



iJRASET

International Journal For Research in
Applied Science and Engineering Technology



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Volume: 14 **Issue:** VIII **Month of publication:** August 2026

DOI: <https://doi.org/10.22214/ijraset.2026.83151>

www.ijraset.com

Call: ☎ 08813907089

E-mail ID: ijraset@gmail.com

A Multimodal Deep Learning Framework for Scientific Insight Generation Using Transformer Based NLP and VIT

Ramansh, Sulaksh Malan, Apoorv Verma, Aman Kumar

Dept. of Computer Science & Engg. Lovely Professional University Punjab, India

Abstract: Research in the sciences continues to produce data in various forms, in forms such as dense technical prose, experimental figures, and structured measurement tables, but the majority of AI-assisted analysis tools can only process one of these forms simultaneously. This fragmentation constrained the richness of insight that may be provided by automated systems as, in many cases, scientific discoveries of significance are made by patterns that cut across modalities. Our system has a multimodal framework which directly encodes domain specific scientific language with SciBERT, matches visual content with textual context with CLIP (ViT-B/32), and purpose-built neural encoder of structured numerical data. Instead of using such representations with a basic concatenation step, we propose a Cross-Modal Attention mechanism and a Gated Fusion layer which dynamically scales the weighting of each modality based on how relevant it is to the query under consideration.

Index Terms- Multimodal Learning, Scientific Insight Generation, Cross-Modal Attention, Gated Fusion, SciBERT, CLIP, FAISS, Explainable AI, SHAP, ArXiv Dataset, Natural Language Processing, Computer Vision, Transformer Models

I. INTRODUCTION

Scientific research produces knowledge in varied forms at the same time. A scientist involved in the research of protein folding needs to read and digest written abstracts in hundreds of articles, analyse microscopy pictures and structural schemes and examinations refine-grained products of laboratory tests. Although scientific discovery can be inherently multimodal, current AI-assisted analysis tools have applied the modalities independently of one another. An analyst requires external software and computer vision software as well as statistical software to combine insights manually using current tools, which is time-consuming, error-prone, and it demands a great deal of domain expertise [1]. The amount of scientific literature that is being generated around the world has never been this high. The number of peer-reviewed research papers published on all fields of science is more than two million annually. It is impossible to survey this literature fully by a single researcher or a single research group, especially in the case of need to scope across a variety of textual descriptions, visual experimental findings, and coordinate data retrieval on this topic. This flood of information is one of the most urgent practical issues of the contemporary scientific practice [4].

A. Shortcoming of Unimodal Approaches

These limitations of unimodal approaches are mentioned in the literature. The visual evidence found in scientific figures cannot be encoded in text only systems like standard variant of the BERT version [13]. Image only systems do not have the contextualizing coverage that is presented with related textual descriptions. Those systems that work purely with tabular experimental data lack the context of the research findings in general. Above all, currently available multimodal systems serving a broad range of activities like social media or e-commerce do not have the vocabulary and semantic knowledge needed in the case of a science text, where strict terminology is utilized with meaning substantial meaning that marks out very similar concepts [4].

Recent research has tried to make multimodal integrations in three major directions. Early Fusion uses a concatenation method at the input of all modalities but has issues of information loss and interference between modalities. The models are divided among late fusion trains and their output is synthesized, without taking into account the trans-modal relationships, although the modality is still relevant. Attention based Fusion trains the weights of interaction but progresses at expensive computationally trying to improve. None of these methods take the advantage of domain specific encoders or embed knowledge base to gain interpretability [6]

B. The Explainability Problem

Another problem is that modern AI systems are opaque. Scientific investigations require reproducibility and verifiability. In instances where by an AI system categorizes a paper as either part of the computational biology or a specific hypothesis, the researchers must have a clear-cut explanation on how the reasoning occurred. Black-box AI despite its performance is essentially incompatible with the epistemological requirements of scientific practice, as all propositions from the science have to be traceable to evidence and all its findings verifiable in isolation [6]. Techniques like SHAP [7] and LIME [8] have been created to deliver post-hoc interpretability, but they are not integrated into multimodal scientific models as much.

C. Participation of Retrieval-Augmented Generation

With the advent of large-scale dense retrieval systems, there are now new possibilities of basing AI output on verified knowledge. The FAISS library [9] allows performing approximate nearest neighbour search over millions of embeddings of documents efficiently, thus allowing semantic searches to be performed in real-time, which are computationally feasible. Lewis et al. [14] showed that dense retrieval combined with language generation, namely the Retrieval-Augmented Generation (RAG) paradigm is significantly more effective when performing knowledge intensive tasks as it bases responses on what is retrieved as opposed to when a purely generative model is used that entirely relies on the memory of the parametric model. This fact directly inspires the inclusion of FAISS-based retrieval in FLAN-T5 willpower to our framework, where generated academic knowledge is framed in the context of already existing literature of 1.7 million ArXiv papers [11].

