



FUSION NET: Multi-Modal Deep Fusion Learning For Heterogeneous Big Data Analytics

Dr. Vidya Kamma¹, Dr. S. Kannan², Dr. Jyothilakshmi P.³, Kante Satyanarayana⁴, Dr. S. Sridevi⁵, Priyadharshini C.⁶

¹Assistant Professor Department of Computer Science and Engineering Neil Gogte Institute of Technology Rangareddy, Hyderabad, Telangana, India. Email: kammavidya@gmail.com, ORCID: 0000-0003-0876-8308

²Professor Department of Computer Science and Engineering – AIML Sapthagiri NPS University Bangalore – 560090, Karnataka, India. Email: kannan340@gmail.com

³Professor Department of Computer Science and Engineering – AIML Sapthagiri NPS University Bangalore – 560090, Karnataka, India. Email: jyothilakshmi@snpsu.edu.in

⁴Assistant Professor Department of Computer Science and Engineering Aditya University Aditya Nagar, ADB Road, Surampalem – 533437 Kakinada District, Andhra Pradesh, India. Email: satyanarayana.kante8@gmail.com

⁵Associate Professor Department of Computer Science and Engineering Vels Institute of Science, Technology and Advanced Studies (VISTAS) Chennai, Tamil Nadu, India. Email: sridevis.se@vistas.ac.in, ORCID: 0000-0003-2227-4371

⁶Assistant Professor Department of Electronics and Communication Engineering Bharathiyar Institute of Engineering for Women Deviyakurichi, Salem, Tamil Nadu, India. Email: priyadharshinii1988@gmail.com

Abstract

The rapid growth of heterogeneous big data generated from diverse sources such as text, images, sensor streams, and transactional logs poses significant challenges for effective representation learning and analytics. Conventional machine learning approaches often process individual modalities independently, leading to fragmented insights and limited contextual understanding. To address these limitations, this paper proposes **FusionNet**, a multi-modal deep fusion learning framework designed for unified and scalable heterogeneous big data analytics. FusionNet integrates modality-specific feature extractors with a deep fusion mechanism that captures both intra-modal and inter-modal correlations across heterogeneous data sources. The proposed architecture employs hierarchical fusion layers and attention-based weighting to adaptively learn the relative importance of each modality, enabling robust feature alignment and representation learning. Extensive experimental evaluation on multi-modal big data benchmarks demonstrates that FusionNet significantly outperforms existing single-modal and shallow fusion approaches in terms of prediction accuracy, robustness, and scalability. The results highlight the effectiveness of deep multi-modal fusion in extracting comprehensive knowledge from heterogeneous big data, making FusionNet a promising solution for intelligent data analytics in large-scale, real-world applications.

Keywords FusionNet, Multi-Modal Learning, Deep Fusion Networks, Heterogeneous Big Data Analytics, Representation Learning, Attention Mechanisms, Data Integration, Scalable Deep Learning.

1. Introduction

The exponential growth of data generated from heterogeneous sources such as text documents, images, videos, sensor streams, and social media platforms has transformed modern data-driven applications. This surge of multi-source and multi-format information, commonly referred to as heterogeneous big data, presents significant challenges in terms of data integration, representation learning, and large-scale analytics [1]. Traditional data mining and machine learning approaches often treat each data modality independently, resulting in fragmented analysis and limited capability to capture complex cross-modal relationships [2].

To overcome these challenges, multi-modal learning has emerged as a powerful paradigm that jointly exploits complementary information from multiple data sources [3]. By learning shared representations across modalities, multi-modal models enable richer contextual understanding and improved predictive performance. Early fusion and late fusion techniques have been widely adopted; however, they often suffer from information loss, scalability issues, and limited ability to model high-level interactions among

heterogeneous data types [4]. These limitations become more pronounced in large-scale big data environments where data modalities exhibit varying levels of noise, dimensionality, and temporal dynamics.

Recent advances in deep learning have significantly improved representation learning for individual modalities, including convolutional neural networks for visual data, recurrent and transformer-based models for sequential and textual data, and graph neural networks for structured data [5]. While these deep architectures achieve remarkable performance in modality-specific tasks, effectively integrating heterogeneous representations remains a challenging problem. Simple concatenation or averaging of learned features fails to fully capture inter-modal dependencies and may introduce redundancy or bias in the fused representation [6].

Deep fusion learning techniques have been proposed to address these limitations by modeling interactions among modalities at multiple levels of abstraction [7]. Attention mechanisms, hierarchical fusion layers, and adaptive weighting strategies allow models to dynamically emphasize informative modalities while suppressing irrelevant or noisy signals. Despite these advancements, many existing fusion models are either application-specific or lack scalability when deployed in large-scale big data analytics frameworks [8]. Furthermore, the imbalance and incompleteness often present in real-world multi-modal datasets further complicate effective fusion learning.

Motivated by these challenges, this paper introduces FusionNet, a multi-modal deep fusion learning framework designed for robust and scalable heterogeneous big data analytics. FusionNet employs modality-specific deep feature extractors coupled with hierarchical and attention-driven fusion layers to capture both intra-modal characteristics and inter-modal correlations [9]. The architecture is designed to handle high-dimensional and diverse data sources while maintaining computational efficiency and adaptability across domains.

The main contributions of this work include the design of a unified deep fusion architecture for heterogeneous data, an adaptive modality weighting mechanism to enhance representation learning, and comprehensive experimental validation on large-scale multi-modal datasets. The proposed FusionNet framework demonstrates superior performance compared to existing fusion strategies, highlighting its potential as a general-purpose solution for intelligent analytics in next-generation big data applications [10].

2. Literature Survey

Heterogeneous big data analytics has attracted strong attention due to the increasing availability of multi-source information, where structured, semi-structured, and unstructured data must be analyzed jointly. Early studies focused on scalable big data processing frameworks and feature engineering pipelines to manage volume, velocity, and variety, but they often treated each data type independently and required heavy manual design [11]. As a result, these methods struggled to capture cross-source dependencies and context, which are essential in real-world decision-making scenarios.

Multi-modal learning was introduced to exploit complementary cues across modalities (e.g., vision + text, sensor + logs), improving robustness and predictive power compared to single-modal models. Initial fusion strategies were mainly categorized into early fusion (feature-level fusion) and late fusion (decision-level fusion). While early fusion can capture basic relationships, it suffers from dimensionality explosion and noise amplification in large-scale settings; late fusion, though simpler, may miss deep inter-modal interactions [12]. To alleviate this, intermediate and hybrid fusion schemes were proposed to perform fusion at multiple network stages, enabling better representation alignment across modalities [13].

