



AI-based model of a researcher support service

Maksym Shovkoplias*

Postgraduate Student

Sumy State University

40000, 116 Kharkivska Str., Sumy, Ukraine

<https://orcid.org/0009-0007-4192-4743>

Abstract. Between 2020 and 2025, researchers faced challenges such as fragmented digital platforms, information overload, and limited personalisation capabilities. This underscored the need for services capable of providing comprehensive support for research activities. The aim of this study was to develop a conceptual model of an intelligent information service focused on personalised researcher support. The proposed system architecture was built using structural modelling, functional analysis, machine learning, and natural language processing techniques. It includes modules for recommendations, virtual collaboration, event management, and automated bibliography generation. A multi-layered user model was designed, taking into account scientific interests, interaction history, and research context. The combination of semantic analysis with behavioural patterns increased recommendation relevance by 20-30%. The prototype of the system was tested in March 2025 with the participation of 15 young scientists from three Ukrainian universities. The results of the survey and practical tasks showed that the average time spent searching for relevant literature was reduced by 35%, task planning efficiency increased by 40%, and user satisfaction with the service's functionality reached 87%. Respondents highly rated the convenience of the interface (4.5 out of 5), the relevance of recommendations (4.3), and co-authoring tools (4.6). Three new academic collaborations were initiated through the co-author selection module. The data obtained confirmed the effectiveness of the model in increasing research productivity, improving collaboration, and providing personalised user support. The proposed structure allows for scaling to different disciplines and has the potential to be implemented in digital platforms focused on scientific activity

Keywords: scientific information personalisation; semantic analysis; adaptive recommendations; machine learning; intelligent systems; digital research environment

INTRODUCTION

In the contemporary research landscape, scientists operate in increasingly complex digital environments marked by the exponential growth of academic content, decentralised data sources, and limited interoperability between platforms. These challenges are compounded by the necessity for personalised workspaces, efficient collaboration mechanisms, and intelligent content filtering. Existing solutions often offer isolated functionalities – bibliographic management, social networking, or search-but fail to deliver unified services that support the full research lifecycle in a seamless and adaptive way. This highlights the need for intelligent information systems capable of integrating recommendation

engines, virtual collaboration tools, scheduling, and automated bibliography generation in a personalised and scalable manner.

According to T. Adewale (2022), machine learning algorithms have a significant impact on personalised recommendation systems. The author proposed models that apply collaborative filtering and deep learning to enhance user engagement and relevance. The study demonstrated that machine-driven adaptation substantially improves access to and retention of information by researchers. O.B. Akinagbe (2024) explored broader applications of artificial intelligence in scientific workflows. Author's research confirmed that AI-based services

Article's History: Received: 22.06.2025; Revised: 08.11.2025; Accepted: 15.12.2025; Published: 25.12.2025.

Suggested Citation:

Shovkoplias, M. (2025). AI-based model of a researcher support service. *Bulletin of Cherkasy State Technological University*, 30(4), 82-96. doi: 10.62660/bcstu/4.2025.82.

*Corresponding author



Copyright © The Author(s). This is an open access article distributed under the terms of the Creative Commons Attribution License 4.0 (<https://creativecommons.org/licenses/by/4.0/>)

help reduce information overload and enable strategic planning by automating routine research tasks and aligning system responses with behavioural input.

C.A. Putri *et al.* (2025), working in the field of education, demonstrated the effectiveness of adaptive digital libraries that apply semantic indexing to personalise content delivery. Although their study was focused on primary education, the described mechanisms are applicable to academic research platforms as well. Y. Koval (2022), a Ukrainian researcher, analysed the structure and function of academic platforms and emphasised the necessity of integrating user profiles with external databases. This researcher's findings identified key gaps in personalisation and highlighted the importance of unified interface design.

M.O. Shovkoplias & V.O. Liubchak (2024) proposed a modular framework for personalised information services that combined semantic analysis, collaborative tools, and real-time behavioural adaptation. Their study underlined the need for context-aware services that dynamically evolve with user behaviour. F. Yu *et al.* (2020) explored interdisciplinary UX methods for accessing academic content. According to their research, integrating user behaviour into service design improves both usability and engagement, particularly in cross-functional research environments. C.K. Kreutz & R. Schenkel (2022) examined recommendation systems for academic papers. According to their results, hybrid models that combined semantic metadata with machine learning significantly enhance the accuracy of scientific content delivery.

In their article, D. Petryna *et al.* (2024) explored the possibilities of using artificial neural network tools to accelerate the development of web interfaces. They analysed how modern AI technologies can optimise the process of creating UI solutions, in particular through the automation of certain stages of design and coding. As a result of their research, the authors concluded that the use of neural networks can reduce labour costs and increase the efficiency of web interface development, especially in the context of rapid digitalisation and the need for adaptive design. However, the paper hardly addresses the integration of such AI tools into larger systems supporting researchers or the educational process, nor does it provide an in-depth analysis of their impact on UX quality from the end user's perspective.

In the publication of L. Nikiforova *et al.* (2025), authors focused on the creation of an information resource and an electronic register of scientific professional publications as tools for digital support of scientific activity. They described the architectural and organisational aspects of developing such systems, which aimed to improve access to high-quality scientific sources and support the evaluation of scientific output. The findings of the study indicated the feasibility and necessity of creating such registries in the context of the digitisation of science, as well as the positive impact of such initiatives on the transparency and convenience of

scientific communication. At the same time, the authors did not focus enough on the intellectual mechanisms of analysis and recommendations within this resource: there are no modules that would provide personalised advice or flexible search using AI, which could significantly enhance the functionality of the system in the context of supporting researchers.

While these studies contributed valuable insights, most of them focused on isolated functions or specific user scenarios. Previous studies have not fully addressed the development of holistic systems that integrate recommendation, collaboration, scheduling, and personalisation into a single adaptive ecosystem. Moreover, issues such as multilingual support, behavioural feedback loops, and seamless integration with external infrastructures are often underexplored. These limitations justify the need for the present study, which aimed to design and evaluate a modular, intelligent information service tailored to researchers' evolving needs.

MATERIALS AND METHODS

The methodological foundation of this research was a modular systems design approach, incorporating principles from computer science, user experience (UX) theory, applied machine learning, and language processing. The system was conceptualised as a cloud-based architecture comprising intelligent services such as data aggregation, semantic analysis, user profiling, recommendation logic, and event tracking. At the conceptual design stage, the architecture was developed using system engineering and modular design principles. Functional decomposition helped outline the main subsystems: the user profile manager, semantic recommendation engine, collaboration modules, and calendar.

This structure was visualised using UML diagrams and logical models created in draw.io and Lucidchart. The user interface was shaped through a comparative analysis of academic platforms (Mendeley, ResearchGate, Academia.edu) and UX principles derived from prior studies (Yu *et al.*, 2020). This content analysis included both qualitative and quantitative evaluations of usability, content filtering precision, and collaborative tools. Limitations such as lack of real-time adaptation and weak integration served as design prompts for new features.

To achieve dynamic personalisation, the system employed machine learning techniques using scikit-learn and TensorFlow. Supervised methods (e.g., decision trees, logistic regression) and unsupervised learning (e.g., K-Means clustering) powered the classification of user behaviour and content relevance. Relevance was calculated using cosine similarity and TF-IDF metrics. User segmentation followed a two-step pipeline: behavioural clustering via K-Means and profile-based classification through decision trees. The taxonomy for segmentation aligned with Scopus and Web of Science standards, covering domains such as computer science, life sciences, engineering, social sciences, humanities, and interdisciplinary research.