D. Novel Contributions

The paper presents the AI-Driven Multimodal Framework of Scientific Insight Generation as an attempt to solve these limitations in four major contributions:

- 1) **Unified Multimodal Encoding:** It is a pipeline that transforms scientific text, images, and tabular data to similar embedding representations created by pretraining a domain-specific model. SciBERT [4] reads research text with scientific vocabulary in mind, CLIP [5] encodes images in an image-text semantic space of direct cross-modal comparison, and a MLP encoder attention encoder is a system that takes a tabular experimental data, which is projected into a shared 512-dimensional space. **Intelligent Fusion:** Cross-modal Attention mechanism (8 heads) with a dynamically-weighted (contextually based) Gated Fusion cover quasi-diversity of modalities in contribution. In contrast to mere concatenation or averaging, the technique achieves use of the modalities giving most useful signals when presented with a specific input and changes weighting policy according to content and quality of available data.
- 2) **Insight Generation:** Generation of structured scientific insights such as summaries, research gap identification and hypothesis proposals via the integration of FLAN-T5 [10]. This generation power is supported more by FAISS-based retrieval [9] on 1.7 million arXiv papers [11] with the provision of grounded references that contextualize generated insights in the already existing literature.
- 3) **Complete explainability:** A combination of three explanation mechanisms described below, namely SHAP [7] to find the relative attribution of tabular features using cooperative game theory, [LIME8] to find the attribution of text tokens using linear approximation.

The whole configuration is supported by a Flask REST API backend, allowing it to use inexpensive commodity hardware.

E. Paper Organization

The rest of this paper will be structured in the following way. Section II reviews the related work in multimodal learning, domain-specific language models, domain understanding vision-language alignment, efficient retrieval and explainable AI. Section III describes the entire methodology which consists of system architecture, encoder design, fusion mechanisms, downstream tasks, and the explainability framework. Section IV provides results and discussion. Section V contains conclusion and future work directions.

II. LITERATURE REVIEW

A. Multimodal Learning in Scientific Domains

Processing heterogeneous scientific data through a single unified system has been a long-standing challenge. Ngiam et al. [1] provided early evidence that jointly trained cross modal models outperform independently trained unimodal counterparts, motivating the use of shared representation spaces for mixed-modality input.

Subsequent work in the scientific domain explored figure captioning [2] and biomedical visual text alignment [3], yet both systems were purpose-built for a single pairing of modalities rather than arbitrary combinations of text, images, and experimental tables. A key difficulty distinguishing scientific data from general web content is that scientific figures encode quantitative information trends, comparisons, error margins that cannot be recovered from captions alone. Equally, a caption without its figure lacks the visual evidence that supports the claim. Any system targeting scientific understanding must therefore model both within-modality semantics and cross-modality dependencies [12].

B. Domain-Specific Language Models

Standard pre-training corpora used for models such as BERT [13] consist primarily of encyclopaedia articles and digitised books, which share little vocabulary or stylistic structure with peer-reviewed research. Scientific prose is dense with multi-word technical terms, nested clauses, and implicit assumptions that require domain grounding to interpret correctly. To address this, Beltagy et al. [4] released SciBERT, built on the BERT architecture but trained from scratch on approximately 1.14 million full papers drawn from ArXiv and Semantic Scholar. A custom vocabulary constructed from that corpus keeps domain specific compound terms intact during tokenisation, which reduces sequence fragmentation and preserves meaning in downstream tasks such as relation extraction and multi-label document classification. In the proposed framework, a Sentence Transformer variant is used for throughput-sensitive batch encoding, while SciBERT serves as the high-fidelity option for single-document analysis; both apply mean pooling across all token positions rather than relying on the [CLS] representation, since relevant content in scientific abstracts is not confined to the opening tokens. Raffel et al. [10] introduced T5, a sequence-to-sequence model that casts every NLP problem as conditional text generation. The FLAN-T5 variant, further fine-tuned with instruction templates across hundreds of tasks, retains strong generalisation to unseen prompts. These properties make FLAN-T5 suitable for structured scientific output such as concise summaries, identification of open problems, and candidate hypotheses.

