



Asian Research Association

INTERNATIONAL RESEARCH JOURNAL OF MULTIDISCIPLINARY TECHNOVATION



Dual Cloud - Sturnus Optimized Intent-BERT with Randomized Secure Anonymization based Bibliographic Network Recommender for Citation Recommendation System

V.A. Nurjahan ^{a,*}, S. Jancy ^a

^a Department of Computer Science Engineering, Sathyabama Institute of Science and Technology (Deemed to be University), Jeppiaar Nagar, Rajiv Gandhi Salai, Chennai-600119, Tamil Nadu, India.

* Corresponding Author Email: noorji84@gmail.com

DOI: <https://doi.org/10.54392/irjmt2654>

Received: 03-01-2026; Revised: 06-07-2026; Accepted: 20-07-2026; Published: 02-09-2026



Abstract: Citation recommendation (CR) systems for manuscripts employ automated systems to identify and suggest relevant scientific publications for effectively citing specific text passages; however, they face challenges related to the promotion of unreliable sources and the privacy of sensitive data. To address these limitations, a novel Dual Cloud Sturnus Optimized Intent Learning BERT with Randomized-Adversarial Secure Anonymization (SIL-BERT) based Bibliographic Network Recommender (BNR) is proposed in this study to enable secure citation recommendation. Among these issues, shilling attacks are particularly important because malicious actors exploit the mathematical foundations of collaborative filtering and manipulate the algorithms. To address this issue, a novel Sturnus Optimized Self-Intent BERT (SS-IBERT) is employed and it effectively reduces artificial inflation caused by algorithmic rhetoric and thereby mitigates the suppression of emerging researchers. Moreover, domain-specific interactions combined with re-identification through linkable attributes generate a distinctive chronological activity fingerprint, which allows anonymized citation logs to be cross-referenced with public metadata. Therefore, to deal with this, a Randomized-Style Adversarial Anonymization (R-SAA) is used, and it effectively suppresses the Small-N Behavioral Vulnerability (S-NBV) that creates elevated re-identification risks. The experimental results demonstrate that the proposed Dual Cloud SIL-BERT achieves a citation-recommendation accuracy of approximately 99%.

Keywords: Citation Recommendation, Natural Language Processing, Bi-directional Encoder Representations from Transformers, Deep Learning and Cloud-based security.

1. Introduction

In contemporary academic discourse, Recommendation Systems (RS) play a crucial role in offering personalized services tailored to user interests, within which Citation Recommendation (CR) functions as a pivotal component. CR serves as an effective tool for substantiating scholarly contributions with pertinent and impactful references. Recommendation systems in scientific literature greatly assist researchers in identifying recently published papers in their domains and discovering existing gaps and trends. CR is the process of providing appropriate citations for a specific manuscript, and in order to achieve this, CR uses different features of the cited paper. These features encompass text content, metadata, and contextual information surrounding the citation, while the metadata include the title, authors, abstract, publication venue, and publication date. Recent studies emphasize that modern citation recommendation systems must integrate both semantic understanding and interaction-

based signals to effectively capture complex citation behaviors and improve recommendation accuracy. In addition, some of the common techniques used to recommend research papers are Collaborative Filtering (CF), Citation Analysis (CA) and Content-Based Filtering (CBF) [1-5]. CBF techniques are commonly used for CR, but they are highly prone to cold-start issues which led to the introduction of deep-learning-based CR models, which use rhetorical-zone embeddings for classifying the rhetorical zones existing in the research article [6]. Furthermore, recent advancements in recommender systems highlight a transition from traditional approaches to transformer-based and data-driven models, enabling improved scalability and recommendation performance across complex domains [7].

Moreover, to capture intricate attributes across diverse entities, a random walk-based Knowledge Graph (KG) method is employed, and it effectively fine-tunes the language model with the generated text, but faces

limitations in data textualization and may fail to preserve necessary information [8].

Recent CR systems have achieved great success in semantic understanding and recommendation accuracy. However, existing approaches are still vulnerable to many security threats, including shilling attacks, data poisoning, adversarial manipulation, sybil attacks, citation spam, membership inference, and model inversion attacks [9]. Such vulnerabilities distort recommendation rankings, reveal sensitive information about researchers, and undermine the fairness and integrity of scholarly recommendation systems [10]. Moreover, the existing privacy-preserving schemes focus primarily on data confidentiality rather than attack resilience, scalable cloud deployment, and robust anonymous recommendation [11]. Despite its effectiveness, CR is also vulnerable to privacy leakage, which leads to data theft and attacks, and, to address this issue, intelligent recommendation systems have been incorporated in the CR. To ensure data confidentiality and user privacy, Federated Learning (FL) with differential privacy is introduced in [12], which uses a multi-party computing protocol to guarantee user privacy. Further, to ensure precise comprehension and context-aware responses, Natural Language Processing models like BERT (Bi-directional Encoder Representations from Transformers) and GPT (Generative Pre-trained Transformers) are employed [13]. Further, in [14], a recommender system known as Multi-embedding Fusion Network for Recommendation (MFNR) has been introduced, using a multi-embedding approach to capture and represent the item and user features in the review texts. However, the performance degrades as only certain specific metrics are considered. Moreover, a fine-tuned BERT model has been presented in [15] to extract API comments, secure code preprocessing, and anonymize tokenization. This ensures code confidentiality, but faces challenges with the metaprogramming constructs.

Further, to enhance the AI (Artificial Intelligence) analytics, without compromising data confidentiality, a Privacy Aware Semantic Intelligence (PASI) has been introduced in [16]. This effectively aids in performing complex semantic relationship extraction from the encrypted data, while supporting baseline queries, but does not support real-time streaming of the encrypted data. Then, to provide a task-specific trust framework in cloud and big-data applications, a Zero-Trust security model is presented in [17], and focuses on the sources of breaches, while facing issues during the processing of data. Meanwhile, in [18], to address data scarcity using a cloud-based platform, the General Data Protection Regulation (GDPR) is defined, and it ensures privacy compliance, but does not adequately address the needs of practitioners and scholars across different sectors.

Therefore, the preceding discussion indicates that the use of the KG method encounters issues with

textualization of data and also may fail to preserve necessary information, while the use of BERT and GPT involves performance degradation as only certain specific metrics are considered. Meanwhile, the use of MFNR considers only specific metrics for the analysis, while fine-tuned BERT encounters issues with metaprogramming constructs. Likewise, the PASI supports baseline queries, but does not support real-time streaming of the encrypted data, while the Zero-Trust security model and the GDPR face issues with data processing and also lag in concentration towards practitioners and scholars across different sectors.

1.1 Main Contributions of the Research

To address these issues and improve CR security, a novel technique has been proposed in this research work, and the main contributions of this study are as follows.

- To address the challenges of handling large-scale data and to attain better privacy, bias control, and scalability in CR, a Dual Cloud BNR is employed.
- To address shilling attacks, a novel Sturnus Optimized Self-Intent BERT is employed, and it effectively reduces artificial inflation caused by algorithmic rhetoric, which further lessens the suppression of emerging researchers.
- To address the challenge posed by distinctive chronological activity fingerprints, a Randomized-Style Adversarial Anonymization is used, and it effectively suppresses the S-NBV that creates elevated re-identification risks.

1.2 Organization of the work

Following these contributions, a review of conventional citation recommendation systems and the methodologies used for attaining effective performance and security, along with the motivation of the work, are discussed in Section 2. Meanwhile, the proposed research work and the various techniques used for attaining secure citation recommendations are discussed in Section 3. Moreover, the results obtained in this study and the comparative analysis with prior techniques are discussed in Section 4. Further, the study is concluded in Section 5, along with the directions for future research.

2. Background and Related Works

In current research, the use of CR helps identify semantically pertinent references that accelerate the process of literature retrieval and also help in improving the quality of the research work. To support this objective, a detailed literature survey was conducted in this section, along with the strengths and limitations. The

discussion is organized into methodological categories and also emphasizes a comparative and analytical understanding of existing approaches

2.1 Deep Learning and Semantic-Based Methods

Ma *et al.*, [19] introduced a multi-task learning model for CR and for argumentative zoning classification, which includes annotation of citing sentences, followed by citation of papers in PubMed dataset. Initially, based on the argumentative zoning types, a corpus of labelled queries is developed, followed by multi-task learning. Then, to generate effective predictions on the test data, an agreement estimation is carried out, and the results showed that the CR performs better when considering argumentative information of the sentences cited. This demonstrates that incorporating contextual features improves recommendation accuracy, particularly in ranking-based evaluation metrics. However, a specific quantity of well-labelled training corpus is necessary, to achieve better citation recommendation.

Nacar & Koubaa research work [20], a Matryoshka embedding-learning framework for training nested embedding models, involving the conversion of different sentence similarity datasets, enables effective comparison of these models over other various dimensions. The defined model demonstrated strong performance in identifying the semantic nuances that are unique in the language, and also outperforms conventional models by about 20-25%. With this, the defined work highlights the need for the adaptation of NLP models in specific scenarios for attaining optimal performance, and this highlights the importance of hierarchical semantic representations in improving similarity modeling. However, the defined technique also indicates the need for improved models for better performance.

2.2 Content-Based and Classical Methods

In the recent research in [21], a study was published in Array that proposed SRS, a scientific publication recommendation system that assists in efficient literature discovery. The proposed framework is a two-step hybrid model, in which candidate papers are first retrieved with content-based filtering based on TF-IDF vectorization and cosine similarity on textual metadata, and then a lightweight bibliometric re-ranking is then performed using the year of publication and the number of citations. The system also promotes transparency through the use of visual analytics and rank-conscious evaluation measures like NDCG, MAP, and MRR. This finding indicates that traditional feature-based methods can achieve reasonable effectiveness in early-stage retrieval tasks. Experimental results demonstrate its effectiveness in retrieving relevant

literature and reducing manual effort during early-stage exploration. Nevertheless, the model is mainly based on conventional feature engineering methods and lacks deep semantic modeling, graph-based relationships, and optimization-based learning, which restricts its capacity to model intricate citation dynamics in large-scale scholarly settings.

Velkumar & Chokkalingam [22] discussed the use of SVM for CR, which enabled effective identification of the research papers. The defined work employs a stochastic matching pattern scheme for excluding the function words, to develop a structured library. Then, with the utilization of a probabilistic text normalization algorithm, the canonicalization of documents is achieved. Then, with the use of Support Vector Machine (SVM), classification is carried out, and it effectively assists in isolating the intricate patterns. Finally, it ranks the citations based on the recommendations and also computes its effectiveness by evaluating the performance metrics. Nevertheless, the exploration process should be improved using more advanced techniques. However, such classical models generally exhibit limitations in capturing complex contextual relationships.

2.3 Graph-Based and Robust Recommendation Models

Gao *et al.*, [23] addresses the major problems of shilling attacks in citation recommendation systems. These attacks involve the manipulation of recommendation results by artificially inflated citation behaviors, which substantially affects the credibility of scholarly recommendations. To address this issue, the authors introduce RSA-CR, a powerful hybrid citation recommendation system that incorporates historical collaboration patterns, important citation relationships, and content-based constraints. The model builds a two-layer academic graph and uses a novel dumbbell inductive learning mechanism to produce trustworthy author and paper embeddings. The experimental findings on various academic datasets demonstrate better performance and robustness than the baseline models. This demonstrates improved robustness and reliability compared to conventional models. Nevertheless, the framework emphasizes robust graph-based representation learning, and lacks semantic intent modeling, optimization strategies, and privacy-preserving mechanisms, which are also significant factors to consider in next-generation citation recommendation systems.

A trusted graph-based collaborative filtering recommendation system is introduced in [24], integrating graph relationships with collaborative filtering to generate reliable recommendations. To enhance both the robustness and reliability of the recommendation process, the reliability of each model is assessed using a trust metric; however, static weights limit adaptability

in dynamic environments. Moreover, the collaborative filtering component uses a service-user matrix to assess similarity and predict the ratings. Meanwhile, the graph-based components develop a trust network that facilitates the trust propagation. Further, the recommendations generated are used for deriving indirect trust, which is then used for leveraging the weighted trust computation. The defined work provides dynamic support to the recommendation system, but it relies heavily on pre-defined rules and static weights.

2.4 Hybrid and Transformer-Based Model

In [25] BeLightRec (BERT-based Light-weight RS) is based on LGCN (Light Graph Convolutional Network) to propagate the collaborative signals and to identify the feature vectors of users and the citations, while it was enhanced with semantic similarity between the title and description of the works, which typically leads to better performance in ranking metrics such as precision and NDCG. The LGCN involves a former matrix that records the interaction amid the titles and the works, while the matrix B records the semantic similarity between each. To reduce the effects of differing measurement scales. The LGCN effectively propagates the signals and also iterates for a specified number of steps before detecting the vector characteristics. Besides, careful modifications are necessary to adjust the weights of similarity measures.

A recommendation model that combines the BERT and Term Frequency-Inverse Document Frequency (TF-IDF) is presented in [26] to extract robust textual features, and Random Forest (RF) for enhancing the precision of the recommendation. A curated dataset is taken into account and is subjected to pre-processing. The defined work effectively calculates the semantic as well as the lexical similarity with the aid of cosine similarity and fuzzy matching. This further enhances the performance of the defined work in cold-start scenarios. The analysis showed that the proposed approach improves system performance, supports user-friendly recommendations and real-time deployment, and semantic embeddings significantly enhance early-stage recommendation accuracy but are limited in improving the adaptability of the recommendations.

A recent work in [27] presents a hybrid citation recommendation model that simultaneously learns sequential collaborative cues and textual semantics. The model captures interaction patterns based on publication indices and at the same time derive semantic representations based on textual content, thereby addressing the natural discrepancy between structural identifiers and textual characteristics. To bridge this gap, a hybrid alignment mechanism is proposed to learn discriminative representations in the presence of dynamic preference changes. Large-scale experiments on several benchmark datasets indicate that SCTRec substantially achieves higher recommendation accuracy

than current methods. Nevertheless, the framework primarily focuses on representation alignment and lacks the resilience to adversarial manipulation, optimization-based learning approaches, and privacy-preserving solutions, which are essential to secure and scalable citation recommendation systems.

2.5 Survey and Analytical Studies

A recent study in [28] presents a comprehensive, graph-centric review of citation recommendation systems, highlighting the growing importance of bibliographic network modelling for scholarly recommendation tasks. The study categorizes existing approaches into unipartite, bipartite, and k-partite graph formulations, capturing complex relationships among papers, authors, venues, and research concepts. This finding indicates that current CR systems are still evolving toward more comprehensive evaluation and modeling strategies. Furthermore, it critically identifies numerous limitations in current models, including the insufficient utilization of heterogeneous relationships, a lack of temporal adaptability, and limited consideration of user-centric factors such as novelty and diversity. In addition, challenges related to cross-domain generalization and inconsistent evaluation practices are highlighted. These observations underline the necessity for advanced citation recommendation frameworks that integrate rich semantic modeling, adaptive learning mechanisms, and scalable architectures to improve both accuracy and practical applicability.

2.6 Summary of Literature Review

The detailed review shows that the use of a multi-task learning model in [19] necessitates a specific quantity of well-labelled training corpus. Following this, the Matryoshka embedding learning in [20] highlights the need for enhanced models to achieve better semantic representation performance. The work in [21] further restricts its capacity to model intricate citation dynamics in large-scale scholarly settings. In a similar direction, the utilization of SVM for CR in [22] faces challenges related to cross-domain generalization and inconsistent evaluation practices. The RSA-CR model in [23] improves robustness against shilling attacks but still lacks semantic intent modeling, optimization strategies, and privacy-preserving mechanisms.

With regard to graph and hybrid approaches, the BeLightRec model in [24] relies heavily on pre-defined rules and static weights, limiting adaptive learning capability. Likewise, the combined BERT and TF-IDF-based model in [25] requires careful modifications to improve adaptability in recommendation performance. The model in [26] also shows limitations in adapting to evolving user preferences despite leveraging sequential and semantic signals. Furthermore, the SCTRec model

in [27] focuses mainly on representation alignment but lacks resilience to adversarial manipulation, optimization-based learning, and privacy-preserving mechanisms. In addition, the Trusted Graph-based collaborative filtering system in [28] relies on static weighting strategies, which reduces flexibility in dynamic environments. Hence, a novel, intelligent and secure framework is essential for effective citation recommendation.

Table 1 reveals several critical limitations in existing citation recommendation approaches. Most prior methods, including BERT-based, graph-based, and hybrid models, primarily focus on either semantic similarity, structural relationships, or shallow feature fusion, but fail to simultaneously capture fine-grained intent understanding, robust representation learning, and security-aware recommendation. Additionally, many approaches suffer from limited adaptability due to static weighting mechanisms, lack of optimization-driven

learning, and absence of adversarial robustness, thereby reducing their effectiveness in large-scale and dynamic scholarly environments. To address these challenges, the proposed SIL-BERT framework is designed as a unified solution that integrates contextual BERT-based embeddings with Intent Contrastive Learning for improved discriminative capability. Furthermore, Self-Attenuation and embedding alignment optimization enhance representation stability, while randomized masking and style adversarial anonymization improve robustness against noisy and manipulated inputs. In addition, the dual cloud secure recommendation layer introduces privacy-preserving and secure processing, which is largely absent in existing studies. Therefore, SIL-BERT effectively addresses the identified research gaps by jointly improving semantic understanding, optimization efficiency, robustness, and security in citation recommendation tasks.

