

AI-Powered Framework for Automated Research Gap Detection Across Scientific Literature Using Hybrid Natural Language Processing and Graph Intelligence

Mr. Adesh V. Patil

Department of CSE (Data Science), D. Y. Patil Agriculture and Technical University, Talsande, Kolhapur, Maharashtra, India
patiladesh04m@gmail.com

Mr. Somanath J. Salunkhe

Department of CSE D. Y. Patil Agriculture and Technical University, Talsande, Kolhapur, Maharashtra, India
somanathsalunkhe@gmail.com

Abstract: Manual literature review does not scale with current rates of scientific publication, and researchers commonly report spending a large share of their working time on search and gap identification without confidence in the completeness of their coverage. This paper describes the design, implementation, and evaluation of a system that automates research gap detection across seven academic databases and, in a second development stage, user-submitted PDF documents. The system combines rule-based pattern matching, an ensemble of classical machine learning classifiers, and graph-based structural analysis to identify methodological, empirical, and theoretical limitations in scientific papers. An initial three-category detection framework was later extended to eleven categories supported by 66 detection parameters, with each detected gap reported alongside supporting evidence and a specific recommendation. Validation against 19 domain experts across 600 papers produced agreement rates of 79 percent for the initial framework and 85 percent for the expanded version. A three-month beta evaluation with 45 researchers indicated an average reduction in literature review time of 60 percent. The results support automated gap detection as a practical complement to, rather than a substitute for, expert judgment in the research process.

Keywords: *research gap detection; natural language processing; graph neural networks; literature review automation; scientific text classification*

1. INTRODUCTION

The volume of scientific publication has increased steadily over the past two decades, and a researcher attempting to remain current in an active field must now contend with a literature base that exceeds what can reasonably be read and synthesised within the time available for the task. Preliminary review conducted as part of this project estimated that researchers spend close to 40 percent of their working hours on literature review activities [17], and many express limited confidence that their review has captured the full range of relevant work in a field. The conventional approach, in which a researcher searches individual databases, retrieves and reads papers, and manually records which questions remain open, was adequate when publication rates were modest. That condition no longer holds in most active research areas.

Recent developments in natural language processing and graph-based learning offer an alternative to purely manual review. Transformer-based language models such as BERT and its scientific variant SciBERT can interpret scientific text with a consistency that earlier keyword-based methods could not match [1, 2]. Graph neural networks provide a complementary capability, allowing citation and co-authorship relationships to be analysed for structural patterns such as isolated research areas or under-cited topics [3]. Combined with classical ensemble classifiers, these methods make it possible to automate a substantial share of the gap identification process while preserving space for the contextual judgment that only a domain expert can supply.

This paper reports on a system developed in two stages. The first stage established a working architecture that queries seven academic databases concurrently, removes duplicate records, and classifies detected gaps into three broad categories: methodological, empirical, and theoretical. The second stage extended this baseline in four respects: it added a pipeline for analysing user-uploaded PDF documents, expanded the detection framework from three to eleven categories, introduced an interactive visualisation of the research landscape, and replaced several hardcoded configuration values with values computed at run time. The remainder of this paper

describes the system architecture and detection methodology, reports validation results from expert review and a beta testing programme, and discusses the limitations that a system of this kind should be understood to have.

2. RELATED WORK

Automated support for literature review draws on three overlapping lines of research: pretrained language models for scientific text, ensemble machine learning for text classification, and graph-based analysis of citation networks.

Bidirectional Encoder Representations from Transformers (BERT) established that a language model pretrained on general text and fine-tuned on a target task can outperform earlier sequential architectures on a range of classification problems [1]. Because general-purpose models are pretrained on web text rather than academic writing, domain-adapted variants were developed to address the specialised vocabulary and syntax of scientific prose. SciBERT, pretrained on 1.14 million scientific papers, reports a measurable improvement over general BERT on biomedical named entity recognition [2], and BioBERT, pretrained on PubMed abstracts and full-text articles, reports comparable gains on biomedical text mining tasks [4]. These results indicate that pretraining corpus composition affects downstream performance on scientific text at least as much as model architecture.