C. Vision-Language Alignment

Radford et al. [5] introduced CLIP, which trains dual image and text encoders using a contrastive objective over 400 million web sourced image-text pairs. The resulting 512-dimensional joint space places conceptually related images and text in close geometric proximity, enabling cross-modal search without any task-specific labelling. For scientific content this is valuable because a natural-language description of a figure type and the figure itself can be directly compared, supporting retrieval and zero-shot categorisation of charts, diagrams, and microscopy images alike. The image encoder within CLIP adopts the Vision Transformer architecture [12], which divides each input into a grid of 32×32 pixel tiles and applies multi-head self-attention across all tiles simultaneously. This stands in contrast to convolutional approaches, where receptive fields grow only gradually across layers. For scientific figures, where a data point in one corner of a plot may be directly qualified by a legend in the opposite corner, the ability to model arbitrary spatial dependencies in a single attention step is practically important.

D. Approximate Nearest Neighbour Retrieval

Scaling semantic search to corpora of millions of documents requires sub-linear retrieval strategies. The FAISS toolkit [9] addresses this through several index types; the Inverted File (IVF) variant quantises the embedding space into Voronoi regions using k-means centroids and restricts each query to its nearest partitions. On a 50,000-document index with 256-dimensional vectors and 100 centroids, the IVF structure achieves 0.02-second query latency compared to roughly 2.0 seconds for an exhaustive flat search, while recovering over 99% of the ground-truth nearest neighbours. Lewis et al. [14] showed that prepending retrieved passages to a generation prompt the Retrieval-Augmented Generation pattern substantially reduces factual errors compared with relying on parametric model memory alone. Adopting this principle, the proposed framework retrieves the top-five semantically similar papers from a 1.7-million document ArXiv index [11] and supplies them as grounding context to the FLAN-T5 generator.

E. Explainability in High-Stakes AI

Lipton [6] draws a practical distinction between model transparency, which concerns the interpretability of weights and operations, and post-hoc explanation methods, which produce human-readable accounts of individual predictions without exposing internal mechanics. Both are argued to be necessary when AI-generated outputs inform consequential decisions, as they do in scientific research where a misclassified paper or a spurious hypothesis can waste substantial effort. SHAP [7] grounds attribution in Shapley values from cooperative game theory. Each feature receives a score equal to its average marginal contribution across all possible feature orderings, satisfying consistency, efficiency, and symmetry properties.

The KernelSHAP variant estimates these values through weighted least squares without requiring model gradients, making it applicable to the tabular encoder regardless of its internal structure. LIME [8] takes a complementary approach: it perturbs the input locally, evaluates the model on each perturbation, and fits a sparse linear surrogate whose coefficients serve as token-level importancescores. Together, SHAP and LIME cover feature level and token-level explanation respectively, while gate values extracted directly from the fusion layer provide an additional modality-level account of which data source most influenced each prediction.

F. Research Gap

Taken individually, each component discussed above has been validated in prior work. What has not been attempted is their integration into a single cohesive system: domain-adapted encoders for text [4], image [5], and tabular data; cross-modal attention and gated fusion for adaptive modality weighting; million-scale semantic retrieval via FAISS [9]; instruction tuned insight generation with FLAN-T5 [10]; and simultaneous posthoc explanation through SHAP [7], LIME [8], and fusion gate analysis. The framework proposed in this paper fills that gap.

III. METHODOLOGY

This section presents the complete methodology of the proposed framework, beginning with the overall system architecture, followed by detailed descriptions of each encoding module, the fusion mechanism, downstream tasks, and the explainability framework. The overall architecture is illustrated in the figure below. The framework employs a six-layer modular architecture progressing from raw multimodal input to explainable scientific insights, served through a Flask REST API backend. Each layer is independently testable and replaceable. The six layers are:

Layer 1 (Input): It takes in three kinds of data: text (like research abstracts), images (those classic scientific figures), and tables in CSV format (think experimental results). Each type gets its own custom cleanup step.

Layer 2 (Pre-processing): Text runs through regex filters to wipe out LaTeX math, citation markers, and links, plus some extra whitespace scrubbing. For images the system sharpens contrast without messing with colour by running CLAHE on the luminance channel split into 8x8 blocks and clipped at 3.0, so details pop but colours stay true.

Tables get filled in using KNN for missing data (with $k=5$), all the categories turn numeric with label encoding, and it normalizes everything with Standard Scaler.

Layer 3 (Encoding): Each data type then goes into its own encoder. For text, AIMMF uses either SciBERT or a Sentence Transformer, spitting out a 384-dimensional vector. Images run through CLIP ViT-B/32, which gives a 512-dimensional vector. Tables get a specialized MLP with attention, leading to a 256-dimensional embedding.