Table 1. Comparative analysis of existing methods and identified research gaps

Ref	Method / Technique	Key Strength	Research Gap
[19]	Multi-task learning with argumentative zoning	Utilizes citation context effectively	Dependence on large labelled datasets limits scalability
[20]	Nested embedding learning	Captures semantic relationships	Limited generalization across different contexts
[21]	Hybrid TF-IDF + bibliometric ranking	Effective and transparent retrieval	Lacks deep learning, graph modeling, and optimization
[22]	SVM-based approach	Strong classification capability	Inability to model complex and dynamic relationships
[23]	Robust graph-based citation recommendation (RSA-CR)	Resists shilling attacks and improves reliability	No transformer-based semantics, optimization, or privacy integration
[24]	Graph-based collaborative filtering	Utilizes graph-based relationships	Lack of adaptability due to static weighting mechanisms
[25]	BeLightRec (BERT + LGCN)	Combines semantic and collaborative features	Difficulty in integrating semantic and collaborative features effectively
[26]	BERT + RF hybrid model	Hybrid feature extraction approach	Limited adaptability to evolving user preferences
[27]	Sequential collaborative + text semantic alignment (SCTRec)	Reduces ID–text semantic gap using hybrid alignment	Does not model adversarial robustness or incorporate optimization-driven parameter tuning; lacks privacy-aware or anonymization strategies for secure recommendation
[28]	Graph-based citation recommendation using multi-relational networks	Provides unified modeling of bibliographic relationships	Lacks semantic modeling, optimization, and security mechanisms

3. Motivation of the research work

In CR systems, the main challenges stem from privacy concerns, trade-offs between security measures and the recommendation accuracy, and malicious data manipulation. Among these issues, shilling attacks play a key role, in which malicious attackers exploit the mathematical foundations of collaborative filtering and manipulate the algorithms. Further, a falsified shilling profile (forged researcher data) is injected, and it contains synthetic researcher data generated through the fabrication of publication records, citations, and the authorship patterns. With this, the attackers cause the system to infer a higher degree of similarity between genuine researchers and fake profiles. This results in Algorithmic Rhetoric, which, in this study, refers to the intentional manipulation of citation recommendation outcomes through biased contextual signals or artificially generated interactions, leading to artificial inflation of citation counts, h-index, and impact metrics of a specific author or a paper, thereby decreasing diversity and suppressing the unconventional research.

Moreover, Embedding Anonymization (EA) serves as the primary stage for securing the CR system, but it is vulnerable to security challenges. The EA transforms the text into vectorized embeddings, so as to safeguard the personal identities and sustain the semantic meaning effectively, but remains vulnerable to re-identification via Linkable attributes. Because human behaviour is often more distinctive than technical identifiers (i.e.) even domain-specific interactions like visiting the niche repository, viewing the abstract, and saving the paper) generate a distinctive chronological activity fingerprint. Moreover, individuals can be identified with 80% accuracy with a short sequence of 10 interactions and this uniqueness permits the anonymized citation logs to get cross-referenced with the public metadata, while exposing private research interests and unpublished project trajectories. This creates Small-N Behavioural Vulnerability (S-NBV) that creates elevated re-identification risks, and even anonymized citation embeddings may be traced back to individual researchers.

4. Proposed Research Work

In general, the CR system relies on embedding-based similarity measures but overlooks the effects of behavioural and stylistic biases, which result in privacy risks [29]. Therefore, to address these limitations, a novel Dual Cloud Bibliographic Network Recommender (BNR) known as “Sturnus Optimized Intent Learning BERT with Randomized-Adversarial Secure Anonymization (SIL-BERT)” framework is proposed in this research for effective Citation Recommendation.

Initially, to address the issues with Shilling attacks that exploit the mathematical foundations of collaborative filtering [30] and manipulate the algorithms,

a novel Sturnus Optimized Self-Intent BERT (SS-IBERT) is employed. The SS-IBERT incorporates a BERT-based transformer Encoder, Self-Attenuated Intent Contrast Learning and Sturnus Vulgaris Optimization Algorithm. First, the pre-processed manuscript is fed to the BERT-based Transformer Encoder, which produces the contextual embeddings that attain the semantic meaning from the text. Then, for reducing the issues with Algorithmic Rhetoric and to improve the semantic fidelity, a Sturnus Optimized Self-Attenuated Intent-Contrastive Learning (SOSA-ICL) is introduced. The Intent-Contrastive Learning (ICL) effectively trains the embeddings, so that the segments that possess similar research intent are brought closer, while the segments with distinct intents are pushed apart. This thereby emphasises the semantic intent with respect to the superficial stylistic cues. Further, using Self-Attenuation mechanism (SA), the signals that are not relevant and noisy signals are down weighted, while safeguarding the framework from overfitting to sentence structures, author-specific writing style and the rhetorical patterns. Then, for optimizing the embedding space, a Sturnus Vulgaris Escape Optimization Algorithm (SVE-OA) is employed, which significantly tunes the embedding alignment, contrastive learning parameters and the attention weights. This effectively ensures that the embeddings represent the semantic intent significantly, while also reducing stylistic bias. Hence, the combination of these techniques with BERT produces embeddings that are semantically effective and less sensitive to the rhetorical biases.

Then, to safeguard researchers' privacy and to address the issues with re-identification via linkable attributes, a novel Randomized-Style Adversarial Anonymization (R-SAA) is employed, which incorporates Randomized Masking Mechanism and Style Adversarial Anonymization. The Randomized Masking Mechanism (RMM) probabilistically suppresses embedding dimension subsets that are highly correlated with stylistic uniqueness, namely topic rarity, intersection sequencing and writing cadence. Moreover, as RMM functions in a randomized manner, the defined technique prevents adversaries from exploiting particular dimensions for reconstruction, which further impedes the reconstruction of behavioural fingerprints from small interaction sequences. Further to improve the privacy protection, a Style-Adversarial Anonymization (SAA) is employed and it excludes the behaviour based or latent author features using the style and temporal network, while safeguarding the core semantic intent. In addition, during training, the SAA uses an embedding generator for preserving the semantic intent, also to increase the error in the discriminator. This, thereby forces the embedding space to reduce the dimensions, which thereby prevents the encoding of stylistic and temporal cues that function as linkable attributes over the datasets. Hence, the integration of these techniques reduces the identifiability of sparsely represented

level embeddings are anonymized with the R-SAA mechanism and get securely encrypted, which addresses the issues with re-identification risks and S-NBV. Furthermore, these preserved embeddings are fed to an untrusted environment (Cloud B), in order to compute the similarity and to retrieve the candidate citation. Then, based on this, the final CR is provided to the user, while ensuring the confidentiality and semantic accuracy of the manuscript till the end of the process. Thus, Cloud A handles pre-processing, embedding refinement, and anonymization, while Cloud B performs only similarity matching on protected vectors.

The steps involved in the proposed Dual Cloud BNR with SIL-BERT are depicted in Figure 2, in which the process starts with uploading the manuscript and is then followed by text pre-processing and embedding generation. Then, embedding anonymization and transfer are performed and then computation of the similarity index along with CR is carried out. Finally, the merging of the results and visualization of the outcomes are performed. Each step produces a structured output (tokens, embeddings, anonymized vectors, similarity scores) that directly feeds into the next, ensuring traceability.

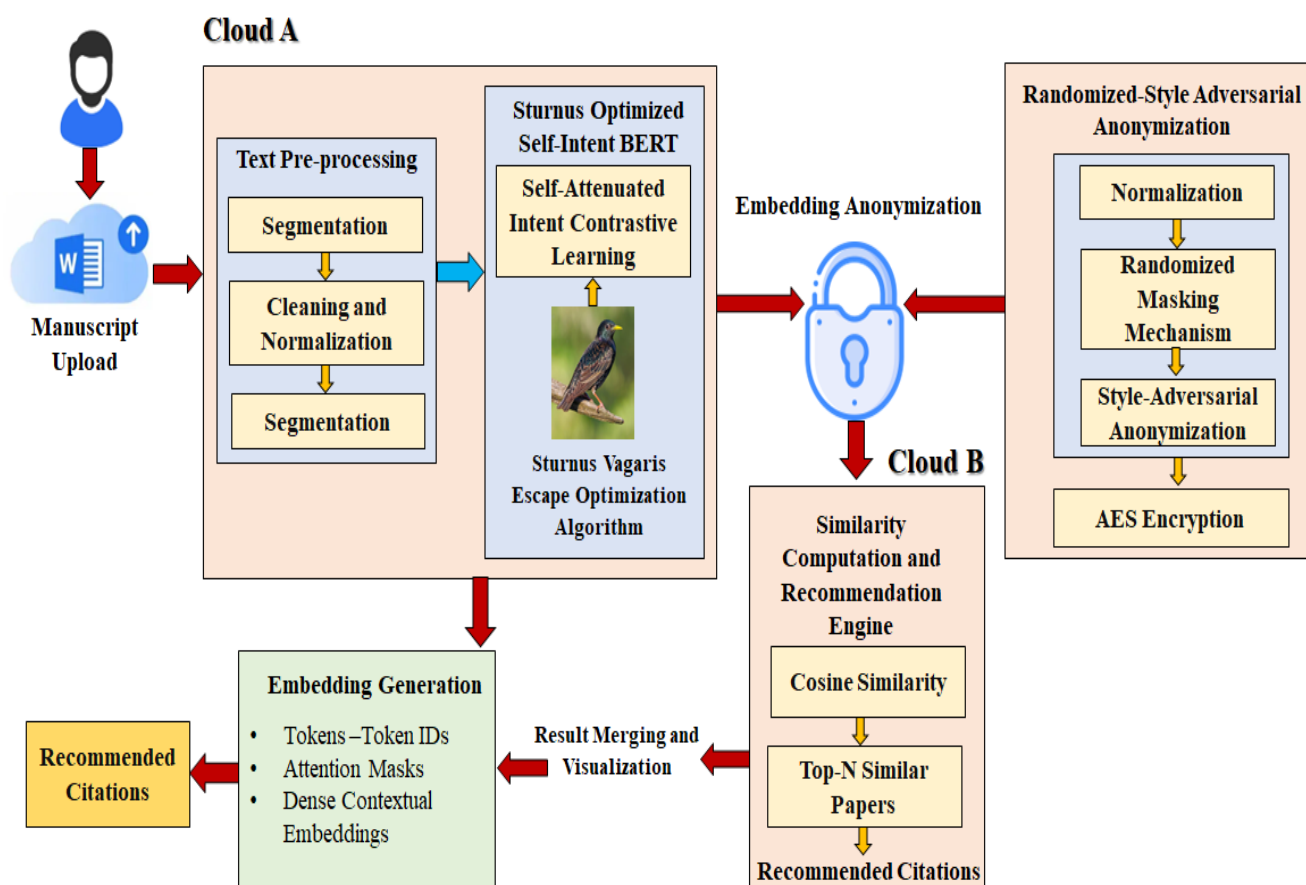


Figure 1. Overall Structure of the Proposed Dual Cloud SIL-BERT based BNR.

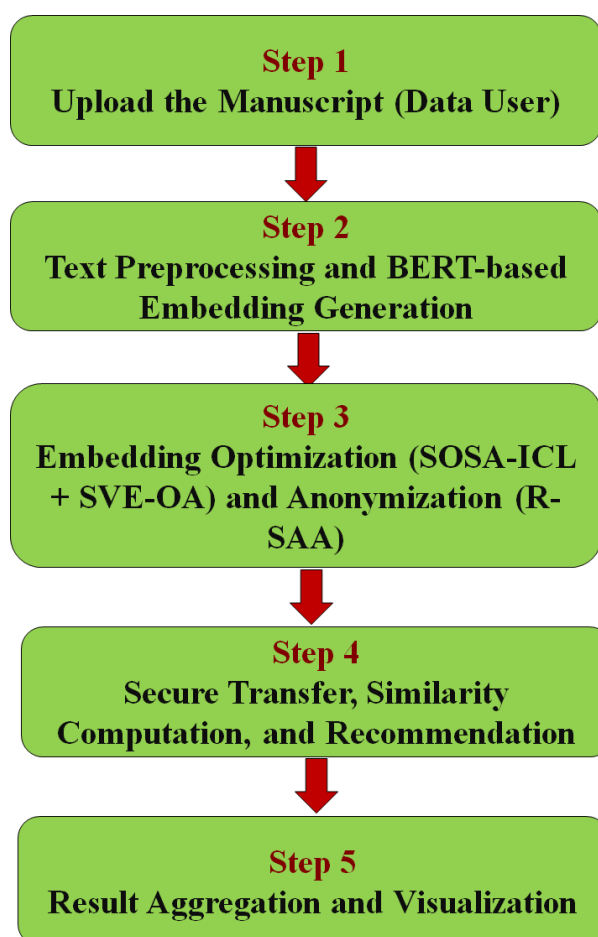


Figure 2. Steps involved in the proposed Dual Cloud BNR with SIL-BERT.

4.1 Text pre-processing

Primarily, the unpublished manuscript prepared by the researcher is uploaded in a trusted environment of Cloud A, in which the text is pre-processed. In this stage, noise and clutter present in the document are removed, and the document is then organized to isolate the title, abstract, references, along with exclusion of equations and tables, while standardizing the text.

- **Segmentation:** *In this stage, the manuscript gets sub-divided into logical units namely sentences and sections.*
- **Normalization:** The text is cleaned and standardized by removing noise, irrelevant symbols, and formatting inconsistencies.
- **Tokenization:** This stage includes tokenization, which is a Pre-BERT stage during which the cleaned text gets prepared to be fed to the proposed SIL-BERT model. The defined stage involves a Word Piece tokenizer that splits the words into sub-word units, and then the special tokens like [CLS] and [SEP] are inserted for identifying the boundaries of the document, while permitting the BERT to generate precise semantic embeddings of the uploaded manuscript.

In this, first, the raw text is converted into a readable form using BERT, in which the words or sub-words are mapped to token IDs based on the current vocabulary of the BERT, so as to identify all the tokens uniquely. In addition, attention masks are generated to identify valid tokens and distinguish them from padding tokens, while the process is mainly carried out to ensure that BERT computes semantically essential text portions and also to maintain contextual integrity, in case of padded sequences. The pre-processing module delivers token IDs plus attention masks. These become the starting point for the BERT-based Transformer Encoder

4.2 Generation of Contextual Embeddings using Sturnus Optimized Self-Intent BERT (SS-IBERT)

Taking the token IDs and attention masks from the pre-processing stage, the SS-IBERT module now generates dense vector representations that capture the semantic meaning of the manuscript. After pre-processing the manuscript, it is fed to the SS-IBERT, where S-IBERT serves as the core embedding generation component of the overall SIL-BERT framework, to produce the dense vector representations or the embeddings. These embeddings assist in storing the semantic meaning of the manuscript, while

permitting effective similarity analysis and also CR. The defined SS-IBERT involves a BERT-based transformer Encoder, Self-Attenuated Intent Contrast Learning and Sturnus Vulgaris Escape Optimization Algorithm (SVE-OA). The proposed SS-IBERT provides embeddings that are effective and also less sensitive to the algorithmic rhetorical biases. The SS-IBERT outputs contextual embeddings, which are then passed to the SOSA-ICL module for bias mitigation.

4.2.1 BERT-based Transformer Encoder

The first component inside SS-IBERT is the BERT encoder. It receives the token IDs and attention masks and produces token-level contextual embeddings. In order to generate the contextual embeddings that incorporate the semantic meaning of the text, the BERT-based Transformer Encoder is employed. The BERT represents a language representation model, which is based on a transformer architecture that surpasses in distinct NLP tasks [31]. The bi-directional training and the masked language modelling involved enable effective understanding with the context in various directions, resulting in superior performance. In order to train the BERT, masked Language Modelling (MLM) and Next Sentence Prediction (NSP) are employed, and it enhances the sentence and word-level comprehension. With this, the BERT highly assists in providing a rich context of bi-directional representation, which is effective for CR.

The general architecture of BERT are illustrated in Figure 3, in which the core of the BERT involves multiple stacked transformer encoder layers. Every single encoder layer incorporates a multi-head self-attention to capture the dependencies amid the words, a Feed-Forward neural Network that applies a non-linear transformation for improving the contextual understanding and then layer normalization with residual connections, which stabilizes the training, also safeguards the gradient flow. In BERT framework, every single token ID is mapped to a vector with the aid of embedding layer and every single token involves three components, namely token embedding, segment embedding and position embedding. The final embedding for a specific token is represented as given in equation (1) [29].

$$e_i = e_{tok}(t_i) + e_{seg}(s_i) + e_{pos}(p_i) \quad (1)$$

Where, e_i denotes the final embeddings, $e_{tok}(t_i)$ specifies the token embedding, $e_{seg}(s_i)$ specifies the segment embedding and $e_{pos}(p_i)$ denotes the position embedding. This token representation is then fed to the stack of multi-layer bi-directional transformer encoder blocks, in which every single layer involves the Updation of token embeddings with self-attention, permitting it to encounter all the other tokens in both backward and forward direction. Each token gives out three matrices like Query (Q), Key (K) and Value (V).

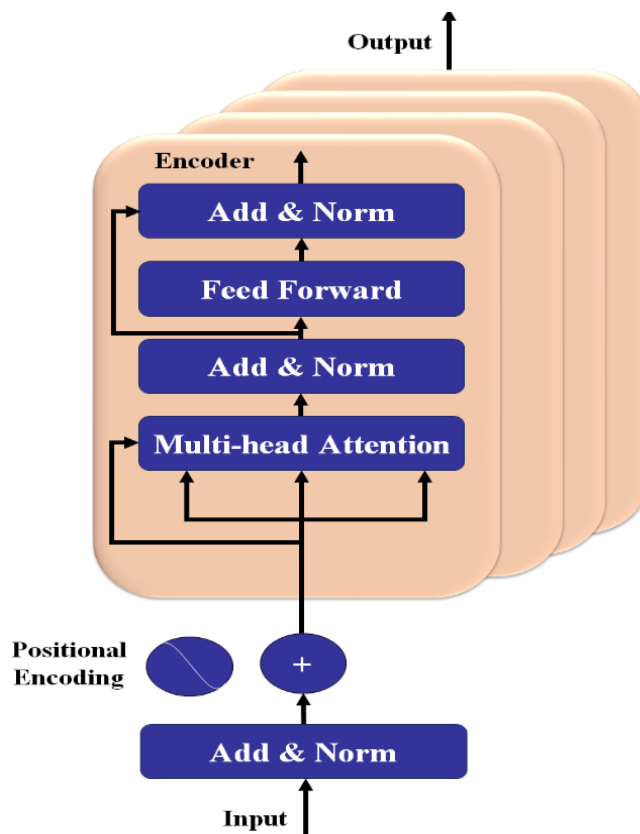


Figure 3. General architecture of BERT.