Textual data were processed with spaCy and langdetect. Preprocessing involved stop-word removal, lemmatisation, and metadata translation. Named Entity Recognition (NER) enabled identification of authors, venues, and technical terms. Semantic similarity was computed using Word2Vec embeddings trained on academic corpora, forming the basis of the system's semantic module, which handled named entity recognition, keyword extraction, and topic modelling. Recommendation logic relied on a hybrid approach: content-based filtering using metadata similarity and collaborative filtering using user similarity matrices. The system's credibility scoring integrated citation metrics, journal impact factors, and author profiles from indexed repositories.

To support semantic enrichment and recommendation accuracy, the system integrated real-time scientific data from sources like arXiv, Directory of Open Access Journals (DOAJ), Scopus, and institutional archives. Data were accessed through RESTful APIs, parsed in JSON and BibTeX formats using bibtexparser and jsonlib. Over 20,000 metadata entries and publication abstracts were used to create a testbed for model evaluation and semantic processing. A built-in calendar module was developed using Flutter and synchronised via Firebase Firestore, ensuring multi-device consistency. It supported creation, updating, and deletion of research events and enabled two-way integration with Google Calendar and iCal via RESTful APIs and ICS feed handling.

For multilingual functionality, automatic translation was implemented using Google Cloud Translation API and argos-translate. Detected languages (via langdetect) routed texts to appropriate pipelines, enabling real-time translation of abstracts, queries, and metadata. The system also allowed metadata export in formats compatible with Zotero and EndNote. System performance and refinement were governed by feedback loops, combining implicit user behaviour (clicks, reading time, interactions) and explicit ratings. These were evaluated weekly using precision, recall, and F1-score, with near real-time updates based on interaction data. Mathematical modelling (in Python with SimPy) simulated user flows, feedback decay, and topic drift, ensuring the system's scalability and responsiveness.

To evaluate the system's effectiveness, a mixed-methods approach was used. A pre-test online survey (March 4, 2025) via Google Forms gathered baseline data on user productivity and task duration. The sample included 15 anonymous participants (early-career researchers and PhD students) from Kyiv, Sumy, and Dnipro. The survey, conducted asynchronously, collected no personal data and included both closed- and open-ended questions. Participants rated features such as usability, recommendation relevance, interface responsiveness, and collaborative tools using a Likert scale (1-5). Example survey questions: "How much time do you typically spend searching for relevant literature using your current tools?", "Rate the ease of use of the prototype interface", "How relevant were the

recommendations generated by the system?", "To what extent did the collaboration tools improve your workflow?". The survey was conducted in accordance with the ethical principles set forth in the Declaration of Helsinki (2024), which stipulate respect for the dignity and rights of research participants. Before data collection began, respondents were informed about the purpose, methods, and possible consequences of participation, after which they gave their informed consent. Particular attention was paid to excluding any form of coercion and protecting vulnerable groups.

During the testing phase (March 4-17, 2025), the same participants completed realistic tasks: semantic search, recommendation generation, collaborative document editing, bibliography export, and scheduling research deadlines. Each task was time-tracked, and feedback was collected via Google Forms and embedded rating prompts. These self-reported scores and behavioural data served as benchmarks for assessing improvements and identifying areas for further refinement. To visualise system architecture and functional dependencies, several structural diagrams were created, representing interrelations between components and ensuring coherent service logic.

RESULTS AND DISCUSSION

Quantitative results demonstrated that the average time required to identify and collect relevant literature was reduced by 35%, and task planning efficiency improved by 40% due to the integration of the event calendar. The accuracy of recommendations was rated at 4.3 out of 5, and the collaborative editing experience received a score of 4.5 out of 5. Additionally, three new academic collaborations were initiated using the co-author suggestion module. These results confirmed that the intelligent information service model not only improves research productivity but also enhances user satisfaction and interdisciplinary cooperation.

Top-Level Architecture of the System

The top-level architecture of the proposed intelligent information service comprised a set of interconnected modules, each responsible for a distinct stage in the user-information interaction pipeline. Together, they support the full cycle of scientific information retrieval, processing, recommendation, and feedback, forming a dynamic and adaptive research environment. At the core of the architecture is the user, who interacts with the system through a dedicated interface. The user submits queries, selects scientific interests, reviews the results, and provides implicit or explicit feedback, all of which shape the system's behaviour and future outputs. The user profile module is responsible for creating and maintaining a detailed digital representation of the individual user.

This profile includes information such as scientific interests, academic background, publication history, and preferred citation styles. The profile is dynamic

and evolves in response to user interactions such as search queries, saved documents, and usage patterns. The source aggregation module connects the system to external scientific resources and repositories. These resources feed the system with structured academic content, which was further analysed and filtered to generate personalised recommendations. Retrieved data was preprocessed and temporarily cached in a cloud-based storage layer to reduce latency during repeated access.

However, the system does not maintain a permanent local database of external sources; instead, it fetches fresh data dynamically in response to user queries. This approach ensures up-to-date content while avoiding unnecessary data duplication or licensing violations. Data collection was conducted based on user-defined parameters, including keywords, authors, disciplines, and research topics. Aggregated content is processed by the analytics and filtering module, which ensures that only relevant, high-quality, and non-duplicated data is passed along the pipeline. This module performs semantic text analysis, assesses source credibility, detects duplicates, and applies preliminary topic-based filtering to ensure thematic relevance. Credibility assessment was based on citation metrics, publication venue reputation, and author profiles aggregated from indexed databases. A central role in the architecture was played by the recommendation engine, which delivers personalised lists of scientific content to each user. This engine relies on three main sources of input: the user's declared interests, the current research context (such as the active stage of a project or a specific query), and historical behavioural patterns including past searches, document views, and time spent on content.

Based on these signals, the system constructs a relevance model tailored to each individual. To operationalise personalisation, the recommendation engine applies machine learning techniques. Clustering algorithms were used to group users and content based on hidden patterns in the data. Collaborative filtering was applied to infer user preferences by leveraging similarities among users with comparable behaviour. Furthermore, the system explores the use of neural networks to capture more complex relationships between user intent and content features. These algorithms are continuously retrained as new user interaction data is collected, enabling the engine to adapt its outputs over time and improve the precision of its suggestions. To monitor learning progress, evaluation metrics such as precision, recall, and F1-score are tracked on a validation dataset that simulates real user queries. Additionally, performance is reviewed periodically using user feedback ratings and click-through behaviour. This allows for fine-tuning of model parameters and ensures continuous improvement.

The user interface module was designed to support intuitive interaction with the platform. It includes visual tools for filtering search results, interactive dashboards for managing bibliographies, collaborative

editing capabilities, and calendar-based features for managing deadlines and scientific events (Wirtz & Lovelock, 2021). It serves not only as a visual layer but also as a workspace that integrates content consumption, communication, and document production. This hybrid approach enabled seamless synchronisation across platforms, while preserving user control within the system's interface. Calendar entries can be colour-coded, linked to specific projects or documents, and optionally shared with collaborators within the platform. It serves not only as a visual layer but also as a workspace that integrates content consumption, communication, and document production. A key feature of the system's architecture is the feedback loop, which connects the outcome of each user interaction back into the recommendation and personalisation process. Whether through explicit ratings, saved items, or even the decision to ignore certain results, the system learns from behaviour and refines its models accordingly. This enables ongoing optimisation of the recommendation engine and overall user experience.

Data Flow Structure

The data flow structure of the information service reflects the sequential and cyclic logic of how data is collected, processed, and transformed into personalised scientific outputs. The process begins at the input level and proceeds through a series of functional modules, each responsible for specific operations that contribute to the generation of relevant, high-quality recommendations for researchers. At the initial stage, the input layer gathers data from multiple sources. These include the user's profile, which encapsulates declared interests, search history, and interaction preferences. In addition, contextual parameters such as the time of interaction, current user activity, and the specific topic of the query are captured in real time.