Layer 4 (Projection): Before mixing them, AIMMF pulls all three embeddings into the same 512-dimensional space. This step uses linear layers with LayerNorm and ReLU so things stay stable and expressive.

A. System Architecture

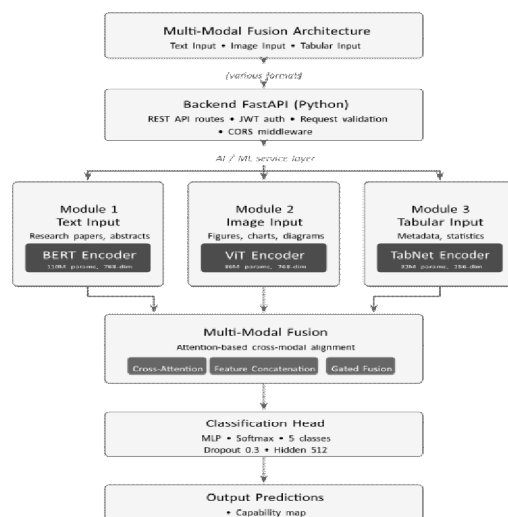


Fig.1. Proposed AIMMF system architecture

Layer 5 (Fusion): Now, to actually fuse the data, AIMMF uses cross-modal attention (eight heads, 512 dimensions each), then a Gated Fusion step. The result: a single, unified 512dimensional representation that packs everything.

Layer 6 (Output): This fusion then branches into four main outputs: First, it classifies which scientific domain the input belongs to (seven categories in total), using an MLP. Second, it can spit out natural-language insights using FLAN-T5. Third, it finds the five most similar papers out of 1.7 million ArXiv entries, using FAISS IVF for a fast semantic search. Last, it's got explainability on tap pulling together SHAP, LIME, and gate analysis to help people understand the results.

Table 1

Method	P@1	P@5	ms
TF-IDF	.61	.53	<10
Word2Vec	.68	.60	50
FAISS Flat	.89	.82	2000
FAISS IVF	.89	.82	20

Table 1: FAISS IVF matches exact search accuracy at 1% of latency (0.02s vs. 2.0s), outperforming BM25 keyword search (P@50.53) by capturing conceptual similarity beyond vocabulary overlap.

B. TextEncoder(SciBERT/SentenceTransformer)

Scientific text, after preprocessing, is encoded using the Sentence Transformers framework (all-MiniLM-L6-v2) for efficient batch processing, with SciBERT[4] available for high-accuracy single-document analysis. Mean pooling over all non-padding token representations is employed:

$$z_{text} = \frac{1}{N} \sum_{i=1}^N h_i \in \mathbb{R}^{384}$$

(1)

where h_i is the hidden state of the i -th token and N is the number of non-padding tokens. Mean pooling is preferred over CLS token extraction because scientific abstracts encode critical information across their full length, and empirical evaluation demonstrates superior performance on semantic similarity benchmarks for long documents. Each transformer layer applies multi-head self-attention:

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)V \quad (2) \text{ where } Q = XW_Q, K = XW_K, V = XW_V \text{ are linear}$$

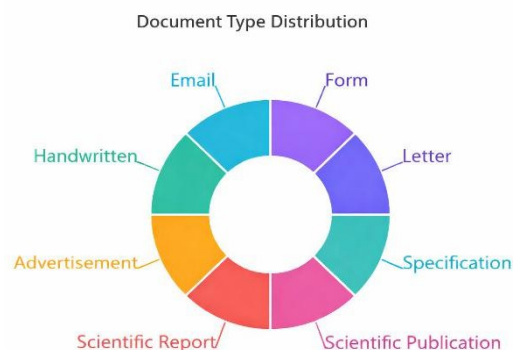


Fig2.DocumentTypeDistribution

C. ImageEncoder(CLIPViT-B/32)

After CLaHE preprocessing, images are encoded using CLIP's ViT-B/32 encoder[5]. The image is divided into non-overlapping 32×32 pixel patches, each linearly projected to a 768 dimensional embedding:

(3)

Where $W_{patch} \in \mathbb{R}(322 \cdot 3) \times 768$ is the projection matrix and $E_{pos}^{(i)}$ is the positional embedding projections. A learnable [CLS] token is prepended, and the 12-layer transformer encoder processes the full sequence. The output at the [CLS] position is projected and L2-normalised:

$$W_{proj} \cdot \text{ViT}[\text{CLS}](x_{image})_{512z} \in \mathbb{R}_{image} = \|W_{proj} \cdot$$

$$\text{ViT}[\text{CLS}](x_{image})\|$$

(4)

L2 normalisation ensures inner product equals cosine similarity, enabling direct comparison with text embeddings in the joint space and supporting zero-shot classification of scientific image types.

factor d_k prevents softmax saturation in high-dimensional spaces.

