

LLM-Powered Framework to Explore Summarized Aggregated Multimedia Vertical Web Search Results

Muhammad Wajeeh Uz Zaman¹, Umer Rashid¹, Abdur Rehman Khan²

¹Department of Computer Sciences, Quaid-i-Azam University Islamabad, Pakistan.

²School of Computer Science, Centre for Data Science, Queensland University of Technology, Brisbane, Australia

*Correspondence: umerrashid@qau.edu.pk

Citation | Zaman. M. W. U, Rashid. U, Khan. A. R, "LLM-Powered Framework to Explore Summarized Aggregated Multimedia Vertical Web Search Results", IJIST, Special Issue pp 01-11, April 2026

DOI | <https://doi.org/10.33411/IJIST/1777>

Received | March 10, 2026 **Revised** | April 15, 2026 **Accepted** | April 22, 2026 **Published** | April 26, 2026.

The exponential growth of multimedia content has shifted users' information-seeking behavior from lookup-based to exploratory search. To aid exploration, search engines adopted two prominent approaches: presenting results in verticals (web images, videos, news) and integrating Generative AI (GenAI) to enable rapid comprehension. However, integrations like GenAI primarily focus on lookup search by providing basic text summaries of top results, which also hinder users' ability to explore information through multimedia. Consequently, users make additional navigation efforts (clicking, scrolling, switching verticals), hindering information exploration. In this approach, we propose a framework that summarizes vertical search results into comprehensive documents. The framework is powered by a large language model (LLM) that extracts topics from search results and groups semantically similar multimedia results across verticals into unified topic-based summaries. This unified interaction reduces users' navigation effort and increases interest in exploration. We evaluated our approach using ROUGE across three domains (Movies, Music, and Sports) and conducted a system usability study with 31 participants, using the Bing search engine as a baseline. The proposed system achieved an average ROUGE F1 at: $R-1 = 0.67 \pm 0.15$, $R-2 = 0.26 \pm 0.17$, $R-L = 0.60 \pm 0.20$. The navigation efforts were significantly reduced in terms of clicks (21.5 vs. 30.8, $p < 0.01$), scrolls (74.6 vs. 218.4, $p < 0.001$), and vertical switches (0 vs. 3.7). The average system usability score was reported at 88%, significantly higher than baseline (77%, $p < 0.05$). These results confirm that our framework reduces exploratory navigation efforts while maintaining high user satisfaction.

Keywords: Aggregated Search, LLMs, Multimedia, Vertical Search, Exploratory Search



Introduction:

The proliferation of multimedia content online transformed how users browse, interact with, and process heterogeneous information on the web. Users rely on web content to satisfy their diverse information needs [1]. The content on the web consists of diverse, heterogeneous multimedia objects assembled from different sources [2]. The rapid development of Information and Communication Technologies has led to the exponential growth of heterogeneous information on the web [3]. Humans are explorers by nature [4] and seek information to satisfy their information needs [5]. Users use vertical web search engines to browse this heterogeneous information from different sources [6].

As a result, user interest in browsing multimedia content has increased, driven by the intention to explore information [7]. To ease information navigation, advancements in Artificial Intelligence (AI) and Machine Learning (ML) have significantly impacted search engines, leading to a paradigm shift [8]. Search engines are integrating AI techniques, such as GenAI, to remain competitive in the search engine market [9]. However, current GenAI integrations remain limited to a single document. Exploratory search requires browsing multiple, diverse yet conceptually similar multimedia documents. As a result, a significant increase in users' effort during exploratory search has been observed [10].

The growth of multimedia content has led to changes in users' information-seeking behavior. Users' information-seeking patterns have shifted toward exploratory rather than lookup search [11]. Exploratory search is non-linear, requiring information foraging, accumulation, and decision-making [12]. However, search engines present heterogeneous information in linear verticals [13]. Users struggle to navigate multiple links and manually aggregate information across different verticals, which leaves them feeling overwhelmed [13]. Users perform additional navigation (scrolling, switching among verticals), resulting in a loss of exploration context. This vertical presentation of search results affects users' overall search context [6].

Recent advances in GenAI have improved search engines [14]. Proprietary Generative Large Language Models (LLMs), such as Google's Gemini and Microsoft's Bing Copilot, are being effectively integrated into search engines [15]. This integration of GenAI provides a quick summary of the search results. However, the vertical presentation forces users to switch among multiple disconnected results. The linear presentation of search results affects users' contextual understanding and their interest in exploring the results [16]. Recent studies have proposed different techniques to aid exploration through query formulation [11] and grouping of search results [13]. However, these vertical presentations require significant navigation efforts to explore search results.

Objectives:

This study aims to reduce user navigation effort and context switching in exploratory vertical search. We proposed a framework that uses LLMs to extract semantic topics from aggregated results and to create comprehensive, topic-based multimedia documents. The framework was evaluated using ROUGE and users' usability metrics. This study reflects recent developments (2021–2026) in exploratory search and generative AI to ensure its relevance. We compiled and studied literature reflecting the rapid evolution of generative AI, exploratory search, users' naturalistic search behavior, and vertical search engine technologies.