In order to attain the value vectors, the BERT multiplies all the embeddings using a learned value weight matrix. Moreover, the attention mechanism involved in the BERT permits every single word in the sentence to concentrate on other words that give out equivalent context of meaning, while the self-attention weights are calculated as given in equation (2) [32].

$$AT(Q_{at}, K_{at}, V_{at}) = \text{Softmax}\left(\frac{Q_{at} K_{at}^T}{\sqrt{D_v}}\right) V_{at} \quad (2)$$

In which the Q_{at} , K_{at} and V_{at} denotes the query vector, key vector and the value vector, while the D_v specifies the dimensionality of the key vectors. The application of the SoftMax function transforms the raw scores into attention weights. Further, the vector attained is fed as input to the Feed Forward Network (FFN), which assists in refining and improving the representation of the tokens attained from the attention. The FFN is a two-layer network and the process involved is given as shown in equation (3) [32].

$$F_N = \text{Max}(0, xw_1 + b_1)w_2 + b_2 \quad (3)$$

Here, the w_1 and w_2 denotes the learned weights of the matrices, while b_1 and b_2 denotes the bias terms. Further, the FFN uses learned projection along with non-linearity and a second projection to generate the refined vector. Then, to ensure about preserving the information and stability, the residual connection and layer normalization is carried out. For maintaining the contextual flow and enhancing the gradient flow, the

defined phase involves the addition of original inputs back to the FFN output by the residual connection. Further, for standardizing the combined output, layer normalization is performed, which balances the activations and also stabilizes the training across the layers. Therefore, this aids in stabilizing the gradients and preserving the prior information, while the process is repeated to attain deep contextual understanding. Hence, with this, the contextual embeddings with semantic meaning of the text are generated and then to enhance the semantic fidelity, SOSA-ICL is employed. The BERT encoder sends token-level contextual embeddings directly into the SOSA-ICL module for intent-contrastive learning.

4.2.2 Sturnus Optimized Self-Attenuated Intent-Contrastive Learning (SOSA-ICL)

Next, the token-level embeddings from BERT flow into the SOSA-ICL module, which refines them by pulling together similar research intents and pushing apart dissimilar ones, while also tuning the embedding space using the Sturnus Vulgaris optimization algorithm. For mitigating the complexities with Algorithmic Rhetoric and to enhance the semantic fidelity, the SOSA-ICL is employed, and it includes techniques like Self-attenuated Intent Contrastive Learning (SA-ICL) and Sturnus Vulgaris Escape Optimization algorithm. The Structure of Sturnus Optimized Self-Attenuated Intent-Contrastive Learning are illustrated in Figure 4.

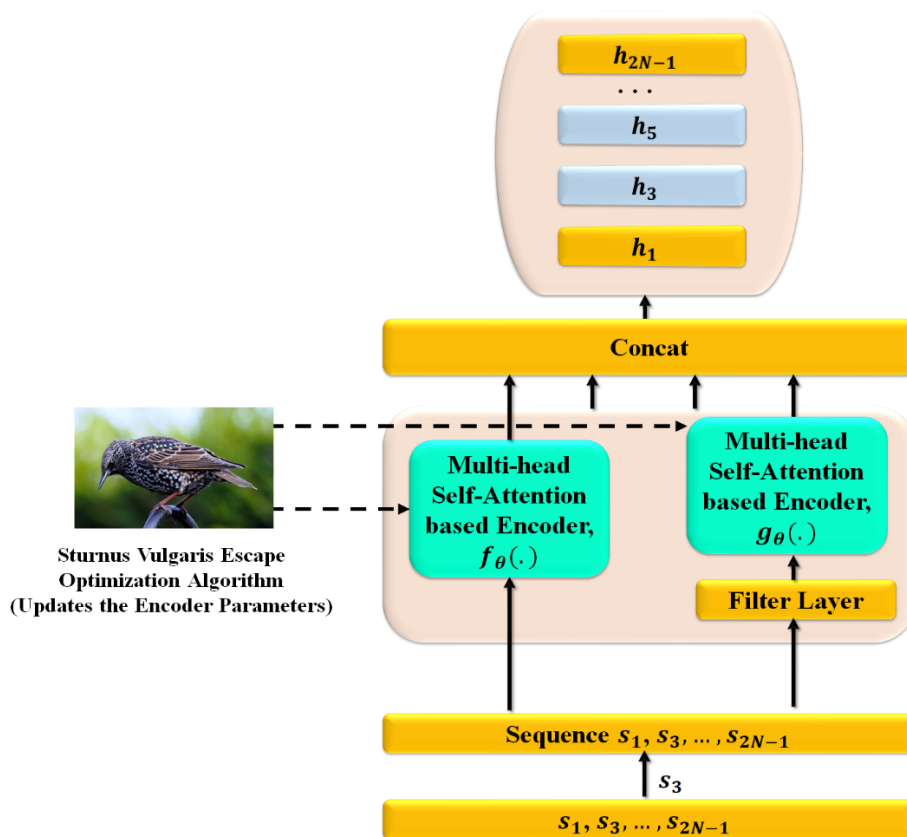


Figure 4. Structure of SOSA-ICL for CR.

4.2.2.1 Self-attenuated Intent-Contrastive Learning

The Self-Attenuated Intent Contrastive Learning (SA-ICL) is a specific representation of embeddings that reflect the document intent above the surface-level semantics. The SA-ICL mainly trains the model to pull together to present the same input and to push apart that represents various intents, with the aid of contrastive losses. The input encoding of SA-ICL involves the encoded texts from BERT and an intent-aware self-attention layer, in which the final embeddings from it are considered as intents. The attention is mainly used for re-weighting the token representation based on the relevance to the intent, which gives out a refined sequence that further brings fixed-length intent embeddings. This acts as the input for the SA-ICL, in which both the positive and negative pair construction is performed. The positive pair construction includes citations belonging to same paper, documents referred to for the same purpose and the query with same paraphrases. Meanwhile, the negative pair construction includes texts that possess the same keywords, but differ in intent and also the random samples attained from other intents. This pair construction is carried out using the Contrastive Objective loss function, which is specified as shown in equation (4) [33, 34].

$$COL = -\log \frac{\exp(SF(E_A, E_{ps})/\tau)}{\sum_n \exp(SF(E_A, E_{ns})/\tau)} \quad (4)$$

In which the E_A specifies the self-attention improvised intent embeddings of the anchor sample, the positive samples are represented as E_{ps} , while the negative samples are denoted by E_{ns} and the temperature parameter is given as τ . Hence, the

encoder gets fine-tuned, which thereby emphasizes the semantic intent in terms of superficial stylistic cues and then for optimizing the SA-ICL, an SVE-OA is employed.

4.2.2.2 Sturnus Vulgaris Escape Optimization Algorithm (SVE-OA)

In order to optimize the embedding space, the SVE-OA is employed and it effectively tunes the embedding alignment, contrastive learning parameters and the attention weights. The SVE-OA is a bio-inspired optimisation algorithm and is inspired with the collective escape behaviour of *Sturnus Vulgaris*, during the time of predation by *Falco peregrinus*. Moreover, the SV involves effective escape strategies that involves high altitude escape strategy, cordon line strategy and eave escape strategy.

The steps involved in Sturnus Vulgaris Escape Optimization Algorithm are illustrated in Figure 5. The algorithm is initialized by defining three individuals, like the best, the leader and the Falco individuals. The strategy used to generate the population in terms of tuning the embedding space is expressed as shown in equation (5) [35].

$$P_{i,j} = LB_i + rand() \times (UB_i - LB_i) \quad (5)$$

$rand()iUB_iLB_i$ denotes a function that generates random numbers uniformly distributed between 0 and 1; i denotes the ordinal number of each individual in the population; and LB and UB denote the dimensional vectors representing, respectively, the lower and upper bounds of the values assumed by the positions of the individuals.

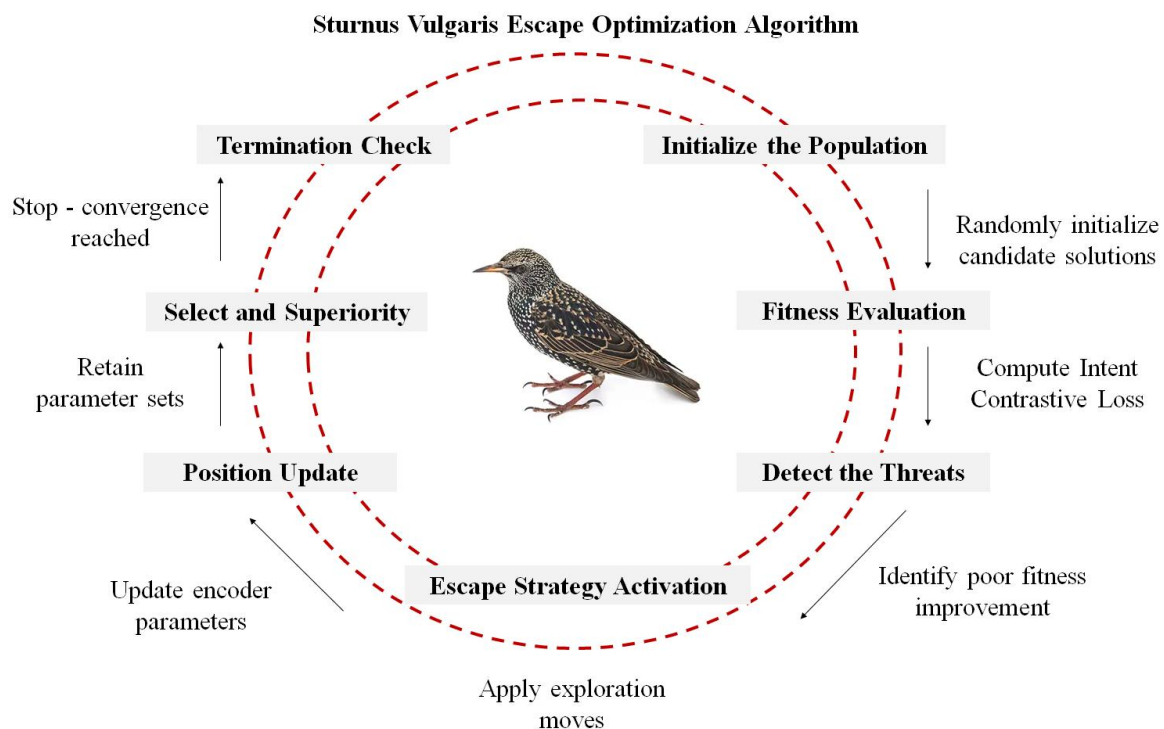


Figure 5. Steps involved in Sturnus Vulgaris Escape Optimization Algorithm.

Exploration Phase: This stage involves strategies like high altitude escape and wave escape strategy, in which the escape and the scattering behavior of the starlings are used to identify the new positions in the search space, while eliminating the premature convergence and is specified in equation (6) [35].

$$P_i^{t+1} = P_i^t + \lambda \cdot \mathcal{N}(0, I) + rad_1(P_{rand} - P_i^t) \quad (6)$$

$\lambda \cdot \mathcal{N}(0, I) \mathcal{N}(0, I) rad_1$ Here, the exploration strength coefficient is denoted by, the gaussian random perturbation is specified as, the randomly chosen agent is given by and is considered to random coefficients. This, thereby, assists in enhancing the diversity, while escaping local optima and extending the search space.

Exploitation Phase: This stage involves the cordon line strategy as well as the wave escape strategy, and it limits the premature maturation of the individuals. Moreover, this effectively captures the cohesive flocking behaviour, while helps the individuals to attain the better solutions and is represented as given in equation (7) [35].

$$P_i^{t+1} = P_i^t + rad_2(P_{Best} - P_i^t) + rad_3(P_{neigh} - P_i^t) \quad (7)$$

$P_{Best} P_{neigh}$ In which the specifies the global best solution and denotes the best local neighbor. This, thereby, enhances the convergence speed and then refines the cluster boundaries, also making the same-intent embeddings robust. Moreover, the adaptive switching amid the exploration and exploitation phase is balanced with the time varying regulation parameter and is represented as in equation (8) [35].

$$P_i^{t+1} = \eta(t) P_i^{Explore} + 1 - \eta(t) P_i^{Exploit} \quad (8)$$

Hence, with the exploration phase, inter-intent margins are identified and then at the exploitation phase, intra-intent variance is reduced, while the embedding gets stabilized by its geometry.

Therefore, the integration of SIL-BERT based techniques gives out semantically significant embeddings, which are substantially less sensitive to algorithmic rhetorical biases, After SOSA-ICL, the refined intent embeddings are aggregated into a fixed-dimensional intent vector, which then moves to the anonymization pipeline. Then, for preserving researchers' privacy, also for reducing the complexities with re-identification via linkable attributes, an R-SAA technique is employed.

4.3 Aggregation of Embeddings

The SOSA-ICL module produces refined token-level intent embeddings. These are now aggregated into a single fixed-dimensional vector that represents the overall research intent of the entire manuscript. The token-level contextual representations are generated and then refined with the aid of the SOSA-ICL module,

which are then aggregated, to develop a single fixed-dimensional vector, giving out the overall intent of the input text. In this step, the variable-length token sequence is transformed into comparable and compact embeddings that are efficient for upcoming tasks. With this attention weight-based aggregation, the tokens most relevant to the intent contribute effectively to the final depiction. The aggregation step creates one intent vector per manuscript, which is immediately forwarded to the anonymization pipeline.

4.4 Embedding Anonymization and Encryption

Once the intent vector is obtained, it enters the privacy-preserving stage. First, the vector is normalized, then anonymized via R-SAA, and finally encrypted before leaving Cloud A. Then, to preserve the sensitive information along with the utility, the defined embedding anonymization is applied to the aggregated intent embeddings, before transferring it. In order to eliminate the complexities with re-identification via linkable attributes a novel anonymization pipeline is employed, which includes Normalization and Randomized-Style Adversarial Anonymization (R-SAA). The R-SAA includes techniques like Randomized Masking Mechanism and Style Adversarial Anonymization.

4.4.1 Normalization of Embeddings

The first step in the anonymization pipeline is normalization. The aggregated intent vector is scaled to unit length, removing magnitude-based leakage while preserving directional semantics. For ensuring stability and scale consistency across all the samples, Normalization of embeddings is performed. The original embedding magnitudes involve leakage of information related to document length, confidence-levels and writing style. This necessitates normalization, in which each embedding is normalized to unit length, while excludes the traces of the magnitudes. This allows the embeddings to be transferred from Cloud A to Cloud B in a safer way. Moreover, Cloud B carries out the similarity check with cosine similarity, which relies on the vector direction and hence the normalization also ensures about the reduction of bias in the similarity measures across the document. This further stabilizes the comparisons performed amid the papers from various sources and corpora.

$et \in \mathbb{R}^d$ For every single token embedding, , that denotes the continuous valued embedding vector along with the real valued numbers and embedding dimensionality, which is generated by BERT, the normalization process rescales it into unit length using the formula specified in equation (9) [36].

$$\hat{et} \in \frac{et}{\|et\|_2}, \text{ where } \|et\|_2 = \sqrt{\sum_{b=1}^d et_b^2} \quad (9)$$

Hence, with this, the embedding attains unit length, while also preserves the semantic direction with security. Further, to address the complexities with re-identification, a R-SAA technique is utilized. Finally, the normalization produces unit-length embeddings, which are then processed by the R-SAA module.

4.4.2 Randomized-Style Adversarial Anonymization (R-SAA)

After normalization, the unit-length embedding enters the core anonymization stage. R-SAA applies randomized masking to suppress stylistic dimensions and adversarial learning to remove behavioural fingerprints. Even though the learned embeddings give out better stylistic and semantic information, it is exposed to higher-level of privacy risks due to re-identification via linkable attributes. Therefore, to limit this a novel Randomized-Style Adversarial Anonymization is used, while it includes techniques like Randomized Masking Mechanism and Style Adversarial Anonymization.

et $\in \mathbb{R}^d$ $\tilde{et}$ **Randomized Masking Mechanism (RMM)**: To further anonymize the embeddings, the RMM technique is employed and it effectively suppresses the dimensions that encode the stylistic or semantic information. rather than perturbing the dimensions uniformly, the RMM randomly masks the embedding components. This thereby decreases the identity leakage risk or intent risk, while safeguards the semantic efficacy. If *et* is the aggregated intent embedding, then the masked embedding is represented using equation (10) [37, 38].

$$\tilde{et} = e \odot m \quad (10)$$

mm $\in \{0, 1\}^d$ Here, specifies the binary masking vector, where, which is a binary masking vector. Hence, the proposed RMM effectively prevents adversaries from inferring the masked dimensions, for reconstruction. This thereby limits the reconstruction attacks and fine-grained stylistic cues.

Style Adversarial Anonymization (SAA): Then, for enhancing the privacy protection, SAA [39] is employed and it eliminates the behaviour-based features. As RMM effectively suppresses the features, the SAA aims at implicit stylistic signals, which remains embedded and for identifying these condition attributes, an adversarial discriminator is employed. The adversarial representation of the defined work involves training of embedding generator and a style discriminator in opposition. The former assists to identify the style-related attributes, while the latter is optimized to prevent the predictions. Therefore, with this process, the generator produces style-invariant embeddings, which are unhelpful, in terms of stylistic features, but remains significant for intent representations. Further, to improve confidentiality, encryption is performed. R-SAA

completes its work by passing anonymized embeddings to the AES-256 encryption stage.

4.4.3 AES Encryption and Protection of Data during Transfer

Finally, the anonymized embeddings are encrypted with AES-256. This turns them into ciphertext, ensuring that even if intercepted, they cannot be traced back to the original manuscript. To ensure the integrity and confidentiality of data at the time of transmission, a symmetric key encryption mechanism known as Advanced Encryption Standard (AES) is used [40]. Before transferring the embeddings from Cloud A to Cloud B, the generated intent embeddings and corresponding metadata are encrypted with the aid of AES, while also ensures the protection against interception, tampering and from unauthorized access. The AES functions with 128 bits of fixed size data and assists a key length of about 128, 192 and 256 bits, while gives out strong resistance over the brute-force and also the cryptographic attacks. In the proposed work, the Cloud A produces a secure key (K-AES) which is known only to the enclave of Cloud B. Every single normalized embedding vector, in terms of ciphertext is represented as given in equation (11) [41].

$$CPT = AES_{KY_{AES}}(\tilde{et}) \quad (11)$$

$\tilde{et} \xrightarrow{AES_{KY_{AES}}} CPT$ Here, the normalized embedding vector is given as, denotes the encryption function with key and CPT is the ciphertext attained. Hence, with the defined technique, even if the data are intercepted during transmission, the original manuscript cannot be reconstructed or linked back the original manuscript, thereby maintaining the integrity and confidentiality throughout the process. Then for identifying the similarity matching, the anonymized embeddings are fed to Cloud B. Encryption turns the embeddings into ciphertext, which is then transmitted to Cloud B for similarity analysis.