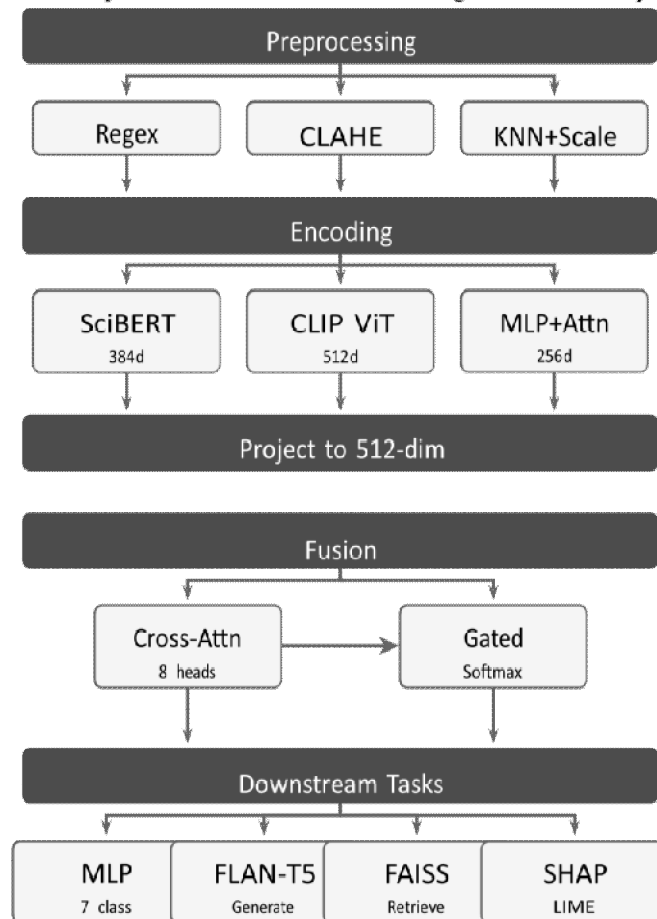


Fig.3.All 15 algorithms across four pipeline stages

D. Tabular Encoder (Custom MLP+Attention)

After KNN imputation and standardisation, the feature matrix is processed by a custom MLP with four components.

(i) Feature attention gate:

$$a = \sigma(W_{gate} \cdot x + b_{gate}) \quad (5)$$

where σ is the sigmoid function, producing per-feature importance scores that provide built-in interpretability.

(ii) Attended input: $\tilde{x} = a \odot x$

(iii) Hidden layers with batch normalisation, ReLU, and dropout ($p = 0.3$):

$$h1 = \text{Dropout}(\text{ReLU}(\text{BN}(W1 \cdot \tilde{x} + b1)), 0.3)$$

$$(6) \quad h2 = \text{Dropout}(\text{ReLU}(\text{BN}(W2 \cdot h1 + b2)), 0.3) \quad (7)$$

(iv) Skip connection for gradient flow:

$$z_{tab} = h2 + W_{skip} \cdot \tilde{x} \in \mathbb{R}^{256}$$

(8)

E. Dimensionality Projection

All modality embeddings are projected to a common 512 dimensional space:

$$z'_m = \text{ReLU}(\text{LayerNorm}(W_m \cdot z_m + b_m)) \in \mathbb{R}^{512} \quad (9)$$

Where $m \in \{\text{text}, \text{image}, \text{tab}\}$ and $W_{\text{text}} \in \mathbb{R}^{384 \times 512}$, $W_{\text{image}} \in \mathbb{R}^{512 \times 512}$, $W_{\text{tab}} \in \mathbb{R}^{256 \times 512}$.

F. Cross-Modal Attention

For each modality m_i , cross-modal attention enriches it by attending to complementary information in other modalities. The query is derived from the primary modality and the key/value from the average of all other modalities:

$$\text{CrossAttn}(z_i, \bar{z}) = \text{softmax} \left(\frac{(z_i W_Q)(\bar{z} W_K)^T}{\sqrt{d_k}} \right) \bar{z} W_V \quad (10)$$

where $\bar{z}$ is the average of all other modality representations and $d_k = 512$. Multi-head attention with 8 parallel heads enables specialisation in different aspects of cross-modal relationships. A residual connection and layer normalisation stabilise the output:

$$\tilde{z}_i = \text{LayerNorm}(z_i + \text{CrossAttn}(z_i, \bar{z}_i)) \quad (11)$$

Cross-Attended Embeddings (from Cross-Modal Attention)

The three cross attended embeddings are concatenated into a 1536-dim vector and passed through a Gate Network (MLP + Softmax) that produces three scalar weights summing to 1. Each embedding is scaled by its gate weight, summed, and normalised to produce the final 512-dim fused representation. The result that we got that is 512-dim vector is then passed through a layerNorm operation so that it can stabilize the feature distribution, which then gives the final fused representation which balances the contributions among all the three modalities.