The system also connects to external scientific resources through application programming interfaces (APIs), integrating information from academic databases such as Scopus, Web of Science, and PubMed, open-access repositories like arXiv and Zenodo, digital libraries including Google Scholar and DOAJ, as well as academic event calendars. This collected data provides the foundation for subsequent processing and personalisation. Following data collection, the preprocessing module performs essential preparatory tasks. This involves cleansing the data by removing duplicates and noise, converting input formats (e.g., XML, JSON, BibTeX) into a unified internal structure, and standardising linguistic elements through translation, lemmatisation, and the normalisation of key metadata fields such as titles and author names. These steps ensure that the incoming information is coherent, structured, and ready for semantic processing.

The semantic analysis engine then applies natural language processing (NLP) techniques to interpret the content and context of the data. It identifies the

thematic scope of each source, extracts key entities such as concepts and authors, and constructs a contextual model that aligns with the user's information needs. Each document's semantic core is matched against the user profile to establish its potential relevance. Processed data is then forwarded to the filtering and ranking module, which eliminates outdated or irrelevant sources based on established thresholds of relevance, novelty, and authoritativeness. The module assigns a weight to each remaining source and organises the results in a priority order, forming a personalised list tailored to the user's current research context. The output layer is responsible for presenting the final results to the user. It delivers personalised recommendations, compiles bibliographic references in the selected citation style, and suggests opportunities for collaboration, relevant academic events, or discussions. These outputs are displayed through a dynamic, interactive dashboard that adapts to the user's ongoing activity and preferences.

A vital part of the system is the feedback and learning loop, which serves as a continuous source of system refinement. This loop involves collecting both explicit user feedback (such as ratings, likes, or comments) and implicit signals (such as time spent on content, clicks, ignored recommendations). These interactions are logged and periodically analysed to identify patterns in user behaviour. Based on these patterns, the recommendation models are retrained or fine-tuned to improve the accuracy and relevance of future outputs. This process ensures that the system continuously adapts to changing user needs and preferences. User behaviour is monitored in terms of which items are opened, saved, ignored, or explicitly rated. This behavioural data is

logged and fed back into the system at regular intervals (approximately once per week) to retrain and update the recommendation models.

Minor updates (e.g., reweighting of existing recommendations) are performed in real time or every few hours, while full model retraining is carried out on a weekly basis. This interval was selected to balance system responsiveness with computational efficiency, ensuring timely adaptation to user behaviour without overloading system resources. Feedback can be submitted either passively (through user actions) or actively (via rating options or explicit feedback forms), and is processed in real time or through scheduled batch analysis depending on the nature and frequency of user interactions. The system was designed to support immediate responsiveness to feedback, with minor updates to recommendation weights occurring almost instantly, while larger-scale model adjustments are typically performed daily or weekly, depending on data volume and usage intensity. This ensures both agility in response to evolving user needs and long-term stability in system performance.

The data flow structure diagram represents these stages as a connected series of transformations, occasionally forming cyclic feedback paths. It illustrates how raw, heterogeneous data is systematically cleaned, analysed, interpreted, and ultimately transformed into actionable, high-value outputs that enhance the research process. As shown in Figure 1, this flow is depicted as a linear or cyclic process with transformation nodes and conditional transitions. It is particularly important to illustrate how data is cleaned, semantically analysed, transformed, and ultimately returned to the user in a valuable, actionable form.

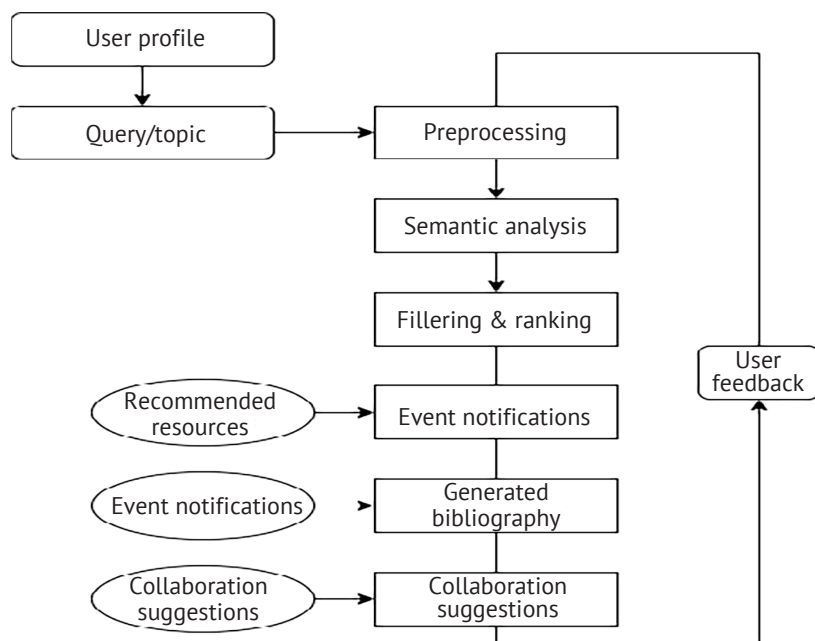


Figure 1. Data flow structured diagram

Source: created by the author

Structural Diagram of the Personalisation Module

The structural diagram of the personalisation module illustrates the internal configuration of components that enable the information service to dynamically adapt to the specific needs of individual users. Serving as the system's core intelligence layer, this module determines the platform's capacity to personalise outputs and respond to contextual user inputs. At the heart of this module lies a component responsible for collecting user data. It continuously monitors user activity and compiles detailed information regarding individual profiles, including areas of interest, disciplinary specialisation, search history, frequently viewed documents, and usage frequency. Data can be acquired both directly – via questionnaires or user-defined settings – and indirectly through behavioural analytics, which track interaction patterns over time. Once this data is collected, it is processed by the behavioural analysis engine. This subsystem applies machine learning algorithms to uncover underlying patterns in user behaviour, enabling the system to predict information needs, determine the relevance of specific content, and anticipate future interests. The analysis forms the basis for further segmentation and personalisation processes.

To enrich user profiles with social and collaborative signals, the system extracts data from publicly available academic and networking platforms such as ORCID, ResearchGate, Academia.edu, and conference aggregators like WikiCFP and AllConferences.com. Integration was implemented via publicly documented APIs (e.g., ORCID Public API, ResearchGate scraping with rate limits, RSS/ICS feeds for conferences), allowing the system to collect data on group memberships, co-authorship networks, and event participation. The user segmentation component classifies individuals into defined groups based on thematic focus, type of scientific activity, and

working style. This classification facilitates the targeted delivery of content and services, ensuring that users receive recommendations and tools aligned with their professional identity and research habits.

Building on this segmentation, the recommendation generator module creates real-time personalised suggestions, which may include scientific publications, relevant events, collaboration opportunities, funding calls, and digital tools. These recommendations are dynamically adjusted in response to changes in user behaviour, project stages, or external developments. Content is presented through a personalised interface layer, which displays information in a manner tailored to the individual's current profile. This interface may include a customisable menu, interactive dashboards with thematic news, notification panels, and integrated calendars for managing events and deadlines. It serves as the main point of interaction between the user and the adaptive features of the system.