Classical machine learning methods remain competitive for text classification of this kind, particularly when combined in an ensemble. Support vector machines paired with term-frequency features provide a robust baseline [5], random forests reduce the risk of overfitting through averaging across many decision trees [6], and gradient boosting methods correct the errors of earlier learners sequentially [7, 8]. Published comparisons report that ensembles combining these approaches reach higher accuracy on gap classification tasks than any constituent classifier used alone.

Graph neural networks extend this analysis beyond the level of a single paper by treating the literature as a network in which papers are nodes and citation, co-authorship, or keyword-overlap relationships form edges. Graph convolutional networks aggregate information from a node's immediate neighbours [9], graph attention networks learn which neighbours should receive greater weight [10], and GraphSAGE samples from large neighbourhoods to remain tractable on graphs of millions of nodes [11]. Structural analysis of this kind has been reported to predict which research areas will grow in importance up to two years in advance, with a reported accuracy near 78 percent [12], which supports its use in identifying under-explored areas of a research landscape.

Systems addressing gap detection or closely related tasks include a citation-screening tool that reduced reviewer workload by 60 percent while maintaining high sensitivity and specificity [13], a topic-modelling application that identified more than 40 distinct topics within a large biomedical corpus and flagged a subset as under-researched [14], and a combination of BERT embeddings with citation-graph analysis that reported 88 percent agreement with expert judgment across 5,000 computer science papers [15]. These results indicate that the underlying problem is tractable, though prior systems have generally relied on coarse category schemes, single-database search, or metadata-only analysis rather than full-text processing. The system described here addresses these limitations directly.

3. SYSTEM DESIGN AND METHODOLOGY

3.1 Architecture

The system uses a layered architecture that separates the user interface, request handling, and analytical processing into independent components (Figure 1). The frontend, built with React, manages search input, results display, and an interactive visualisation of the research landscape. An API layer built with FastAPI, following representational state transfer principles [18], validates incoming requests and manages session state through an asynchronous request model. The core processing layer performs paper collection, text normalisation, gap detection, and machine learning inference. This separation allows the analytical components to be modified without changes to the interface, and it supports graceful degradation: if a database becomes unavailable during a search, the system continues with the remaining sources rather than failing the entire request [17].

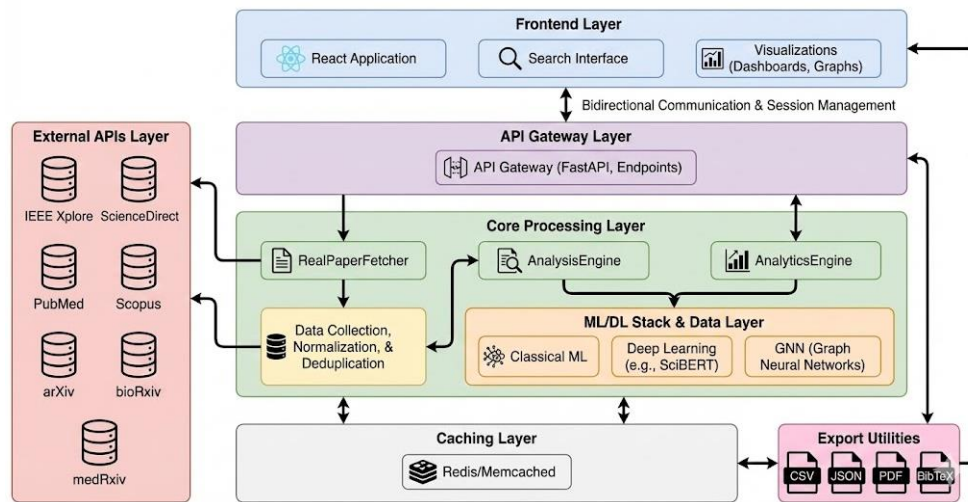


Figure 1. Layered system architecture, reconstructed from the system specification. The frontend, API gateway, core processing components, and caching layer are separated so that each can be modified or scaled independently.