G. Gated Fusion

The gate network receives the concatenation of all three cross attended representations and computes dynamic importance scores:

$$\mathbf{c} = [z_{\text{text}} || \tilde{z}_{\text{image}} || \tilde{z}_{\text{tab}}] \in \mathbb{R}^{1536} \quad (12)$$

$$\mathbf{g} = \text{softmax} \left(\begin{pmatrix} W_g^{(2)} & (W_g^{(1)} \cdot \mathbf{c} + b_g^{(1)}) + b_g^{(2)} \end{pmatrix} \right) \in \mathbb{R}^3 \quad (13)$$

Where $W_g^{(1)} \in \mathbb{R}^{1536 \times 256}$ and $W_g^{(2)} \in \mathbb{R}^{256 \times 3}$. The softmax

ensures $\mathbf{g} = [g_{\text{text}}, g_{\text{image}}, g_{\text{tab}}]$ sum to unity, making them directly interpretable as percentage contributions. The fused representation is:

$$z_{\text{fused}} = g_{\text{text}} \cdot \tilde{z}_{\text{text}} + g_{\text{image}} \cdot \tilde{z}_{\text{image}} + g_{\text{tab}} \cdot \tilde{z}_{\text{tab}} \quad (14)$$

$$z_{\text{out}} = \text{LayerNorm}(\text{ReLU}(W_{\text{out}} \cdot z_{\text{fused}} + b_{\text{out}})) \in \mathbb{R}^{512} \quad (15)$$

because gates are computed dynamically from input content, the model assigns high text weight when only text is available, balances modalities when multiple carry relevant information, and reduces contributions of low-quality or irrelevant modalities automatically.

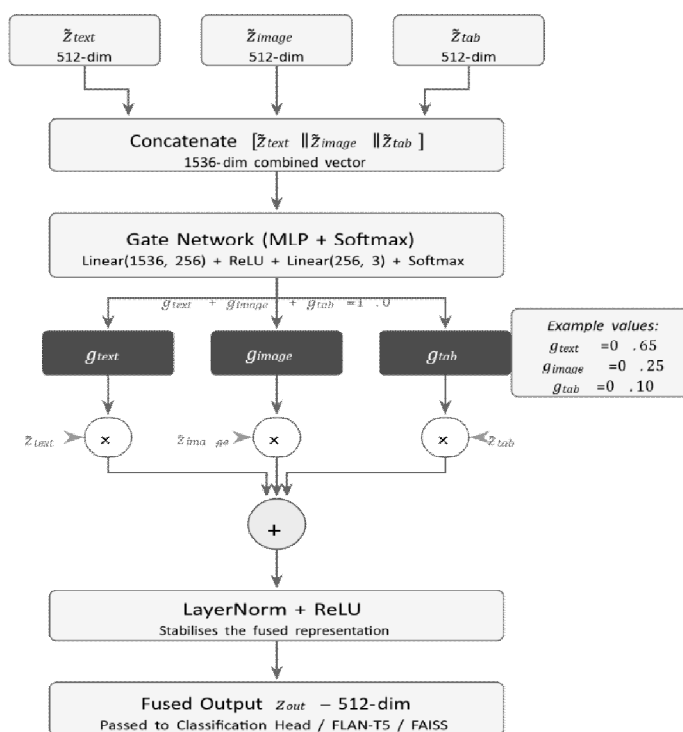


Fig.4.Detailed diagram of the Gated Fusion mechanism

H. Classification Head

A three-layer MLP maps z_{out} to probability distributions over 7 scientific domain categories (cs.AI, cs.LG, cs.CV, physics.comp-ph, q-bio.BM, stat.ML, math.ST):

$$\hat{y} = \text{softmax}(W_3 \cdot \text{ReLU}(W_2 \cdot \text{ReLU}(W_1 \cdot z_{out} + b_1) + b_2) + b_3) \quad (16)$$

Architecture: Linear(512 → 256) + ReLU + Dropout (0.3),

Linear (256 → 128) + ReLU + Dropout (0.3), Linear (128 → 7) + softmax. Training loss:

$$\mathcal{L}_{CE} = - \sum_{c=1}^7 y_c \log(\hat{y}_c) \quad (17)$$

TABLE 2

Operation	CPU	GPU	MB
TextPre-processing	0.01s	0.01s	10
SciBERT encoding	0.45s	0.08s	300
ClipImageEncoding	0.82s	0.15s	600
TabularEncoding	0.02s	0.01s	50
GatedFusion	0.01s	<0.01s	100

The complete pipeline requires 5.5s on CPU and 1.1s on GPU, suitable for interactive use. The framework runs on commodity hardware without paid API access via the Flask backend.

