

# Evolution of Multimodal Deep Learning Architectures in Computational Biology: A Software Engineering and Big Data Perspective

Mr. Ninad Janardan Dani  
Research Scholar,

School of Computational Sciences, Swami Ramanand Teerth Marathwada University, Nanded

Dr. Nilesh K. Deshmukh  
Associate Professor,

School of Computational Sciences, Swami Ramanand Teerth Marathwada University, Nanded

**Abstract** - The intersection of Big Data and computational biology has catalyzed a paradigm shift in how complex biological systems are modeled. As genomic, proteomic, and structural databases expand exponentially, the software engineering challenge has transitioned from data acquisition to algorithmic scalability. This paper provides a comprehensive algorithmic survey of the evolution of deep learning architectures in computational biology, evaluating them strictly through the lens of software engineering, computational complexity, and system architecture. We trace the structural migration from single-modality encoders—such as 1D Convolutional Neural Networks (CNNs) and Graph Neural Networks (GNNs)—to modern multimodal Transformer architectures. Furthermore, we critically analyze the hardware bottlenecks introduced by overparameterized networks, specifically evaluating the quadratic time and space complexity inherent in standard attention mechanisms. By examining these architectures as software pipelines, this survey highlights the imperative for "Green AI" and resource-optimized computing. We conclude by outlining best practices in Machine Learning Operations (MLOps) and pipeline reproducibility, arguing that future advancements in bioinformatics depend on optimized multi-run validation frameworks and efficient memory footprints rather than brute-force parameter scaling.

**Keywords** - Deep Learning Architectures; Software Engineering; Big Data Analytics; Computational Complexity; Multimodal Fusion; Graph Neural Networks; Machine Learning Operations (MLOps).

## I. INTRODUCTION

### A. The Big Data Bottleneck in Bioinformatics

Over the last decade, computational biology has transformed into a premier domain of Big Data analytics. High-throughput sequencing technologies, cryo-electron microscopy, and combinatorial chemical libraries have generated petabytes of unstructured biological data. Databases such as the Protein Data Bank (PDB) and PubChem require highly scalable software infrastructures for storage, retrieval, and processing [1] [2]. Consequently, the primary barrier in modern bioinformatics is no longer data scarcity, but algorithmic throughput and software architecture optimization.

From a software engineering perspective, modeling biomolecular interactions requires processing highly heterogeneous data structures: 1D sequential strings (e.g., SMILES and FASTA formats), 2D topological graphs (molecular connectivity), and 3D spatial point clouds (atomic coordinates). To handle this multimodal data, deep learning (DL) architectures have grown increasingly complex. However, designing predictive algorithms that operate efficiently at this scale demands careful management of computational time complexity, GPU memory allocation, and inference latency [3].

### B. The Evolution of Predictive Software Architectures

Early predictive software in computational biology relied on classical machine learning models, such as Random Forests and Support Vector Machines (SVMs), built on top of handcrafted molecular descriptors [4]. While computationally lightweight, these models suffered from severe feature-engineering bottlenecks.

The advent of representation learning automated the feature-extraction pipeline. Convolutional Neural Networks (CNNs) and Recurrent Neural Networks (RNNs) became the standard for processing 1D sequential biological data [5]. As the field recognized the limitations of linear representations for non-Euclidean chemical structures, Graph Neural Networks (GNNs) emerged to model 2D molecular topologies [6]. Most recently, the empirical success of self-attention mechanisms in natural language processing has prompted the widespread adoption of multi-billion-parameter Transformer architectures to model long-range biological dependencies [7] [8].

### C. The Software Engineering Challenge: Scaling vs. Efficiency

While these architectural advancements have achieved remarkable accuracy on standardized benchmarks, they have introduced severe systems engineering challenges. Modern multimodal architectures frequently suffer from extreme parameter bloat. The unchecked scaling of deep learning models has resulted in exponentially increasing carbon footprints and prohibitive hardware requirements [9].