4.5 Similarity Analysis and Citation Retrieval at Cloud B

Cloud B receives the encrypted embeddings, decrypts them inside its secure enclave, and then performs similarity matching. The result is a ranked list of recommended citations sent back to the user. For identifying the relevant citations of the query or manuscript, the Similarity Analysis (SA) is carried out and every single input text is given as a fixed-dimensional intent embeddings, while captures the semantic content, along with the existing intents. Moreover, in proposed work, Cloud B does not process the text; instead it works only with the numerical vectors. The steps involved in the defined process are as follows.

- **Step 1 Receiving the Decrypted Embeddings:** Cloud B decrypts the embeddings inside its

secure enclave, while each row provides a higher dimensional vector that aids in capturing the contextual meaning of every word in the manuscript.

- **Step 2 Similarity Computation:** In this step, the angular proximity of the optimized embedding space is measured and is also invariant to the magnitude of the pre-processing that makes it well suited for the intent-based comparison. For computing the similarity, Cosine similarity is employed and is represented as given in equation (12) [42].

$$CS(a, b) = \frac{a \cdot b}{\|a\| \times \|b\|} \quad (12)$$

Here, a and b denotes the embedding corresponding to the uploaded manuscript and the other paper in the database and $\|a\| \times \|b\|$ specifies the vector magnitudes.

- **Step 3 Ranking and the Top-N Retrieval:** Cloud B repeats the similarity computation for all the papers in the database and then ranks the papers in descending order with respect to the similarity scores, thereby producing relevant citation recommendations. The citation retrieval based on the query embeddings is represented as given in equation (13) [42].

$$Rk(c) = Sort_{des}(sim(a, b)) \quad (13)$$

Hence, from this, the Top-N citations that possess higher similarity values are chosen as the most relevant references, while the defined technique ensures that the recommended citation is not only related semantically, but also is intent-consistent with the uploaded manuscript. This also ensures that the recommended process is highly secure and does not expose any sensitive information related to the paper or the author. Cloud B delivers a ranked list of candidate citations, which serves as the final output to the user.

4.6 Overall Technical flow of the proposed SIL-BERT pipeline

The proposed framework uses a sequential processing pipeline in which each module operates on the output generated by the previous stage.

- **Cloud A – Pre-processing:** The uploaded manuscript is cleaned, segmented, normalized, and tokenized using the Word Piece tokenizer. The output consists of token IDs along with attention masks.
- **BERT-based Transformer Encoder:** The token sequence is passed to the BERT model, which generates contextualised embedding vectors for each token.
- **SOSA-ICL:** These embeddings are processed using the Self-Attenuated Intent Contrastive Learning mechanism, where semantically similar representations are grouped together and dissimilar ones are separated. The learning process is guided by a contrastive loss function. In addition, the Sturnus Vulgaris Escape Optimization (SVE-OA) strategy is employed to refine model parameters (attention weights, contrastive parameters, embedding alignment), reducing the need for manual tuning.
- **Aggregation:** The refined embeddings are aggregated into a fixed-dimensional vector representing the overall intent of the manuscript.
- **R-SAA Anonymisation:** The aggregated vector is normalised and processed through a randomised masking mechanism, followed by style-based anonymisation to reduce identifiable patterns.
- **AES Encryption:** The anonymised representation is encrypted using a symmetric encryption scheme before transmission.
- **Transfer to Cloud B:** The encrypted data is transmitted to a secondary cloud environment for further processing.
- **Similarity & Retrieval:** The representation is decrypted, and cosine similarity is computed against stored embeddings to identify relevant documents. The top-N citations are then returned as recommendations. Each stage produces an output that serves as the input to the subsequent module, enabling a clearly traceable workflow from input manuscript to final citation output.

Table 2. Input-output relationships between modules

Step	Module	Input	Output
1	Pre-processing	Raw Manuscript	Token IDs + Masks
2	BERT Encoder	Token IDs + Masks	Contextual Embeddings (768-D)
3	SOSA-ICL + SVE-OA	Contextual Embeddings	Refined Intent Vector
4	R-SAA Anonymization	Intent Vector	Anonymized Vector
5	AES-256 Encryption	Anonymized Vector	Encrypted Data
6	Similarity & Retrieval	Decrypted Vector	Ranked Citations (Top-N)

The input-output relationship of each processing module is summarized in Table 2.

captures the meaning of the words, word order and the sentence identity.

4.7 Illustrative Example

Assume, the researcher uploads, "Graph Neural Networks (GNNs) can capture the relationship between papers in citation graphs. These models help improve citation recommendation".

Step 1: After text pre-processing, it becomes "Graph Neural Networks can capture the relationships between papers in citation graphs. These models help improve the citation recommendation accuracy".

Tokens: [CLS], graph, Neural Networks, can, capture, the relationships, between, papers, in, citation, graphs, These, models, help, improve, the, citation, recommendation, accuracy, [SEP].

Token IDs: [CLS] = 101, graph = 4689, neural = 3138, networks = 2301..., [SEP] = 102. $X = [101, 4689, 3138, 2301, 2064, 5382, 5315, 4623, 2090, 2470, 4842, 1999, 8296, 5251, 102]$. With this the paragraph gets converted into a sequence of 21 tokens.

Step 2: In Embedding generation with SIL-BERT, the paragraph is given as a 21×21 matrix and it

Table 3. Derived final 5-D embedding vectors

Component	5-D Values
Token Embedding	[0.22, -0.15, 0.58, 0.41, 0.05]
Segment Embedding	[0.00, 0.00, 0.00, 0.00, 0.00]
Position Embedding	[0.01, -0.02, 0.00, 0.03, 0.01]
Final (Sum)	[0.23, -0.17, 0.58, 0.44, 0.06]

Tables 3 and 4 illustrate the derived 5D embedding vectors and the representation of all Tokens with 5-D values. These values are illustrative representations computed using the BERT-base architecture as described in [31, 32]. Hence, these embeddings form a 16,128-feature matrix and the snapshot of the embedding matrix that show the sample 5-D representation of tokens are as follows.

The BERT-base model includes 12 layers that effectively refine the word meanings and after all the layers, the citation captures its relationship with "papers and the "recommendation" in the research context. The final embeddings obtained are presented below.

Table 4. Representation of all Tokens with 5-D values

Token	Token Embedding (E_{token})	Segment Embedding ($E_{segment}$)	Position Embedding ($E_{position}$)	Final Embedding ($E_i = sum$)
Graph	[0.10, -0.20, 0.55, 0.33, 0.05]	[0, 0, 0, 0, 0]	[0.02, -0.01, 0.02, 0.02, 0.00]	[0.12, -0.19, 0.54, 0.35, 0.05]
Neural	[0.24, -0.30, 0.42, 0.37, 0.08]	[0, 0, 0, 0, 0]	[0.03, -0.02, 0.01, 0.01, 0.02]	[0.27, -0.32, 0.42, 0.38, 0.10]
Networks	[0.18, -0.25, 0.61, 0.29, 0.04]	[0, 0, 0, 0, 0]	[0.04, 0.00, 0.02, 0.00, -0.01]	[0.22, -0.25, 0.63, 0.29, 0.03]
can	[0.09, -0.18, 0.40, 0.28, 0.02]	[0, 0, 0, 0, 0]	[0.01, 0.00, 0.00, 0.02, 0.01]	[0.10, -0.18, 0.40, 0.30, 0.03]
capture	[0.12, -0.23, 0.46, 0.31, 0.06]	[0, 0, 0, 0, 0]	[0.02, -0.01, 0.01, 0.02, 0.01]	[0.14, -0.24, 0.47, 0.32, 0.07]
relationships	[0.20, -0.19, 0.50, 0.33, 0.03]	[0, 0, 0, 0, 0]	[0.03, 0.01, -0.01, 0.01, 0.02]	[0.23, -0.18, 0.49, 0.33, 0.05]
papers	[0.15, -0.22, 0.52, 0.34, 0.04]	[0, 0, 0, 0, 0]	[0.02, 0.00, 0.01, 0.02, 0.01]	[0.17, -0.22, 0.53, 0.35, 0.05]
citation	[0.22, -0.15, 0.58, 0.41, 0.05]	[0, 0, 0, 0, 0]	[0.01, -0.02, 0.00, 0.03, 0.01]	[0.23, -0.17, 0.58, 0.44, 0.06]
recommendation	[0.16, -0.26, 0.48, 0.28, 0.05]	[0, 0, 0, 0, 0]	[0.04, 0.01, 0.02, 0.01, 0.01]	[0.20, -0.27, 0.50, 0.29, 0.06]
accuracy	[0.13, -0.21, 0.55, 0.31, 0.04]	[0, 0, 0, 0, 0]	[0.03, -0.01, 0.00, 0.03, 0.00]	[0.16, -0.21, 0.55, 0.33, 0.04]

[0.12, -0.19, 0.54, 0.35, 0.05...],	←	"graph"
[0.27, -0.32, 0.42, 0.38, 0.10 ...],	←	"neural"
[0.22, -0.25, 0.63, 0.29, 0.03...],	←	"networks"
[0.10, -0.18, 0.40, 0.30, 0.03 ...],	←	"can"
[0.14, -0.24, 0.47, 0.32, 0.07...],	←	"capture"
[0.23, -0.18, 0.49, 0.33, 0.05 ...],	←	"relationships"
[0.17, -0.22, 0.53, 0.35, 0.05...],	←	"papers"
[0.23, -0.17, 0.58, 0.44, 0.06, ...],	←	"citation"
[0.20, -0.27, 0.50, 0.29, 0.06...],	←	"recommendation"
[0.16, -0.21, 0.55, 0.33, 0.04...],	←	"accuracy"

"graph"	→	[0.30, -0.15, 0.61, 0.39, ...]
"citation"	→	[0.28, -0.08, 0.55, 0.25, ...]
"recommendation"	→	[0.42, -0.17, 0.74, 0.48, ...]

Step 3: R-SAA based Embedding Anonymization includes normalization of the contextual embeddings attained from the BERT, $et = [0.28, -0.08, 0.55, 0.25, \dots]$ and is given as,

$$et = \frac{1}{0.67} [0.28, -0.08, 0.55, 0.25, \dots] \approx [0.418, -0.119, 0.821, 0.373, \dots]$$

Further, AES Encryption is carried out, in which the normalized semantic embeddings represent the contextual meaning of every single token, (i.e.),

$$\text{"citation"} \rightarrow [0.418, -0.119, 0.821, 0.373, \dots]$$

Therefore, the embeddings get securely encrypted with the aid of AES-256 algorithm and further it is fed to Cloud B. As Cloud B's secure enclave recognizes the symmetric encryption key, it transforms the embeddings into ciphertext, while the resultant encrypted data appears in terms of random hexagonal values specified as,

$$\text{"citation"} \rightarrow \text{"8F3A2B7E90D1A4C3E2"}$$

Hence, with this, the encryption process ensures that even if data gets interrupted with Algorithmic Rhetoric issues and S-NBV, the proposed Dual Cloud SIL-BERT based BNR model effectively prevents the reconstruction or link back to the original manuscript, thereby maintaining the integrity and security of the manuscript throughout the entire transfer.

5. Results and Discussion

This research work introduces a novel Dual Cloud Sturnus Optimized Intent Learning BERT with Randomized-Adversarial Secure Anonymization for secure CR. Initially, manuscript uploaded in Cloud A gets pre-processed and is then improved with SA-ICL. Further, anonymization is performed using R-SAA and is then fed to Cloud B for calculating the similarity index. In order to analyze the effectiveness of the defined work, it

is implemented in Python platform and the performance metrics attained are evaluated. In addition, to demonstrate the significance of the defined work, it is compared with conventional techniques.

5.1 Dataset Description

For identifying the effectiveness of the proposed work, it is analyzed with DBLP dataset [43], which is a major source used for generation and testing. The defined dataset involves largest collection of bibliographic records and it includes millions of research papers along with the metadata. This incorporates abstracts, titles, authors, publication venue and also the citation relationship. Moreover, the citation network permits the model to learn about the contextual and semantic relationship amid the research work. The defined dataset includes 6.6 million of publications involving conferences, journals and workshops. Then, for producing semantic embeddings from this massive dataset, the defined system employs batch based embedding generation strategy. Moreover, the papers are kept in small batches precisely to about 1000 records at the place of loading the overall dataset inside the memory. Meanwhile, the title, as well as abstract of every single paper gets tokenized and encoded using the SIL-BERT model, so as to generate the dense 768-dimensional semantic vectors that reflect the topical and contextual similarity of the research paper. Then, the embeddings resulted are stored as compressed NumPy (.npz) files, along with the paper ID and the vector representation.

5.2 System Configuration

This section involves the system configuration used for implementing the Dual Cloud SIL-BERT based BNR system, while the configurations involved is given below.

Software Used	: Python
OS	: Windows 10 (64-bit)
Processor	: Intel i5
RAM	: 8GB

5.3 Implementation and Training Details

The proposed SIL-BERT framework is built upon the BERT-based (BERT-based-uncased) architecture with 12 transformer layers, 768 hidden dimensions, and 12 attention heads. The model is fine-tuned on the DBLP dataset for 100 epochs using a batch size of 32. The AdamW optimizer is employed with an initial learning rate of 2×10^{-5} , which is decayed by a factor of 0.9 every 10 epochs. For intent-aware learning, the SOSA-ICL module utilizes contrastive loss, where the temperature parameter (τ) is set to 0.1 based on empirical tuning. The Sturnus-based optimization

algorithm is applied with an exploration coefficient (λ) of 0.8 over 50 iterations to refine embedding alignment. The example sentence used in Section 4.6 is processed using the trained model to demonstrate the embedding workflow; however, the numerical values shown in Tables 3 and 4 are simplified projections of the actual 768-dimensional embeddings for interpretability. In the implemented system, embeddings are generated in full high-dimensional space and then reduced only for visualization purposes. For privacy preservation, embedding-level anonymization and encryption are implemented using standard cryptographic libraries with AES-256 encryption

5.4 Simulation Results

In the defined research work, a Dual Cloud SIL-BERT based BNR is proposed for CR system, in which the manuscript uploaded initially encounters text-pre-processing that formats and transforms the paper into a readable format. Further, for addressing the challenges with Shillings attacks, SS-IBERT is employed. The SA-ICL trains the embeddings, so that the embeddings with similar and distinct research intent are segmented [33, 34].

The top 20 and 50 words in DBLP manuscript along with Word Cloud of DBLP manuscript are illustrated in Figure 6(a, b). In order to identify the high-

frequency words that possess the essential information of the task, the top 20 and 50 words are identified. The smaller feature set with 20 words provides quick computation, while the larger dataset with 50 words captures richer semantic content. By analyzing the model with various word counts assists in identifying the sensitivity of the model. Moreover, Figure 6(c) presents the most frequent as well as semantically salient terms, while provides an intuitive overview of dominant search themes. The words with large font show the high occurrence and this act as the initial qualitative validation, assisting in embedding generation. This thereby excludes the issues with S-NBV thereby decreasing the risks of re-identification.

The intent-strength analysis with Normal/Anomalous Text Segments and Vector Magnitude are illustrated in Figure 7 (a) and (b). This analysis is mainly performed to identify the strength of the individual text segments that specifies the citation intent. In Figure 7 (a), each point denotes the citation index with intent strength score, and reflects the semantic focus degree and the necessary citation captured by the Self tuned BERT model. Moreover, the anomalous segments act as outliers, giving out high intent strength due to rhetorical overemphasis. Hence, this significantly compresses the algorithmic rhetoric, while ensures that only genuine intent contributes in embedding aggregation and in upcoming CRs.

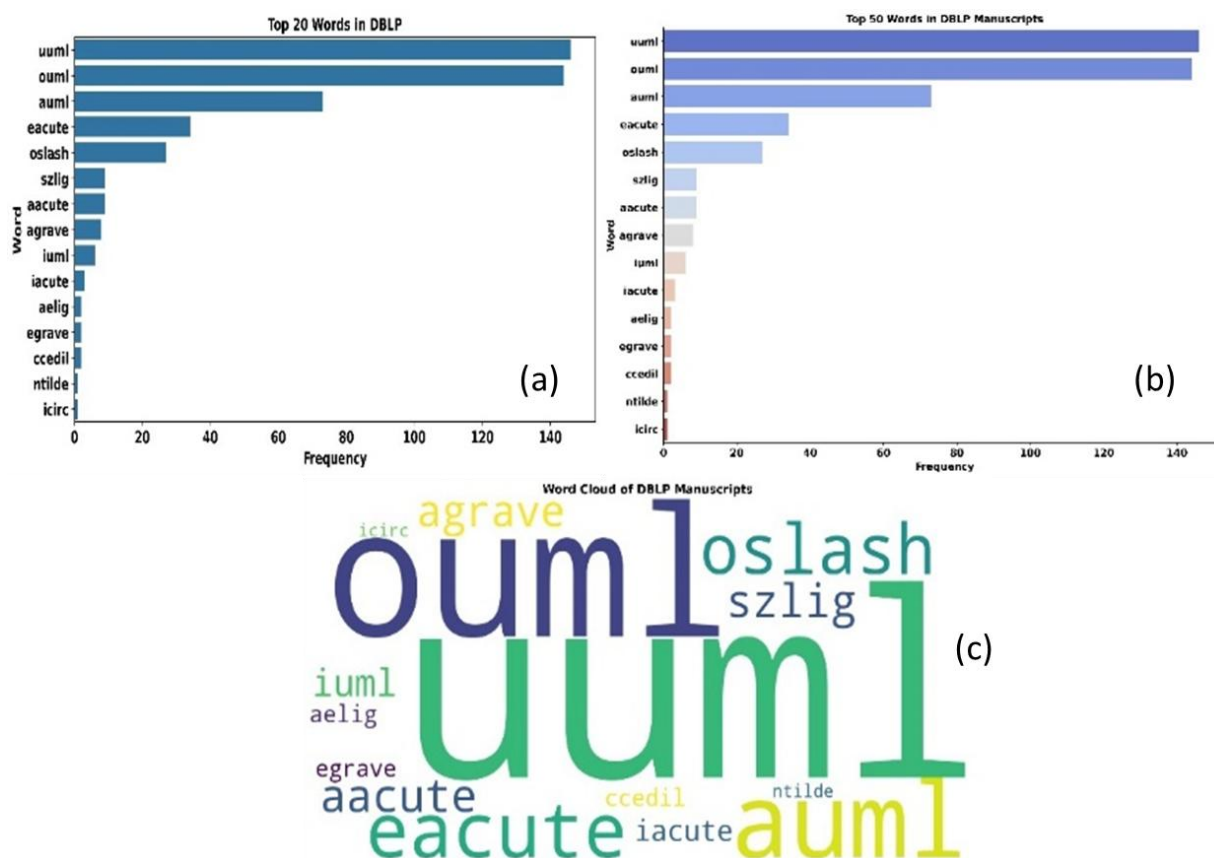


Figure 6. (a) Top 20 Words in DBLP, (b) Top 50 Words in DBLP manuscript and (c) Word Cloud of DBLP manuscript.