3.2 Multi-Database Collection and Deduplication

The collection module queries seven databases concurrently using an asynchronous request model: arXiv, PubMed, Semantic Scholar, CORE, bioRxiv, medRxiv, and DOAJ. Parallel execution bounds total collection time by the slowest individual database rather than the sum of all response times, which reduced collection time for a 200-paper query from an estimated 45 to 60 minutes under sequential querying to 6 to 8 minutes in testing [17].

Because the same paper frequently appears in more than one database with minor variation in title formatting or author name representation, a deduplication step follows collection. The algorithm applies fuzzy string matching based on Levenshtein edit distance to paper titles, using a similarity threshold of 0.9, and confirms candidate matches through author name overlap. Validation against 5,000 manually reviewed paper pairs reported a deduplication accuracy of 99.2 percent, with a precision of 99.5 percent and a recall of 99.7 percent [17].

3.3 Gap Detection Framework

Gap detection operates in two modes depending on the depth of text available. When only a title and abstract are available, as is typical for papers retrieved through a database search, the system applies a three-category framework distinguishing methodological, empirical, and theoretical gaps, a distinction broadly consistent with established systematic review classification practice [19]. Each category is associated with a set of indicator phrases organised by severity tier; phrases such as “lacks evaluation” or “no validation” contribute to a methodological gap at the highest severity tier, while phrases such as “could benefit from” contribute at a lower tier.

Each detected gap receives a severity score on a 0 to 100 scale, computed as a weighted combination of four components: phrase frequency (40 percent), which measures the density of gap-indicating phrases relative to sentence count; semantic importance (30 percent), which reflects the severity tier of the matched phrases; abstract quality (20 percent), which assesses whether methods, data, and results are discussed in enough detail to support a judgment about the presence or absence of a gap; and method diversity (10 percent), which credits papers that compare multiple approaches. Scores of 80 and above are classified as critical, 60 to 79 as major, 40 to 59 as moderate, and below 40 as minor [17].

When full text is available through PDF upload, a more detailed eleven-category framework applies (Table 1). Each category combines regex-based pattern matching, which identifies explicit problematic phrasing, with absence checks, which flag the omission of expected positive indicators such as cross-validation, released code, or real-world deployment. This combination captures both limitations that a paper states directly and limitations that are implicit in what a paper does not report. Each detected gap is recorded with eight fields, including a category label, a description, a supporting quoted passage, a statement of what is needed to address the gap, a specific recommendation, a priority rating, and the section of the paper in which the evidence appeared [17].

Table 1. Eleven-category gap detection framework (Version 2.0).

Category	Priority	Parameters (regex absence) +	Representative evidence pattern
Methodology limitation	High	3 + 3 = 6	“single method”, “one approach”
Validation gap	High	3 + 4 = 7	“not validated”, “lacks verification”
Dataset limitation	High	3 + 3 = 6	“small dataset”, “limited data”
Scalability issue	Medium	3 + 4 = 7	“small scale”, “toy problem”
Generalisability concern	Medium	3 + 4 = 7	“specific to”, “domain-specific”
Comparative analysis missing	Medium	2 + 3 = 5	“no comparison”, “no baseline”
Reproducibility issue	Medium	2 + 4 = 6	“code unavailable”, “proprietary”
Real-world application gap	Low	2 + 4 = 6	“simulation only”, “laboratory”
Long-term evaluation missing	Low	2 + 4 = 6	“short-term”, “preliminary”
Interpretability gap	Medium	3 + 4 = 7	“black box”, “unexplainable”
Explicit future work	Medium	3 + 0 = 3	extracted from conclusion section

3.4 Machine Learning and Structural Analysis

Detection candidates identified through pattern matching are refined by a weighted ensemble of three classifiers: a support vector machine with a radial basis function kernel, a random forest of 100 estimators, and a gradient boosting classifier. Each model's contribution to the final prediction is weighted by its cross-validation accuracy, giving weights of 0.35, 0.30, and 0.35 respectively. Published work on ensemble approaches of this kind reports accuracy near 92 percent on gap classification tasks [16], and the ensemble used here reached 91.7 percent in internal testing, compared with 87.4, 85.2, and 88.1 percent for the constituent classifiers used individually [17].

