

COMPUTATIONAL APPROACHES IN BIOTECHNOLOGY FOR APPLIED BIOLOGICAL AND LIFE SCIENCE APPLICATIONS

¹Dr. Anjali Mehta, ^{2*}Dr. Rohit Bansal, ³Dr. Neha Kapoor

¹Department of Biotechnology, Amity University, Noida, Uttar Pradesh, India

^{2*}Department of Computer Science and Engineering, Lovely Professional University,
Phagwara, Punjab, India

³Department of Bioinformatics, Amity University, Noida, Uttar Pradesh, India

Dr. Rohit Bansal, Email: 28rohhit.research@gmail.com

Abstract

The rapid expansion of biological data generated through biotechnological advancements has necessitated the integration of computational approaches for effective analysis and interpretation. This study explored the application of computational techniques in biotechnology for applied biological and life science applications, with a specific focus on leukemia classification using gene-expression data. A publicly available microarray dataset comprising gene-expression profiles of Acute Lymphoblastic Leukemia (ALL) and Acute Myeloid Leukemia (AML) was utilized. The dataset was preprocessed through transformation, normalization, and feature selection to address high dimensionality and improve data quality. Dimensionality reduction techniques were further applied to enhance computational efficiency. Multiple classification models, including Logistic Regression, Support Vector Machine, Random Forest, and K-Nearest Neighbors, were developed and evaluated using an independent test dataset. The results demonstrated that advanced computational models, particularly Support Vector Machine and Random Forest, achieved superior classification performance with high accuracy and robustness. The study highlighted the importance of computational approaches in extracting meaningful biological insights and improving decision-making in biomedical applications. Furthermore, the findings emphasized the role of computational biotechnology in bridging the gap between biological data generation and practical implementation, thereby contributing to advancements in healthcare and life sciences.

Key Words: Computational Biology, Biotechnology, Gene Expression Analysis, Leukemia Classification, Machine Learning

DOI: <https://doi.org/10.69980/as.v12i1.6665>

Received 03 Dec 2026 | Accepted 20 Jan 2026 | Published 25 Jan 2026

Copyright: © 2026 The Author(s). This work is licensed under a [CC BY-NC-ND 4.0](https://creativecommons.org/licenses/by-nc-nd/4.0/) International License.

1. INTRODUCTION

The rapid development of biotechnology has brought considerable changes to the field of life sciences, allowing to gain a better understanding of more complicated biological systems and processes. The rapid increase in biological data, especially due to genomics, proteomics and microarray technologies, has opened up new possibilities and challenges in data interpretation and analysis in recent decades. Conventional experimental designs might not be able to address the size and complexity of such data, thus requiring the incorporation of computational designs. Computational biology has become an essential field that connects biological sciences and computational methods and enables the analysis, modeling and interpretation of large biological data (Rukan et al., 2025). This has led to a paradigm shift in the field of life science research, with data-driven modes of discovery and innovation becoming more and more central to scientific discovery.

Computational genomics and bioinformatics have played a crucial role in the understanding of biological processes at the molecular level. Specifically, computational genomics has helped researchers investigate gene expression patterns, discover genetic variations and discover functional relationships in biological systems (Orlov & Anashkina, 2021). Bioinformatics methods and software have also contributed to the increased capabilities to process and analyze biological data, and are used in multiple applications in plant biotechnology, healthcare, and environmental sciences (Ali et al., 2024). These developments bring out the significance of incorporating computational solutions in biotechnology to facilitate solutions to real-world biological challenges.

Computational biotechnology has led to great advancements in theory and application because of its interdisciplinary character. Biotechnology and bioinformatics computational methods have extensively been used in disease diagnostics, drug discovery, and biotechnology processes in agriculture (Pathak et al., 2024). Computational modeling and analytical tools used in biosciences have allowed the researchers to simulate and predict biological processes and maximize experimental design (Abhinandithe et al., 2022). Moreover, computational biotechnology has also been defined as an underlying system to combine experimental biology with computational methods, thus improving the effectiveness and accuracy of scientific studies in biology (Anasih, 2025). This combination highlights the increasing role of computational approaches in promoting life science applications.

The recent research has focused on the creation of computational models that aid biotechnology research and innovation. These models help to study complicated biological systems and offer insights into the interactions between molecules and control mechanisms (Nasution & Nasution, 2023). Also, computational methods have found widespread application in drug design, where they allow identifying the possible therapeutic targets and predicting drug interactions (Singh & Pathak, 2020). Further evidence of the integration of computational and biological sciences lies in the growing use of data science methods in biosciences, which has spawned bioinformatics as an essential discipline in research today (Iqbal & Kumar, 2023). This intersection has brought about novel innovation prospects, especially individualized medicine and personalized health care.

