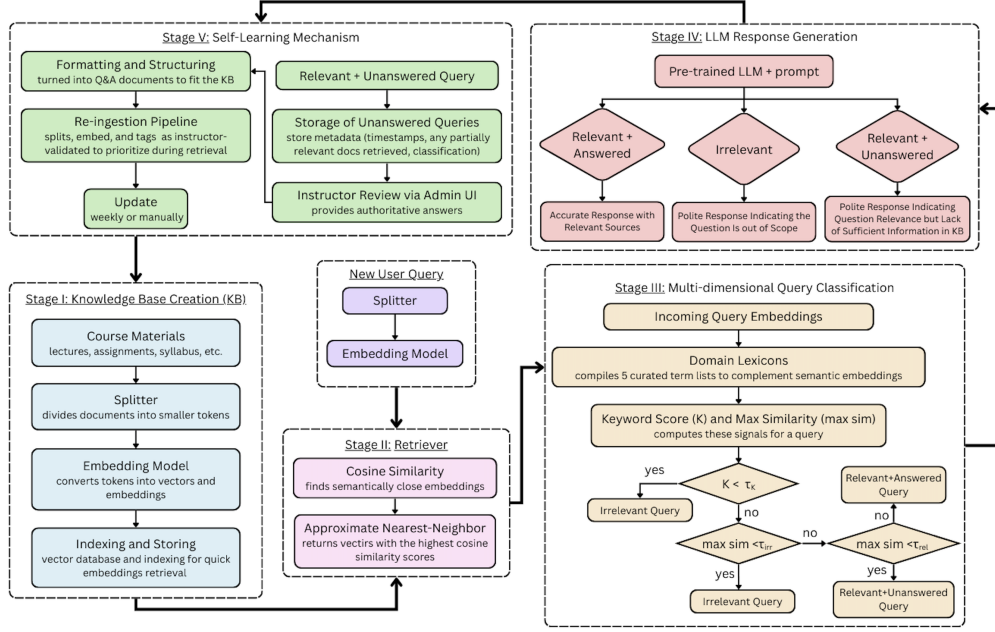


Figure 1 RAGent’s high-level overview on the five core operational stages – initial document ingestion and knowledge base creation (stage I), retrieval processing (stage II), multi-dimensional classification for a new query (stage III), LLM response generation (stage IV), and the self-learning feedback loop (stage V). This diagram abstracts away implementation details to focus on the essential system components and flow.

3.1 System Architecture and Operational Workflow

RAGent’s knowledge base construction ingests diverse course materials through format-specific loaders. Documents are preprocessed and segmented into semantically coherent chunks using recursive character text splitting with overlap. These chunks retain metadata connections to their source documents and are transformed into vector representations indexed for efficient retrieval. When a student submits a query, RAGent normalizes and encodes it, then performs an approximate nearest-neighbor search to identify relevant content fragments. Similarity between query embedding q and document chunk embedding d_i is quantified using cosine similarity:

$$\text{sim}(q, d_i) = \frac{q \cdot d_i}{\|q\| \cdot \|d_i\|} = \frac{\sum_{j=1}^n q_j \cdot (d_i)_j}{\sqrt{\sum_{j=1}^n q_j^2} \cdot \sqrt{\sum_{j=1}^n (d_i)_j^2}} \quad (1)$$



■ **Figure 2** RAGent’s detailed architecture, including the multi-dimensional query-classification module and self-learning feedback loop – where instructor input iteratively sharpens classification boundaries and expands the knowledge base.

A key innovation is RAGent’s query classification mechanism that determines the nature of student inquiries using terminology lists and similarity thresholds. The classification assigns queries to one of three categories:

$$f(q) = \begin{cases} \text{irrelevant,} & \text{if } K(q) < \tau_K \text{ and } \max_i \text{sim}(q, d_i) < \tau_{\text{irr}} \\ \text{relevant-unknown,} & \text{if } K(q) \geq \tau_K \text{ and } \max_i \text{sim}(q, d_i) < \tau_{\text{rel}} \\ \text{relevant-known,} & \text{if } \max_i \text{sim}(q, d_i) \geq \tau_{\text{rel}} \end{cases} \quad (2)$$