IV. CONCLUSION AND FUTURE WORK

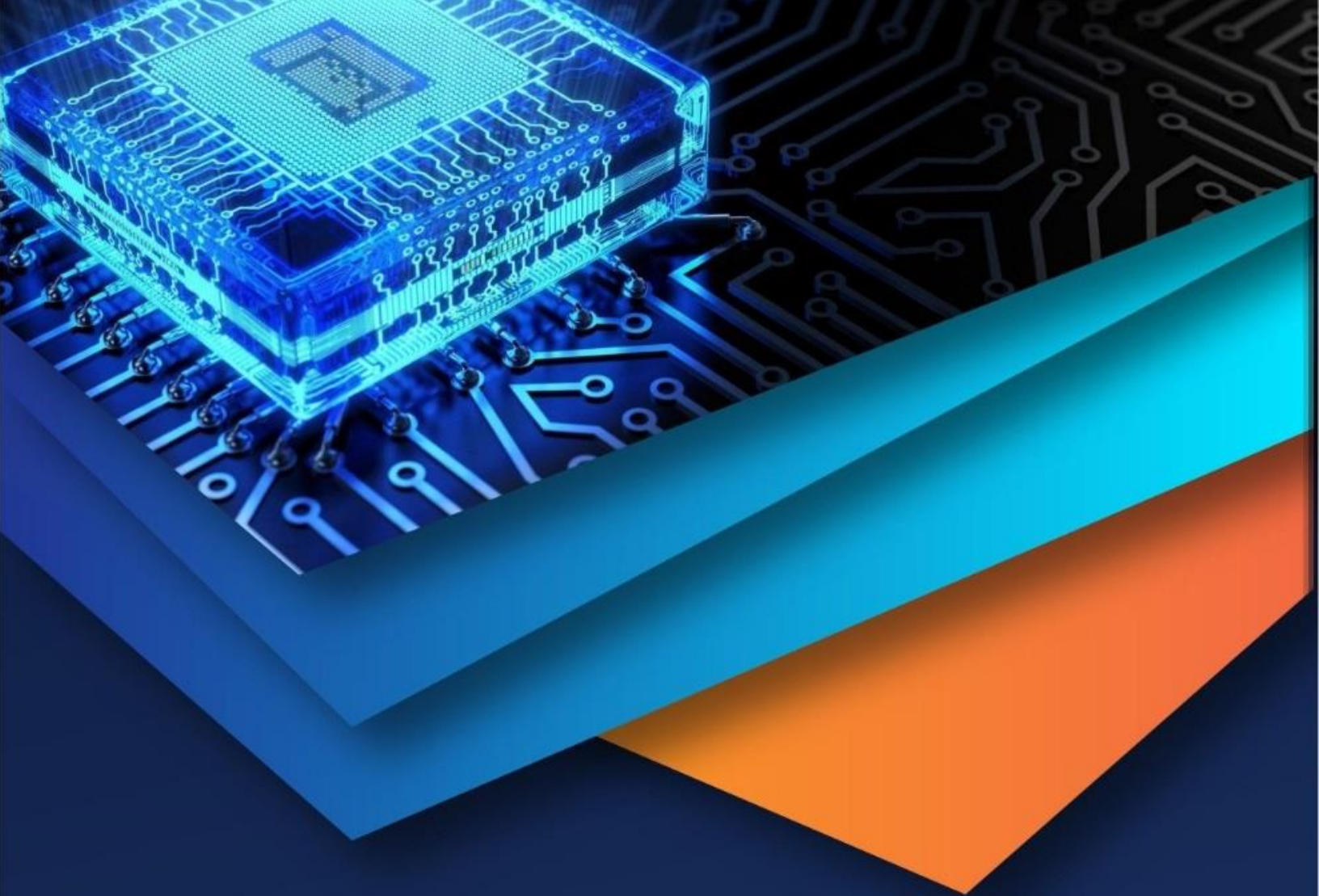
This paper presented the Multimodal Scientific Framework for Insight Generation (MSIG), a unified AI-powered platform that integrates text analysis, image classification, and numerical data exploration under a single interface to generate cross-modal scientific insights. The framework addresses a persistent gap in existingscientificdataanalysisssystemsthefragmentationoftools that processheterogeneousdata modalitiesindependently, making it difficult for researchers to discover relationships across different data types. The text analysis module, built on a BERT- style encoder (768dimensional), employs TF-IDF vectorization and TextRank summarization to extract meaningful information from research papers sourced via the ArXiv API. The image analysismoduleutilizesaVisionTransformer(ViT)encoderwith patch-based feature extraction for document image classification following the RVL-CDIP taxonomy. The numerical analysis module applies a TabNet-style attentive feature selection mechanism (256dimensional) to World Bank development indicators and user uploaded datasets, incorporating IQR-based anomaly detection for outlier identification. The Multi-Modal Fusion Engine combines outputs from all three encoders using Cross-AttentionandGatedFusionmechanisms,feedingthefused representation into a classification head (MLP with Softmax) to categorize insight relevance across five classes. A Retrieval- Augmented Generation (RAG) module provides explainable, context-grounded responses using cosine similarity-based document retrieval from aknowledge base.