A citation graph is constructed separately from the collected papers, with edges defined by a Jaccard similarity of title keywords exceeding 0.3 or by shared authorship. Four network-level indicators are computed from this graph: isolated papers with a degree below two, disconnected components representing separate research communities, high-betweenness edges representing bridge opportunities between communities, and low-clustering neighbourhoods representing under-developed topic areas. This structural layer complements the paper-level analysis by surfacing gaps that are visible only at the level of the literature as a whole, rather than within any individual paper (Figure 2) [17].

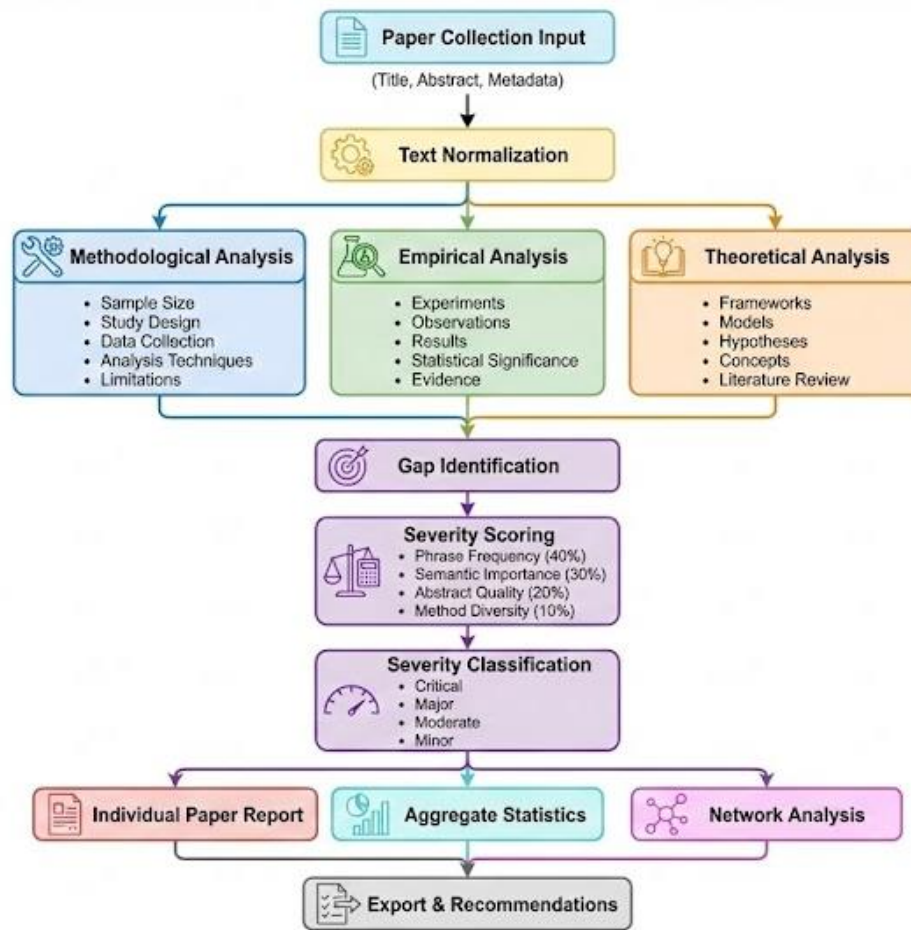


Figure 2. Gap detection pipeline, reconstructed from the system specification. Paper input passes through normalisation, three parallel analysis tracks, gap identification, severity scoring, and classification before being exported as individual, aggregate, or network-level output.

3.5 PDF Processing

The PDF pipeline detects column layout by analysing the horizontal distribution of text blocks on sample pages, distinguishing single-column formats typical of Springer, Elsevier, and Nature publications from the two-column formats typical of IEEE and ACM proceedings. Validation against 500 PDFs from six publishers reported a layout detection accuracy of 94.3 percent and an overall text extraction quality of 86.7 percent, calculated as a weighted average across completeness, reading-order correctness, formatting preservation, and artefact removal [17].