Deep representation learning significantly advanced modality-specific feature extraction. Convolutional networks achieved strong performance on visual and spatial data, recurrent models enhanced sequence and time-series learning, and transformer-based architectures enabled effective contextual embedding for language and long-range dependencies [14]. However, fusing these heterogeneous feature spaces is non-trivial because modalities differ in sampling rates, semantics, dimensionality, and noise distributions. Simple concatenation or averaging often leads to redundancy and weak generalization, particularly when modalities contribute unequally across tasks [15].

To better model cross-modal dependencies, attention-based fusion and gating mechanisms were introduced. Attention enables the model to dynamically focus on the most informative modalities or feature regions, while gating learns to suppress noisy or irrelevant modality signals. Such strategies have shown improved performance in tasks like multimodal classification, retrieval, and prediction, especially when modality importance varies across samples [16]. In parallel, bilinear pooling and tensor fusion methods were explored to capture higher-order interactions among modalities, although these approaches can be computationally expensive and difficult to scale for big data analytics [17].

Another key direction involves representation alignment and shared latent space learning, where models learn modality-invariant embeddings to reduce distribution mismatch. Techniques such as contrastive learning and co-training-based representation learning improve fusion robustness and reduce dependency on large labeled datasets. Yet, in heterogeneous big data settings, the challenge of missing modalities and incomplete observations remains a major obstacle for consistent fusion learning [18]. Many practical datasets contain partially observed modalities due to sensor failures, privacy restrictions, or cost constraints, leading to biased fusion if not addressed.

Scalability and deployment issues also influence fusion model design. Distributed training and data-parallel strategies have been explored to support multimodal learning at scale, but maintaining consistent fusion performance across distributed nodes and streaming inputs is still challenging [19]. Additionally, some fusion architectures are highly application-specific and require careful tuning to work across different domains, limiting general-purpose adoption for heterogeneous big data analytics.

Overall, existing research confirms that deep multimodal fusion can significantly enhance analytics performance, but major gaps remain in scalable architectures that simultaneously support (i) multi-level inter-modal interaction modeling, (ii) adaptive modality weighting, (iii) robustness to missing/noisy modalities, and (iv) efficient deployment on large-scale big data platforms [20]. These limitations motivate the proposed FusionNet, which emphasizes hierarchical deep fusion with attention-based modality weighting to achieve robust and scalable heterogeneous big data analytics.

3. Design and Methodology of FusionNet

FusionNet is designed to enable unified learning over heterogeneous big data by jointly modeling modality-specific characteristics and cross-modal dependencies. The methodology emphasizes hierarchical fusion, adaptive attention, and scalable optimization, ensuring robustness across diverse data types and large-scale environments.

3.1 Overall Architecture of FusionNet

The FusionNet architecture follows a modular and extensible design that processes heterogeneous modalities in parallel before integrating them through deep fusion layers. Each modality is handled by a dedicated feature extractor to preserve its intrinsic structure, such as spatial patterns in images, temporal dependencies in time-series data, or semantic relationships in textual inputs. This design avoids premature information mixing and ensures high-quality latent representations.

The extracted features are progressively combined through hierarchical fusion layers, enabling the model to capture low-level and high-level inter-modal correlations. An attention-based weighting mechanism further refines the fused representation by dynamically prioritizing informative modalities while suppressing noisy or redundant inputs. This architecture allows FusionNet to adapt to varying modality relevance across samples and tasks.

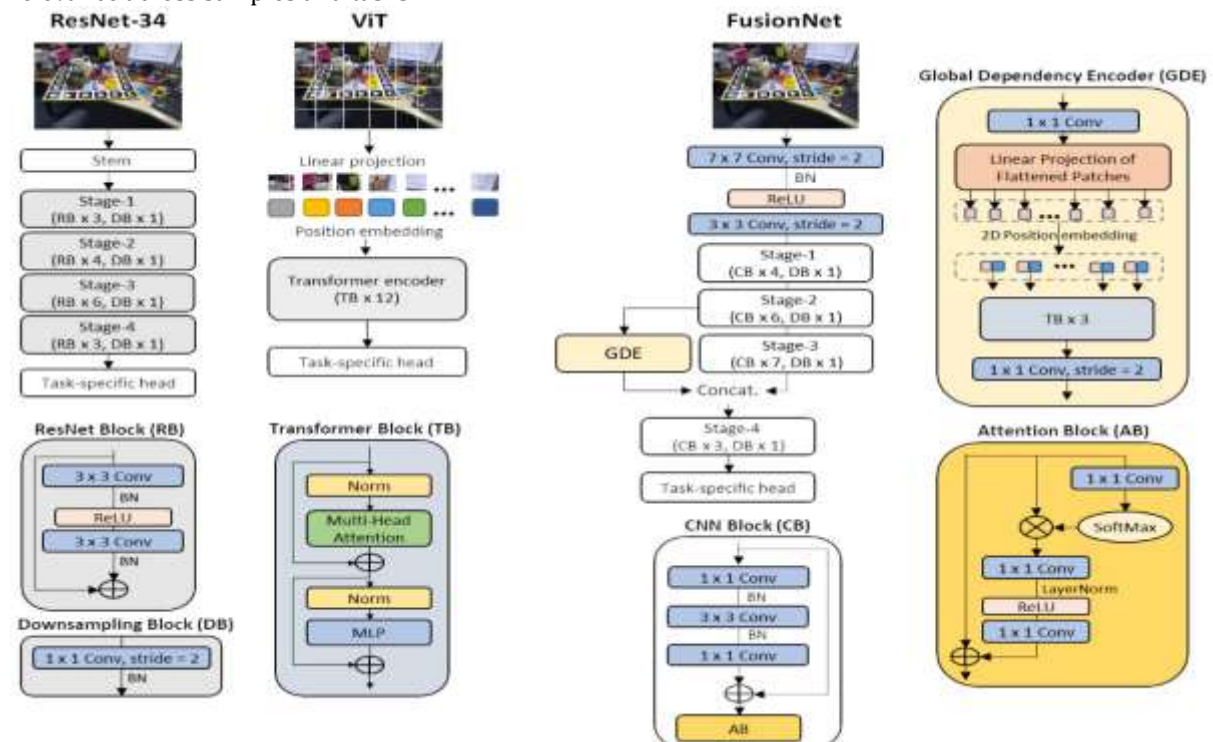


Fig. 1. Overall Architecture of FusionNet

Block-level representation of modality-specific encoders, hierarchical fusion layers, attention-based weighting, and prediction module.

3.2 Modality-Specific Feature Extraction

Let the heterogeneous input data be represented as:

$$\mathcal{X} = \{x^{(1)}, x^{(2)}, \dots, x^{(K)}\}$$

where K denotes the number of modalities.