Novelty and Research Contributions:

This study makes a novel contribution to the fields of Multimedia, Information Retrieval, Artificial Intelligence, and human-computer interaction by integrating LLMs into multimedia vertical search engines to enhance users' exploratory search behavior. The framework aggregates and summarizes multimedia search results into comprehensive, topic-based documents. Grouping semantically related results across verticals reduces excessive navigation effort and context switching, enabling more efficient exploratory search.

The rest of the paper is organized as follows. The Materials and Methods section details the design and implementation of the proposed framework. The Results and Discussion section provides a comprehensive analysis of the results and findings. The conclusion section concludes the paper and outlines future research directions.

Materials and Methods:

Given the issues with search engines, we have proposed a framework that uses an LLM to extract topics from the metadata (titles and descriptions) of vertical search results and generate summarized multimedia documents. The summarized documents, organized by topic, allow users to efficiently explore and satisfy their information needs. Our framework comprises five layers and four components. The component-based diagram of our proposed framework is presented in Figure 1.

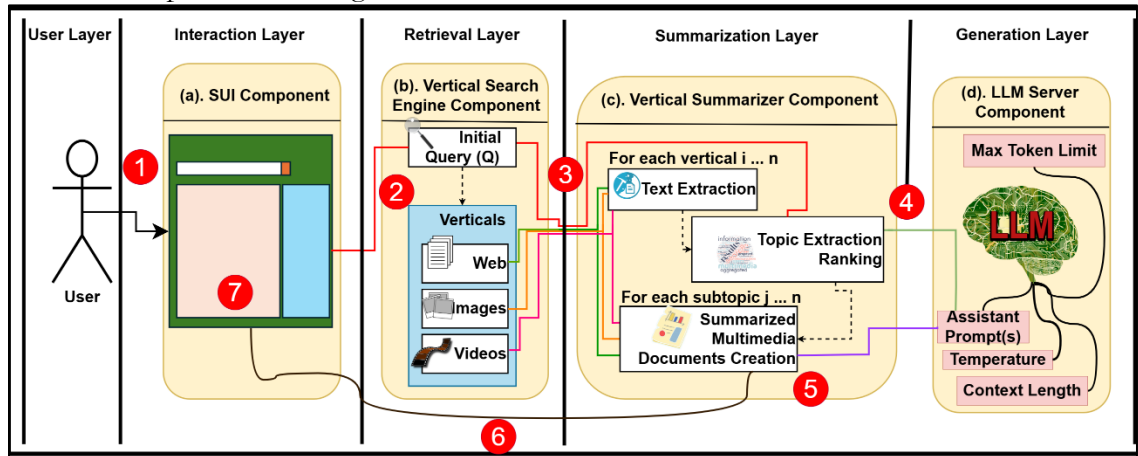


Figure 1. Component-based approach of our proposed four-layered framework. Workflow of the proposed LLM-based summarization framework. Steps: (1) User inputs query; (2) Fetch vertical results (web, images, videos) from search engine; (3) Aggregation of metadata from all verticals; (4) topic extraction and ranking using LLM; (5) Compose multimedia documents; (6) Display documents in SUI grid; (7) User explores details. (a). SUI Component (b). Vertical Search Engine Component (c). Vertical Summarization Component (d). LLM Server Component

Search Engine User Interface (SUI) Component:

The first component of our framework is the SUI component. The SUI component provides a user interface for entering queries and exploring summarized search results. User queries are forwarded to the summarizer component, which retrieves summarized multimedia documents for non-linear presentation. The non-linear presentation of comprehensive multimedia documents allows users to quickly scan summarized documents by topic. Users can click a topic-specific document to view detailed information in the details panel that addresses their information needs, as presented in Figure 2.

Vertical Search Engine Component:

The second component of our framework is the vertical search engine. This component acts as a black-box system that takes a user query as input and produces structured multimedia results across verticals (web, images, and videos) as output. The summarizer component later on utilizes the corresponding output. For a given user query (q), we convert both the query and each retrieved document into numerical vectors using Term Frequency-Inverse Document Frequency (TF-IDF) weighting. TF-IDF weighting reflects the importance of a word in a document within a collection. The weighting balances term frequency with the term's frequency across all documents.