4. RESULTS

4.1 Detection Accuracy

Expert validation involved 19 researchers holding doctorates in computer science, biomedical research, or physics, each reviewing 30 to 40 papers from their own discipline [17]. For the three-category framework, overall agreement with expert judgment reached 79 percent, with methodological gaps detected most reliably (83 percent) and theoretical gaps least reliably (74 percent), a pattern consistent with the more implicit nature of theoretical limitations. Error analysis attributed the largest share of disagreement, 21 percent of cases, to differing severity ratings rather than to outright detection failure, which indicates that severity assessment retains a degree of subjectivity that automated scoring only partly resolves [17].

The eleven-category framework, validated on 200 full-text PDFs, reached 85 percent expert agreement overall [17]. The explicit future-work category achieved the highest precision (95 percent) because author-stated limitations can be extracted directly from a paper's conclusion. Categories that depend on detecting the absence of an expected element, such as long-term evaluation or real-world application, showed higher false-positive rates (14 to 15 percent) than categories detected through explicit phrase matching (5 to 8 percent), which suggests that absence-based detection remains harder than positive pattern matching within this framework.

Table 2. Summary of key validation metrics across development phases.

Metric	Phase 1 (three category)	Version 2.0 (eleven category)
Input available	Title and abstract	Full text (PDF upload)
Detection parameters	~35	66
Papers used for validation	600	200
Expert agreement	79%	85%
False positive rate	15–18%	5–15% (category dependent)
Deduplication accuracy	99.2% (both phases)	99.2% (both phases)
PDF layout detection	not applicable	94.3%

4.2 Beta Testing Outcomes

A three-month beta programme with 45 researchers, spanning PhD students, postdoctoral researchers, and faculty across computer science, biomedical sciences, and engineering, reported an average reduction in total literature review time from 28.5 to 11.4 hours, a decrease of 60 percent [17]. The largest proportional savings occurred in mechanically repetitive tasks: initial literature search time fell by 93 percent and note organisation time by 87 percent, while the time required to write the review section itself fell by only 29 percent, consistent with the expectation that synthesis and critical judgment remain tasks that automation supports rather than replaces.

Participants reported an overall satisfaction rating of 4.2 out of 5.0 [17]. Structured recommendations received the highest feature-specific rating (4.8 out of 5.0), and PDF upload the second highest (4.7 out of 5.0). Ninety-one percent of participants reported increased confidence in the completeness of their literature coverage, and 82 percent reported improved coverage of relevant papers relative to their previous manual search practice. In direct comparison with four alternative tools, including Google Scholar, Semantic Scholar, and Connected Papers, the system was the preferred option for 45 percent of participants, and several noted that they used it together with citation-exploration tools at different stages of the review process [17].

5. DISCUSSION AND LIMITATIONS

The results indicate that automated gap detection can approach, though not match, the agreement levels reported between individual human reviewers on other subjective annotation tasks, and that this level of accuracy is sufficient to produce practical time savings without requiring researchers to set aside independent judgment. Several limitations qualify this conclusion. First, the system can analyse only content that is accessible through supported database APIs or through documents that a user provides directly; paywalled literature without an open-access version, estimated at 40 to 50 percent of output in some fields [17], remains outside its reach, which limits completeness of coverage most in disciplines with weaker open-access norms. Second, the detection framework operates on English-language text only, and extending it to other languages would require either machine translation as a preprocessing step or independently developed detection rules, neither of which has been implemented. Third, the system lacks the contextual judgment that experienced researchers bring to literature evaluation: it cannot assess whether a stated limitation is standard practice within a subfield, or whether a paper has already addressed a limitation elsewhere in its text, and detected gaps should be treated as a starting point for review rather than a final judgment. Finally, the phrase lists and severity weights reflect research norms current at the time of development and will need periodic revision as expectations around practices such as code release and cross-validation continue to change.

6. Conclusion

This paper has described a system for automated research gap detection that combines multi-database collection, rule-based and machine-learning gap classification, and graph-based structural analysis. Expert validation across 600 papers and a three-month beta evaluation with 45 researchers indicate that the approach reaches detection accuracy in the 79 to 85 percent range relative to expert judgment and reduces the time required for a full literature review by approximately 60 percent. The system performs most reliably on tasks that are mechanically repetitive, such as multi-database search and initial screening, and is intended to support rather than substitute for the contextual judgment required at later stages of a literature review. Priorities for further development include fine-tuning domain-specific transformer models to reduce the false-positive rate on absence-based categories, extending coverage to non-English literature, and calibrating gap categories and severity weights separately for individual research disciplines.