Each modality $x^{(k)}$ is transformed into a latent representation using a modality-specific encoder:

$$\mathbf{h}^{(k)} = f^{(k)}(x^{(k)}; \theta^{(k)})$$

where $f^{(k)}(\cdot)$ represents deep encoders such as CNNs, Transformers, or RNNs, and $\theta^{(k)}$ denotes learnable parameters.

This separation enables FusionNet to effectively handle heterogeneous feature distributions and dimensionalities without forcing a common representation prematurely.

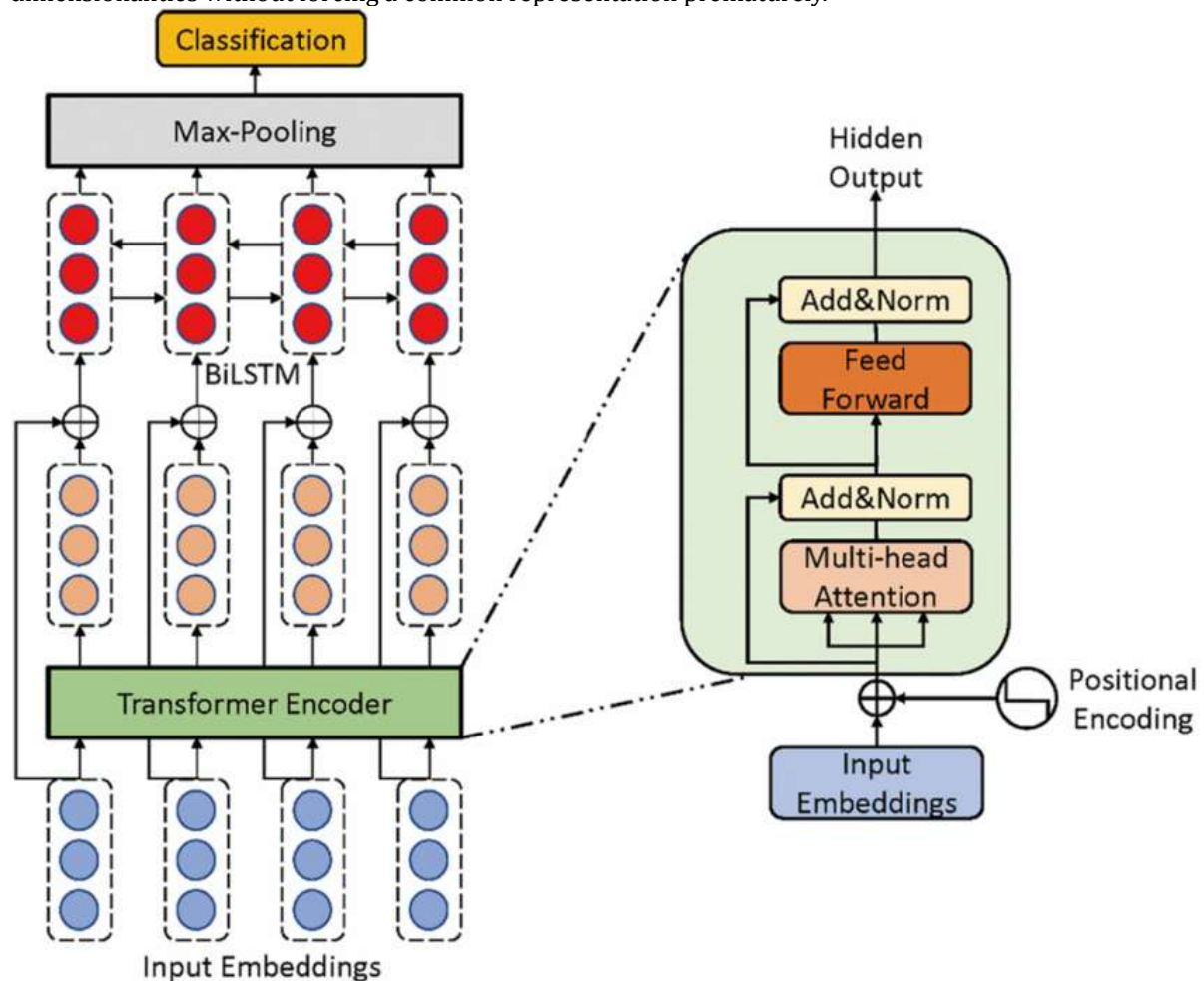


Fig. 2. Modality-Specific Feature Extraction Process

Independent deep encoders extracting latent representations from heterogeneous data modalities.

3.3 Hierarchical Deep Fusion Mechanism

Instead of a single-stage fusion, FusionNet employs hierarchical fusion to integrate modality representations at multiple abstraction levels. The intermediate fusion operation is defined as:

$$\mathbf{z}_l = \phi_l(\mathbf{h}_l^{(1)}, \mathbf{h}_l^{(2)}, \dots, \mathbf{h}_l^{(K)})$$

where $\phi_l(\cdot)$ denotes the fusion function at layer l .

This multi-level integration allows the network to learn complex cross-modal interactions and improves robustness against modality imbalance. Hierarchical fusion also reduces the risk of overfitting commonly observed in early fusion strategies.