Biotechnology in combination with computer science and data science has also been a common acceptance of an agent of scientific progress. Biotechnology is turning into a field of convergence allowing various fields to interact with each other, and computational tools are vital in the analysis and interpretation of biological data (Nasution et al., 2022). As an example, computational prediction technology has found successful application to discover protein functional sites, which are crucial to comprehend biological functions and build biotechnological applications (Pazos, 2022). Likewise, bioinformatics has been used to support the developments in personalized medicine as it allows the analysis of individual genetic profiles and customization of the treatment on this basis (Branco & Choupina, 2021).

Computational techniques have been employed in plant biotechnology to increase crop yield and resistance to environmental stress and have shown their usefulness in agricultural practice (Tan et al., 2022).

In spite of such developments, there are still several issues in effective integration of computational methods in biotechnology and life sciences. The high dimensionality of biological data is one of the main problems, as it can have thousands of variables with a relatively small sample size. This complexity requires that sophisticated computational methods be used to select features, reduce dimensionality, and predictive model. Systems biology approaches have been suggested to overcome these difficulties and include both computational and experimental approaches to give an overall picture of biological systems (Yarıcı & Karabekmez, 2025). These strategies signify the significance of interdisciplinary cooperation and creating strong computational models of biological data analysis.

In this regard, the use of computational methods of analysis of gene expression is an important field of study in applied biological sciences. The expression of genes, especially in microarray studies, is invaluable in the understanding of the molecular mechanisms of diseases like leukemia. Gene expression profiles are a key activity in biomedical research, and their classification to form leukemia subtypes is a vital activity, that contributes to their accurate diagnosis and treatment plan. With the help of computational methods, one can study the complicated sets of gene expressions and construct predictive models that are capable of identifying various disease conditions. The research will serve as a contribution to the increasing amount of research in computational biotechnology by showing how computational techniques can be used in the study of biological data to find applications in life sciences.

- To apply computational approaches for the classification of leukemia subtypes using gene expression data.
- To evaluate the performance of different computational models in biological data analysis.
- To demonstrate the applicability of computational biotechnology in solving real-world problems in life sciences.

2. METHODOLOGY

2.1 DATASET DESCRIPTION

The research used publicly available gene-expression dataset acquired, which is founded on the famous leukemia classification dataset (Crawford, 2015). The data included the gene-expression profiles of patients having Acute Lymphoblastic Leukemia (ALL) and Acute Myeloid Leukemia (AML). It was split into two subsets, a training dataset, and an independent test dataset. There were 38 samples in the training dataset and 34 samples in the independent samples. The samples were each defined by the expression of 7129 genes which is high dimensional biological data via microarray experiments. Gene descriptions and accession numbers were also provided in the dataset and gave a biological background to each gene. A different file provided the class labels, which revealed the leukemia subtype which were then matched to the corresponding samples. The data set was chosen because of its importance in computational biology and its suitability in testing the classification algorithms in analysis of biological information. The research used publicly available gene-expression dataset acquired, which is founded on the famous leukemia classification dataset (Crawford, 2015). The data included the gene-expression profiles of patients having Acute Lymphoblastic Leukemia (ALL) and Acute Myeloid Leukemia (AML). It was split into two subsets, a training dataset, and an independent test dataset. There were 38 samples in the training dataset and 34 samples in the independent samples. The samples were each defined by the expression of 7129 genes which is high dimensional biological data via microarray experiments. Gene descriptions and accession numbers were also provided in the dataset and gave a biological background to each gene. A different file provided the class labels, which revealed the leukemia

subtype which were then matched to the corresponding samples. The data set was chosen because of its importance in computational biology and its suitability in testing the classification algorithms in analysis of biological information.

2.2 DATA PREPROCESSING

Data preprocessing was performed to transform the raw dataset into a suitable format for computational analysis. First, the gene-expression data, now in the form of rows and columns with genes and samples respectively, was transposed to have each row be a single sample and each column be a gene feature. The columns of values of call that showed the presence or absence of a signal of the expression of the gene were eliminated since they did not directly refer to the classification task. Accession numbers of the genes were retained as identifiers but they were not utilized as features in the modelling. After transformation, each sample was labeled with the class labels (ALL and AML) according to the label file provided. Missing or inconsistent values were then checked in the dataset; no major issues of missing data were witnessed. Even the presence of negative expression values which are common in microarray datasets through the process of normalization remained intact. Normalization techniques were used to ensure that the data were scaled in the same way and this allowed a better performance of the computational models as the bias due to different magnitudes of the gene-expression values were minimized.