where $K(q)$ is the keyword matching score and τ_K , τ_{irr} , and τ_{rel} are empirically determined thresholds (2, 0.65, and 0.78, respectively). Response generation employs context-aware prompting strategies tailored to query classification. For “relevant-known” queries, the RAGent incorporates retrieved context with attribution; for “relevant-unknown” queries, it acknowledges knowledge gaps while providing partially relevant context; and for “irrelevant” queries, it redirects students while suggesting reformulation. The framework varies temperature (0.3-0.7) based on query type to maintain educational value. RAGent’s self-learning mechanism enables improvement by capturing “relevant-unknown” queries for expert review. Instructors provide authoritative answers through an administrative interface, and these curated pairs are processed and indexed. The knowledge base updates according to:

$$KB_{t+1} = KB_t \cup \{(q_i, a_i, e(q_i), e(a_i)) | q_i \in Q_t\} \quad (3)$$

where Q_t represents unanswered queries collected during period t , and $e(q_i)$ and $e(a_i)$ are embedding vectors of questions and answers. This feedback loop enhances RAGent’s capabilities, creating an adaptive educational tool that evolves alongside the course.

4 Implementation

RAGent’s implementation integrates several specialized libraries: Langchain for orchestrating RAG components, FAISS for efficient similarity search, OpenAI’s text-embedding-ada-002 model, GPT-4o for response generation, and Streamlit for the user interface. A Postgres database stores relevant but unanswered questions for instructor review. This stack creates a modular architecture where components can be independently optimized as educational needs evolve or new AI capabilities emerge. Source materials are parsed and segmented using recursive character text splitting, producing coherent chunks of approximately 1000 characters with 200-character overlap to preserve context across boundaries. Each chunk maintains metadata links to its source file, page, and offset. We embed chunks with text-embedding-ada-002 (1536-dimensional vectors) and store them in a FAISS index optimized for high-dimensional nearest-neighbor search. When a student submits a question, the classification module assigns one of three relevance labels. For in-scope queries, RAGent performs cosine-similarity search and selects $K \in [3, 7]$ chunks based on local embedding density. The retrieved context, classification metadata, and original query are injected into a GPT-4o prompt template that enforces citation formatting; the LLM then generates responses with inline source attributions. Queries labeled relevant but unanswered are stored in Postgres for instructor review. Verified answers are embedded and merged into the FAISS index during periodic ingestion jobs, expanding coverage without redeployment. The Streamlit interface provides real-time interaction for students, presenting answers with collapsible source excerpts. An administration dashboard allows teaching staff to monitor unanswered queries, author responses, and trigger re-ingestion when needed.

5 Experimental Design and Evaluation

We evaluated RAGent through a dual-pronged approach combining controlled quantitative assessment with ecological user testing involving 125 undergraduate data science students at the University of Maryland. Our controlled evaluation employed a stratified test corpus with balanced representation of in-scope and out-of-scope queries across conceptual understanding, technical implementation, and administrative concerns. We incrementally tested with sample sizes of 100, 200, 300, and 400 queries to assess performance stability and scalability. For each query, we recorded classification decisions, retrieval metrics, and response characteristics, comparing outputs against the ground truth labels using standard information retrieval metrics: accuracy, precision, recall, F1-score, and error rates. The ecological evaluation engaged students in structured interaction sessions through a Streamlit-based interface. Participants posed diverse questions (administrative, conceptual/technical, and deliberately out-of-scope) and completed post-interaction feedback forms capturing both quantitative ratings and qualitative impressions. Metrics included response correctness, classification precision, system reliability, user satisfaction (10-point scale), adoption willingness (Yes/No/Maybe), and qualitative feedback. For all evaluations, RAGent operated with its complete architecture using authentic course materials, with experiments conducted on standardized configurations to ensure reproducibility. This combined approach provided both rigorous performance measurements and insights into authentic user interactions, enabling comprehensive assessment of RAGent’s educational efficacy and implementation readiness. The controlled experiments demonstrated RAGent’s ability to distinguish between relevant and irrelevant queries with high accuracy, while the ecological testing revealed strong user acceptance and identified improvement opportunities. The dual methodology allowed us to verify that performance characteristics observed in controlled settings translated effectively to real-world educational

contexts, providing validation of both technical capabilities and practical utility. Through this evaluation framework, we were able to systematically assess RAGent’s performance across multiple dimensions critical to educational technology: technical accuracy, appropriate boundary recognition, response quality, user experience, and pedagogical alignment. These findings informed both our understanding of the RAGent’s current capabilities and our roadmap for future enhancements to further improve its educational value.