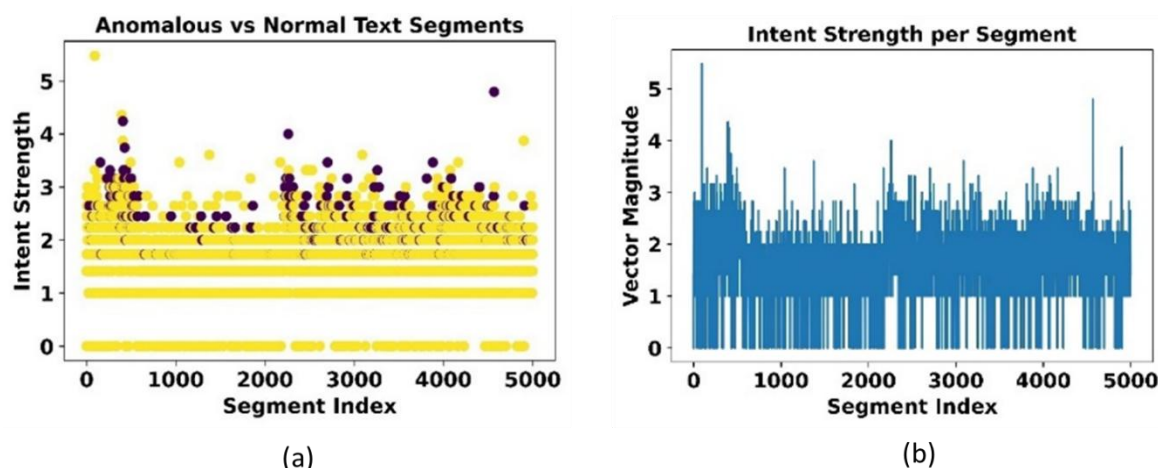


Figure 7. (a) Intent Strength analysis with **(b)** Normal/Anomalous Text Segments and Vector Magnitude.

Meanwhile, Figure 7 (b) denotes the discriminations in intent strength over successive text segments of the manuscript. Every single segment index corresponds to a citation index and a higher vector magnitude of about 5.6 at 200 segment index gives out the stronger intent strength. This shows the potential intent concentration zones, thereby ensures that any type of sharp spikes is realized, with which the bias in CR is eliminated.

5.5 Performance Analysis

The proposed Dual Cloud SIL-BERT based BNR is analyzed with standard retrieval and ranking techniques, so as to assess the accuracy of CR and also the quality of ranking. The proposed method is evaluated in terms of Precision, Accuracy, Recall, MRR (Mean Reciprocal Rank), MPP (Mean Precision at Position), F1-Measure, NDCG (Normalized Discounted Cumulative Gain) and MSE (Mean Square Error).

The accuracy of proposed SS-IBERT based BNR for CR is shown in Figure 8, which shows that, at epoch of 20, the accuracy is about 87%, while it increases and reaches a value of about 99% at an epoch of 100. This clearly demonstrates the ability of the system to retrieve the semantically correct citations. The use of SIL-BERT produces context-aware embeddings, which captures the intent nuances, [31, 32] which thereby decreases the bias existing in any of the specific portions, while also enhances the relevance of the recommendation system.

Figure 9 depicts the Precision of proposed SS-IBERT based BNR and it shows that as the epoch expands to 20, the precision increases to 0.78, while it further rises and reaches a precision value of about 0.87. The involvement of self-tuned embedding generation aligns the recommendations with the relevant manuscripts, while it decreases the issues with stylistic similar papers. This integration better assists in

retrieving the most relevant papers, enhancing the precision.

The Recall of proposed SS-IBERT based BNR are illustrated in Figure 10, which denotes the steep rise of recall value from 0.81 to 0.95 from epoch 20 to epoch 100. This shows that the recall value increases as the number of epochs increases. As the proposed work involves semantic contextual embeddings, less obvious but better relevant papers are taken into account, apart from considering the highly cited papers. Hence, these approaches make the relevance coverage better, while also improves the recall.

Figure 11 illustrates the F1-Measure of proposed SS-IBERT based BNR, from which the results show that there is a steady rise in the F1-Measure value, as the number of epoch's increases. At an epoch of 20, the F1-Measure value is about 0.76, while at an epoch of 100, the F1-Measure is about 0.89. The F1-Measure, which is an efficient metric brings balance amid the precision and the recall, [44, 45] while provides a balanced measure of system performance. The ICL highly ensures about the system in attaining the nuanced semantic relevance, while excludes the dependance over the highly cited papers.

The MAP of proposed SS-IBERT based BNR are illustrated in Figure 12, in which the increase in epoch increases the MAP value. The MAP is about 0.80 at an epoch of 20, while it increases and attains a value of about 0.89 at an epoch of 100. This demonstrates the improvement in MAP with the application of Optimized BERT. The re-ranking consistency involved, gives out effective enhancement in the MAP, while also provides well-ordered and more appropriate citation recommendations.

Figure 13, depicts the MRR of proposed SS-IBERT based BNR and the plot shows that the MRR has a steady rise from epoch 20 to epoch 100, while the concerned MRR values are about 0.80 and 0.95.

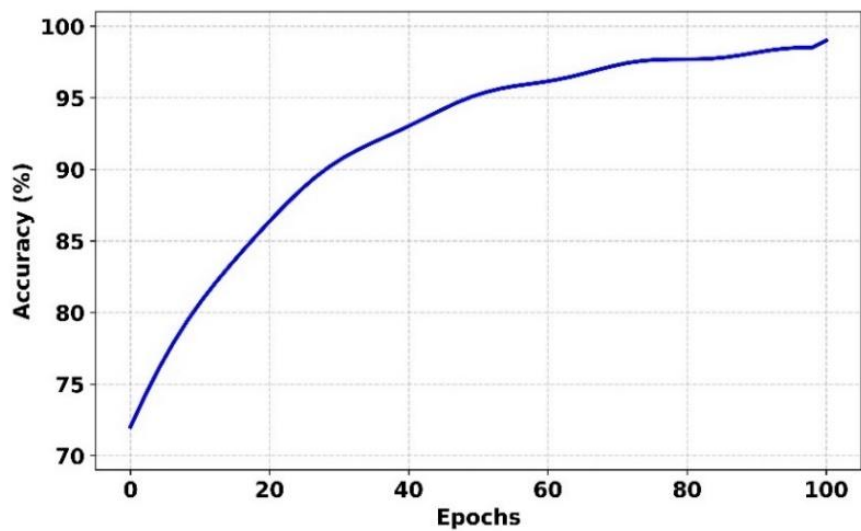


Figure 8. Accuracy of proposed SS-IBERT based BNR.

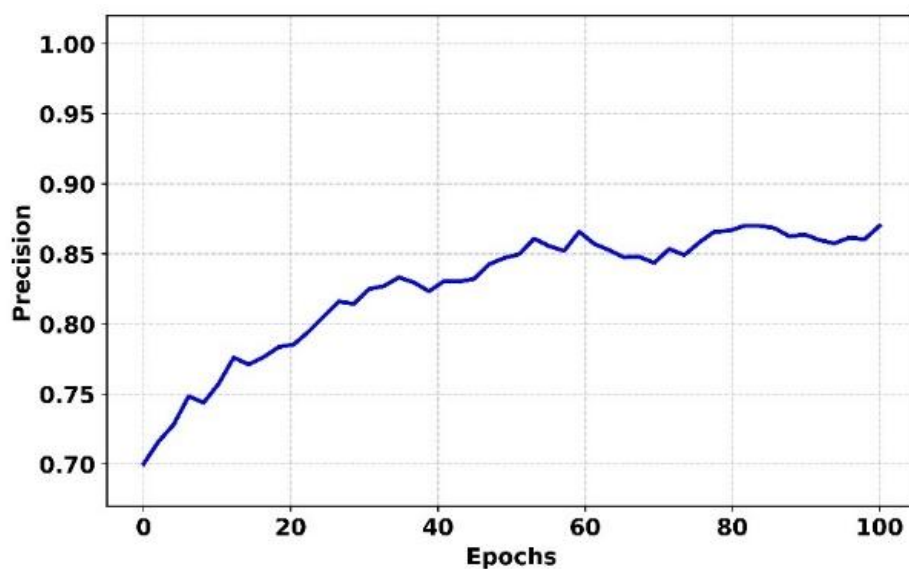


Figure 9. Precision of proposed SS-IBERT based BNR.

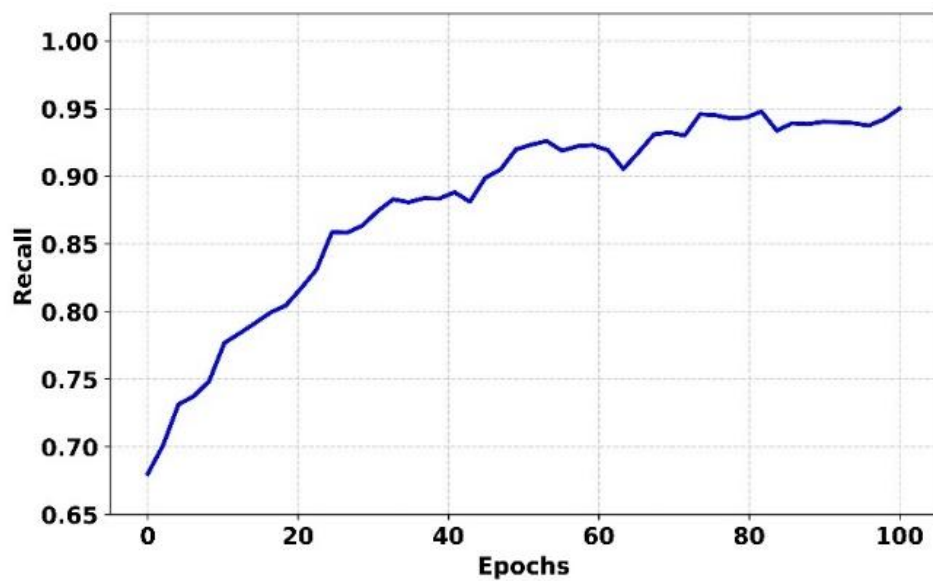


Figure 10. Recall of proposed SS-IBERT based BNR.

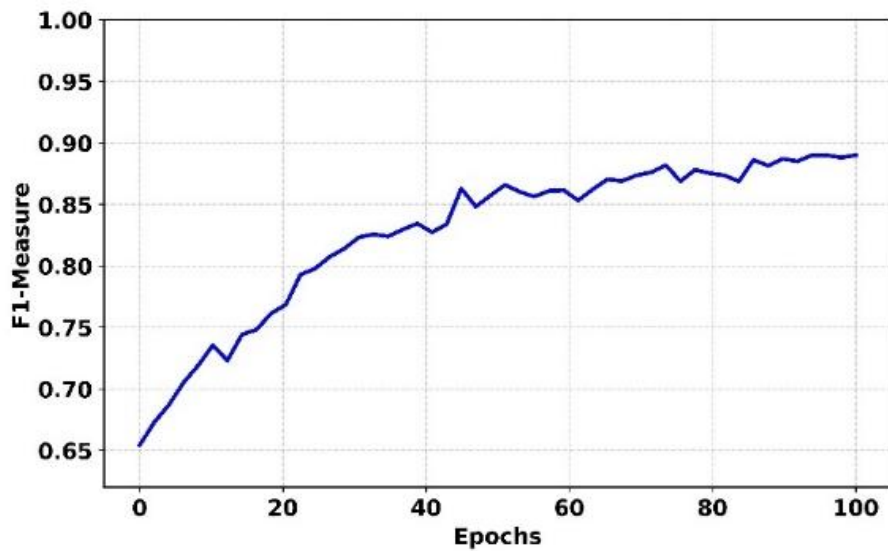


Figure 11. F1-Measure of proposed SS-IBERT based BNR.

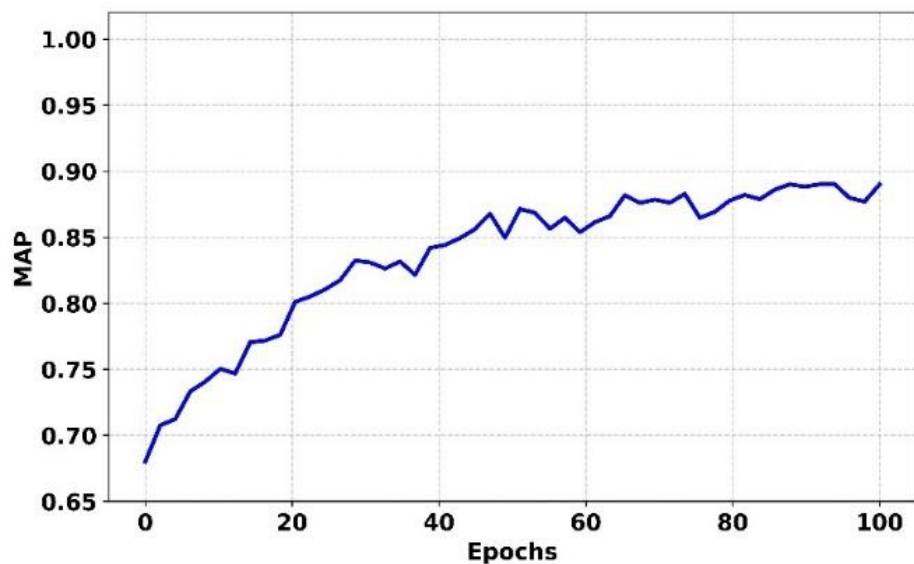


Figure 12. MAP of proposed SS-IBERT based BNR.

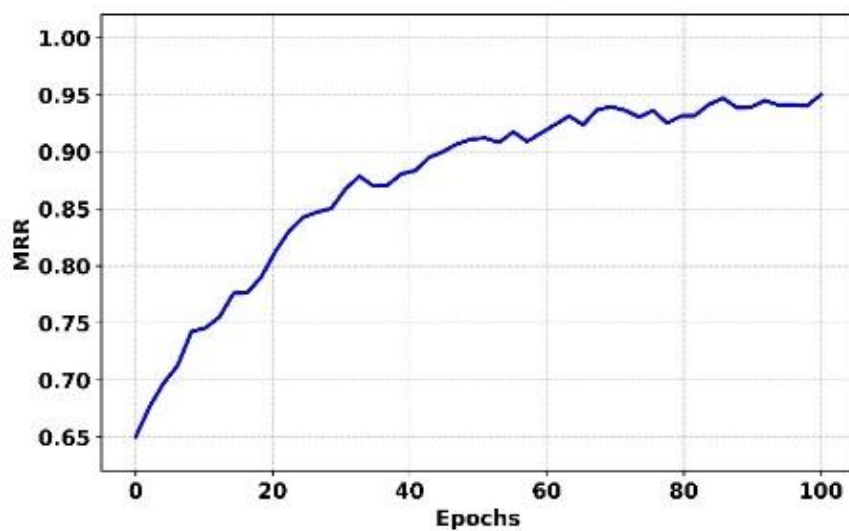


Figure 13. MRR of proposed SS-IBERT based BNR.

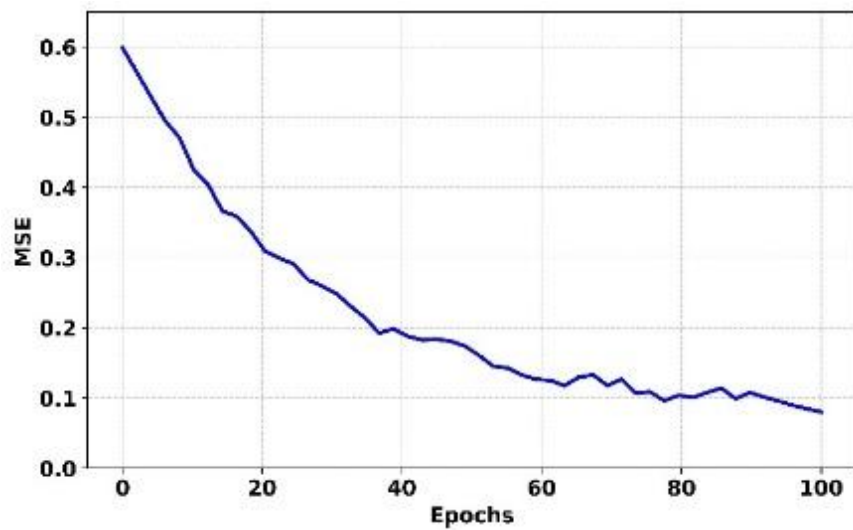


Figure 14. MSE of proposed SS-IBERT based BNR.