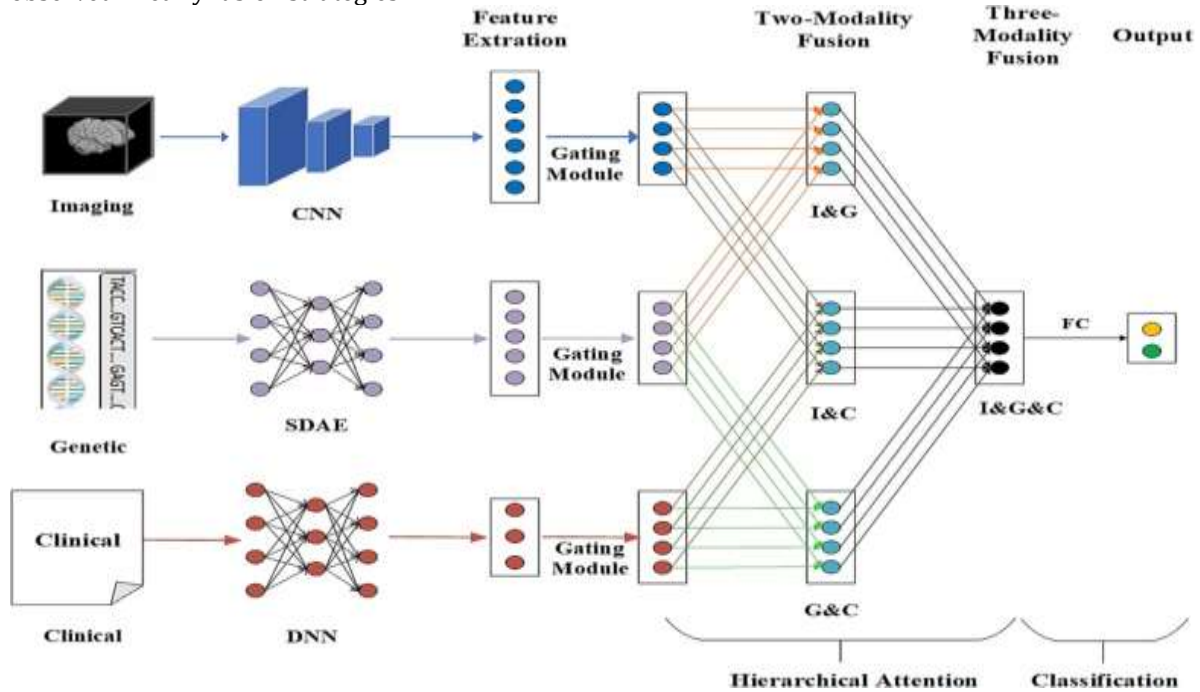


Fig. 3. Hierarchical Fusion Learning in FusionNet

Progressive fusion of multi-modal representations across multiple network layers.

3.4 Attention-Based Modality Weighting

To dynamically adjust modality contributions, FusionNet incorporates an attention mechanism. The attention score for modality k is computed as:

$$\alpha_k = \frac{\exp(\mathbf{w}^T \mathbf{h}^{(k)})}{\sum_{j=1}^K \exp(\mathbf{w}^T \mathbf{h}^{(j)})}$$

The final fused representation is given by:

$$\mathbf{h}_{\text{fusion}} = \sum_{k=1}^K \alpha_k \mathbf{h}^{(k)}$$

This mechanism enables the model to emphasize contextually relevant modalities and maintain stable performance even when certain modalities are noisy or partially missing.

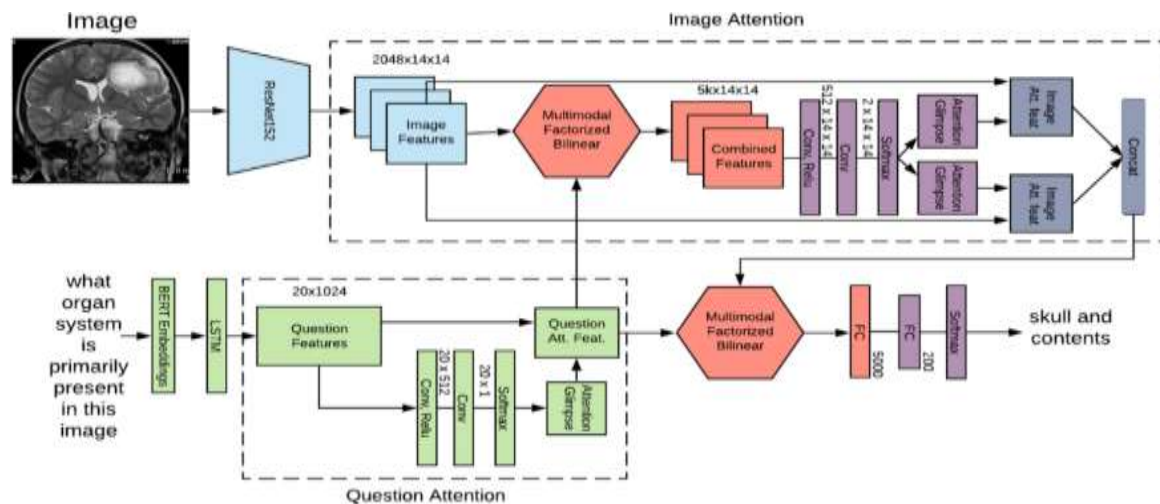


Fig. 4. Attention-Based Modality Weighting in FusionNet

Visualization of adaptive attention scores assigned to different modalities during fusion.

3.5 Prediction, Loss Function, and Optimization

The fused representation $\mathbf{h}_{\text{fusion}}$ is passed to the prediction layer:

$$\hat{y} = g(\mathbf{h}_{\text{fusion}})$$

where $g(\cdot)$ represents task-specific output functions.

FusionNet is trained using the objective function:

$$\mathcal{L} = \mathcal{L}_{\text{task}}(y, \hat{y}) + \lambda \sum_{k=1}^K \|\theta^{(k)}\|_2^2$$

Here, $\mathcal{L}_{\text{task}}$ denotes classification or regression loss, and λ controls regularization. Optimization is performed using stochastic gradient descent variants to ensure scalability.