6 Results and Discussion

Our evaluation of RAGent produced both quantitative performance metrics and qualitative user insights that collectively demonstrate the framework’s effectiveness in educational settings.

6.1 Quantitative Performance Analysis

RAGent demonstrated robust classification capabilities across progressively larger question sets, showing consistent improvement as corpus size increased – suggesting effective generalization rather than overfitting to specific question types [4].

■ **Table 1** The classifier steadily improves on every metric – accuracy rises from 95.0% to 96.5%, precision and recall both edge upward, and the balanced F1 score increases from 0.974 to 0.982 – while the Type II error rate (false negatives) is cut by more than half, dropping from 0.021 to 0.010. The trend indicates that RAGent generalizes well to a broader and more diverse set of queries rather than overfitting to smaller samples.

Question Set Size	Accuracy	Precision	Recall	F1 Score	Type II Error
100 questions	0.950	0.969	0.979	0.974	0.021
200 questions	0.955	0.974	0.979	0.977	0.021
300 questions	0.960	0.973	0.986	0.980	0.014
400 questions	0.965	0.975	0.990	0.982	0.010

As shown in Table 1, accuracy improved from 95.0% to 96.5% as the test corpus expanded, while Type II errors decreased from 0.021 to 0.010 – a 52.4% reduction particularly significant in educational contexts where missing relevant questions could impact learning outcomes.

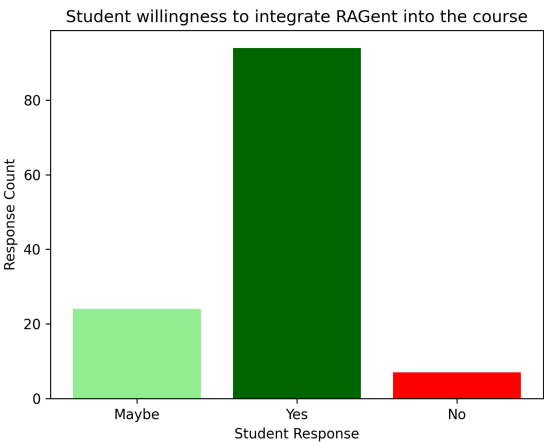
The in-situ deployment with 125 student participants validated RAGent’s effectiveness in authentic settings. RAGent successfully differentiated between course-relevant and irrelevant questions with an average of 1.9 incorrect responses per session. System reliability was exceptional (100% crash-free), with a mean satisfaction rating of 7.85/10 demonstrating strong user acceptance.

6.2 Qualitative Assessment and User Experience

Student responses regarding potential course integration revealed strong support for RAGent, with 94.4% of participants (94 “Yes” and 24 “Maybe” responses) expressing positive sentiment toward integration into their course workflow. These results show that the system is capable of handling routine queries that would otherwise require direct instructor intervention. The high reliability of the system along with immediate responses addresses the temporal mismatch between student question generation and staff availability, particularly beneficial during peak assignment periods and outside traditional office hours.

■ **Table 2** Quantitative results from a RAGent demonstration with 125 undergraduates. RAGent ran crash-free for all sessions, achieved a mean user-satisfaction score of 7.85/10, and averaged 1.9 incorrect answers per session, indicating solid reliability and generally positive student reception in an authentic instructional setting.

Metric	Value
Total participants	125
Average satisfaction rating	78.5%
Minimum rating (1–10 scale)	5
Maximum rating (1–10 scale)	10
Average incorrect answers per session	1.9
System reliability (crash-free sessions)	100%



■ **Figure 3** Distribution of 125 student responses to the question “Would you like this tool included in the course?”. 94 students answered “Yes” and 24 answered “Maybe” and provided some valuable feedback, while only 7 chose “No”. Combined, the affirmative responses represent 94.4% of the cohort, indicating overwhelming interest in adopting RAGent as a course resource.

Participant feedback included actionable recommendations for UI modifications, multimodal capabilities extension, and enhanced search functionality for exam preparation. The predominance of positive terms in student feedback suggests RAGent’s strong appreciation, with favorable impressions of the framework’s usefulness for data science education.

6.3 Ethical and Privacy Considerations

RAGent’s deployment in educational settings requires careful attention to ethical and privacy considerations. All student interactions with the system were conducted under informed instruction with anonymized data collection for research purposes. Query logs were stripped of personally identifiable information to protect student privacy.