For software engineers deploying high-throughput virtual screening (HTVS) pipelines, inference latency is a critical metric. A model that achieves a marginal fractional improvement in accuracy but requires ten times the GPU memory and compute time represents a failure in software optimization. Therefore, analyzing these models requires looking beyond statistical accuracy to evaluate their algorithmic efficiency, big-O complexity, and pipeline reproducibility.

This review systematically deconstructs the evolution of deep learning algorithms in computational biology through an engineering lens, assessing the computational trade-offs of 1D, 2D, and multimodal fusion architectures.

## II. ALGORITHMIC COMPLEXITY OF SINGLE-MODALITY ENCODERS

### A. 1D Sequence Encoders: CNNs and RNNs

In the context of computational biology, 1D sequential data—such as Simplified Molecular-Input Line-Entry System (SMILES) strings for chemical ligands or FASTA amino acid sequences for target proteins—are traditionally processed using 1D Convolutional Neural Networks (CNNs) or Recurrent Neural Networks (RNNs).

From a computational complexity standpoint, a standard 1D CNN layer applies a filter of kernel size  $k$  across a sequence of length  $L$ . Assuming the input and output channel dimensions are both  $d$ , the time complexity of evaluating a single 1D convolutional layer is bounded by  $O(L \cdot k \cdot d^2)$  [5]. Because the convolution operation is inherently stateless across the spatial dimension, the matrix multiplications can be parallelized efficiently across GPU Streaming Multiprocessors (SMs). However, CNNs struggle to capture distant biochemical dependencies, such as allosteric binding sites situated far apart on a primary protein sequence but folded closely together in 3D space.

Conversely, RNNs—specifically Long Short-Term Memory (LSTM) networks—were adopted to model these long-range sequential dependencies. While the time complexity of a single LSTM step is  $O(d^2)$ , processing an entire sequence of length  $L$  requires  $O(L \cdot d^2)$  operations. The critical software engineering bottleneck of RNNs lies in their strict sequential execution constraint. Unlike CNNs, step  $t$  cannot be computed until step  $t - 1$  is resolved, resulting in  $O(L)$  sequential operations that severely throttle parallel hardware utilization. This makes RNNs highly inefficient for high-throughput screening pipelines dealing with long macromolecules [4].

### B. 2D Topological Encoders: Graph Neural Networks (GNNs)

To overcome the limitations of 1D linear representations, the field shifted toward Graph Neural Networks (GNNs) to map the 2D topological connectivity of molecular atoms and bonds. A molecule is represented computationally as an undirected

graph  $G = (V, E)$ , where  $|V|$  is the number of atoms (nodes) and  $|E|$  is the number of chemical bonds (edges).

The standard Graph Convolutional Network (GCN) updates node representations via a neighborhood aggregation (message passing) scheme defined as  $H^{(l+1)} = \sigma(\tilde{A}H^{(l)}W^{(l)})$ , where  $\tilde{A}$  is the normalized adjacency matrix,  $H^{(l)}$  represents the node features at layer  $l$ , and  $W^{(l)}$  is the trainable weight matrix [6].

The algorithmic time complexity of a GCN layer is dictated by sparse matrix multiplication. Transforming the node features requires  $O(|V| \cdot d^2)$  operations, and propagating these features across the graph topology requires  $O(|E| \cdot d)$  operations, resulting in a total time complexity of  $O(|V| \cdot d^2 + |E| \cdot d)$  [6]. While computationally efficient for small-molecule ligands, extending GNNs to massive protein graphs introduces severe memory-access bottlenecks. Because graph topology is highly irregular, memory fetching on GPUs becomes non-contiguous, leading to hardware cache misses and suboptimal memory bandwidth utilization compared to the dense matrix operations of CNNs.

## III. THE SHIFT TO MULTIMODAL FUSION AND TRANSFORMERS

### A. Multimodal Integration Strategies

As the biological complexity of predictive modeling increased, architectures evolved from single-modality networks into multimodal fusion systems. These software pipelines are designed to ingest heterogenous data structures simultaneously—such as pairing a 2D GNN ligand encoder with a 1D CNN protein sequence encoder.