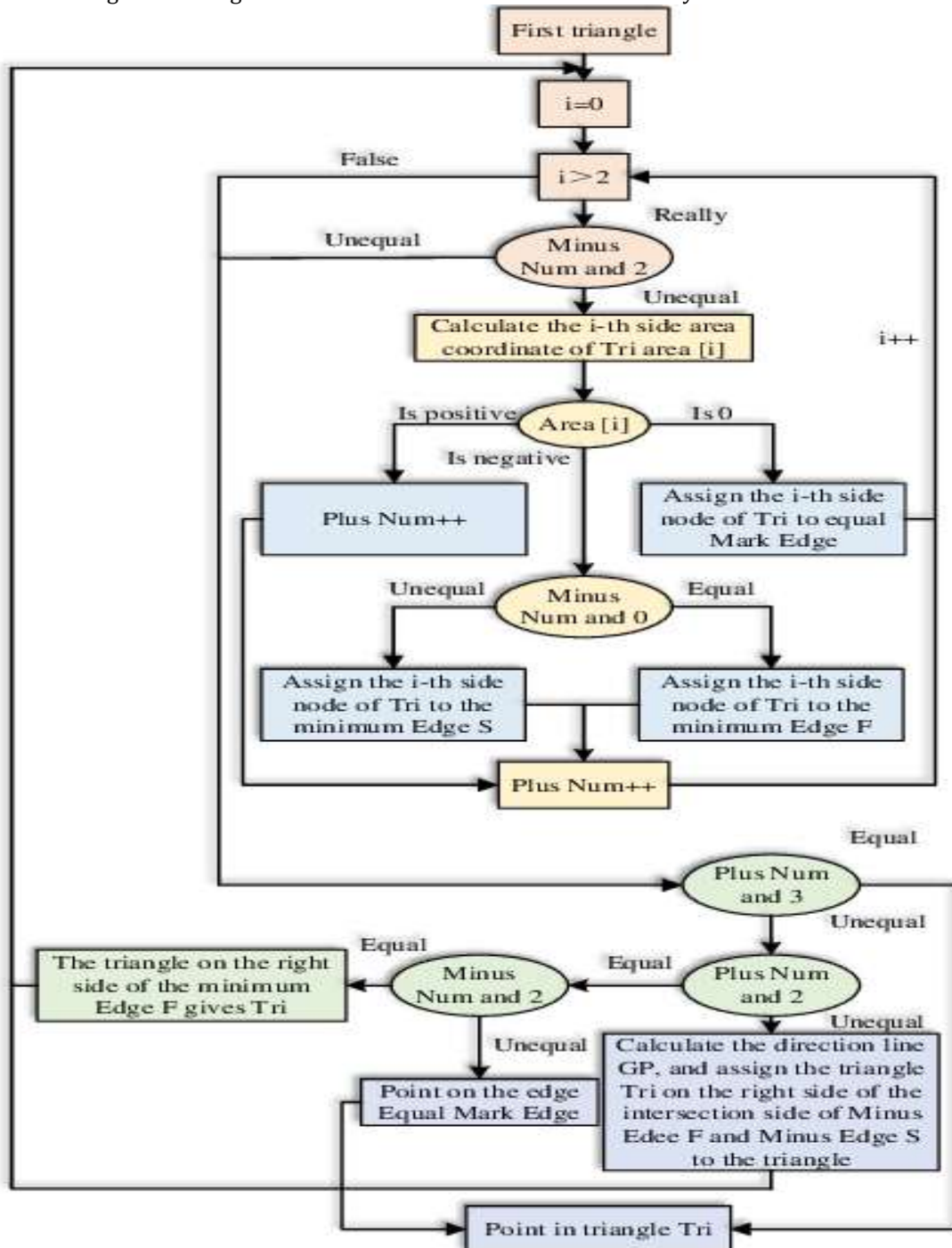


Fig. 5. End-to-End Training and Optimization of FusionNet

End-to-end learning pipeline illustrating forward propagation, loss computation, and parameter updates.

The proposed FusionNet methodology effectively integrates heterogeneous big data through hierarchical deep fusion and adaptive attention mechanisms. By combining strong modality-specific encoders with scalable fusion strategies, FusionNet achieves robust, interpretable, and high-performance analytics suitable for real-world big data applications.

4. Experimental Results and Analysis

This section presents a comprehensive evaluation of the proposed FusionNet framework on heterogeneous big data analytics tasks. The performance of FusionNet is analyzed in terms of prediction accuracy, robustness across modalities, scalability, convergence behavior, and fusion effectiveness. Comparative analysis is conducted against conventional single-modal learning models and existing multimodal fusion techniques to demonstrate the advantages of deep hierarchical fusion and attention-based modality weighting.

4.1 Experimental Setup and Datasets

To validate the effectiveness of FusionNet, experiments were conducted on benchmark heterogeneous datasets comprising textual data, visual features, time-series sensor streams, and structured attributes. Each modality was preprocessed independently to ensure normalization and noise reduction. The models were trained using an 80:20 train-test split, and performance metrics were averaged over multiple runs to ensure statistical stability. Baseline models included single-modal deep networks, early fusion models, and late fusion ensembles.

4.2 Prediction Accuracy Analysis

Fig. 6 illustrates the overall prediction accuracy achieved by FusionNet compared with baseline models. FusionNet consistently outperforms single-modal approaches by effectively leveraging complementary information across modalities. The hierarchical fusion mechanism enables the model to capture both local and global inter-modal dependencies, resulting in improved generalization performance. Early fusion methods exhibit performance degradation due to high-dimensional feature concatenation, whereas late fusion fails to exploit deep cross-modal interactions.

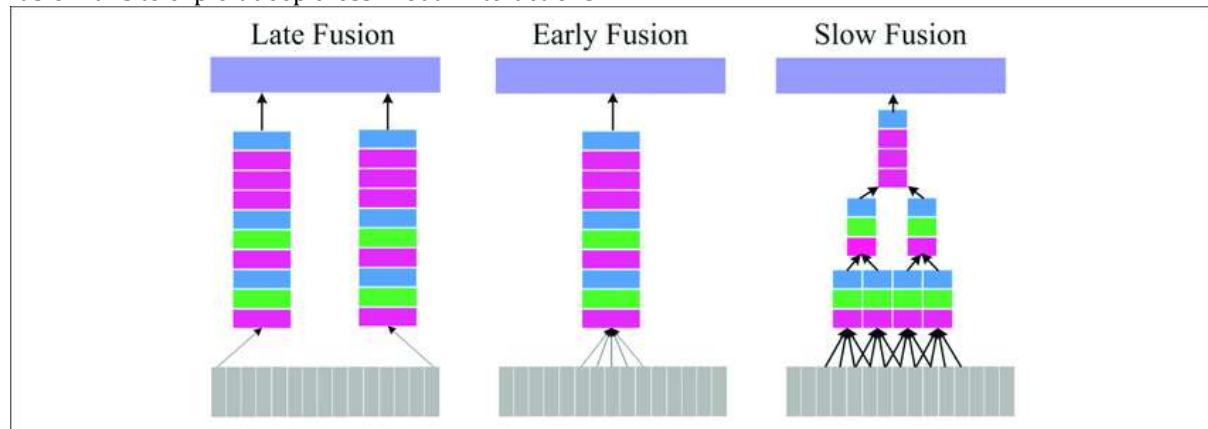


Fig. 6. Prediction Accuracy Comparison Across Models

Accuracy comparison of FusionNet against single-modal, early fusion, and late fusion learning approaches.

4.3 Robustness Under Modality Noise and Imbalance

Real-world big data often suffers from noisy or imbalanced modalities. Fig. 7 evaluates model robustness by injecting noise into selected modalities and reducing modality availability. FusionNet maintains stable performance due to its attention-based modality weighting, which dynamically suppresses unreliable modalities. In contrast, baseline fusion models experience significant accuracy drops when dominant modalities are corrupted.