$\vec{q} \in R^d$ represents the TF-IDF vector of the user query, where d is the dimensionality of the vocabulary (number of unique terms across all documents). $\vec{d}_i \in R^d$ is a TF-IDF vector

of the i -th retrieved document from a vertical $V \in \{Web, Image, Video\}$. n represents the total number of documents retrieved across all verticals, and $R_V = \{d_1, d_2, \dots, d_k\}$: The set of *top* – k ranked results for vertical V , where k is a fixed number (5 results from each vertical)

The equation for cosine similarity is:

$$\text{sim}(\vec{q}, \vec{d_i}) = \frac{\vec{q} \cdot \vec{d_i}}{\|\vec{q}\| \|\vec{d_i}\|}$$

where $\cdot$ denotes the dot product and $\|\cdot\|$ the Euclidean norm. The result ranges from 0 (orthogonal, no similarity) to 1 (identical direction). For each vertical V , we compute $\text{sim}(\vec{q}, \vec{d_i})$ For all retrieved documents, sort them in descending order, and select the top k :

$$R_V = \text{Top}_k(\text{sim}(q, d_i))$$

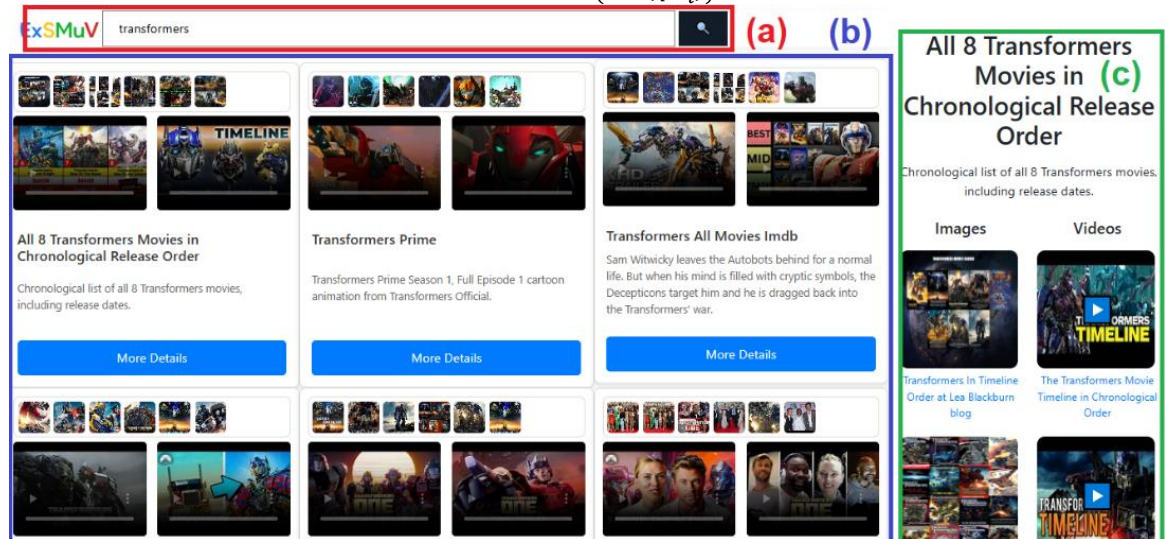


Figure 2. The SUI implementation of our proposed framework. (a). Query input panel (b). Summarized multimedia documents panel (c). Detailed Search Results Panel.

LLM Server Component:

The LLM server component takes system prompts as input and provides structured output to the summarizer component. The prompt and metadata for verticals are used to create topics and generate titles and descriptions for each summarized multimedia document. Let $M = \{(t_i, d_i)\}_{i=1}^n$ represent the combined set of metadata, where t_i is the title and d_i is the description of the i^{th} search result. The metadata is aggregated across all verticals (web, image, video). The input to the LLM for subtopic generation is constructed by concatenating all text, which is defined as:

$$x = t_1 \oplus d_1 \oplus t_2 \oplus d_2 \oplus \dots \oplus t_n \oplus d_n$$

Where $\oplus$ denotes string concatenation. The concatenated input x is a single text string that contains all titles and descriptions from the aggregated metadata. This allows the LLM to see the entire set of search results when extracting topics. We developed two functions that utilize LLM to generate topics and metadata for summarized documents. *ftopics* is a function that identifies and creates topics and *fmeta* then creates condensed information to represent multimedia document summaries. Both functions follow a controlled prompt design to ensure extractive behavior and avoid hallucination. The outputs are validated indirectly via ROUGE-based comparisons against expert-generated references.

Let $f_{topics}: X \rightarrow Tk$ be the subtopic mining function implemented using an LLM. It takes the input space X (all possible metadata strings) and produces a list of k high-level topics $T = \{\tau_1, \tau_2, \dots, \tau_k\}$. The number k is not fixed a priori. The LLM determines an appropriate number based on the metadata's diversity. Keyword-based phrase extraction avoids hallucination. The subtopic mining function can be defined as:

$$T = f_{topics}(x) = \{\tau_1, \tau_2, \dots, \tau_k\}, \text{ where } k = 5$$

For each summarized multimedia document (MM document), we first select a set of m ranked snippets:

$$S = \{(t_j, d_j)\}_{j=1}^m$$

The LLM is prompted to generate a concise title $\mathbf{f}$ and a representative description $\mathbf{d}$ that generalize over semantic content in $\mathbf{S}$. Formally:

$$f_{meta}: S \rightarrow X \times X$$

Summarizer Component:

The summarizer component is the core component responsible for creating comprehensive multimedia documents using the LLM Server Component.