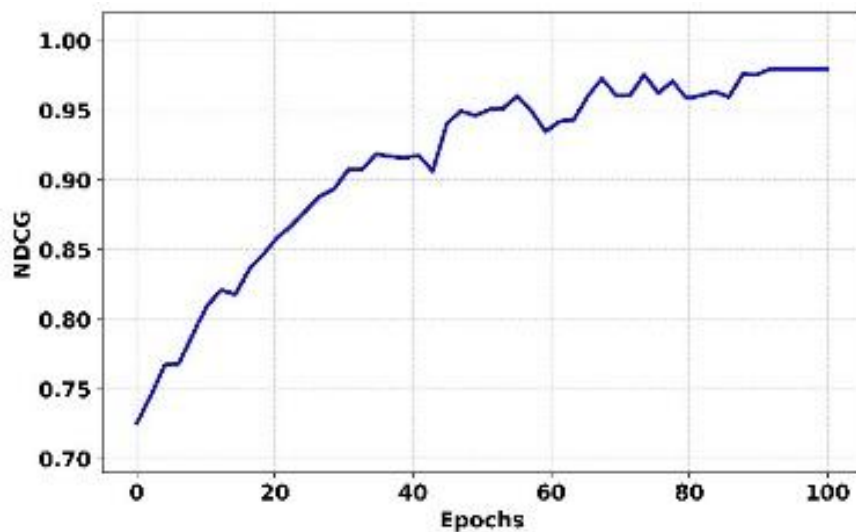


Figure 15. NDCG of proposed SS-IBERT based BNR.

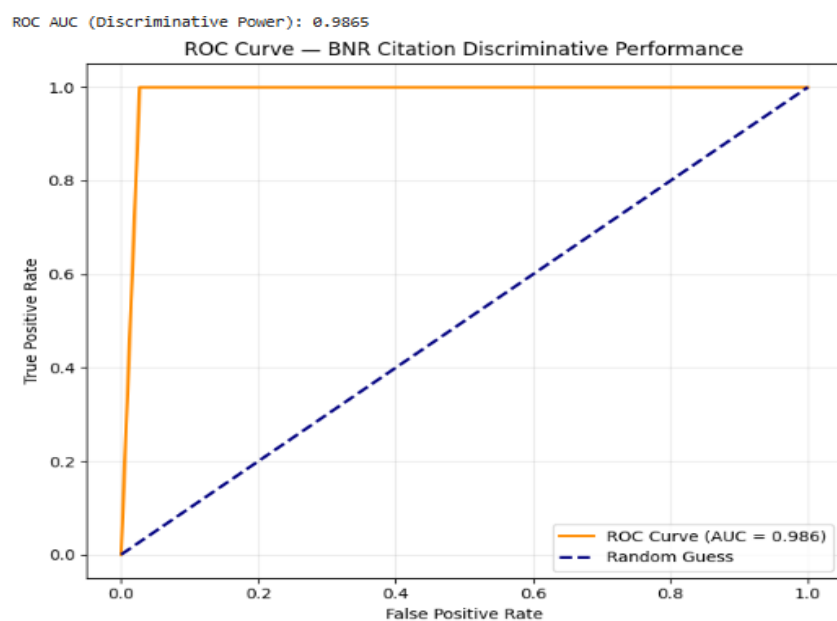


Figure 16. ROC-AUC Evaluation.

The MRR specifically defines the measure of first relevant citation that has to appear in the recommended list and as the proposed work uses Optimized ICL the MRR gets improved and also balances with the re-ranking with refinement of the top positions that results in quick and precise attainment of citation.

The MSE of proposed SS-IBERT based BNR, are illustrated in Figure 14, from which the results show that the use of novel SIL-BERT quantifies the deviations amid the ground-truth relevance and the predicted citation. The plot depicts that at an epoch of 20, the MSE is about 0.3, while it decreases and reaches a value of about 0.08 at an epoch of 100. From this it is concluded that the proposed work achieves a greater reduction in MSE with higher CR.

The NDCG of proposed SS-IBERT based BNR is represented in Figure 15 and it shows that NDCG increases in NDCG as the number of epoch's increases. At an epoch of 20, there is a rise to 0.86, while it further rises and then steadily reaches an NDCG value of about 0.98. The NDCG not only identifies the relevance, but also notes the ranking of the citations. [28, 44] The defined SIL-BERT based ranking in the proposed work, promotes meaningful and non-redundant citations, to align the recommendation and the manuscript, and thereby enhances the NDCG of the system.

For analyzing the discriminative ability of the proposed SIL-BERT model, an ROC evaluation is carried out and is shown in Figure 16, in which every single citation is assigned with a rank-based similarity score, also the system is tested in terms of its ability to distinguish relevant and irrelevant citations. Moreover, the ROC curve provides a trade-off amid the True Positive rate (TPR) and the False Positive Rate (FPR).

The higher AUC indicates that the defined model has the capability to differentiate TPR and FPR, while it can differentiate the various text data existing in the dataset.

5.6 Comparative Analysis of the Proposed and Existing Models

To demonstrate the significance of the proposed SS-IBERT based BNR for CR, a comparative analysis is carried out amid the conventional techniques like MSCN [22], GS [22], RRMF [22], CCF [23], Col [23], SB [23], SVM [23], TF-BERT [27], PAC [27], CF [27], HCR [27], BM25, SciBERT-NTX, SPECTER, and SPECTER2 [46] with respect to Accuracy, Precision, Recall, F1-Score, MAP, MSE and MRR, NDCG, Hit Rate (HR).

The Accuracy Comparison of proposed SS-IBERT based BNR are illustrated in Figure 17 with conventional techniques like TF-BERT, PAC, CF and HCR. From the chart, it is noted that the proposed work achieves an accuracy of about 99%, while the other techniques like TF-BERT, PAC, CF and HCR possess an accuracy of about 91.25%, 91%, 81.9% and 51.10%. [27] Hence, it is confirmed that the proposed SIL-BERT achieves higher accuracy than that of the conventional techniques. As SS-IBERT captures the fine-grained contextual intent effectively, the accuracy improves.

Figure 18 illustrates the Precision Comparison of proposed SS-IBERT based BNR with that of the existing CR techniques like COI, CB and SVM. It is identified from the plot that compared to conventional techniques; the proposed work provides better precision of about 0.87 than that of the conventional techniques like COI, CB and SVM that possessed a precision of 0.5, 0.79 and 0.77.

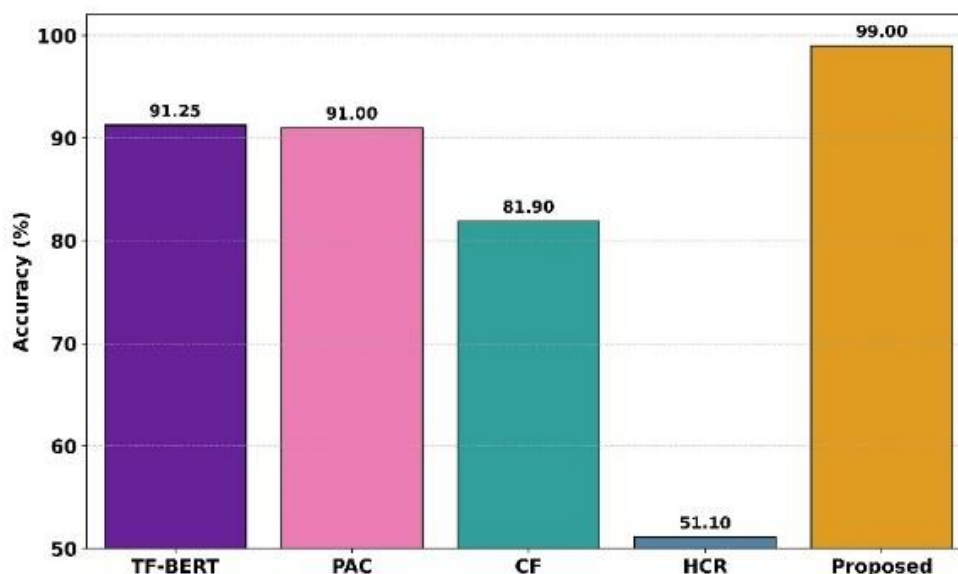


Figure 17. Accuracy Comparison of proposed SS-IBERT based BNR.

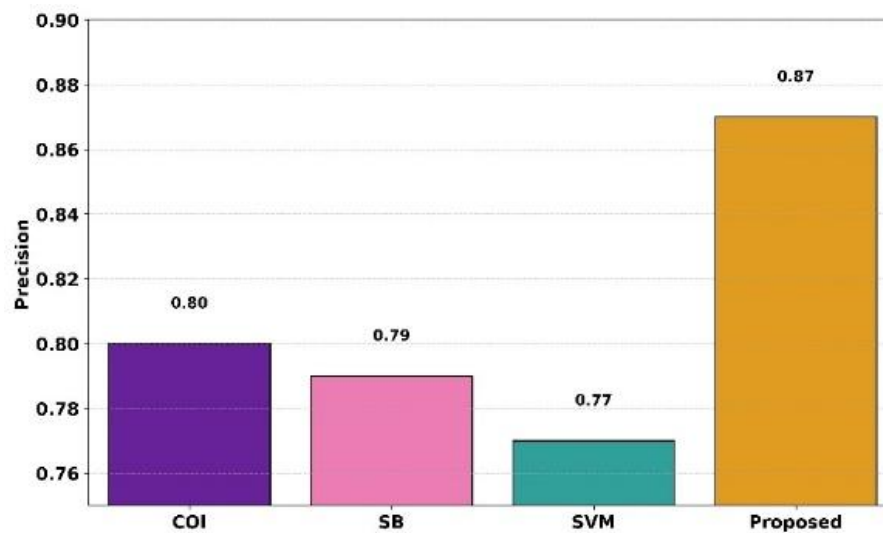


Figure 18. Precision Comparison of proposed SS-IBERT based BNR.

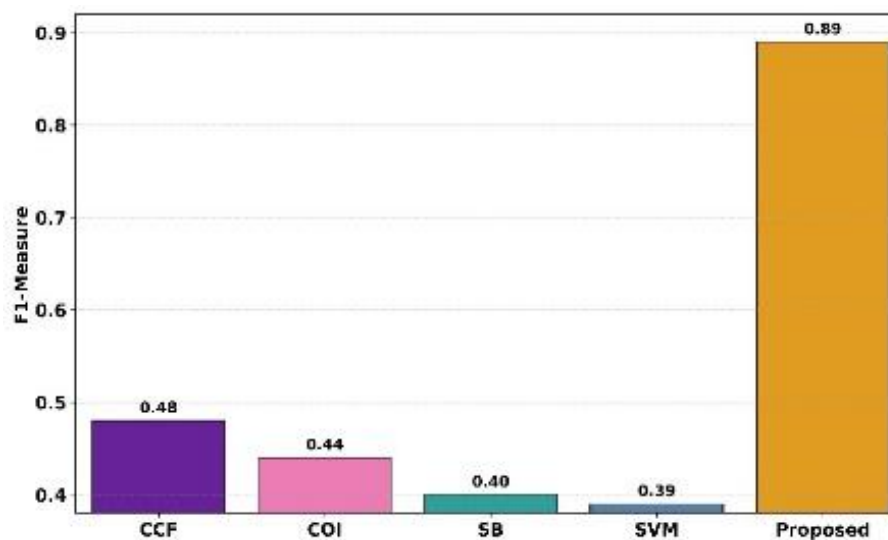


Figure 19. F1-Measure Comparison of proposed SS-IBERT based BNR.

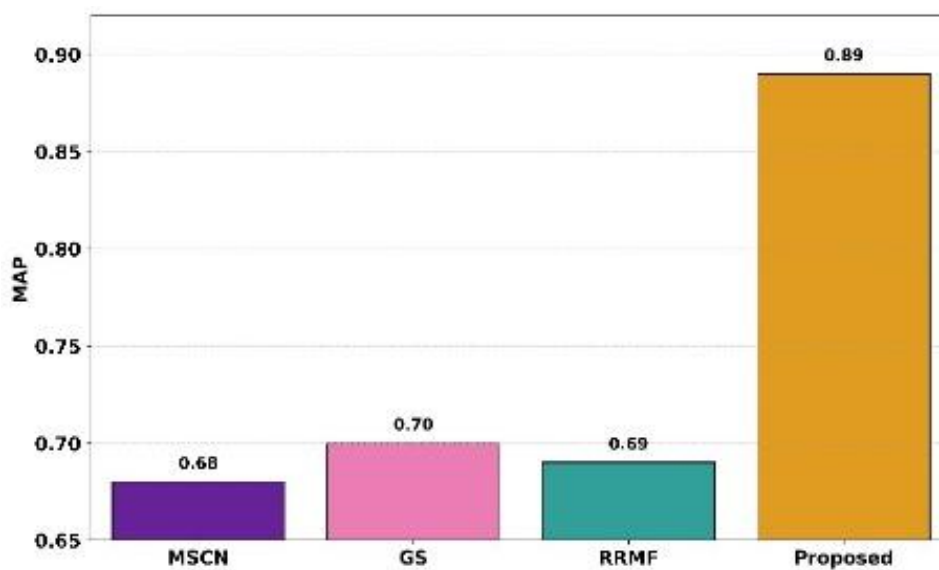


Figure 20. MAP Comparison of proposed SS-IBERT based BNR.

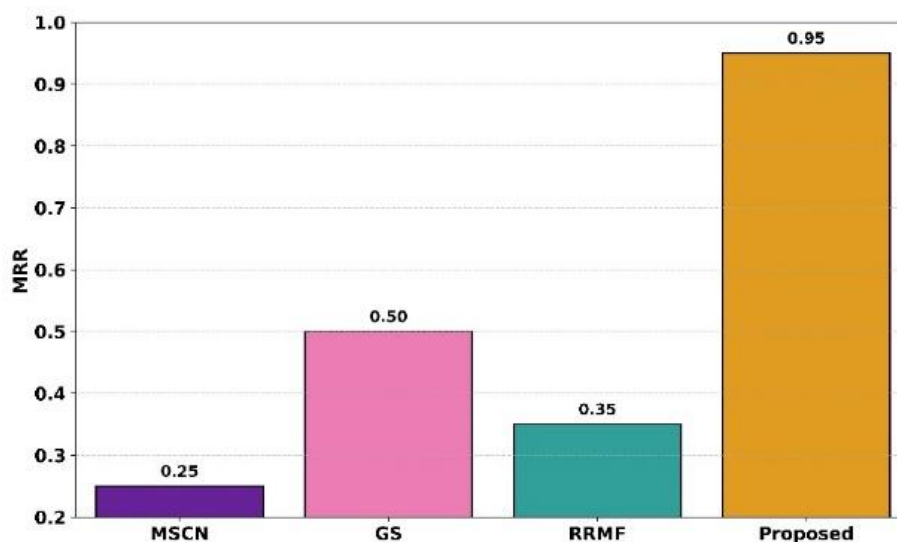


Figure 21. MRR Comparison of proposed SS-IBERT based BNR.

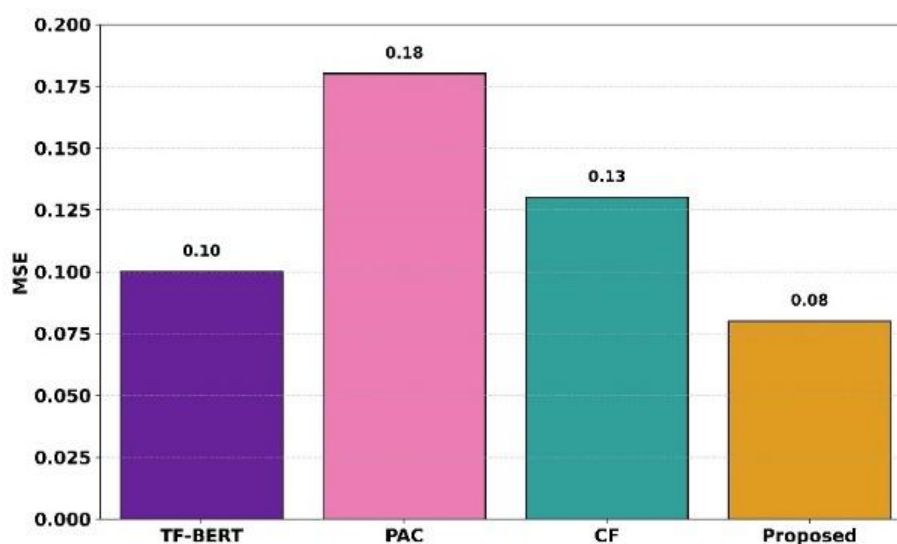


Figure 22. MSE Comparison of proposed SS-IBERT based BNR.

Therefore, it is clear that the defined work gives out better precision than that of the conventional techniques. The generation of intent focused contextual embedding, significantly decreases the issues with rhetorical noise and also the popularity bias, enabling better accuracy.

The F1-Measure Comparison of proposed SS-IBERT based BNR with that of the conventional techniques like CCF, COI, CB and SVM is depicted in Figure 19. The comparison results show that the conventional techniques like CCF, COI, CB and SVM possess an F1-Measure value of about 0.48, 0.44, 0.40 and 0.39, while the proposed SIL-BERT gives out an F1-score value of about 0.89. Hence, with this the results show that the proposed model has a higher F1-score than that of the conventional techniques. Since the defined BERT involves effective learning of intent discriminative contextual embedding, better aligns

citations with the relevant papers, and thus enhances the F1-Score.

Figure 20 depicts the MAP Comparison of proposed SS-IBERT based BNR and the existing CR techniques like MSCN, GS and RRMF. From the plot, it is concluded that the use of SIL-BERT in the defined CR, achieves a MAP of about 0.89, while the other techniques namely MSCN, GS and RRMF have MAP values of about 0.68, 0.70 and 0.49. This thereby depicts that the proposed method achieves better MAP than that of the conventional techniques. The learning, as well as the balanced embedding aggregation involved effectively decreases noise along with the section bias, while improving the MAP of the model.

The MRR Comparison of proposed SS-IBERT based BNR with that of the other techniques namely MSCN, GS and RRMF are illustrated in Figure 21, from which it is noted that the defined work achieves an MRR of approximately 0.95, while other conventional

techniques like MSCN, GS and RRMF provide MRR values of about 0.25, 0.50 and 0.35 respectively. Moreover, the results show that the proposed work has a higher MRR than other conventional techniques like MSCN, GS and RRMF. The inclusion of re-ranking with similarity retrieval enhances the ranking accuracy, which further improves the MRR.