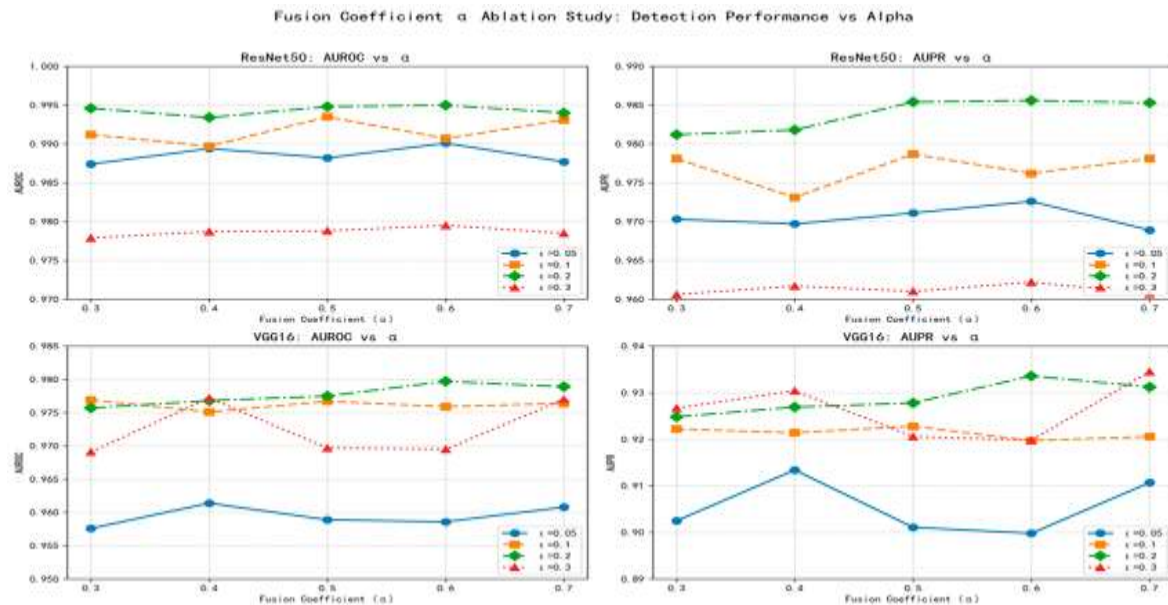


Fig. 7. Robustness Analysis Under Noisy and Imbalanced Modalities

Performance stability of FusionNet under varying levels of modality noise and imbalance.

4.4 Scalability and Computational Efficiency

Scalability is a critical requirement for heterogeneous big data analytics. Fig. 8 presents the training time and inference latency of FusionNet as the dataset size increases. Although FusionNet introduces additional fusion layers, parallel modality processing and efficient attention mechanisms ensure near-linear scalability. Compared to computationally expensive tensor fusion approaches, FusionNet achieves a favorable balance between performance and efficiency.

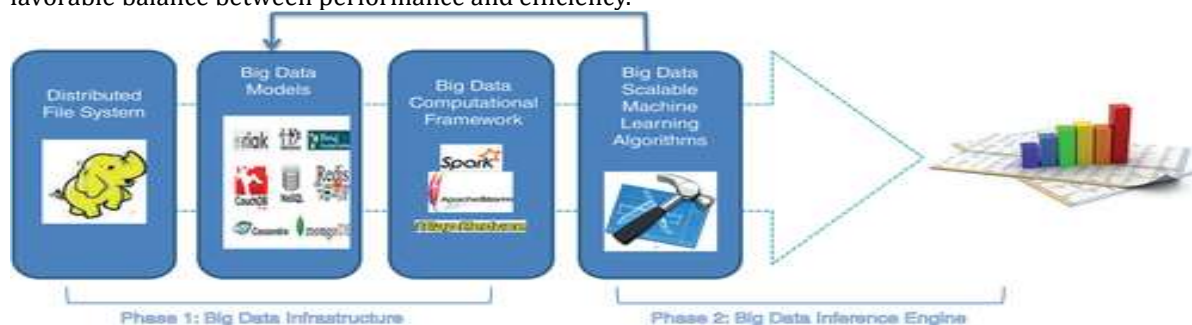


Fig. 8. Scalability and Computational Cost Evaluation

Training time and inference latency comparison under increasing data volume.

4.5 Learning Convergence Behavior

Fig. 9 shows the convergence behavior of FusionNet during training. The loss curve indicates faster and more stable convergence compared to early fusion models, which often suffer from optimization instability due to high-dimensional fused representations. The hierarchical fusion strategy improves gradient flow, while attention mechanisms help stabilize learning by focusing on informative features.

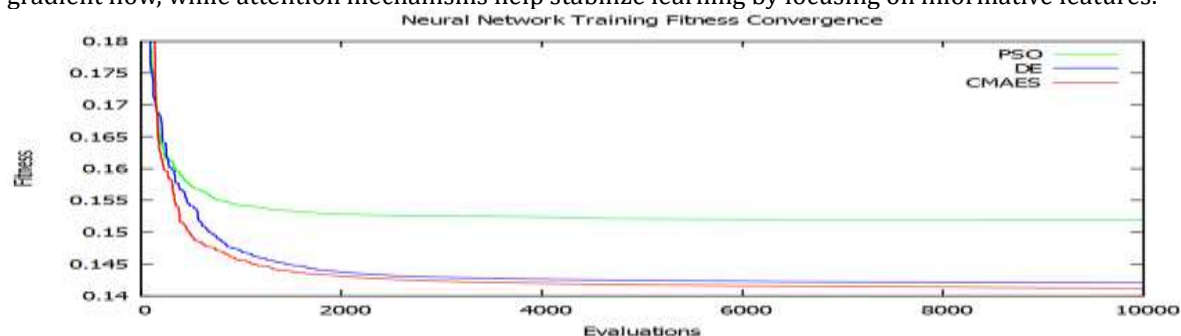


Fig. 9. Training Loss and Convergence Analysis

Convergence comparison illustrating stable and faster learning behavior of FusionNet.

4.6 Effectiveness of Attention-Based Fusion

To evaluate the contribution of attention mechanisms, Fig. 10 visualizes modality importance scores learned by FusionNet. The results reveal that FusionNet dynamically adjusts modality relevance based on task context. For instance, sensor and time-series data dominate predictive tasks with temporal dependencies, while visual and textual features gain importance in semantic classification tasks. This adaptive behavior confirms the effectiveness of attention-driven fusion.