Input Aggregation:

Let for each vertical v_i , let $R_i = \{(t_{ij}, d_{ij})\}_{j=1}^{n_i}$ be the list of search results with metadata (titles t and descriptions d). Let $\mathbf{M}$ be the combined metadata (titles and descriptions) from all verticals represented as:

$$M = \cup_{i=1}^m R_i$$

Topic Generation and Expansion:

The metadata $\mathbf{M}$ is passed to the LLM Server, producing a set of subtopics:

$$T = f_{topics}(M) = \{\tau_1, \tau_2, \dots, \tau_k\}$$

This set is concatenated with similar queries. $Q' = \{q'_1, q'_2, \dots, q'_l\}$ and query suggestions $S = \{s_1, s_2, \dots, s_r\}$ retrieved from the baseline search engine to produce an extended candidate topic set. The extended topic set can be represented as:

$$T = T \cup Q' \cup S$$

Topic Ranking:

Each candidate topic $\tau \in T$ is scored against the original user query q using cosine similarity in an embedding space. Let $E: \text{text} \rightarrow R^d$ be an embedding function that maps text to a d -dimensional dense vector. We used TF-IDF vectors for consistency. The similarity is:

$$\cos(E\vec{q}, E\vec{\tau}) = \frac{(E\vec{q} \cdot E\vec{\tau})}{(\|E\vec{q}\| \|E\vec{\tau}\|)}$$

This ranking ensures that the most query-relevant topics appear first in the user interface. The final ordered list is:

$$T^* = \text{Sort}_{\downarrow}(T, \cos(E(q), E(\tau)))$$

Each topic $\tau \in T$ is scored using cosine similarity $\cos(q, \tau)$, where q and τ are vector embeddings. Let $E: T \cup Q \rightarrow R^d$ be an embedding function that maps text to a d -dimensional vector. Where $\text{Sort}_{\downarrow}$ sorts topics in descending order of similarity.

Multimedia Document Construction:

For each top-ranked topic $\tau^* \in T^*$, we collect all search results across all verticals whose metadata (title and description) are semantically similar to the topic. The similarity threshold $\theta \in [0, 1]$ controls the strictness of relevance. We set $\theta = 0.65$ empirically. The threshold was determined in our pilot studies. The set of results $\text{topic } \tau^*$ is:

$$D_{\tau^*} = \cup_{i=1}^m \left\{ \cos(E(\tau^*), E(t_{ij} \oplus d_{ij})) \geq \theta \right\}$$

where θ is a relevance threshold such that $\theta \in [0,1]$.

Where R_i the set of search results from vertical i (web, image, or video), t_{ij} , d_{ij} is metadata (title and description) of the j -th result in vertical i , $\oplus$ represents string concatenation, E is the same embedding function used for topic ranking, and θ shows relevance threshold (a hyperparameter that we set to 0.65 in this study).

Metadata Summarization:

From each D_τ^* , a representative set of k results is selected. The top k results are selected based on relevance scores. The metadata is passed to the LLM Server to generate a summarized title $\hat{t}$ and description $\hat{d}$:

$$(\hat{t}, \hat{d}) = f_{meta}(D_\tau)$$

Final Output:

Each multimedia document is represented as:

$$M_\tau^* = (\hat{t}, \hat{d}, D_\tau^*)$$

The full output to the SERP consists of the set: $\{M_{\tau_1^*}, M_{\tau_2^*}, \dots, M_{\tau_p^*}\}$ where p is the number of selected topics. This enables exploration of comprehensive, summarized multimedia documents.

Framework Implementation:

We used various tools and techniques to implement each component of our framework. The Search User Interface of our proposed framework is presented in Figure 2. The SUI component provides a user interface built with HTML, CSS, and JS, a query box for user input, and panels for exploring summarized search results. The user's query from this component is sent to the vertical search engine component for further processing. We used PHP to communicate between backend components. We used UniServer [17] to run PHP code and serve pages to users for querying and exploring search results. The vertical search engine component is based on Bing. Bing and Google are popular vertical search engines that present multimedia search results in linearly ranked lists [18]. This component receives the user's query from the SUI component and sends it to the Bing search engine to retrieve multimedia vertical search results. We used the Bing Search API on the backend and Selenium WebDriver [19] to fetch metadata for vertical search results.

The Summarizer component fetches vertical search results and metadata from the vertical search engine component, communicates with the LLM server to extract topics from the metadata, ranks each topic-query pair using cosine similarity, and creates multimedia documents for each topic by ranking vertical search results against each topic. The summarizer component sends a prompt along with vertical search results metadata to the LLM server component, which retrieves results as a structured output. The composed, summarized multimedia documents are then sent to the SUI component for nonlinear presentation and exploration. The LLM Server Component processes prompts and returns structured responses to the Summarizer component. The prompts are used to extract keywords/topics from the metadata of vertical search results, or to generate a title/description for a single metadata item. We have used LM Studio [20] as an LLM Server Environment component to communicate with the summarizer component, process prompts, and return responses. We used lightweight LLM (LM2 [21]) in our study due to its strict prompt following, lower hallucination rate, and small inference time. The model parameters used in the approach are presented in Table 1.