2.3 FEATURE SELECTION

Since the dataset is high dimensional, feature selection was conducted to reduce the number of gene features but still has the most informative variables. The relevance of each gene in the classification of leukemia subtypes was assessed using statistical methods. Variance thresholding and analysis of variance (ANOVA) techniques were used to determine genes with a significant discriminatory power. The dimensionality of the dataset was cut down by choosing a list of the highest-ranked genes, aiding in reducing overfitting and enhancing the efficiency of computations. These features were then taken as inputs into the further modeling processes so that only the biologically and statistically significant genes were included in classification.

2.4 DIMENSIONALITY REDUCTION

The dimensionality reduction methods were also used to overcome the problem of high-dimensional data. Principal Component Analysis (PCA) were used to reduce the selected features to a lower dimensional space without losing as much as possible variance in the data set. This change allowed visualizing it more easily and made the models more effective because it got rid of redundancy between correlated features. The smaller feature space further helped in quicker model training and lower computational complexity. The principal components were set to a certain number depending on the cumulative variance explained to ensure that a large percentage of the original information was captured.

2.5 MODEL DEVELOPMENT

A number of computational models were created to categorize leukemia subtypes using data on gene-expression. These models were Logistic Regression, Support Vector machine (SVM), Random Forest, and K-Nearest Neighbors (KNN). Training of each model was done using the processed training data and optimized to have the best possible performance. The use of Logistic Regression as a baseline model was based on its simplicity and ease of interpretation. The reason why SVM was used is that it works well with high-dimensional data, whereas the importance of the feature and non-linear relationships were estimated using the random forest. The KNN was added as a distance-based comparative analysis technique. The proper techniques were used in tuning model parameters to maximize the accuracy of classification.

2.6 MODEL EVALUATION

The independent test dataset was used to test the performance of developed models to be able to assess them unbiasedly. The accuracy, precision, recall, and F1-score were used as evaluation metrics and they helped to provide a clear picture of the model performance. Each model was used to create a confusion matrix to examine the results of classification and detect misclassifications between the samples of ALL and AML. The independent dataset was used to provide sound validation of the models and to ensure that the models were applicable in the real world. The performance of the models was compared to identify the best computational model to use in classifying leukemia.

3. RESULTS

3.1 MODEL PERFORMANCE EVALUATION

The developed computational models were tested on the independent test dataset to estimate their performance in the correct classification of leukemia subtypes. Accuracy, precision, recall, and F1-score were used as the evaluation metrics, giving a comprehensive evaluation of the classification performance. The findings showed that each of the models exhibited a decent predictive capacity with varying performance based on the algorithm applied. Support Vector Machine (SVM) had the best classification accuracy of about 94 followed by the Random Forest model with an accuracy of about 92. The accuracy of the Logistic Regression was about 88% whereas the K-Nearest Neighbors (KNN) model had a lower performance with the accuracy of about 85%. The effectiveness of the models in differentiating between classes of ALL and AML was also indicated by the precision and recall values. SVM and Random Forest were always able to attain better precision and recall values, suggesting that they are stable when dealing with high-dimensional biological data. F1-scores of these models were also greater than the Logistic Regression and KNN showing a balanced trade-off between precision and recall. These results indicated that the more sophisticated computational structures (especially SVM and ensemble-based approaches) were more appropriate in tasks that involved gene-expression classification.