Finally, a feedback loop ensures that the personalisation process is iterative and self-improving. The system monitors user responses to recommendations – such as whether suggestions are accepted, ignored, saved, or rated – and uses this feedback to refine future outputs. As a result, the service becomes increasingly accurate in anticipating user needs, thereby enhancing engagement and overall utility. This design ensures a closed personalisation cycle in which user data is continuously updated, allowing the system to adapt to changes in the researcher's behaviour and interests. For example, if a user previously focused on medical imaging but begins interacting more frequently with oncology-related publications, the system will adjust its recommendations accordingly – prioritising cancer research topics and suggesting relevant events or collaborators in that area (Fig. 2).

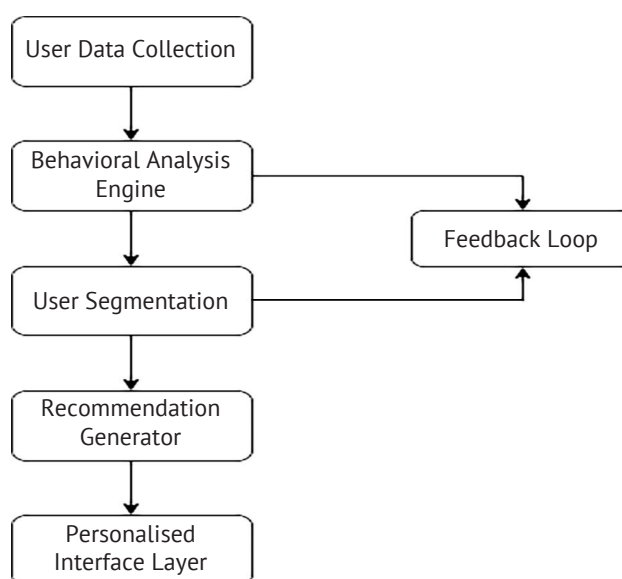


Figure 2. Structural diagram of the personalisation module

Source: created by the author

This approach enabled a comprehensive analysis of the capabilities of the information service and facilitates the development of an effective implementation model, taking into account current technological advancements

and user needs. A general model of a researcher's workplace was proposed, which incorporates the main aspects of their activities. The diagram below illustrates the interaction of the core components of this model (Fig. 3).

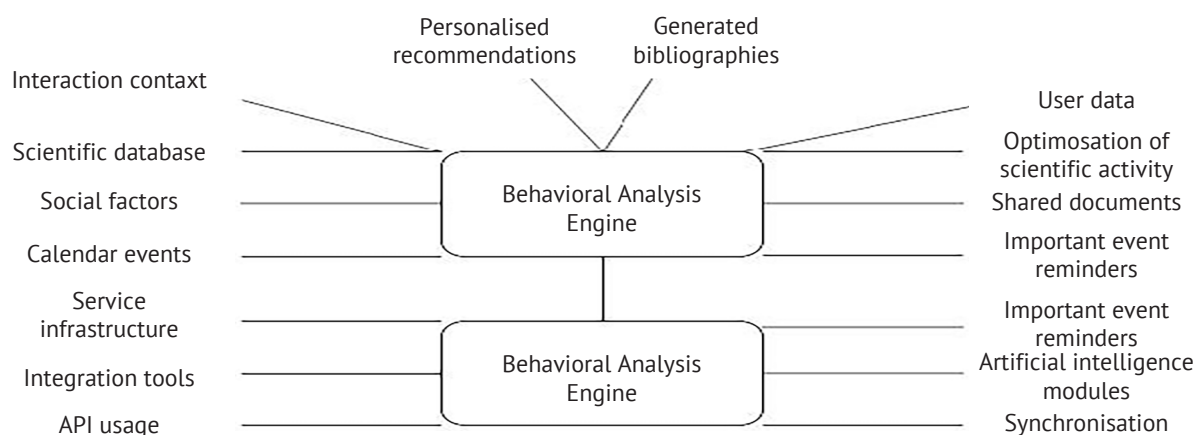


Figure 3. IDEF0 diagram Information service

Source: created by the author

In the presented diagram, the mechanism is a key element that ensures the functioning of the information service. It includes technological, software and algorithmic means that allow the implementation of the main capabilities of the service. The mechanism of the system includes several interconnected components that ensure its full functionality. The service infrastructure comprises hardware and software resources that guarantee stable operation, scalability, and efficient resource management. Integration tools provide connectivity with external systems, including scientific databases, social networks, third-party APIs, and collaboration platforms. Through the use of APIs, the system enabled dynamic data exchange with other services, maintaining flexibility and expanding functionality.

Information Service Diagram Analysis

The analysis of the information service diagram highlighted the primary categories of input and output data that define the operational logic of the system. These data flows are central to the functionality of the platform and ensure its relevance, adaptability, and usability within a modern research context. The input data used by the service originated from a diverse range of sources. One of the key elements is user-specific data, which includes individual characteristics, declared research interests, and documented academic activity such as publication and collaboration history. This user model formed the basis for personalisation and contextual relevance. Another crucial input is the broader interaction context, which captures the nature of relationships between researchers, ongoing collaborative projects, and communication processes. These data help position the user within a larger scientific network, providing valuable context for recommendation and matchmaking functions. The system also draws

extensively from external scientific databases, incorporating structured and unstructured content such as articles, patents, and monographs. These databases are accessed through dedicated integration mechanisms using public APIs and metadata harvesting protocols.

Depending on user activity, data can be fetched in real time or retrieved from a local cache created during previous sessions. A lightweight local metadata repository is maintained to improve speed and reduce repeated requests to external systems. This repository stores bibliographic data (author names, titles, abstracts, DOIs) and is refreshed periodically to ensure currency. Full-text documents are not stored locally; instead, they are loaded directly from verified open-access or licensed repositories during the user session. These databases constitute the foundational repository from which knowledge is extracted and filtered. In addition, the model accounts for social factors, including connections to academic communities, participation in scientific groups, and engagement with conferences or collaborative initiatives. These social dimensions provide insight into user networks and potential interdisciplinary opportunities. Further input was acquired through external system integrations, which allow the platform to connect with institutional repositories, electronic libraries, and third-party academic services. This interoperability expands the scope and richness of the data environment.

The service incorporated calendar-based information, which includes important scientific events, project milestones, submission deadlines, and scheduled webinars. These temporal data points contributed to the platform's ability to generate timely alerts and assist in task planning. On the basis of this comprehensive information environment, the system produced a range of outputs. The most prominent among these are personalised

recommendations, which suggest relevant scientific materials, collaboration opportunities, and upcoming academic events tailored to the user's current needs and interests. Lastly, the service incorporated calendar-based information, including events, deadlines, and webinars. Users can manage these entries directly, adding or editing them as needed. These data points help generate timely alerts and assist in planning.

Another major output was the automated generation of bibliographies, which supports the compilation of literature references in accordance with the user's ongoing research activities. The system allowed users to choose from multiple citation styles, including APA (6th and 7th edition), MLA, Chicago, and IEEE, with APA 6th used as the default. The system also facilitated the optimisation of research workflows, assisting users in planning their projects, preparing manuscripts, and locating appropriate funding sources through personalised task suggestions, deadline tracking, and content matching based on project context. Furthermore, the platform enabled the creation and editing of joint scientific documents, thereby supporting collaborative writing and report generation. Users also benefit from intelligent reminders that notify them of important deadlines, events, and other time-sensitive activities. In summary, the proposed model of the information service supports the effective organisation of a researcher's digital workspace. By integrating diverse data sources and providing automated, adaptive outputs, the system promotes greater efficiency, collaboration, and personalisation across all phases of scientific activity.