Figure 22 illustrates the MSE Comparison of proposed SS-IBERT based BNR with that of the conventional techniques like TF-BERT, PAC and CF. the results show that the PAC exhibits higher MSE of about 0.18, while the techniques like CF and TF-BERT show an MSE of about 0.13 and 0.10. It is also noted that compared to these conventional techniques the proposed SIL-BERT achieves a lower MSE of about 0.08. Because randomized Anonymization is involved; the MSE at task level is lower than other techniques, while it maximizes the reconstruction error for author-based inference, which thereby balances the privacy and the utility.

The NDCG Comparison of proposed SS-IBERT based BNR and the prior CR techniques like MSCN, GS and RRMF are illustrated in Figure 23, from which the results show that the proposed SIL-BERT has better NDCG than that of the conventional techniques like MSCN, GS and RRMF. Because the self-tuned SIL-BERT produces higher intent aware embeddings, the semantic match enhances and hence the combined re-ranking process, effectively improves the NDCG.

The Comparative Analysis of proposed SS-IBERT based BNR with SOTA techniques like TF-IDF, Word2Vec, Doc2Vec and SPECTER in terms of Precision, Recall and F1@10 are illustrated in Figure 24. From the plot, the results show that the proposed dual cloud SS-IBERT based BNR performs better than all the specified conventional techniques. As the proposed technique provides balanced and relevant citations,

effectively captures the semantic intent as well as reduces the bias existing in the top-ranked recommendations.

The graphical analysis of the proposed system are illustrated in Figure 25 and compares the response time of computing amid the Single Cloud and the Dual Cloud of BNR architecture. From the analysis, it is concluded that the response time differs with 100 queries, while it is noted that the Dual Cloud system provides a faster response than that of the Single Cloud. The response-time difference is approximately 15-20 ms. Hence, from this it is concluded that the dual cloud provides higher responsiveness, while also reduces processing delays.

Figure 26 compares the performance of various information retrieval methods based on the Hit Rate (HR) metric at two different thresholds such as HR@10 and HR@20. The proposed method significantly outperforms all baseline models, achieving a near-perfect hit rate of approximately 0.967 at HR@10 and 0.975 at HR@20. In contrast, the baseline methods such as BM25, SciBERT-NTX, SPECTER, and SPECTER2 show considerably lower performance, with their HR@10 values ranging between roughly 0.55 and 0.65. Among the baselines, SciBERT-NTX appears to be the most competitive, while SPECTER exhibits the lowest accuracy. Across all methods, there is a consistent trend where performance improves when moving from HR@10 to HR@20, though the proposed model maintains a dominant lead in both categories, suggesting superior retrieval precision and recall.

Figure 27 presents the analysis of Average Response Time with respect to Single and Dual cloud. From the chart, it is noted that the single-cloud architecture gives a response time of about 13.2ms, while the Dual cloud has an average response time of approximately 4.1ms.

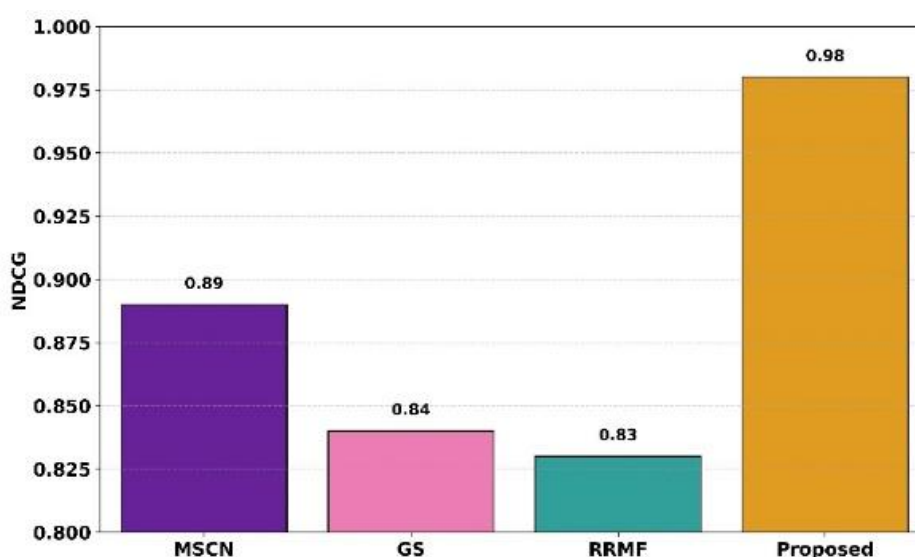


Figure 23. NDCG Comparison of proposed SS-IBERT based BNR.

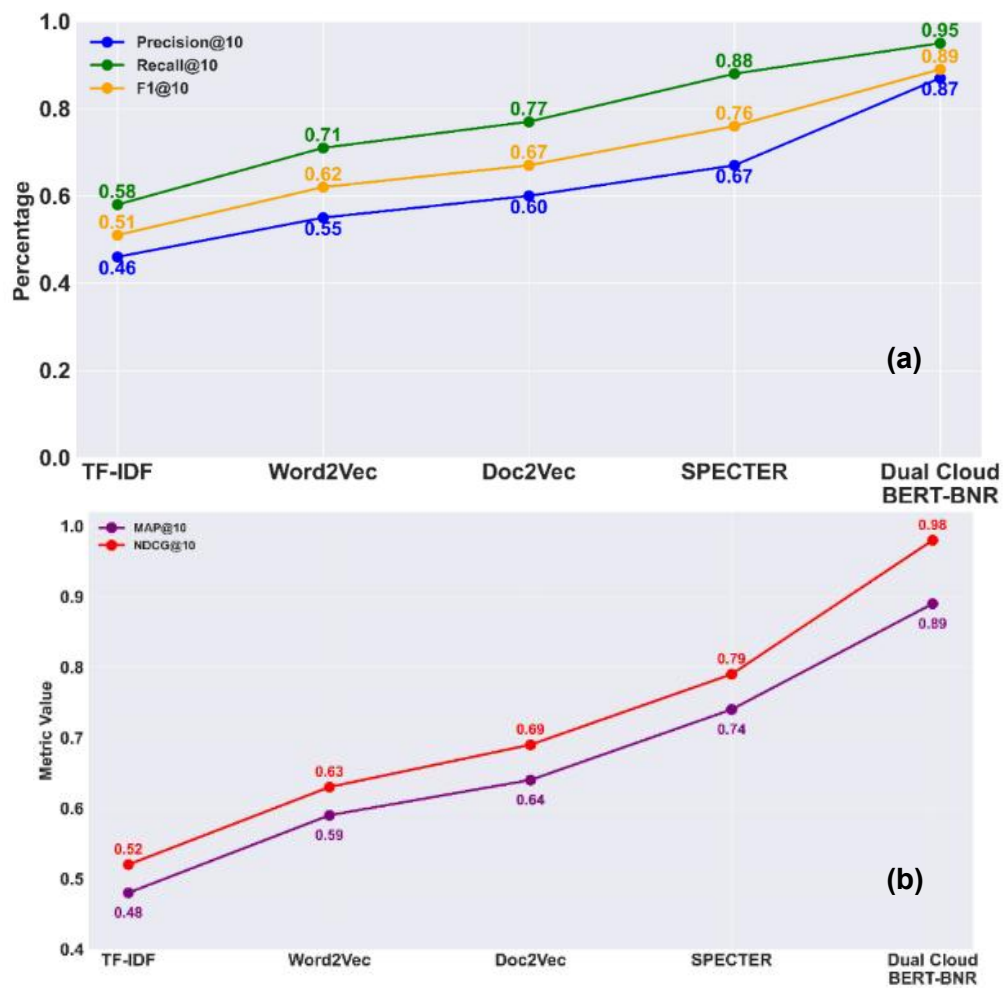


Figure 24. Comparative Analysis of proposed SS-IBERT based BNR with SOTA techniques in terms of (a) Precision, Recall and F1@10, (b) MAP@10 and NDCG@10.



Figure 25. Graphical Analysis.

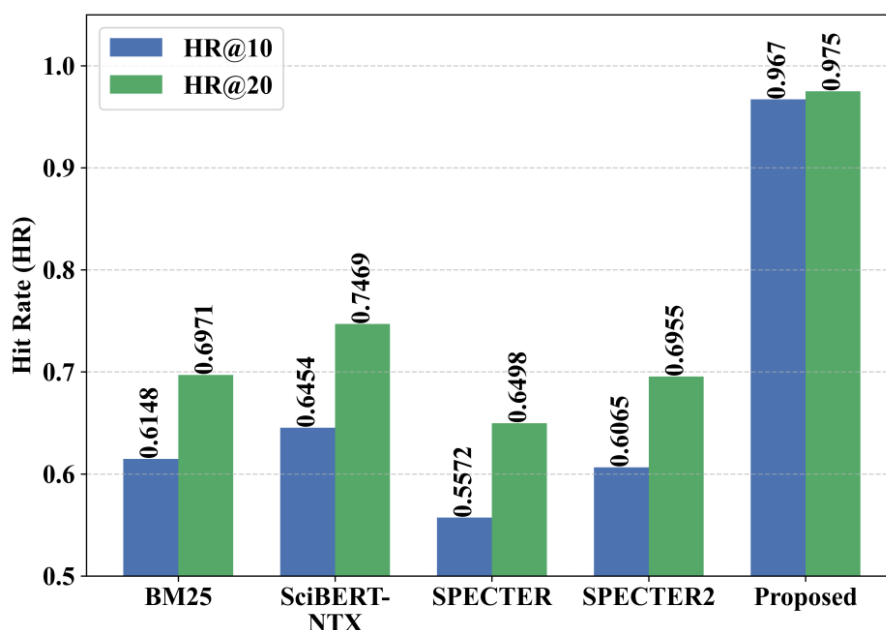


Figure 26. Comparison of Hit rate with existing models.

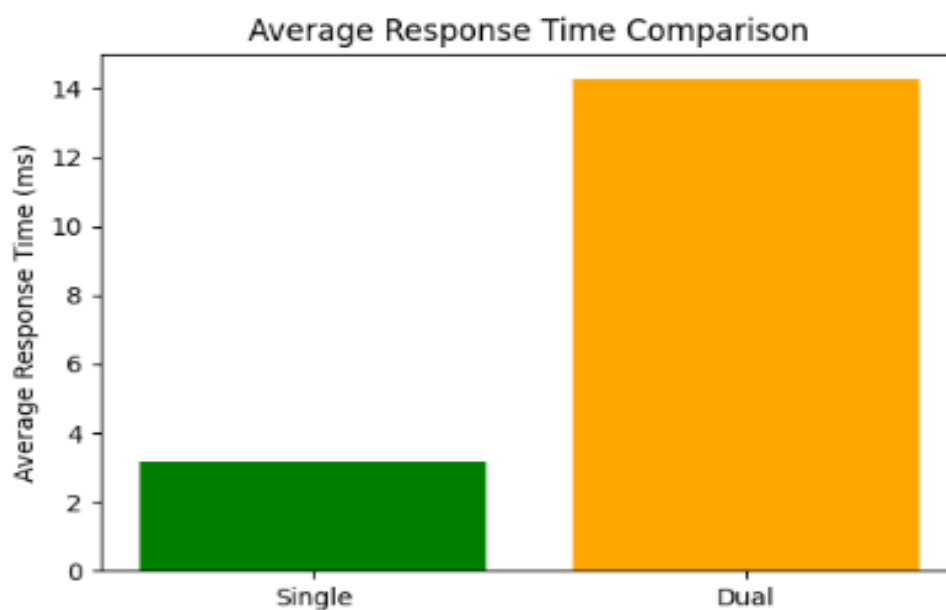


Figure 27. Graphical Analysis.

This confirms that the Dual Cloud BNR system significantly decreases the latency and also assists in optimizing the computation efficiency, while providing faster query processing, along with system responsiveness than that of the single cloud structure.

5.7 Ablation Study

To evaluate the contribution and effectiveness of each component in the proposed SIL-BERT framework, an ablation study is conducted. In this analysis, key modules are selectively removed or modified to evaluate their individual impact on overall performance. This helps in understanding the significance of each design element, particularly in improving semantic representation, reducing formal

bias, and ensuring privacy preservation in the citation recommendation process. The corresponding results are presented in Table 5.

The results of the ablation study clearly show that every element of the SIL-BERT framework has a positive impact on overall performance. When individual modules are removed, a gradual decline in performance is seen in all metrics. Specifically, the removal of the BERT encoder decreases the accuracy from 99.00% to 96.12%, which demonstrates its significant influence on the contextual representation. Likewise, when ICL, SA, and optimization modules are omitted, the accuracy decreases to 96.85, 97.21, and 97.68, respectively, confirming their incremental contribution to model learning and stability.

Table 5. Ablation study of the proposed SIL-BERT framework

Configuration	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
Without BERT Encoder (Contextual Embedding Backbone)	96.12	83.45	88.60	85.94
Without Intent Contrastive Learning (ICL)	96.85	84.12	90.05	86.98
Without Self-Attenuation (SA) Mechanism	97.21	84.96	91.22	87.97
Without Sturnus Optimization (SVE-OA)	97.68	85.40	92.10	88.62
Without Embedding Alignment Optimization	97.95	85.78	92.85	89.17
Without Randomized Masking Mechanism (RRM)	98.12	86.05	93.40	89.58
Without Style Adversarial Anonymization (SAA)	98.35	86.32	93.85	89.93
Without Dual Cloud Secure Recommendation Layer	98.72	86.75	94.35	90.39
Without Adversarial Robustness Handling	98.93	86.98	94.78	90.72
Full Proposed Model (SIL-BERT)	99.00	87.00	95.00	89.00

Table 6. Re-identification Success Rate and Utility Comparison

Method	Re-identification Success Rate (%)	NDCG@10
No anonymization	78	0.98
Gaussian noise ($\sigma = 0.5$)	34	0.82
Proposed R-SAA	11.5	0.95

Security-oriented features like SAA and adversarial robustness also contribute to preserving higher performance, as observed in the reduction of precision and recall when removed. Overall, the full SIL-BERT model achieves the best performance (99.00% accuracy, 87.00% precision, 95.00% recall, 89.00% F1-score), which demonstrates the effectiveness of the combination of all modules.

5.8 Privacy Validation – Re-identification Resistance

In order to test the privacy-preserving property of the proposed R-SAA mechanism, a re-identification attack was simulated in a realistic adversarial environment. The adversary, with access to publicly available metadata (e.g., author identities, publication years, and co-authorship networks), was exposed to 10,000 anonymized citation embeddings generated using R-SAA.

A trained Random Forest classifier (100 trees) was used to associate anonymized representations with original authors, using non-anonymized embeddings and metadata. The success rate of re-identification was calculated using Equation (13):

$$R_{id} = \frac{N_{correct}}{N_{total}} \quad (13)$$

As shown in Table 6, the absence of anonymization results in a high re-identification success rate of 78%, indicating significant privacy risk. The proposed R-SAA method substantially reduces this rate to 11.5%, demonstrating strong resistance to re-identification attacks. In comparison, Gaussian noise perturbation achieves a success rate of 34%, confirming the superior effectiveness of R-SAA. In terms of utility, the model achieves an NDCG@10 score of 0.98 without anonymization. With R-SAA, the score slightly decreases to 0.95 (approximately 3% reduction), whereas Gaussian noise reduces it to 0.82, indicating a considerable degradation in recommendation quality.

Overall, these results demonstrate that the proposed R-SAA framework effectively reduces re-identification risk while preserving high recommendation performance, achieving a balanced trade-off between privacy protection and utility.

6. Discussion

From the performance analysis and comparative analysis with the conventional techniques

like MSCN, GS, RRMF, CCF, Col, SB, SVM, TF-BERT, PAC, CF, HCR, TF-IDF, Word2Vec, Doc2Vec and SPECTER, with respect to Accuracy, Precision, Recall, F1-Score, MAP, MSE and MRR, NDCG, the results show that the proposed work attains an Accuracy of about 99%, Precision of about 0.87, Recall of about 0.95, F1-Score of about 0.89, MAP of about 0.89, MRR of about 0.95, MSE of about 0.08 and NDCG of about 0.98. [22, 23, 27] Moreover, the comparison between single and dual cloud shows better response with the proposed work, in terms of dual-cloud performance [17], improving the CR. Besides, the experimental setup is designed to evaluate the accuracy of the proposed SIL-BERT model on the citation recommendation task, formulated as a ranking problem where the system retrieves the most relevant citation for a given query text from the DBLP dataset. Top-N accuracy was used as the evaluation metric, indicating whether the ground-truth citation appears at rank 1. The DBLP dataset (6.6 million publications) is split into 80% training (5.28M), 10% validation (660K), and 10% testing (660K). From the test set, 10,000 citation queries are constructed, each paired with its true cited paper. For evaluation, each query is matched against a pool of 100 candidates (1 correct and 99 negative samples). Cosine similarity is used to rank candidates based on embedding similarity, and a prediction is correct when the true citation ranked first. Under this setup, the proposed model achieves 99% Top-1 accuracy (9,900/10,000 correct predictions) and 98.5% validation accuracy, showing strong generalization. All experiments were run with fixed random seeds to ensure reproducibility and consistent evaluation.

6.1 Optimization Objective, Variables, Stopping Criteria, and Computational Cost

$O(N \cdot d \cdot k)$ SVE-OA optimizes the embedding space by reducing the contrastive loss (Equation (4)). The objective is to enhance the alignment of intent-aware embeddings so that similar documents are brought closer together and dissimilar documents are separated. The variables that are tuned in the optimization process are the temperature τ in the contrastive loss and the exploration coefficient λ in the escape maneuvers (Equation (6)). These parameters affect the contrastive learning and flocking dynamics of the algorithm. The optimization process stops when either (i) a maximum number of iterations is reached, or (ii) the difference in the contrastive loss between two consecutive iterations falls below a small threshold, suggesting convergence. These conditions were selected empirically on the basis of validation results. The time complexity of SVE-OA is dominated by the computation of the contrastive loss and the update of nearest neighbors for the flocking strategy. The complexity per iteration is estimated as , where N is the number of documents, d is the embedding dimension, and k is the number of nearest neighbors used for local

interactions. This is linear in the number of documents, making it feasible for large bibliographic data like DBLP.