In MLOps pipelines, integrating these distinct data streams typically occurs via early fusion (concatenating raw inputs prior to feature extraction) or late fusion (concatenating high-level latent vectors immediately before a fully connected dense classifier). However, static concatenation fails to dynamically model the cross-modal physical interactions that define biochemical binding. To address this, software architects universally adopted attention mechanisms to actively weight the relevance of specific ligand atoms against specific protein residues during the forward pass.

### B. The Transformer Bottleneck: Quadratic Complexity

The introduction of the Transformer architecture revolutionized multimodal fusion by abandoning convolutions and recurrences entirely in favor of self-attention [7]. Transformers allow a neural network to explicitly learn the mathematical correlation between every token in a sequence or graph simultaneously.

The core computational engine of the Transformer is the scaled dot-product attention:

$$A = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)V$$

Where  $Q$  (Query),  $K$  (Key), and  $V$  (Value) are linear projections of the input embeddings.

However, from an algorithmic scaling perspective, the standard Transformer introduces a catastrophic hardware bottleneck. Computing the  $QK^T$  attention matrix for an input sequence of length  $N$  requires  $O(N^2 \cdot d)$  time complexity and necessitates  $O(N^2)$  space complexity just to store the attention weights in GPU VRAM [7] [8].

In natural language processing,  $N$  typically represents a few hundred sub-word tokens. In computational biology, a single protein sequence can easily exceed  $N = 2000$  amino acids. Fusing a ligand graph with a protein sequence via cross-attention directly compounds this scaling issue. When evaluating batch sizes necessary for stable gradient descent, the  $O(N^2)$  memory footprint rapidly exceeds the physical capacity of standard enterprise GPUs (e.g., NVIDIA A100 or H100), frequently resulting in Out-of-Memory (OOM) failures. This software limitation restricts high-throughput deployment and forces machine learning engineers to drastically reduce batch sizes, implement aggressive gradient accumulation, or manually truncate biological sequences—all of which degrade computational throughput and pipeline stability.

#### IV. HARDWARE BOTTLENECKS AND INFERENCE LATENCY IN VIRTUAL SCREENING

##### A. The Memory Wall in Computational Drug Design

As the bioinformatics community rapidly adopts Transformer architectures initially designed for Large Language Models (LLMs), a significant "memory wall" has emerged. The memory wall refers to the growing disparity between the computational speed of modern GPUs (FLOPs) and the bandwidth/capacity of GPU VRAM (High-Bandwidth Memory - HBM) [9].

In structural biology, encoding the 3D coordinates of a protein-ligand complex requires immense memory overhead. While the time complexity of attention mechanisms is quadratic ( $O(N^2)$ ), the physical space complexity is equally restrictive. For a protein containing 1,500 residues interacting with a 50-atom ligand, storing the intermediate activation tensors and attention maps for a deep multi-layer Transformer easily exceeds the 80GB VRAM limit of a single NVIDIA A100 GPU during backpropagation [8]. This hardware bottleneck forces software engineers into sub-optimal workarounds:

- **Sequence Truncation:** Artificially shortening biological sequences, which destroys long-range structural dependencies and degrades predictive accuracy.
- **Gradient Accumulation:** Processing micro-batches sequentially, which drastically increases the wall-clock time required for model training.

##### B. Inference Latency and the High-Throughput Screening Bottleneck

While training massive models is computationally expensive, inference latency presents a more critical bottleneck in the pharmaceutical pipeline. High-Throughput Virtual Screening (HTVS) pipelines are designed to evaluate libraries containing billions of synthesized or virtual compounds, such as the ZINC database.

If a multi-billion-parameter Transformer requires 0.5 seconds to predict the affinity of a single protein-ligand pair, evaluating a library of  $10^8$  compounds would require roughly 1.5 years of continuous inference on a single GPU. From a software architecture perspective, deploying an algorithm with such high latency fundamentally breaks the HTVS pipeline. Therefore, in the context of predictive drug discovery, an algorithm's value must be measured not merely by its Concordance Index (CI) or Mean Squared Error (MSE), but by its Pareto efficiency—the optimal trade-off between predictive accuracy and inference throughput (predictions per second).