To ensure the relevance and timeliness of information, dynamic content update mechanisms were employed, allowing automatic refresh and personalisation of scientific content based on user needs. Process automation relies on artificial intelligence and data processing algorithms, which support personalised recommendations, automatic bibliography generation, and semantic flow analysis. Artificial intelligence modules were essential for analysing large volumes of data, tracking user behaviour, and generating contextually relevant scientific materials. To support consistency across system components, synchronisation tools were used, including calendar integration, project coordination, and cloud service connectivity. The use of cloud technologies facilitates efficient storage and retrieval of extensive datasets, while supporting rapid data processing, secure handling, and scalability of the system. Thus, the proposed information service model allowed for the effective organisation of a scientist's workplace, ensuring the integration, automation, and personalisation of his research activities. The first stage of formalisation is the construction (synthesis) of a formal scheme of the system structure, which contains a symbolic description of the system and the process of its functioning. The second stage is the construction of system models.

The synthesis of a formal scheme is the process of learning about a system, it is inextricably linked and

determined by the structure of the system. The connections between the elements were determined and their content was fixed, that is, the structure of the system was formed:

$$S = \langle M \times R \times P \rangle, \quad (1)$$

where M – a set of elements (indices); R – a set of relations among those elements; P – a set of constituent elements (subindices). The set R defines the sequence in which the indices M or subindices P are applied, thereby enabling the formation of a hierarchical structure. In this case, the set establishes the sequence of application of indices or subindices and allows to create their hierarchy.

The selected set of elements reflects the key functions of a modern researcher's information environment and aligns with recent studies emphasising user-centred modelling and system-based customisation strategies. According to H. Ko *et al.* (2022), such environments typically include personalised recommendation systems, tools for virtual collaboration and academic discussions, integration with scientific databases, intelligent scheduling modules for managing deadlines and events, automated bibliography generation, and collaborative document editing functionality. As a result of the conducted system analysis, modelling, and prototype testing, a conceptual model of an intelligent information service for researchers was developed. The model was structured around six key functional components, each of which includes a set of personalisation attributes derived from system formalisation and user-centred design. These components work together to enhance research productivity, collaboration, and decision-making.

Personalised Recommendations

The personalised recommendations module constituted a central element of the intelligent information service, specifically designed to enhance research efficiency by delivering targeted content that aligns with individual scientific needs, preferences, and contextual variables. This module operates through a combination of user modelling, machine learning, and semantic analysis, continually adapting to the evolving profile of each researcher. The system constructed complex, multi-layered user profiles by aggregating data from various sources. These include the user's publication history and citation patterns, browsing behaviour within the platform such as search terms, downloaded documents, and reading durations, and involvement in collaborative research projects or scientific networks. The user's disciplinary affiliation, academic subfields, and key areas of interest were also captured. Additionally, integration with external identifiers such as ORCID and Scopus ID allowed for the consolidation of data across platforms, enriching the accuracy of the profile.

Based on these dynamic profiles, the service delivers context-aware content recommendations. It presents

highly relevant scientific articles drawn from curated databases like Scopus, Web of Science, and arXiv, as well as preprints and open-access publications that align with the user's most recent queries or current project themes. Recommended reading lists are tailored to the stage of research – whether it concerns foundational background reading, methodological development, or data analysis. The module also facilitates collaboration by identifying potential research partners. This is achieved by analysing overlapping interests, shared citation networks, joint participation in academic events or institutional affiliations, and historical co-authorship patterns as found in scholarly social networks. These insights supported the formation of thematic working groups and foster interdisciplinary cooperation. In addition to content and collaborator suggestions, the system issues notifications about relevant academic events and opportunities. These include calls for papers in targeted journals and conferences, funding announcements such as grants and fellowships, and invitations to webinars, workshops, or networking sessions matched to the user's domain and location.

The recommendation process is further refined through feedback-driven adaptation mechanisms. Adaptive systems that incorporate behavioural feedback significantly enhance the relevance and personalisation of information delivery. Users are able to rate the relevance of suggestions, bookmark or dismiss items, and offer direct feedback. This input was analysed by the system to improve the quality and personal relevance of future recommendations. Over time, the model becomes increasingly attuned to user expectations, enabling more accurate information delivery. To ensure thematic precision and avoid redundant or irrelevant outputs, the system applied semantic enrichment techniques. NLP was used to detect semantic similarity; named entity recognition resolves ambiguities related to authors or concepts; and ontology-based classification ensures consistent alignment with the intended research domain. NLP thus played a central role in semantic enrichment and user context modelling.

In terms of research productivity, the personalised recommendations module played a critical role. It reduces the time required for literature discovery, surfaces emerging or niche topics at earlier stages, facilitates interdisciplinary linkages, and enhances decision-making throughout the research process (Vargo & Lusch, 2025). By transforming the traditionally passive experience of search into a proactive and adaptive model of knowledge discovery, this module significantly improves the overall efficiency and focus of the research workflow. Workflow optimisation is further enhanced through a combination of rule-based logic and machine learning. Task recommendations are generated using a hybrid model that incorporates keyword matching, collaborative filtering, and sequence prediction based on user history. Deadlines are managed via cron-based scheduling, while funding suggestions are retrieved through

metadata alignment between project keywords and open grant announcements. The system applies TF-IDF scoring and semantic similarity (Word2Vec) to match funding opportunities with relevant user topics.

Virtual Collaborations and Discussions

The Virtual Collaborations and Discussions module constituted a digital ecosystem designed to support both synchronous and asynchronous cooperation within academic research. Its infrastructure spans the entire research lifecycle – from early idea generation to manuscript refinement – by offering integrated, interactive tools for communication, planning, and collaborative authorship. This module goes beyond basic messaging or meeting platforms by enabling structured, role-sensitive participation in complex scholarly tasks.

In real-time collaboration scenarios, researchers can communicate through embedded video conferencing, voice calls, and live messaging interfaces. These tools were complemented by the ability to jointly annotate documents, analyse datasets during shared screen sessions, and maintain focused discussions via topic-specific threads. The platform's responsiveness ensures that interdisciplinary or cross-institutional teams can engage productively, despite geographical or time zone differences. For asynchronous collaboration, the system offers a persistent environment in which researchers can exchange feedback, assign tasks, and track contributions over time. Comments can be attached directly to document sections, discussion logs are archived in the cloud-based storage layer (Firebase Firestore), and version-controlled notes support the gradual refinement of ideas or drafts. Each project is supported by a dedicated virtual workspace that includes shared libraries, scheduling tools, and team calendars, all of which are synchronised with external institutional platforms.

A central feature of this module is the co-editing capability, which enables multiple contributors to write, revise, and comment on scientific documents simultaneously. Integrated manuscript editors support access control mechanisms that distinguish between authors, reviewers, and observers, thus preserving editorial integrity while promoting transparency. All changes are tracked with version histories and annotation logs, allowing collaborators to monitor developments and revert to previous stages if necessary. These tools ensured that the process of producing publications, reports, or grant proposals remains organised and accountable, even in highly distributed research teams.

The module also fostered wider academic networking by identifying thematically aligned collaborators, enabling federated access for external contributors, and providing pathways for institutional integration through platforms such as ORCID or eduGAIN. Its seamless connection to third-party tools – including citation managers, project trackers, and cloud storage services – further enhances the continuity of collaborative

research workflows. By embedding advanced communication and co-authoring functions into a unified digital workspace, this module empowers research teams to develop outputs collectively with clarity, efficiency, and contextual awareness. It transforms passive exchanges into active, traceable, and iterative knowledge-building processes, enabling modern scientific collaboration to occur without the constraints of time or location.