To address potential bias in AI-generated responses, RAGent uses several mitigation strategies: (1) diverse training materials representing multiple perspectives within data science pedagogy, (2) teaching staff review of flagged responses to identify and correct biased outputs, and (3) transparent citation of source materials. The system’s query classification mechanism explicitly acknowledges knowledge limitations rather than generating potentially misleading responses, supporting responsible AI use in educational contexts.



■ **Figure 4** Word-cloud visualization of the 125 short student reactions. Word size corresponds to frequency, with “Good” and “Okay/okay” being most common descriptors, followed by strongly positive terms such as “Helpful”, “Very helpful”, “Amazing”, and “Impressive”. The dominance of affirmative language indicates a broadly favorable first impression of RAGent’s usefulness.

7 Future Work

While RAGent demonstrates considerable promise, our evaluation identified several opportunities for enhancement. Key work directions include knowledge boundaries for specialized queries, performance variance across query types, lack of visual capabilities, cross-topic synthesis challenges, and interface customization. Our enhancement roadmap addresses these constraints through three strategic dimensions: multimodal capabilities for visualizing statistical concepts, enhanced retrieval mechanisms using hierarchical approaches and knowledge graph representations, and adaptive interface customization based on user preferences. Additionally, we plan to implement a temporally-aware “Gradual Knowledge Expansion” that aligns with course progression, where for any given week N , the system will only access content from weeks 1 through $N - 1$ plus evergreen materials. This approach keeps answers temporally relevant, prevents exposure to upcoming concepts, and maintains alignment with instructional pacing. Looking beyond single-course implementations, we are developing a generalized framework with automated ingestion routines, course-agnostic prompt templates, a low-code instructor dashboard, and inter-course knowledge sharing mechanisms. This abstraction would significantly reduce deployment overhead and facilitate institution-wide scaling, transforming RAGent from a course-specific tool to a broadly applicable educational platform that can adapt to diverse disciplinary contexts while maintaining domain-specific relevance. Additional enhancement would include multimodal capabilities to support image and code visualization. Another promising direction is represented by the personalized response generation based on individual learning history and performance patterns, allowing RAGent to adapt explanation complexity and provide targeted resources in regard to each student’s knowledge level. To support broader adoption while maintaining institutional flexibility, we plan to develop comprehensive documentation and implementation guidelines to allow other institutions reproduce RAGent’s architecture using readily available components. This includes detailed technical specifications, configuration templates, and best-practice recommendations for adapting the framework to diverse curricular.

8 Conclusion

This paper has presented RAGent, a specialized RAG-based intelligent agent addressing the unique challenges of university-level data science education. By integrating LLM capabilities with course-specific knowledge retrieval, RAGent enhances academic support accessibility, consistency, and scalability. Our evaluation, combining controlled experiments with ecological testing involving 125 undergraduate students, demonstrated RAGent's effectiveness in distinguishing between relevant and irrelevant queries while providing contextually appropriate responses. The framework showed consistent performance improvements across increasingly diverse question sets (accuracy improving from 95.0% to 96.5%), suggesting robust generalization capabilities. Student feedback was overwhelmingly positive, with 94.4% of participants expressing interest in integrating the tool into their course. RAGent's novel contributions include: (1) a multi-dimensional query classification system that effectively categorizes student questions, (2) a self-learning mechanism that captures unanswered but relevant queries for knowledge base expansion, and (3) a contextually aware response generation approach that maintains alignment with course terminology. These features address the challenges posed by ambiguous student queries and the structured nature of course content. While our evaluation identified certain opportunities for improvement – including knowledge boundaries, performance variance, and lack of multimodal capabilities –these provide clear directions for future development through progressive knowledge expansion, multimodal support, and cross-course generalization. The broader implications of this work extend beyond a single course implementation. RAGent demonstrates how AI-augmented educational tools can bridge gaps in traditional academic support, particularly for diverse student populations with varying technical preparation. As institutions face increasing enrollment pressures, frameworks like RAGent offer a scalable approach to maintaining educational quality while accommodating individual student needs. Our future work should validate RAGent's effectiveness across multiple institutions and diverse course contexts to establish broader generalization and identify context-specific adaptation requirements.