Table 1. Summarization model parameters.

SR#	Model Attribute	Value
1.	Model Name	lfm2-1.2b [21]
2.	Quantization	Q8_0
3.	Size	1.25 GB
4.	Context Length (Set For this Research)	10,000
5.	Temperature (Set For this Research)	0.15
6.	Max Tokens (Set For this Research)	2000

Framework Evaluation:

We evaluated our framework using ROUGE scores and a usability evaluation. In the LLM summarization evaluation. We created summarized multimedia documents for the sub-topics of queries (Transformers, Lewis Capaldi, and Cristiano Ronaldo) across three domains (Movie, Music, and Sports). We shared the topic-based ranked multimedia search results (one of the top summarized multimedia documents from each domain) with the domain expert and asked them to generate a title and description for the summarized multimedia document based on the contents, as well as keywords/topics that represent the document. We compared the outputs from the experts and the LLM used in our proposed approach using the ROUGE measure to evaluate summarization quality.

In the usability evaluation, we conducted a study in which 31 participants of various age groups (average = 34.2 Years, Males = 21, Females = 10) with an average of 4.4 years of experience in using search engines were voluntarily recruited to explore search results for three domains (the exact domains given to experts for evaluation of LLM responses) on a baseline (Bing Search Engine) and a proposed interface. All participants provided written informed consent before participation and were informed of their right to withdraw at any time without penalty. No personally identifiable information was collected. Participants, including students and staff from different university departments, were recruited voluntarily. Participants aged over 20 years with more than one year of search engine experience were recruited.

The study protocol was approved by the Research Ethics and Support Committee (Department of Computer Sciences, Quaid-i-Azam University, Islamabad). Participants were asked to explore search results starting from initial queries of two different domains on both interfaces. Participants were asked to perform exploratory search tasks starting from initial queries. Participants were free to explore search results and refine their queries as they gained domain knowledge. During the task, we recorded their navigation efforts using monitoring tools (What Pulse [22] and Auto Hotkey scripts [23]) for comparison. After exploring the search results on each interface, we asked them to complete the SUS questionnaire to assess usability and effectiveness. Although search results are cached, the time overhead of local LLM inference was not considered in the evaluation.

Results:

The results of the ROUGE evaluation highlighted differences in summarization performance across domains, as presented in Table 2. ROUGE (Recall-Oriented Understudy for Gisting Evaluation) is a set of n-gram overlap metrics to measure the summarization quality: ROUGE-1 measures unigram overlap, ROUGE-2 measures bigram overlap, and ROUGE-L measures the longest common word sequence between the generated summary and references [24]. In our research, the average ROUGE-1, ROUGE-2, and ROUGE-L F1 scores across all domains are 0.53, 0.26, and 0.46, respectively.

Usability Evaluation:

The usability experiment results revealed that user participants made more navigation efforts, link clicks, query reformulations, and switches among verticals to explore search results on baseline search engines. Compared to our proposed approach, participants spent more

time gathering information on the baseline than with our proposed approach. The creation of comprehensive documents makes it convenient for users to explore heterogeneous information. The user satisfaction with the use of interfaces was calculated through System Usability Scale (SUS) scores. Figure 3 presents a comparison of SUS scores. The average score revealed that participants are more satisfied (88%) with exploring summarized search results than with the baseline vertical search engine (77%). The comparison of evaluation measures is shown in Table 3.

Table 2. Empirical Evaluation of composed, summarized multimedia documents.

Domain	Initial Query	ROUGE 1			ROUGE 2			ROUGE L		
		P	R	F1	P	R	F1	P	R	F1
Movies	Transformers	0.44	0.66	0.53	0.05	0.09	0.07	0.44	0.66	0.53
Music	Lewis Capaldi	0.38	0.51	0.43	0.21	0.46	0.29	0.38	0.51	0.43
Sports	Cristiano Ronaldo	0.45	0.91	0.65	0.3	0.63	0.41	0.33	0.66	0.44

Table 3. Effort comparison on Baseline vs Proposed Approach

Metric	Baseline (Mean $\pm$ SD)	Proposed (Mean $\pm$ SD)	t-statistic (df=30)	p-value
Clicks	30.8 $\pm$ 6.2	21.5 $\pm$ 4.7	4.82	< 0.001
Scrolls	218.4 $\pm$ 41.1	74.6 $\pm$ 10.3	11.34	< 0.001
Vertical switches	3.7 $\pm$ 1.4	0 $\pm$ 0	14.76	< 0.001
Avg. re-queries	4.5 $\pm$ 1.6	0 $\pm$ 0	15.68	< 0.001