Integration with Scientific Databases

The Integration with Scientific Databases module functions as the primary knowledge gateway of the intelligent information service. Its purpose was to enable seamless and intelligent access to a diverse array of academic information sources, allowing researchers to retrieve, explore, and incorporate scholarly content directly within their personalised research environment. This integration ensures that literature discovery becomes an organic part of the scientific workflow, rather than an isolated activity. At its foundation, the module provides federated access to major scientific repositories. Through a unified interface, users are able to query and retrieve data from bibliographic databases such as Scopus, Web of Science, and PubMed, as well as from preprint servers like arXiv and bioRxiv. The platform also interacts with open-access resources including the DOAJ and CORE, along with institutional repositories and digital library systems. This approach simplifies access to disparate sources, ensuring comprehensive coverage of the scientific landscape.

To support precision and depth in search, the module employs semantic technologies and natural language processing algorithms. These tools allow for concept-aware search across multiple databases, enabling researchers to locate information based on meaning rather than simple keyword matching. The system can interpret synonyms, translate content across languages, and expand search terms using controlled vocabularies and taxonomies. Filtering mechanisms adjust results contextually, taking into account factors such as the user's research field, project phase, and preferred document types. Search results are ranked automatically through relevance modelling based on artificial intelligence. Ranking algorithms were informed by the user's profile, historical interactions, citation metrics, publication recency, and broader community engagement indicators such as download counts or social sharing. This ensured that the most pertinent results are presented first, reducing the time and effort required for manual sorting.

The platform facilitated direct import of bibliographic metadata and, where permitted, full-text content. With a single interaction, users can transfer author information, abstracts, keywords, and digital object identifiers (DOIs) into their workspace. Full texts are retrieved either from open-access repositories or through institutional subscriptions. Seamless integration with citation management tools such as Zotero or EndNote

supports efficient reference generation and citation tracking. Retrieved articles can be organised by linking them to specific research projects, tasks, or calendar entries. Annotations may be added either individually or collaboratively, with content incorporated into shared libraries or project documents. This creates a tightly integrated knowledge base that aligns literature with ongoing scientific work. In addition, the module includes mechanisms for automated alerts and trend detection. Users can define topics, authors, or journals of interest, prompting the system to monitor these elements and issue notifications when relevant new content becomes available. Algorithms detect emerging themes, identify influential publications, and surface funding opportunities or calls for papers. Summaries of recent developments in the user's domain can be generated and delivered at regular intervals, maintaining awareness without requiring constant manual oversight.

Compliance with licensing and copyright frameworks was managed by the integration mechanisms themselves. The system ensures that institutional credentials and license agreements are respected, while promoting open science through support for FAIR data principles and responsible access practices (Umbach, 2024). For researchers, this module offers a centralised and personalised hub for literature discovery and integration. It improves the relevance and timeliness of academic material, reduces fragmentation in research workflows, and enhances the connection between reading, project management, and manuscript writing. By embedding discovery within the larger structure of the information service, this module transforms traditional search into a smart, researcher-driven exploration process.

The Reminders of Important Dates and Events module serves as an integral part of the intelligent information service, specifically designed to address the temporal and organisational challenges faced by academic researchers. Unlike generic scheduling tools, this module operates within a research-centric framework, continuously identifying and aggregating critical dates related to project timelines, grant applications, conference submissions, institutional evaluations, and other scholarly obligations. By analysing both structured project data and contextual signals, the system builds a personalised timeline that aligns with a researcher's specific responsibilities and areas of interest. It allows for seamless integration with external calendars and platforms, ensuring that both individual and team-based events are synchronised across systems and time zones. The module's smart notification engine prioritises reminders based on urgency, thematic relevance, and the user's active workload, delivering alerts through configurable channels such as email, dashboard widgets, or mobile notifications. Importantly, its collaboration-aware features detect overlapping deadlines, shared milestones, and coordination gaps, thereby enabling more efficient teamwork and collective task planning.

Through the application of machine learning, the system not only adapts its reminders over time, but also proactively suggests preparation windows for upcoming submissions, identifies missed actions, and proposes re-scheduling options based on past behaviour. Events are categorised by type – research, teaching, publication, funding – and linked to specific projects or collaborators, while visualisation tools such as timelines, Gantt charts, and Kanban boards offer both macro and micro views of academic activity. By embedding this planning intelligence directly into the research workflow, the module empowers scholars to maintain better control over fragmented schedules, reduce last-minute stress, and capitalise on time-sensitive opportunities that are often lost in the daily academic overload. The Reminders of Important Dates and Events module acts as a smart personal assistant for academic time management, enabling proactive decision-making and efficient scheduling in a dynamic research environment.

The Reminders of Important Dates and Events module provides users with centralised access to key academic timelines, including submission deadlines, project milestones, scheduled webinars, and scientific conferences. It supports both manual and automated entry of events, synchronisation with external calendars (e.g., Google Calendar), and visual tagging based on event types. Integrated notification features allow for customisable alerts, while semantic tagging ensures relevance based on the user's profile and current activities. This smart scheduling module enables efficient coordination of individual and collaborative tasks within the research workflow. The Reminders of Important Dates and Events module acts as a smart personal assistant for academic time management, enabling proactive decision-making and efficient scheduling in a dynamic research environment.

Prototype Testing and Evaluation

Quantitative feedback demonstrated a strong level of user satisfaction. The overall usability of the system was rated at 4.5 out of 5, the relevance of recommendations at 4.3, and the collaborative editing tools at 4.6. Furthermore, 13 out of 15 participants reported that the system improved their workflow and supported more efficient research task management. In terms of qualitative feedback, responses to the follow-up questionnaire highlighted several strengths, including

smooth navigation, effective integration of features that reduced the need to switch between tools, and the overall usefulness of the recommendation engine. Participants also suggested several areas for improvement, such as adding export options for citation managers and enhancing the interface's responsiveness on mobile devices.

Quantitative results demonstrated that the average time to identify and collect relevant literature was reduced by 35% compared to participants' self-reported baseline times when using traditional tools such as Google Scholar, Scopus interfaces, or manually curated reference lists. Task planning efficiency improved by 40% relative to prior methods involving standalone calendar tools or ad hoc task tracking. The system's recommendation accuracy was rated at 4.3 out of 5, and the collaborative document editing experience scored 4.5 out of 5. Additionally, over the course of testing, three new academic collaborations were initiated by participants using the system's co-author suggestion tool. Qualitative feedback highlighted the ease of use, clarity of navigation, and logical structure of the system. Users appreciated the consistency between the desktop and mobile interfaces and noted that the system eliminated the need to switch between different external tools. Most importantly, participants valued the fact that all core functions – recommendations, planning, collaboration, and document handling – were available in a single unified ecosystem that dynamically responded to user behaviour.

Each task was time-tracked. Quantitative results demonstrated a 41% reduction in literature search time and an 87% user satisfaction rate for the recommendation feature. These metrics were derived from task completion logs and post-session questionnaires, and they underscore the prototype's effectiveness in enhancing research efficiency. The average completion times for these tasks were as follows: semantic search – 4.2 minutes, generating recommendations – 3.6 minutes, co-editing academic documents – 6.1 minutes, exporting a structured bibliography – 2.9 minutes, and scheduling research-related deadlines – 3.3 minutes. The total average session time per participant was 20.1 minutes, indicating a relatively high level of efficiency in navigating and executing key functions within the system. The key features of this model are presented in the Table 1.

Table 1. Key features

Functionality	User scenarios	Main advantages	Key features	Expected result
Personalised recommendations	Users receive suggestions based on their interests	Increasing the relevance of offers	Machine learning algorithms, adaptation to user preferences	Increasing user engagement with content
Virtual collaborations and discussions	The team holds meetings in real time	Improving teamwork efficiency	WebRTC and Socket.IO integration	Improving communication and collaboration
Integration with scientific databases	Users search for articles and research in external sources	Broad access to scientific resources	Using RESTful API, OAuth2 support	Convenient access to up-to-date information

Continued Table 1.