#### V. MLOPS BEST PRACTICES: THE NEED FOR MULTI-RUN REPRODUCIBILITY

##### A. The Vulnerability of Single-Run Testing

A persistent software engineering failure in the deployment of computational biology models is the reliance on single-run validation. Deep neural networks, particularly when processing inherently noisy biological data (e.g.,  $K_i$  or  $K_d$  values from high-throughput wet-lab assays), are highly sensitive to stochastic variables, including weight initialization seeds, mini-batch shuffling, and hardware-specific floating-point execution orders.

When researchers report state-of-the-art results based on a single training run, it is mathematically impossible to distinguish true architectural generalizability from a statistical anomaly driven by a "lucky" random seed [10].

##### B. Establishing Rigorous Engineering Standards

To ensure the reliability of predictive software in critical healthcare domains, the bioinformatics community must adopt stricter Machine Learning Operations (MLOps) protocols. Moving forward, architectural superiority should only be claimed if supported by rigorous statistical averaging. We propose that standard software engineering practices for AI in computational biology must include:

1. **Multi-Run Validation:** Executing a minimum of 10 independent training runs per architectural configuration and reporting the averaged evaluation metrics alongside the standard deviation.
2. **Resource Profiling:** Mandating the publication of computational metrics alongside statistical metrics, including total parameter count, floating-point operations per second (FLOPs), and inference latency.

## VI. CONCLUSION

The evolution of multimodal deep learning architectures has provided unprecedented predictive power in computational biology. However, as this survey demonstrates, the uncritical adoption of massive, overparameterized architectures—such as standard quadratic Transformers—has introduced severe software engineering bottlenecks. The future of AI in drug discovery relies on breaking the memory wall and reducing inference latency. To achieve scalable, "Green AI" virtual screening pipelines, software architects must pivot away from brute-force parameter scaling. The next generation of predictive architectures must focus on resource optimization, lightweight cross-modal attention mechanisms, and uncompromising multi-run statistical reproducibility.

## REFERENCES

- [1] H. M. Berman *et al.*, "The protein data bank," *Nucleic acids research*, vol. 28, no. 1, pp. 235–242, 2000.
- [2] S. Kim *et al.*, "PubChem in 2021: new data content and improved web interfaces," *Nucleic acids research*, vol. 49, no. D1, pp. D1388–D1395, 2021.
- [3] J. Vamathevan *et al.*, "Applications of machine learning in drug discovery and development," *Nature reviews Drug discovery*, vol. 18, no. 6, pp. 463–477, 2019.
- [4] M. Bagherian, E. Sabeti, K. Wang, M. A. Sartor, Z. Nikolovska-Coleska, and K. Najarian, "Machine learning approaches and databases for prediction of drug–target interaction: a survey paper," *Briefings in bioinformatics*, vol. 22, no. 1, pp. 247–269, 2021.
- [5] H. Öztürk, A. Özgür, and E. Ozkirimli, "DeepDTA: deep drug–target binding affinity prediction," *Bioinformatics*, vol. 34, no. 17, pp. i821–i829, 2018.
- [6] T. N. Kipf and M. Welling, "Semi-Supervised Classification with Graph Convolutional Networks," Feb. 22, 2017, *arXiv*: arXiv:1609.02907. doi: 10.48550/arXiv.1609.02907.
- [7] A. Vaswani *et al.*, "Attention is all you need," *Advances in neural information processing systems*, vol. 30, 2017, Accessed: Jun. 12, 2026. [Online]. Available: <https://proceedings.neurips.cc/paper/2017/hash/3f5ee243547dee91fbd053c1c4a845aa-Abstract.html>
- [8] J. Jumper *et al.*, "Highly accurate protein structure prediction with AlphaFold," *nature*, vol. 596, no. 7873, pp. 583–589, 2021.
- [9] R. Schwartz, J. Dodge, N. A. Smith, and O. Etzioni, "Green AI," *Commun. ACM*, vol. 63, no. 12, pp. 54–63, Nov. 2020, doi: 10.1145/3381831.
- [10] T. Pahikkala *et al.*, "Toward more realistic drug–target interaction predictions," *Briefings in bioinformatics*, vol. 16, no. 2, pp. 325–337, 2015.