Discussions:

The expert-based evaluation scores indicate that the proposed extractive approach consistently captures key lexical content. The scores also indicate that summarized multimedia documents are ranked with higher precision. The Sports domain performed best, with an F1 score of 0.65 and high recall (0.91), because the name ‘Cristiano Ronaldo’ appears consistently across sources. The Movies domain performed moderately, with a score of 0.53, because ‘Transformers’ (movie series) is an ambiguous term that affects precision. The Music domain shows a score of 0.42, as the artist’s name is specific, but descriptions vary widely across song lyrics, reviews, and biographies, lowering overlap. Recent literature indicates that in extractive summarization, ROUGE-1 F1 scores range from 0.40 to 0.50, ROUGE-2 from 0.20 to 0.30, and ROUGE-L from 0.40 to 0.50. The average ROUGE-2 is low (0.26) across all domains, as exact bigram matching is rare in exploratory search. The scores are typical for state-of-the-art systems on standard text datasets [25]. Thus, our observed average scores are within or above commonly reported acceptable ranges. In conclusion, the ROUGE scores confirm that the proposed approach performs reliably at the keyword level while exhibiting domain-dependent variations in phrase-level coherence.

Implications:

Our study extends the theory of exploratory search by demonstrating that LLM-based topic extraction can operationalize the concept of information foraging across heterogeneous verticals [12]. Reducing vertical switches during exploratory search provides empirical support for the claim that context switching disrupts exploration [6]. Moreover, the ROUGE scores obtained using a small (1.2B parameter) LLM challenge the assumption that large-scale models are necessary for effective summarization in constrained domains.

The components of the framework can be applied in domains such as academia [26], medical research [27], and e-commerce [28] to assemble topic-based information from diverse sources. The framework can be improved to support the exploration of summarized vertical search results based on user profiles and previous search patterns. In the future, we will conduct studies to reduce perceived search effort [29] and to improve learning during exploratory search [30].

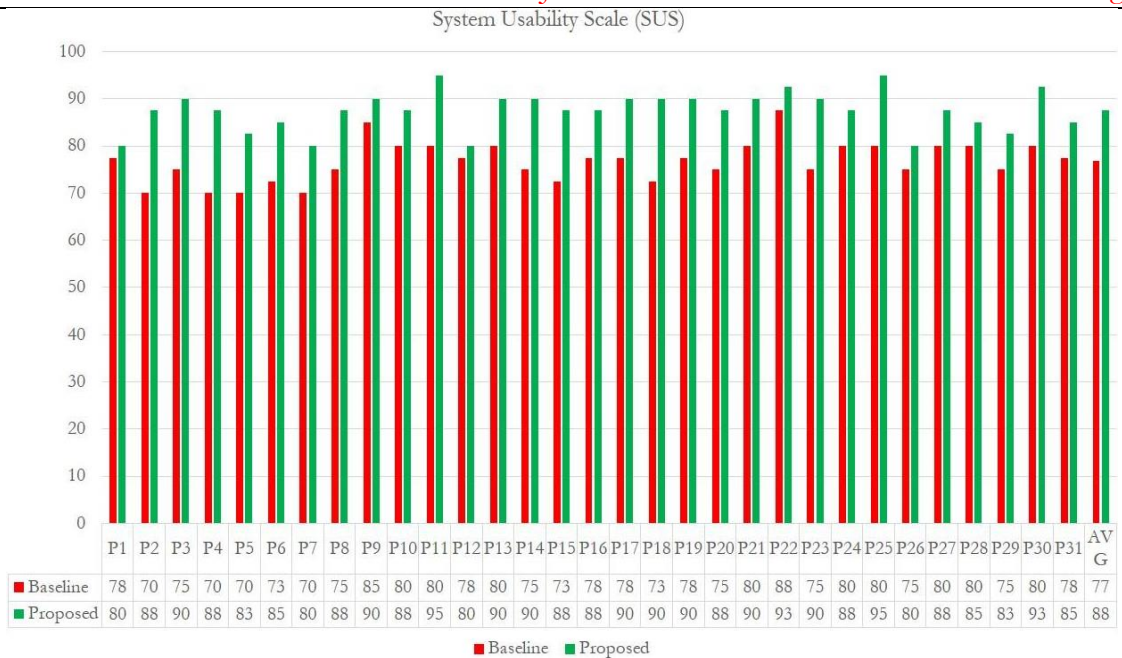


Figure 3. The SUS scores of the participants involved.

Conclusion:

We propose an LLM-powered framework to explore summarized vertical web search results more efficiently, thereby reducing users' navigation effort. The work extends exploratory search theory by operationalizing information foraging across heterogeneous verticals without context switching. The framework provides a blueprint for reducing navigation efforts and improving the exploratory search process. The results revealed that LLM-based aggregation of multimedia search results can generalize across domains and search platforms. The LLMs are general-purpose and rely primarily on metadata rather than domain-specific features. Despite limited resources, we employed a recent lightweight LLM. Moreover, the 31 participants are too few to demonstrate the significance of our proposed approach. The participants recommended improving the SUI to improve the exploration process further. In the future, we plan to conduct large-scale usability studies with diverse demographics and multiple baselines to generalize the findings of this study.