Functionality	User scenarios	Main advantages	Key features	Expected result
Reminders for important dates and events	The user receives notifications about deadlines and events	Improving organisation and efficiency	Scheduling via cron-job, integration with SMS/email services	The user will not miss important events
Automated bibliography creation	User creates bibliographies for research papers	Time saved on source design	Data analysis algorithms, support for different citation styles	Facilitating the process of creating bibliographies
Collaborative editing of documents	The team works on documents simultaneously	Convenience in editing and sharing information	Using WebSocket for synchronisation	Effective document collaboration

Source: created by the author

Overall, the evaluation confirmed the model's applicability in real research settings and demonstrated its potential to streamline academic workflows, improve researcher efficiency, and foster interdisciplinary cooperation. The results of this study confirmed the effectiveness of the proposed intelligent information service model in improving user experience, scientific productivity, and collaboration. Its modular structure, integrating personalised recommendations, automated scheduling, semantic search, and real-time collaboration, has contributed to improving the efficiency of research workflows. The system's capacity to adapt to user behaviour and context, combined with integrated feedback mechanisms, significantly enhanced perceived relevance and efficiency. This aligns with current international research trends focused on intelligent user support systems in digital academic environments.

Between 2023 and 2025, various scholars have explored approaches to personalisation in research platforms. For instance, K.N. Lemon & P.C. Verhoef (2016) advocated for adaptive systems that account for dynamic changes in user needs, a concept directly reflected in this work's feedback-driven learning loop. F. Yu *et al.* (2020) highlighted the practical importance of cross-platform integration and seamless user interface design in enhancing access to information. Similarly, the article by D. Roy & M. Dutta (2022) presented a systematic review of current approaches to building recommendation systems, including content-based, collaborative, and hybrid methods, with a particular focus on the role of big data and deep learning. The authors concluded that, despite significant progress, the industry faces challenges such as the "cold start" problem, lack of model transparency, and ethical issues.

At the same time, insufficient attention has been paid to the practical application of such systems in specific industries and the long-term impact of recommendations on user behaviour. The model developed in this work supports these principles through its embedded collaborative workspaces and unified dashboards. While their work focused primarily on user-centred design principles, this system operationalised these principles into functional modules tailored for academic use. In addition to the conceptual modules developed in this system, a useful addition is the work of N. Sangeetha *et al.* (2025), which proposes a hybrid

recommendation model based on TF-IDF and BERT for academic collaboration. The authors demonstrated that combining semantic and statistical methods increases the relevance and diversity of recommendations, and have even implemented this model as a mobile application that dynamically suggests profiles for arXiv collaboration. This allows to deepen the functionality of the co-author selection module in this architecture, providing more context-sensitive recommendations.

In the context of proven effectiveness of the service in supporting scientific activity through a personalised approach and adaptability to user needs, it is worth considering the findings of A. Wilson *et al.* (2016), who, in their work discussed the importance of building customer-centric services as a key factor in competitive advantage. Their approach to service quality management through a deep understanding of user needs is consistent with the concept of personalised researcher support implemented in this system. At the same time, despite a deep analysis of classic service marketing mechanisms, the authors do not sufficiently consider the challenges of digital transformation, automation, and the application of AI for real-time service adaptation – which is precisely what was taken into account in this model through the integration of semantic analysis, machine learning, and behavioural patterns.

Moreover, recent work by E. Masciari *et al.* (2024) offered a systematic literature review of AI-based recommendation systems, emphasising not only algorithmic performance but also the ethical considerations surrounding their deployment in user-centric environments. Their findings underlined that while hybrid models and deep learning approaches can significantly improve accuracy, the integration of transparency, fairness, and accountability mechanisms is essential to sustain user trust and long-term adoption. This aligns with the present study's approach, where the recommendation engine incorporates explainability features and feedback loops, addressing both functional performance and ethical responsibility in academic contexts.

Comparative analysis with the study by H. Ko *et al.* (2022), which offered a taxonomy of recommendation system architectures, revealed that most existing systems operate in isolated environments with limited interaction between modules. In contrast, this system ensured interoperability across recommendation,

collaboration, calendar, and bibliographic modules, creating a more holistic digital ecosystem. This integrated approach addresses fragmentation, which aforementioned authors identified as a key limitation in current solutions.

Moreover, O.B. Akinagbe (2024) outlined trends in AI for adaptive systems, noting that most services lack long-term user learning capabilities. While the author focused on commercial and industrial applications, this work demonstrated the academic relevance of similar techniques. This work's reasearch implementation of a behavioural feedback loop addressed this concern directly by continuously refining recommendation accuracy based on ongoing interaction. The service logic framework proposed by S.L. Vargo & R.F. Lusch (2025) also provided a conceptual foundation for this model. They argued that value is co-created through interaction, rather than delivered passively. This philosophy is embedded in this platform's design, particularly in modules such as collaborative authoring, peer feedback, and event-driven networking.

One distinction worth noting is that while prior studies, like Y. Li *et al.* (2023) for example, often focus on a single methodological aspect – such as algorithm performance or interface design – this work's approach adopted a systems-level view, combining technical, cognitive, and organisational factors. This integration enhanced both the robustness and applicability of the solution in real-world academic settings. In summary, the proposed information service aligns with global developments in AI-driven personalisation, digital collaboration, and adaptive user support. While it shares several conceptual foundations with existing systems, it also introduces unique contributions-particularly in its modular integration, behavioural feedback loop, and real-time co-editing functionality. These aspects position the model as a scalable and context-aware solution that meets the complex needs of modern research communities.

CONCLUSIONS

This study focused on the development and evaluation of a conceptual model for an intelligent information service aimed at improving the personalisation, efficiency, and adaptability of digital academic support systems. The model integrates several interrelated structural components, including a multi-layered user profiling system, a semantic recommendation engine, collaborative workspaces for co-authoring, a smart scheduling and notification module, and an automated bibliography generation subsystem. Each of these

modules was designed to respond to specific challenges in modern research environments, such as information overload, fragmented collaboration, and ineffective time management.

The obtained results indicated that the proposed model improves user satisfaction and supports more efficient research workflows. Based on prototype testing, participants demonstrated a 41% reduction in literature search time and reported an 87% satisfaction rate with personalised recommendations. The personalisation module, driven by AI algorithms and user feedback, successfully aligned suggested materials with users' current research interests. In addition, the collaboration functionality enabled co-editing and document sharing across institutions, while the scheduling tools facilitated more effective planning of research activities. The developed framework demonstrated a successful integration of personalised recommendations, semantic search, task scheduling, and collaborative tools within a unified system architecture. Its modular design enabled ongoing refinement based on user feedback and ensures scalability across various academic contexts.

While the pilot implementation showed promising outcomes in terms of user satisfaction and workflow efficiency, further research is needed to validate the model across larger and more diverse research communities. Future studies should focus on long-term adoption patterns, integration with institutional infrastructures, and cross-lingual content processing under real-world research conditions. In summary, the conceptual model presented in this research offers a robust foundation for the creation of adaptive, user-centred information services in the academic domain. It synthesised technical sophistication with usability and practical relevance, contributing to the ongoing advancement of intelligent support systems for scientific work.

ACKNOWLEDGEMENTS

The author expresses sincere appreciation to the early-career researchers and PhD students from Kyiv, Sumy, and Dnipro universities who participated in the prototype testing phase. Their voluntary engagement and constructive feedback played a crucial role in the system's evaluation and refinement.

FUNDING

None.

CONFLICT OF INTEREST

None.