References

- 1 Mohamed Aleedy, Eric Atwell, and Souham Meshoul. Using ai chatbots in education: Recent advances, challenges and use case. In *Artificial Intelligence and Sustainable Computing*, pages 661–675. Springer, 2022. doi:10.1007/978-981-19-1653-3_50.
- 2 Bashar Alsafari, Eric Atwell, Alan Walker, and Michael Callaghan. Towards effective teaching assistants: From intent-based chatbots to llm-powered teaching assistants. *Natural Language Processing*, 8:100101, 2024. doi:10.1016/j.nlp.2024.100101.
- 3 Jose Belda-Medina and Veronika Kokoskova. Integrating chatbots in education: Insights from the chatbot–human interaction satisfaction model (chism). *International Journal of Educational Technology in Higher Education*, 20:62, 2023. doi:10.1186/s41239-023-00432-3.
- 4 Matt Gardner, Yoav Artzi, Victoria Basmov, Jonathan Berant, Ben Bogin, Sihao Chen, Pradeep Dasigi, Dheeru Dua, Yanai Elazar, Ananth Gottumukkala, Nitish Gupta, Hannaneh Hajishirzi, Gabriel Ilharco, Daniel Khashabi, Kevin Lin, Jiangming Liu, Nelson F. Liu, Phoebe Mulcaire, Qiang Ning, Sameer Singh, Noah A. Smith, Sanjay Subramanian, Reut Tsarfaty, Eric Wallace, Ally Zhang, and Ben Zhou. Evaluating models' local decision boundaries via contrast sets. In *Findings of the Association for Computational Linguistics: EMNLP 2020*, pages 1307–1323. Association for Computational Linguistics, 2020. doi:10.18653/v1/2020.findings-emnlp.117.
- 5 IBM Client Success Stories. European school of management and business (eude) case study, 2025. Retrieved 2 May 2025. URL: <https://www.ibm.com/case-studies/eude>.

- 6 Chokri Kooli. Chatbots in education and research: A critical examination of ethical implications and solutions. *Sustainability*, 15(7):5614, 2023. doi:10.3390/su15075614.
- 7 Sushil Kumar, Dipesh Paikar, Karthik Sai Vutukuri, Hassan Ali, Sai Rohit Ainala, Arun M. Krishnan, and Yi Zhang. Katzbot: Revolutionizing academic chatbot for enhanced communication. *arXiv preprint*, 2024. arXiv:2410.16385.
- 8 Lasha Labadze, Maya Grigolia, and Lela Machaidze. Role of ai chatbots in education: Systematic literature review. *International Journal of Educational Technology in Higher Education*, 20(1):56, 2023. Published 31 October 2023. doi:10.1186/s41239-023-00426-1.
- 9 Patrick Lewis, Barlas Oguz, Ruty Rinott, Sebastian Riedel, and Veselin Stoyanov. Retrieval-augmented generation for knowledge-intensive nlp tasks. In *Advances in Neural Information Processing Systems*, volume 33, pages 9459–9474, 2020.
- 10 McGraw Hill Higher Education. Aleks: Adaptive learning and assessment platform, 2025. Retrieved 2 May 2025. URL: <https://www.mheducation.com/highered/digital-products/aleks.html>.
- 11 National Academies of Sciences, Engineering, and Medicine. *Data Science for Undergraduates: Opportunities and Options*. National Academies Press, 2018. doi:10.17226/25104.
- 12 OpenAI. Gpt-4 technical report, 2023. URL: <https://openai.com/research/gpt-4>.
- 13 Kanishk Sikka, Shubham Singh, Kashyap Ramesh, et al. Discourse-aware prompt design for education question answering. In *Proceedings of the NAACL 2022 Workshop on Innovative Use of NLP for Building Educational Applications*, 2022.
- 14 Kartik Taneja, Priyanka Maiti, Shashank Kakar, Pooja Guruprasad, Shruti Rao, and Ashok K. Goel. Jill watson: A virtual teaching assistant powered by chatgpt. *arXiv preprint*, 2024. arXiv:2405.11070.
- 15 Dominik Thüs, Sarah Malone, and Roland Brünken. Exploring generative ai in higher education: A rag system to enhance student engagement with scientific literature. *Frontiers in Psychology*, 15:1474892, 2024. doi:10.3389/fpsyg.2024.1474892.
- 16 University of Maryland. Umd virtual agent. <https://ai.umd.edu/resources/services>, 2025. Accessed May 12, 2025.
- 17 Pengcheng Yin, Ziyi Wang, and Graham Neubig. Tabsum: Table summarization with selection, aggregation and generation. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics*, 2021.