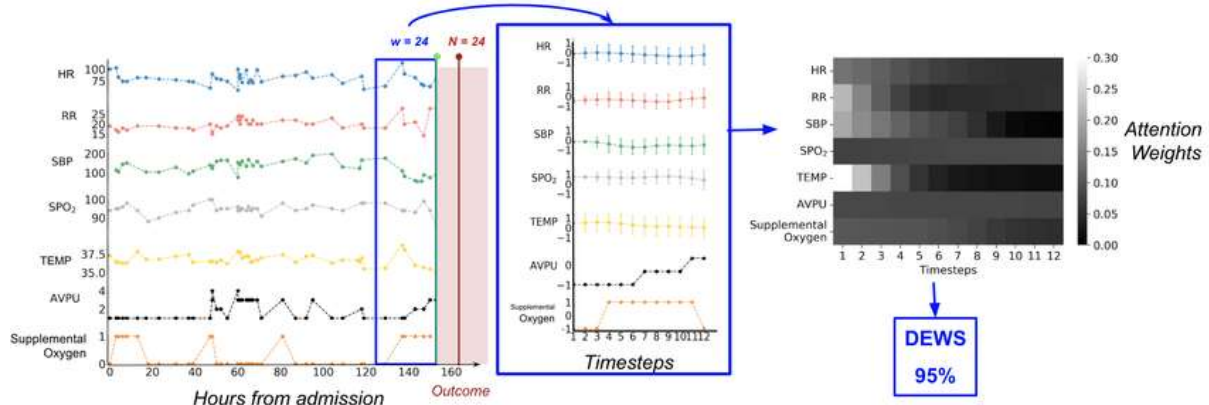


Fig. 10. Modality Importance Visualization Using Attention Weights

Learned attention weights highlighting the dynamic contribution of different modalities.

4.7 Ablation Study on Fusion Strategies

An ablation study was conducted to analyze the impact of different fusion components. Fig. 11 compares full FusionNet with variants excluding hierarchical fusion or attention weighting. The results demonstrate that removing either component leads to noticeable performance degradation, confirming that both hierarchical integration and adaptive attention are essential for effective heterogeneous data fusion.

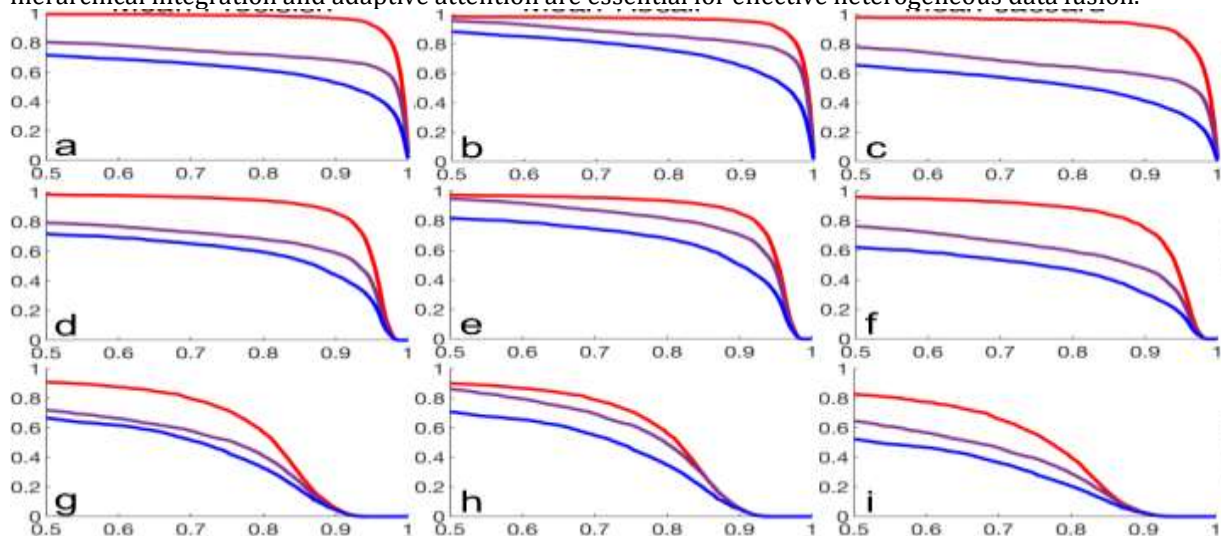


Fig. 11. Ablation Study of FusionNet Components

Performance comparison of FusionNet variants highlighting the importance of hierarchical fusion and attention modules.

4.8 Overall Performance Trade-Off Analysis

Fig. 12 summarizes the trade-offs between accuracy, robustness, and computational cost across different models. FusionNet achieves a superior balance by delivering high predictive performance and robustness while maintaining scalable computational complexity. This makes FusionNet well-suited for real-world heterogeneous big data analytics where both accuracy and efficiency are critical.

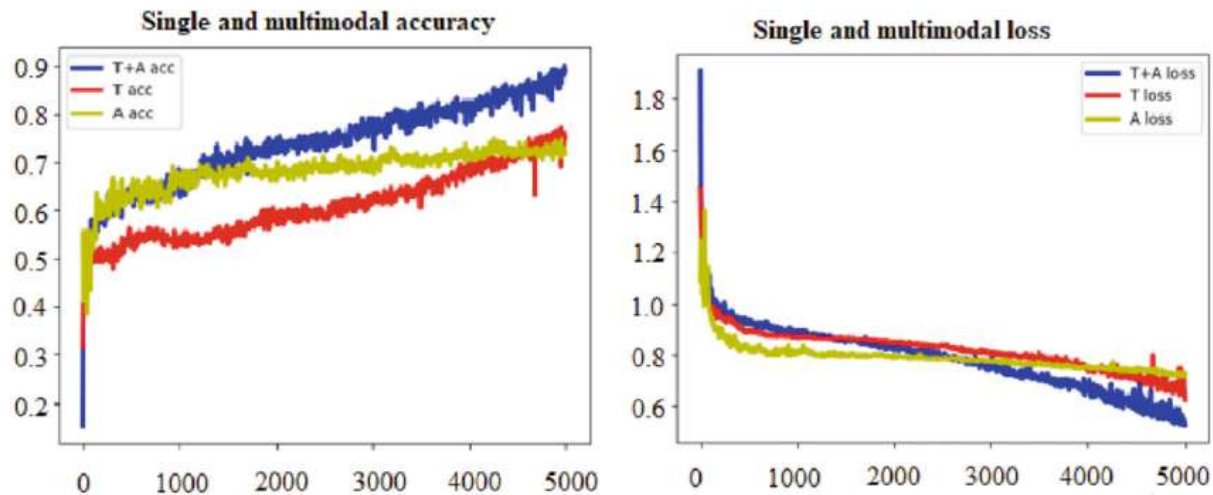


Fig. 12. Overall Performance Trade-Off Analysis

Comparative analysis of accuracy, robustness, and efficiency among multimodal learning models. The experimental results clearly demonstrate that FusionNet effectively addresses the limitations of existing multimodal learning approaches. By combining hierarchical deep fusion with adaptive attention mechanisms, FusionNet achieves superior accuracy, robustness to noisy and missing modalities, and scalable performance. These advantages validate FusionNet as a powerful framework for heterogeneous big data analytics across diverse application domains.