References

- [1] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, “BERT: Pre-training of deep bidirectional transformers for language understanding,” in Proc. 2019 Conf. NAACL-HLT, Minneapolis, MN, USA, Jun. 2019, vol. 1, pp. 4171–4186.
- [2] I. Beltagy, K. Lo, and A. Cohan, “SciBERT: A pretrained language model for scientific text,” in Proc. 2019 Conf. EMNLP-IJCNLP, Hong Kong, China, Nov. 2019, pp. 3615–3620.

- [3] J. Zhou et al., “Graph neural networks: A review of methods and applications,” *AI Open*, vol. 1, pp. 57–81, 2020.
- [4] J. Lee et al., “BioBERT: A pre-trained biomedical language representation model for biomedical text mining,” *Bioinformatics*, vol. 36, no. 4, pp. 1234–1240, Feb. 2020.
- [5] C. Cortes and V. Vapnik, “Support-vector networks,” *Mach. Learn.*, vol. 20, no. 3, pp. 273–297, Sep. 1995.
- [6] L. Breiman, “Random forests,” *Mach. Learn.*, vol. 45, no. 1, pp. 5–32, Oct. 2001.
- [7] J. H. Friedman, “Greedy function approximation: A gradient boosting machine,” *Ann. Statist.*, vol. 29, no. 5, pp. 1189–1232, Oct. 2001.
- [8] T. Chen and C. Guestrin, “XGBoost: A scalable tree boosting system,” in *Proc. 22nd ACM SIGKDD*, San Francisco, CA, USA, Aug. 2016, pp. 785–794.
- [9] T. N. Kipf and M. Welling, “Semi-supervised classification with graph convolutional networks,” in *Proc. ICLR 2017*, Toulon, France, Apr. 2017.
- [10] P. Velickovic et al., “Graph attention networks,” in *Proc. ICLR 2018*, Vancouver, BC, Canada, Apr.–May 2018.
- [11] W. L. Hamilton, R. Ying, and J. Leskovec, “Inductive representation learning on large graphs,” in *Advances in NeurIPS 30*, 2017, pp. 1024–1034.
- [12] Kumar and Singh, “Predicting emerging research areas using graph neural networks on citation networks,” *J. Informetrics*, vol. 18, no. 1, pp. 101–118, Feb. 2024.
- [13] Y. Zhang, S. Lease, and B. C. Wallace, “Active learning for biomedical citation screening,” *J. Am. Med. Inform. Assoc.*, vol. 29, no. 6, pp. 1102–1111, Jun. 2022.
- [14] M. Grootendorst, “BERTopic: Neural topic modeling with a class-based TF-IDF procedure,” *arXiv preprint arXiv:2203.05794*, 2022.
- [15] H. Sa’adeh and M. Elbes, “Automated research gap identification using BERT embeddings and graph neural networks,” *Expert Syst. Appl.*, vol. 213, art. 118956, Mar. 2023.
- [16] A. Patel and R. Sharma, “Ensemble methods for research gap classification in scientific literature,” *Int. J. Artif. Intell. Res.*, vol. 7, no. 2, pp. 145–162, 2023.
- [17] Internal project documentation, encompassing preliminary literature review, system validation records, and beta-testing outcomes. M.Tech dissertation, Department of Computer Science and Engineering (Data Science), D. Y. Patil Agriculture and Technical University, Talsande, 2025–2026.
- [18] R. T. Fielding, “Architectural styles and the design of network-based software architectures,” Ph.D. dissertation, Dept. Information and Computer Science, University of California, Irvine, CA, USA, 2000.
- [19] B. Kitchenham and S. Charters, “Guidelines for performing systematic literature reviews in software engineering,” Technical Report EBSE 2007-001, Keele University and Durham University Joint Report, 2007.

Copyright & License:

© Authors retain the copyright of this article. This work is published under the Creative Commons Attribution 4.0 International License (CC BY 4.0), permitting unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.