Acknowledgement: The study included human participants in its research. Ethical and experimental procedures were reviewed and approved by the Computer Sciences Department at Quaid-i-Azam University in Islamabad, Pakistan. Code available at [GitHub](#).

Author's Contribution: **Muhammad Wajeesh Uz Zaman:** Conceptualization, Data curation, Formal analysis, Investigation, Methodology, Software, Validation, Visualization, Writing - Original Draft, Writing - Review & Editing.

Umer Rashid: Conceptualization, Formal analysis, Methodology, Project administration, Supervision, Validation, Writing - Review & Editing, Resources.

Abdur Rehman Khan: Writing - Review & Editing, Conceptualization, Formal analysis, Supervision, Validation, Methodology.

Conflict of interest: The authors declare no conflict of interest.

References:

- [1] Muhammad Wajeesh Uz Zaman, Umer Rashid, "ExSMuV: [Ex]ploration software for [S]ummarized [Mu]ltimedia [V]ertical search results," *SoftwareX*, vol. 33, 2026, doi: 10.1016/j.softx.2025.102501.
- [2] U. Rashid and A. R. Khan, "Discovery work-space design to explore semantically aggregated graph-cluster spaces of multimedia document results," *Behav. Inf. Technol.*, Mar. 2025, doi: 10.1080/0144929X.2025.2532151.

- [3] Leila Shahrzadi, Ali Mansouri, "Causes, consequences, and strategies to deal with information overload: A scoping review," *Int. J. Inf. Manag. Data Insights*, vol. 4, no. 2, p. 100261, 2024, doi: <https://doi.org/10.1016/j.jjime.2024.100261>.
- [4] Holli Anne Passmore, Ashley N. Krause, "The beyond-human natural world: Providing meaning and making meaning," *Int. J. Environ. Res. Public Health*, vol. 20, no. 12, p. 6170, 2023, [Online]. Available: <https://pubmed.ncbi.nlm.nih.gov/37372757/>
- [5] Timothée Schmude, Laura Koesten, "Information that matters: Exploring information needs of people affected by algorithmic decisions," *Int. J. Hum. Comput. Stud.*, vol. 193, p. 103380, 2025, doi: <https://doi.org/10.1016/j.ijhcs.2024.103380>.
- [6] M. W. U. Zaman, A. R. Khan, and U. Rashid, "Towards Summarization of Aggregated Multimedia Verticals Web Search Results," *18th IEEE Int. Conf. Emerg. Technol. ICET 2023*, pp. 263–268, 2023, doi: 10.1109/ICET59753.2023.10374811.
- [7] S. Naseer, U. Rashid, M. Saddal, and A. R. Khan, "SERB: an interactive exploration mechanism for blind web users to explore web search content," *Univers. Access Inf. Soc.* 2025 243, vol. 24, no. 3, pp. 2335–2350, Feb. 2025, doi: 10.1007/s10209-025-01199-2.
- [8] E. Amer and T. Elboghhdady, "The End of the Search Engine Era and the Rise of Generative AI: A Paradigm Shift in Information Retrieval," *4th Int. Mobile, Intelligent, Ubiquitous Comput. Conf. MIUCC 2024*, pp. 374–379, 2024, doi: 10.1109/MIUCC62295.2024.10783559.
- [9] A. R. Ashraf, T. K. Mackey, and A. Fittler, "Search Engines and Generative Artificial Intelligence Integration: Public Health Risks and Recommendations to Safeguard Consumers Online," *JMIR Public Heal. Surveill.*, vol. 10, no. 1, p. e53086, Mar. 2024, doi: 10.2196/53086.
- [10] Pubali Mukherjee, Varsha Jain, "How lack of choice to opt-out of generative artificial intelligence in traditional search engines drives consumer switching intentions: The mechanism of empowerment," *J. Retail. Consum. Serv.*, vol. 88, p. 104506, 2026, doi: <https://doi.org/10.1016/j.jretconser.2025.104506>.
- [11] Tajmir Khan, Umer Rashid, "End-to-end vertical web search pseudo relevance feedback queries recommendation software," *SoftwareX*, vol. 27, p. 101872, 2024, doi: <https://doi.org/10.1016/j.softx.2024.101872>.
- [12] Y. Liu, C. Qin, X. Ma, J. Chen, H. He, and J. Mao, "Characterising exploratory search tasks: Evidence from different fields," *J. Inf. Sci.*, 2025, doi: 10.1177/01655515251330611.
- [13] Milad Momeni, Orland Hoeber, "A study of search result aggregation approaches for the digital humanities," *J. Assoc. Inf. Sci. Technol.*, 2025, doi: <https://doi.org/10.1002/asi.70006>.
- [14] Soyoung Kim, Randi Priluck, "Consumer Responses to Generative AI Chatbots Versus Search Engines for Product Evaluation," *J. Theor. Appl. Electron. Commer. Res.*, vol. 20, no. 2, p. 93, 2025, doi: <https://doi.org/10.3390/jtaer20020093>.
- [15] Lin Liu, Jiajun Meng, "LLM technologies and information search," *J. Econ. Technol.*, vol. 2, no. 4, pp. 269–277, 2024, doi: 10.1016/j.ject.2024.08.007.
- [16] Nakyoung Kim, Yong Gu Ji, "Exploratory Search with Generative AI: An Empirical Study on the Impact of Interaction Design Strategies on Information Exploration and Cognitive Load," *Int. J. Hum. Comput. Stud.*, vol. 210, p. 103771, 2026, doi: <https://doi.org/10.1016/j.ijhcs.2026.103771>.
- [17] A. Akintoye, "Installing Laravel, Development, and Test Tools Installation," *API Dev. with Laravel*, pp. 21–64, 2025, doi: 10.1007/979-8-8688-1576-8_2.
- [18] Göker Cebeci, Banu Diri, "Disentangling Technical and Content Attributes in Search Engine Ranking: A Comparative Study of Google and Bing," *IEEE Access*, 2026, doi:

- 10.1109/ACCESS.2026.3657977.
- [19] Boni García, Filippo Ricca, "Test automation with selenium: A survey," *Inf. Softw. Technol.*, vol. 194, p. 108077, 2026, doi: <https://doi.org/10.1016/j.infsof.2026.108077>.
 - [20] Shariq Murtuza, "Forensic Implications of Localized AI: Artifact Analysis of Ollama, LM Studio, and llama," *arXiv:2603.23996*, 2026, [Online]. Available: <https://arxiv.org/abs/2603.23996>
 - [21] Alexander Amini, Anna Banaszak, Harold Benoit, "LFM2 Technical Report," *arXiv:2511.23404*, 2025, [Online]. Available: <https://arxiv.org/abs/2511.23404>
 - [22] "WhatPulse 5.0: Mouse Scrolls, Distance, oh my! | WhatPulse Blog." Accessed: Apr. 21, 2026. [Online]. Available: <https://whatpulse.org/blog/whatpulse-5-0-mouse-scrolls-distance-oh-my-ded09dc3e64d>
 - [23] "AutoHotkey." Accessed: Mar. 11, 2026. [Online]. Available: <https://www.autohotkey.com/>
 - [24] Mousumi Akter, Naman Bansal, "Assessing the effectiveness of ROUGE as unbiased metric in Extractive vs. Abstractive summarization techniques," *J. Comput. Sci.*, vol. 87, p. 102571, 2025, doi: <https://doi.org/10.1016/j.jocs.2025.102571>.
 - [25] Haopeng Zhang, Philip S. Yu, "A Systematic Survey of Text Summarization: From Statistical Methods to Large Language Models Haopeng Zhang, Philip S. Yu, Jiawei Zhang," *ACM Comput. Surv.*, 2024, [Online]. Available: <https://arxiv.org/abs/2406.11289>
 - [26] "Can Generative AI Replace Search Engines for Learning? Understanding Student Preferences, Use and Perceived Proficiency in Using AI | TechTrends | Springer Nature Link." Accessed: Apr. 21, 2026. [Online]. Available: <https://link.springer.com/article/10.1007/s11528-025-01095-9>
 - [27] A. M. Garcia-Carmona, M. L. Prieto, E. Puertas, and J. J. Beunza, "Leveraging Large Language Models for Accurate Retrieval of Patient Information From Medical Reports: Systematic Evaluation Study," *JMIR AI*, vol. 4, 2025, doi: 10.2196/68776.
 - [28] Panagiota Papastamoulou, Nikos Antonopoulos, "Artificial Intelligence in E-Commerce: A Comparative Analysis of Best Practices Across Leading Platforms," *Systems*, vol. 13, no. 9, p. 746, 2025, doi: <https://doi.org/10.3390/systems13090746>.
 - [29] Joon Soo Lim, Nalae Hong, "Perceived search overload, generative AI credibility, and comparative usefulness: a channel complementarity approach to health information seeking," *Heal. New Media Res.*, vol. 9, no. 1, pp. 124–145, 2025, doi: <https://doi.org/10.22720/hnmr.2025.00115>.
 - [30] Gusman Lesmana, Indra Prasetya, "Depth Counseling for Students Experiencing Cognitive Crisis in Learning," *J. Educ. Stud. Multidiscip. Approaches*, vol. 6, no. 1, 2026, doi: 10.51383/jesma.2026.142.



Copyright © by authors and 50Sea. This work is licensed under the Creative Commons Attribution 4.0 International License.