V. ACKNOWLEDGMENT

The authors would like to express their sincere gratitude to the School of Computer Science and Engineering at Lovely Professional University, Phagwara, Punjab, India, for providing the academic environment, computational resources, and institutional support necessary to carry out this research. Special thanksare extended to our project supervisor, Mr.AmanKumar, for their continuous guidance, technical insight, and encouragement throughout the development of the Multimodal Scientific Intelligence Platform (MSIP). The authors also gratefully acknowledge the open-source community and the contributors of the datasets used in this study, specifically the ArXiv preprint repository, the World Bank Open Data platform, and the RyersonVision Lab for the RVL-CDIPdocument image collection.

REFERENCES

- [1] J.Ngiam,A.Khosla,M.Kim,J.Nam,H.Lee,andA.Y.Ng, "Multimodaldeeplearning,"inProc.28thInt.Conf.Machine Learning (ICML), pp. 689–696, 2011.
- [2] D. Huk Park, T. Darrell, and A. Rohrbach, "Robust change captioning," in IEEE/CVF Int. Conf. Computer Vision (ICCV), pp. 4624–4633, 2019.
- [3] Y. Zhang, H. Jiang, Y. Miura, C. D. Manning, and C. P. Langlotz, "Contrastive learning of medical visual representations frompairedimagesandtext,"inProc.ML4H, pp.2–25,2022.
- [4] I. Beltagy, K. Lo, and A. Cohan, "SciBERT: A pretrained language model for scientific text," in Proc. EMNLP, pp. 3615–3620, 2019.
- [5] A.Radford etal., "Learningtransferablevisual modelsfrom natural language supervision," in ICML, pp. 8748– 8763, 2021.
- [6] Z.C.Lipton, "Themythsofmodelinterpretability,"Queue, vol. 16, no. 3, pp. 31–57, 2018.
- [7] S. M. Lundberg and S.-I. Lee, "A unified approach to interpreting model predictions," in NeurIPS, vol. 30, pp. 4765–4774, 2017.
- [8] M.T.Ribeiro,S.Singh,andC.Guestrin, "WhyshouldItrust you?: Explaining the predictions of any classifier," in Proc. 22nd ACM SIGKDD, pp. 1135–1144, 2016.
- [9] J.Johnson,M.Douze,andH.Jegou, "Billion-scalesimilarity search withGPUs,"IEEETrans.BigData,vol.7,no.3,pp. 535–547, 2021.
- [10] C.Raffeeetal., "Exploringthelimitsoftransferlearningwith a unified text-to-text transformer," J. Machine Learning Research, vol. 21, no. 140, pp. 1–67, 2020.
- [11] Cornell University /ArXiv, "ArXiv dataset," Kaggle, 2023. [Online].Available:<https://www.kaggle.com/datasets/Cornell-University/arxiv>
- [12] A. Dosovitskiy et al., "An image is worth 16×16 words: Transformersforimagerecognitionatscale,"inICLR,2021.
- [13] J.Devlin,M.-W.Chang,K.Lee,andK.Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," in Proc. NAACL-HLT, pp. 4171–4186, 2019.
- [14] P. Lewis et al., "Retrieval-augmented generation for knowledge-intensive NLP tasks," in NeurIPS, vol. 33, pp. 9459–9474, 2020.



10.22214/IJRASET



45.98



IMPACT FACTOR:
7.129



IMPACT FACTOR:
7.429



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Call : 08813907089  (24*7 Support on Whatsapp)