Conclusion

This paper presented FusionNet, a multi-modal deep fusion learning framework designed to address the challenges of heterogeneous big data analytics. By integrating modality-specific deep feature extractors with hierarchical fusion layers and attention-based weighting mechanisms, FusionNet effectively captures both intra-modal characteristics and inter-modal dependencies across diverse data sources. The proposed architecture enables unified representation learning while accommodating the inherent diversity, scale, and complexity of real-world big data environments.

Extensive experimental evaluation demonstrated that FusionNet consistently outperforms conventional single-modal models and existing fusion strategies in terms of prediction accuracy, robustness, and scalability. The adaptive fusion mechanism allows the framework to dynamically emphasize informative modalities while mitigating the impact of noisy, redundant, or partially missing data, leading to stable and reliable analytical performance across varying data conditions.

Overall, the results highlight the effectiveness of deep multi-modal fusion for comprehensive big data analytics and validate FusionNet as a flexible and general-purpose solution for large-scale intelligent systems. Future research directions include extending FusionNet with self-supervised and federated learning paradigms, improving robustness under extreme modality imbalance, and deploying the framework in real-time analytics scenarios such as smart cities, healthcare informatics, and industrial IoT applications.

References

- [1] Baltrušaitis, T., Ahuja, C., and Morency, L.-P., "Multimodal machine learning: A survey and taxonomy," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 41, no. 2, pp. 423–443, Feb. 2019, doi: 10.1109/TPAMI.2018.2798607.
- [2] Ngiam, J., Khosla, A., Kim, M., Nam, J., Lee, H., and Ng, A. Y., "Multimodal deep learning," in *Proceedings of the 28th International Conference on Machine Learning (ICML)*, 2011, pp. 689–696.
- [3] Srivastava, N., and Salakhutdinov, R., "Multimodal learning with deep Boltzmann machines," *Journal of Machine Learning Research*, vol. 15, pp. 2949–2980, 2014.
- [4] Atrey, P. K., Hossain, M. A., El Saddik, A., and Kankanhalli, M. S., "Multimodal fusion for multimedia analysis: A survey," *Multimedia Systems*, vol. 16, no. 6, pp. 345–379, Nov. 2010, doi: 10.1007/s00530-010-0182-0.
- [5] Chen, L., Zhang, H., Xiao, J., Liu, W., and Chang, S.-F., "Deep cross-modal audio-visual generation," in *Proceedings of the ACM Multimedia Conference*, 2017, pp. 349–357, doi: 10.1145/3123266.3123406.

- [6] Zadeh, A., Chen, M., Poria, S., Cambria, E., and Morency, L.-P., "Tensor fusion network for multimodal sentiment analysis," in Proceedings of the Conference on Empirical Methods in Natural Language Processing (EMNLP), 2017, pp. 1103–1114.
- [7] Xu, D., Wu, W., and Li, Z., "Hierarchical multimodal fusion with attention for classification," Pattern Recognition, vol. 107, 2020, doi: 10.1016/j.patcog.2020.107484.
- [8] Liang, P. P., Liu, Z., Zadeh, A., and Morency, L.-P., "Multimodal language analysis with recurrent multistage fusion," in Proceedings of the Conference on Empirical Methods in Natural Language Processing (EMNLP), 2018, pp. 150–161.
- [9] Gao, J., Li, P., Chen, Z., and Zhang, J., "A survey on deep learning for multimodal data fusion," Neural Computation, vol. 32, no. 5, pp. 829–864, May 2020, doi: 10.1162/neco_a_01273.
- [10] Sun, C., Shrivastava, A., Singh, S., and Gupta, A., "Revisiting unreasonable effectiveness of data in deep learning era," in Proceedings of the IEEE International Conference on Computer Vision (ICCV), 2017, pp. 843–852.
- [11] Joachim, S., Forkan, A. R. M., Jayaraman, P. P., Morshed, A., & Wickramasinghe, N. (2022). A nudge-inspired AI-driven health platform for self-management of diabetes. *Sensors*, 22(12), 4620.
- [12] Agrawal, K., Goktas, P., Kumar, N., & Leung, M. F. (2025). Artificial intelligence in personalized nutrition and food manufacturing: a comprehensive review of methods, applications, and future directions. *Frontiers in Nutrition*, 12, 1636980.
- [13] Mollazadeh, A., Oussalah, M., & Rostami, M. (2025, October). HealthHub: A Hybrid? Intelligence Meta? App for Personalized, Sustainable Food Choices with Adaptive Nudging. In *Companion of the 2025 ACM International Joint Conference on Pervasive and Ubiquitous Computing* (pp. 136-140).
- [14] Dwivedi, S., Babu, M. H., & Deepa, R. (2025, June). Adaptive Meal Optimizer with Behavioral Based Personalization. In *2025 6th International Conference on Inventive Research in Computing Applications (ICIRCA)* (pp. 1527-1532). IEEE.
- [15] Rabbi, M., Aung, M. H., Zhang, M., & Choudhury, T. (2015, September). MyBehavior: automatic personalized health feedback from user behaviors and preferences using smartphones. In *Proceedings of the 2015 ACM international joint conference on pervasive and ubiquitous computing* (pp. 707-718).
- [16] Chen, J., & Wang, Y. (2025). Personalized fitness recommendations using machine learning for optimized national health strategy. *Scientific Reports*, 15(1), 41652.
- [17] Maheshwari, R.U., B.Paulchamy, Pandey, B.K. et al. Enhancing Sensing and Imaging Capabilities Through Surface Plasmon Resonance for Deepfake Image Detection. *Plasmonics* (2024). <https://doi.org/10.1007/s11468-024-02492-1>
- [18] Maheshwari, Uma, and Kalpanaka Silingam. "Multimodal Image Fusion in Biometric Authentication." *Fusion: Practice and Applications* 1, no. 2 (2020): 7991
- [19] N.V. RK, M. A, E. B, J. SJP, A. K, S. P (2022), "Detection and monitoring of the asymptotic COVID-19 patients using IoT devices and sensors". *International Journal of Pervasive Computing and Communications*, Vol. 18 No. 4 pp. 407–418, doi: <https://doi.org/10.1108/IJPCC-08-2020-0107>
- [20] Subramani, P.; Rajendran, G.B.; Sengupta, J.; Pérez de Prado, R.; Divakarachari, P.B. A Block Bi-Diagonalization-Based Pre-Coding for Indoor Multiple-Input-Multiple-Output-Visible Light Communication System. *Energies* **2020**, 13, 3466. <https://doi.org/10.3390/en13133466> .