7. Conclusion and Future Scope

In this research work, an SIL-BERT-based dual-cloud BNR is proposed for effective and secure CR in this research work. Shilling attacks are addressed with the inclusion of SS-IBERT, which effectively decreases the algorithmic rhetoric issue, while the use of R-SAA reduces the challenges with S-NBV, thereby reducing the risk with re-identification. Therefore, with this, the defined SIL-BERT gives out bias-resilient modelling and also ensures embedding-level anonymization, while provides a secure CR system. To analyze the effectiveness of the proposed work, it is implemented in Python platform and the obtained results are compared with conventional CR systems, which clearly demonstrate the significance of the defined work. From the implementation results, the results show that the proposed work attains better Accuracy (99%), Precision (0.87), Recall (0.95), F1-Score (0.89), MAP (0.89), MRR (0.95), MSE (0.08) and NDCG (0.98). Hence, the use of self-tuned BERT with optimization enhances the semantic accuracy and the NDCG, while the novel anonymization suppresses the author-based stylistic influence. Moreover, the Dual Cloud based architecture enhances the security, stability and fairness by effectively reducing bias, thereby improves the citation recommendation. While the proposed framework demonstrates strong performance, the following limitations and practical considerations provide a balanced perspective on its application.

7.1 Limitations

The proposed Dual Cloud SIL-BERT framework demonstrates strong performance; however, several limitations must be acknowledged.

- First, the model is evaluated primarily on the DBLP dataset. As a result, its performance may degrade when applied to emerging or low-resource research domains with limited citation data, which remains a common challenge in citation recommendation systems.
- The incorporation of privacy-preserving mechanisms, including anonymization and encryption, introduces additional computational overhead. Experimental observations indicate an approximate 15–20% increase in processing time compared to non-secure baseline systems, highlighting a trade-off between privacy preservation and system efficiency.
- Then, the current implementation is limited to English-language manuscripts. The model may not generalize effectively to multilingual or

cross-lingual scenarios without additional training on diverse linguistic datasets.

- Finally, the Sturnus-based optimization strategy is tuned specifically for citation recommendation tasks. Its applicability to other domains, such as patent or legal document recommendation, requires further investigation and parameter adaptation.

7.2 Practical Deployment Considerations

In addition to technical limitations, several practical aspects must be considered when deploying the proposed system in real-world environments. The dual-cloud architecture may not align with the data governance policies of certain institutions that require strict control over data storage and processing within a single trusted infrastructure. Additionally, anonymization and encryption at the embedding level, although improving privacy, decrease interpretability, potentially affecting the system's capacity to explain the recommended citations. Furthermore, the requirement to upload unpublished manuscripts, even within a secured environment, may raise concerns in cases involving sensitive or proprietary research data, thereby affecting adoption in certain organizations.

7.3 Future Scope

The defined work can be extended with the inclusion of multi-lingual citation recommendation that aids in interdisciplinary research. Moreover, by combining the user-feedback signals along with the online learning can further enhance the accuracy of CR, also its adaptability.

References

- [1] Y. Zhao, Y. Wang, Y. Liu, X. Cheng, C.C. Aggarwal, T. Derr, Fairness and Diversity in Recommender Systems: A survey. *ACM Transactions on Intelligent Systems and Technology*, 16(1), (2025) 1–28. <https://doi.org/10.1145/3664928>
- [2] H. Djafer, O.E. Hamdane, H.E. Degha, Advancing Context-Aware Paper Citation Recommendation using BERT and CNN Enhanced Multi-Model Deep Learning. *Rawafid for Research and Studies Review*, (2024) 120–138. <https://asjp.cerist.dz/en/article/249503>
- [3] B.D.S.B. Contreras, L.A. Digiampietri, Advances in Scientific Article Recommendation Systems: A Systematic Literature Review. *Revista Brasileira de Computação Aplicada*, 17(3), (2025) 78–86. <http://dx.doi.org/10.5335/rbca.v17i3.16376>
- [4] D.B. Rajesh, A. Kumar, Collaborative Filtering models an Experimental and detailed Comparative Study. *Scientific Reports*, 15(1), (2025) 31667. <https://doi.org/10.1038/s41598-025-15096-4>
- [5] T. Kanwal, T. Amjad, Research Paper Recommendation System based On Multiple features from Citation Network. *Scientometrics*, 129(9), (2024) 5493–5531. <https://doi.org/10.1007/s11192-024-05109-w>
- [6] M.A. Abbas, S. Ajayi, M. Bilal, A. Oyegoke, M. Pasha, H.T. Ali, A Deep Learning Approach for Context-Aware Citation Recommendation using Rhetorical Zone Classification and Similarity to Overcome Cold-Start Problem. *Journal of Ambient Intelligence and Humanized Computing*, 15(1), (2024) 419–433. <https://doi.org/10.1007/s12652-022-03899-6>
- [7] S.V. Oprea, A. Bâra, T. Ghinea, Recommender Systems: Emerging Trends from Four Decades of Research using Bibliometric Analysis and Transformer-Based Models. *Electronics*, 15(4), (2026) 763. <https://doi.org/10.3390/electronics15040763>
- [8] Z. Ali, G. Qi, I. Ullah, A.A. Mohammed, P. Kefalas, K. Muhammad, GLAMOR: Graph-based Language Model embedding for Citation recommendation. In *Proceedings of the 18th ACM Conference on Recommender Systems*, ACM, (2024) 929–933. <https://doi.org/10.1145/3640457.3688171>
- [9] W.N. Chen, P. Grzybowski, Security and privacy Challenges in AI-Powered Library Recommender Systems: A Systematic Literature Review. *Proceedings of the Association for Information Science and Technology*, 62(1), (2025) 1386–1389. <https://doi.org/10.1002/pra2.1412>
- [10] B. Mallik, S. Yasar, A. Shafkat, M.M. Rahman, S. Sharmin, (2025) Privacy and security threats of recommendation systems: Awareness and perspective of users of Bangladesh. In *2025 2nd International Conference on Next-Generation Computing, IoT and Machine Learning (NCIM)*, IEEE, Gazipur, Bangladesh. <https://doi.org/10.1109/NCIM65934.2025.11160260>
- [11] Q. Zhang, P. Yang, J. Yu, H. Wang, X. He, S.M. Yiu, H. Yin, (2025) A survey on point-of-interest Recommendation: Models, Architectures, and Security. *IEEE Transactions on Knowledge and Data Engineering*, 37(6), (2025) 3153–3172. <https://doi.org/10.1109/TKDE.2025.3551292>
- [12] D. Zheng, M.H. Alkawaz, M.G.M. Johar, Privacy Protection and Data Security in Intelligent Recommendation Systems. *Neural Computing and Applications*, 37(34), (2025) 28431–28448. <https://doi.org/10.1007/s00521-025-11189-3>
- [13] A. Raju, C. Raju, Advancing AI-driven customer service with NLP: A novel BERT-based model for automated responses. *Journal of Emerging Technologies and Innovative Research (JETIR)*, 12(1), (2025) h89–h123.

- [14] H. Lim, Q. Li, S. Yang, J. Kim, A BERT-based Multi-Embedding Fusion Method using Review Text for recommendation. *Expert Systems*, 42(5), (2025) e70041. <https://doi.org/10.1111/exsy.70041>
- [15] K.A. Alam, W. Ali, S. Aziz, M. Haroon, M.H. Khan, Z. Ahmad, N. Abbas, (2025) CodeComClassify: Automating Code Comments Classification using BERT-based Language Models. In 2025 IEEE/ACM International Workshop on Natural Language-Based Software Engineering (NLBSE), IEEE, Ottawa, ON, Canada. <https://doi.org/10.1109/NLBSE66842.2025.00017>
- [16] N.R. Gadicharla, Secure Semantic Databases: Integrating Knowledge Management and Encrypted Machine Learning for Intelligent Data Systems. *Journal of Advances in Computational Intelligence Theory (JACIT)*, 8(1), (2026) 61-77. <https://doi.org/10.5281/zenodo.17483458>
- [17] R.K. Vankayalapati, D. Valik, V.K.A.T. Ganti, Zero-trust Security Models for Cloud Data Analytics: Enhancing Privacy in Distributed Systems. *Journal of Artificial Intelligence & Cloud Computing*, 4(1), (2025) 1-8. [https://doi.org/10.47363/JAICC/2025\(4\)415](https://doi.org/10.47363/JAICC/2025(4)415)
- [18] N. Borovits, G. Bardelloni, H. Hashemi, M. Tulsiani, D.A. Tamburri, W.J. van den Heuvel, Addressing Data Scarcity with Synthetic Data: A Secure and GDPR-Compliant Cloud-Based platform. *ACM Transactions on Software Engineering and Methodology*, 35(3), (2025) 1-50. <https://doi.org/10.1145/3702317>
- [19] S. Ma, C. Zhang, H. Zhang, Z. Gao, Citation Recommendation based on Argumentative Zoning of User Queries. *Journal of Informetrics*, 19(1), (2025) 101607. <https://doi.org/10.1016/j.joi.2024.101607>
- [20] O. Nacar, A. Koubaa, Enhancing Semantic Similarity Understanding in Arabic NLP with Nested Embedding Learning. In: Koubaa, A., Ammar, A., Ghouti, L., Boulila, W., Benjdira, B. (eds) *Generative AI and Large Language Models: Opportunities, Challenges, and Applications*. *Studies in Computational Intelligence*, 1214, (2025) 179–216. https://doi.org/10.1007/978-3-031-90573-5_6
- [21] J.W. Sanchez-Obando, N.D. Duque-Méndez, A. Echeverri-Rubio, SRS: A Recommendation System for Scientific Publication. *Array*, 30, (2026) 100802. <https://doi.org/10.1016/j.array.2026.100802>
- [22] K. Velkumar, B. Chokkalingam, Improving Citation Recommendation Accuracy using an SVM-based model. *Journal of Information Systems Engineering and Management*, 10(29s), (2025) 1–12. <https://doi.org/10.52783/jisem.v10i29s.4602>
- [23] X. Gao, Y. Liu, Z. Ma, X. Qiao, H. Li, C. Xu, K. Zhang, J. Cui, RSA-CR: Resisting shilling attacks in Citation Recommendation via Dumbbell Inductive Learning. In *Proceedings of the AAAI Conference on Artificial Intelligence*, 40(42), (2026) 35446–35454. <https://doi.org/10.1609/aaai.v40i42.40854>
- [24] U.R. Saxena, Y. Khatri, R. Kadel, S. Shailendra, A service recommendation model in cloud environment based on trusted graph-based collaborative filtering recommender system. *Network*, 5(3), (2025) 30. <https://doi.org/10.3390/network5030030>
- [25] M.M. Van, T.T. Tran, BeLightRec: A Lightweight Recommender System Enhanced with BERT. In: Thai-Nghe, N., Do, TN., Benferhat, S. (eds) *Intelligent Systems and Data Science. ISDS 2024. Communications in Computer and Information Science*, Springer, Singapore, 2191, (2024) 30–43. https://doi.org/10.1007/978-981-97-9616-8_3
- [26] A.N. Hasoon, S.K. Abdulateef, R.S. Abdulameer, M.L. Shuwandy, An Intelligent Hybrid AI course Recommendation Framework Integrating BERT Embeddings and Random Forest Classification. *Computers*, 14(9), (2025) 353. <https://doi.org/10.3390/computers14090353>
- [27] J. Wang, N. Jia, L. Wu, H. Zhao, L. Wu, Align Sequential Collaborative Signals and Text Semantics for Citation Recommendation: A hybrid Perspective. *ACM Transactions on Knowledge Discovery from Data*, 20(3), (2026) 1–24. <https://doi.org/10.1145/3797954>
- [28] P. Kefalas, Z. Ali, Y. Manolopoulos, A critical Review on Citation Recommendation Systems through a Graph-Centric Approach. *SN Computer Science*, 7(3), (2026) 267. <https://doi.org/10.1007/s42979-026-04848-2>
- [29] B.A. Ojokoh, F.O. Isinkaye, M. Zhang, J.J. Tom, A.J. Gabriel, O. Afolabi, B. Afolabi, Privacy and Security in Recommenders: An analytical review. *Artificial Intelligence Review*, 58(11), (2025) 351. <https://doi.org/10.1007/s10462-025-11333-4>
- [30] G. De Toni, E. Purificato, E. Gomez, B. Lepri, A. Passerini, C. Consonni, You don't bring me Flowers: Mitigating unwanted Recommendations through Conformal Risk Control. In *Proceedings of the 19th ACM Conference on Recommender Systems (RecSys 2025)*, ACM, (2025) 492-502. <https://doi.org/10.1145/3705328.3748054>
- [31] M.H. Al Sarori, M.M.A.E. Sufyan, M. Al Asaly, G.A.A. Al Maamari, BERT and beyond: A Comprehensive Survey of Natural Language Processing Techniques for Information Retrieval. *Journal of Intelligent Communication*, 4(2), (2025) 93–114. <https://doi.org/10.54963/jic.v4i2.1706>
- [32] N.M. Gardazi, A. Daud, M.K. Malik, A. Bukhari, T. Alsahfi, B. Alshemaimri, BERT Applications in

- Natural Language Processing: A review. Artificial Intelligence Review, 58(6), (2025) 1–49. <https://doi.org/10.1007/s10462-025-11196-5>
- [33] T. Chen, S. Kornblith, M. Norouzi, G. Hinton, A simple framework for contrastive learning of visual representations. In Proceedings of the 37th International Conference on Machine Learning (ICML 2020), PMLR, 119, (2020) 1597–1607. <https://proceedings.mlr.press/v119/chen20j.html>
- [34] Y. Zhang, Y. Fan, T. Sheng, A. Wang, Temporal-Aware and Intent Contrastive Learning for Sequential Recommendation. Symmetry, 17(10), (2025) 1634. <https://doi.org/10.3390/sym17101634>
- [35] Y. Liu, Y. Fan, J. Ma, Sturnus vulgaris escape Algorithm and its Application to Mechanical Design. Scientific Reports, 15(1), (2025) 6552. <https://doi.org/10.1038/s41598-025-91270-y>
- [36] S. Abakarim, S. Qassimi, S. Rakrak, Emotion and Sentiment Enriched Decision Transformer for Personalized Recommendations. Scientific Reports, 15(1), (2025) 21020. <https://doi.org/10.1038/s41598-025-06386-y>
- [37] C. Zhang, Z. Hu, X. Xu, Y. Liu, B. Wang, J. Shen, T. Li, Y. Huang, B. Cai, W. Wang, DMRP: Privacy-Preserving Deep Learning Model with Dynamic Masking and Random Permutation. Journal of Information Security and Applications, 89, (2025) 103987. <https://doi.org/10.1016/j.jisa.2025.103987>
- [38] S. Liang, J. Ma, Q. Zhao, T. Chen, X. Lu, S. Ren, C. Zhao, L. Fu, H. Ding, A Sequential Recommendation Method using Contrastive Learning and Wasserstein Self-Attention Mechanism. PeerJ Computer Science, 11, (2025) e2749. <https://doi.org/10.7717/peerj-cs.2749>
- [39] W. Chen, M. Yuan, Z. Zhang, R. Xie, F. Zhuang, D. Wang, R. Liu, FairDgcl: Fairness-Aware Recommendation with Dynamic Graph Contrastive Learning. IEEE Transactions on Knowledge and Data Engineering, 37(5), (2025) 2301–2315. <https://doi.org/10.1109/TKDE.2025.3349821>
- [40] N.C. Nelakuditi, N.K. Namburi, J. Sayyad, D.V. Rudraraju, R. Govindan, P.V. Rao, Secure file operations: Using Advanced Encryption Standard for strong data protection. International Journal of Safety & Security Engineering, 14(3), (2024) 1007–1014. <https://doi.org/10.18280/ijssse.140330>
- [41] P. Selvi, S. Sakthivel, A hybrid ECC-AES Encryption Framework for Secure and Efficient Cloud-based Data Protection. Scientific Reports, 15(1), (2025) 30867. <https://doi.org/10.1038/s41598-025-01315-5>
- [42] R. Onozuka, A. Shimpei, H. Okumura, T. Ogawa, Revealing Hidden Interdisciplinary Collaborations: A Novel Approach using Citation Analysis and Semantic Embeddings. Journal of Emerging Technologies and Innovative Research (JETIR), 12(1), (2025). <https://aisel.aisnet.org/pacis2025/aiandml/aiandml/15/>
- [43] M. Ley, The DBLP Computer Science Bibliography: Evolution, Research Issues, Perspectives. In: Laender, A.H.F., Oliveira, A.L. (eds) String Processing and Information Retrieval. SPIRE 2002. Lecture Notes in Computer Science, Springer, Berlin, Heidelberg, 2476 (2002). https://doi.org/10.1007/3-540-45735-6_1
- [44] S.Y. Lim, N. Hashim, L.L. Thanh, Research Article Recommender Systems: A Comprehensive Review of Models, Approaches and Evaluation Metrics. Journal of Informatics and Web Engineering, 4(3), (2025) 166–190. <https://doi.org/10.33093/jiwe.2025.4.3.10>
- [45] B. Manzanares-Salor, D. Sánchez, P. Lison, Evaluating the Disclosure Risk of Anonymized Documents via a Machine Learning-based re-Identification Attack. Data Mining and Knowledge Discovery, 38(6), (2024) 4040–4075. <https://doi.org/10.1007/s10618-024-01066-3>
- [46] M.T.R. Ratul, Z. Chen, K. Fu, T. Ji, L. Zhang, (2026) MasterSet: A Large-Scale Benchmark for Must-Cite Citation Recommendation in the AI/ML Literature. arXiv preprint arXiv:2604. <https://doi.org/10.48550/arXiv.2604.17680>

Authors Contribution Statement

Both the Authors equally contributed, Read and Approved the Final Version of the Manuscript.

Funding

The authors declare that no funds, grants or any other support were received during the preparation of this manuscript.

Competing Interests

The authors declare that there are no conflicts of interest regarding the publication of this manuscript.

Data Availability

The datasets generated during and/or analysed during the current study are available from the corresponding author on reasonable request.

Has this article screened for similarity?

Yes

About the License

© The Author(s) 2026. The text of this article is open access and licensed under a Creative Commons Attribution 4.0 International License.