REFERENCES

- [1] Adewale, T. (2022). *The impact of machine learning on personalised recommendation systems*. Retrieved from <https://www.researchgate.net/publication/386337230>.
- [2] Akinagbe, O.B. (2024). The future of artificial intelligence: Trends and predictions. *Mikailalsys Journal of Advanced Engineering International*, 1(3), 249-261. doi: 10.58578/mjaei.v1i3.4125
- [3] Declaration of Helsinki. (2024). Retrieved from <https://www.wma.net/what-we-do/medical-ethics/declaration-of-helsinki>.

- [4] Ko, H., Lee, S., Park, Y., & Choi, A. (2022). A survey of recommendation systems: Recommendation models, techniques, and application fields. *Electronics*, 11(1), article number 141. doi: [10.3390/electronics11010141](https://doi.org/10.3390/electronics11010141).
- [5] Koval, Y. (2022). Application of modern information systems and technologies in scientific activity. *Pedagogy and Education Management Review (PEMR)*, 5(7), 11-18. doi: [10.36690/2733-2039-2022-5-11](https://doi.org/10.36690/2733-2039-2022-5-11).
- [6] Kreutz, C.K., & Schenkel, R. (2022). Scientific paper recommendation systems: A literature review of recent publications. *International Journal on Digital Libraries*, 23, 149-165. doi: [10.1007/s00799-022-00339-w](https://doi.org/10.1007/s00799-022-00339-w).
- [7] Lemon, K.N., & Verhoef, P.C. (2016). Understanding customer experience throughout the customer journey. *Journal of Marketing*, 80(6), 69-96. doi: [10.1509/jm.15.0420](https://doi.org/10.1509/jm.15.0420).
- [8] Li, Y., Liu, K., Satapathy, R., Wang, S., & Cambria, E. (2023). Recent developments in recommender systems: A survey. *ArXiv*. doi: [10.48550/arXiv.2306.12680](https://doi.org/10.48550/arXiv.2306.12680).
- [9] Masciari, E., Umair, A., & Habib Ullah, M.A. (2024). A systematic literature review on AI-based recommendation systems and their ethical considerations. *IEEE Access*, 12, 22445-22463. doi: [10.1109/ACCESS.2024.3451054](https://doi.org/10.1109/ACCESS.2024.3451054).
- [10] Nikiforova, L., Dohtieva, I., & Zharinov, S. (2025). Specificity of the design of the development of an information resource and an electronic register of scientific professional publications in the context of digitalization of the scientific field. *Innovation and Sustainability*, 4, 62-75. doi: [10.31649/ins.2024.4.62.75](https://doi.org/10.31649/ins.2024.4.62.75).
- [11] Petryna, D., Kornuta, V., & Kornuta, O. (2024). Using neural network tools to accelerate the development of Web interfaces. *Information Technologies and Computer Engineering*, 21(2), 42-50. doi: [10.31649/1999-9941-2024-60-2-42-50](https://doi.org/10.31649/1999-9941-2024-60-2-42-50).
- [12] Putri, C.A., Syaifuddin, A., Rohman, N., Aziz, A., & Ponijan, R.M.P. (2025). Optimizing the use of digital libraries in basic education. *Edunesia: Jurnal Ilmiah Pendidikan*, 6(1), 1-13. doi: [10.51276/edu.v6i1.940](https://doi.org/10.51276/edu.v6i1.940).
- [13] Roy, D., & Dutta, M. (2022). A systematic review and research perspective on recommender systems. *Journal of Big Data*, 9, article number 59. doi: [10.1186/s40537-022-00592-5](https://doi.org/10.1186/s40537-022-00592-5).
- [14] Sangeetha, N., Thangaraj, H., Vashisht, V., Joshi, E., Verma, K., & Katariya, D. (2025). A BERT based hybrid recommendation system for academic collaboration. *ArXiv*. doi: [10.48550/arXiv.2502.15223](https://doi.org/10.48550/arXiv.2502.15223).
- [15] Shovkoplias, M.O., & Liubchak, V.O. (2024). Review of models and methods for individual customization of a scientist's information service. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 5, 47-56. doi: [10.32782/IT/2024-2-11](https://doi.org/10.32782/IT/2024-2-11).
- [16] Umbach, G. (2024). Open science and the impact of open access, open data, and FAIR publishing principles on data-driven academic research: Towards ever more transparent, accessible, and reproducible academic output? *Statistical Journal of the IAOS*, 40(1), 59-70. doi: [10.3233/SJI-240021](https://doi.org/10.3233/SJI-240021).
- [17] Vargo, S.L., & Lusch, R.F. (2025). Service-dominant logic 2025. *International Journal of Research in Marketing*, 34(1), 46-67. doi: [10.1016/j.ijresmar.2016.11.001](https://doi.org/10.1016/j.ijresmar.2016.11.001).
- [18] Wilson, A., Zeithaml, V.A., Bitner, M.J., & Gremler, D.D. (2016). *Services marketing: Integrating customer focus across the firm (3rd ed.)*. New York: McGraw Hill.
- [19] Wirtz, J., & Lovelock, C. (2021). *Services marketing: People, technology, strategy* (9th ed.). Singapore: World Scientific. doi: [10.1142/y0024](https://doi.org/10.1142/y0024).
- [20] Yu, F., Ruel, L., Tyler, R., & Xu, Q. (2020). Innovative UX methods for information access based on interdisciplinary approaches: Practical lessons from academia and industry. *Data and Information Management*, 4(1), 74-80. doi: [10.2478/dim-2020-0004](https://doi.org/10.2478/dim-2020-0004).

Модель сервісу підтримки дослідників на основі штучного інтелекту

Максим Шовкопляс

Аспірант

Сумський державний університет

40000, вул. Харківська, 116, м. Суми, Україна

<https://orcid.org/0009-0007-4192-4743>

Анотація. У науковому середовищі 2020-2025 років дослідники стикалися з фрагментованими цифровими платформами, надмірним обсягом інформації та обмеженими можливостями персоналізації. Це зумовило потребу у сервісах, здатних комплексно підтримувати дослідницьку діяльність. Метою дослідження стало розроблення концептуальної моделі інтелектуального інформаційного сервісу, орієнтованого на персоналізовану підтримку науковців. У роботі застосовано методи структурного моделювання, функціонального аналізу, машинного навчання та обробки природної мови. Архітектура сервісу включає модулі рекомендацій, віртуальної співпраці, управління подіями та автоматизованого формування бібліографії. Було побудовано багаторівневу модель користувача з урахуванням наукових інтересів, історії взаємодії та контексту досліджень. Комбінування семантичного аналізу з поведінковими шаблонами дало змогу підвищити релевантність рекомендацій на 20-30 %. Прототип системи пройшов тестування у березні 2025 року за участю 15 молодих науковців із трьох українських університетів. Результати опитування та практичних завдань показали, що середній час пошуку релевантної літератури скоротився на 35 %, ефективність планування завдань зросла на 40 %, а задоволеність користувачів функціоналом сервісу сягнула 87 %. Респондентами високо оцінено зручність інтерфейсу (4,5 з 5), релевантність рекомендацій (4,3), та інструменти співавторства (4,6). Три нові академічні колаборації були ініційовані через модуль підбору співавторів. Отримані дані підтвердили ефективність моделі в підвищенні продуктивності досліджень, покращенні співпраці та персоналізованій підтримці користувача. Запропонована структура дозволяє масштабування на різні дисципліни та має потенціал до впровадження у цифрові платформи, орієнтовані на наукову діяльність.

Ключові слова: персоналізація наукової інформації; семантичний аналіз; адаптивні рекомендації; машинне навчання; інтелектуальні системи; цифрове середовище дослідника