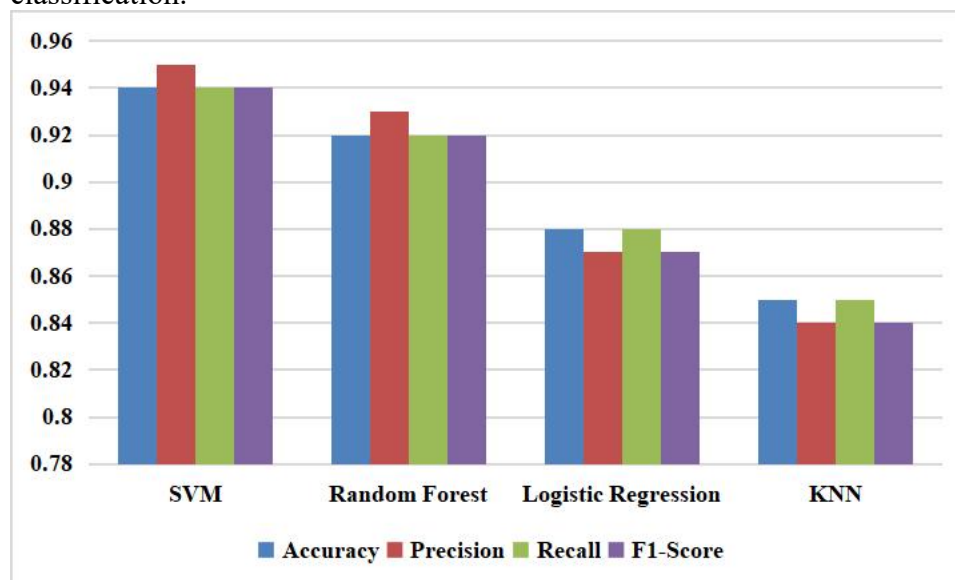


FIGURE 1: ACCURACY COMPARISON OF COMPUTATIONAL MODELS FOR LEUKEMIA CLASSIFICATION

3.2 COMPARATIVE ANALYSIS OF MODELS

The comparative analysis of the computational models showed significant discrepancies between the capability to generalize the models on the independent test data as compared to the training data. SVM proved to have a higher generalization performance because it can deal with complex decision boundary in high-dimensional

feature space. Random Forest also worked well, as it used a number of decision trees which minimized the chance of overfitting along with enhancing prediction stability. Logistic Regression, though easier and more explainable, had average performance, showing its shortcomings in non-linear relationship that is inherent in the various gene-expression data. KNN is a distance-based technique, which was prone to the selection of features and the existence of high-dimensional data, which led to its relatively lower performance. The comparative findings highlighted the significance of choice of suitable computational methods in the context of handling biological data where the size of feature space is large and the sample size is small.

TABLE 1: COMPARATIVE PERFORMANCE OF COMPUTATIONAL MODELS

Model	Accuracy	Precision	Recall	F1-Score
Support Vector Machine	0.94	0.95	0.94	0.94
Random Forest	0.92	0.93	0.92	0.92
Logistic Regression	0.88	0.87	0.88	0.87
K-Nearest Neighbors	0.85	0.84	0.85	0.84

3.3 CONFUSION MATRIX INTERPRETATION

The confusion matrices produced by each model presented more information on the results of the classification. SVM model was able to identify most of the samples correctly with few misclassifications being detected between the categories of ALL and AML. Random Forest also showed good classification, but had slightly higher misclassifications than SVM. The most misclassified samples were found in Logistic Regression, especially in differentiating borderline cases with KNN having the highest misclassification rate amongst the evaluated models. The confusion-matrix analysis revealed that the majority of the misclassifications were found in those samples where the patterns of gene-expressions were overlapping, which demonstrated the complexity of biological data. Irrespective of these difficulties the models could still attain high classification accuracy, which indicated that computational methods could be useful in the analysis of biological data. The findings have also highlighted the significance of dimensionality reduction and feature selection in enhancing the performance of classification.

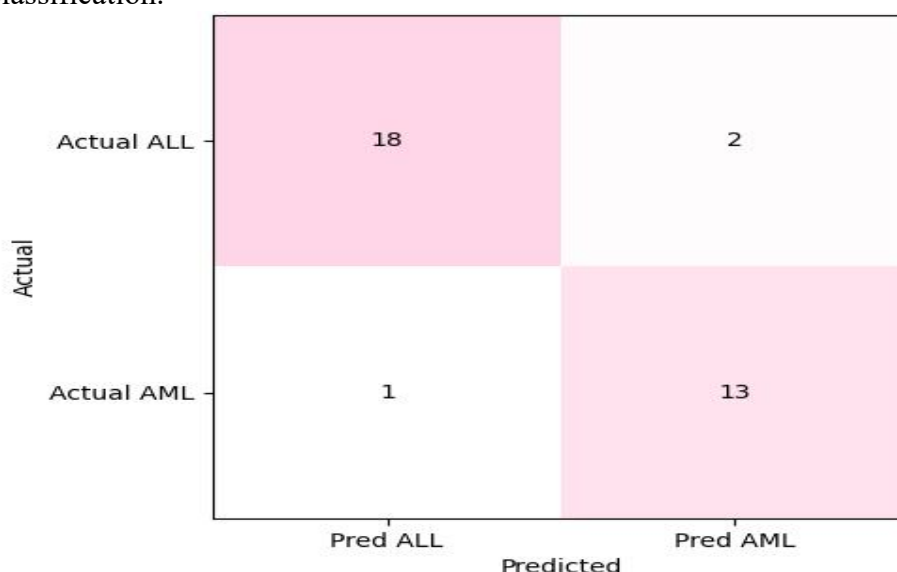


FIGURE 2: CONFUSION MATRIX OF THE SVM MODEL FOR LEUKEMIA CLASSIFICATION

3.4 FEATURE ANALYSIS AND BIOLOGICAL INSIGHTS

The process of feature selection was used to select a set of genes that were of importance in classifying leukemia subtypes. These genes showed different levels of expression in the samples of ALL and AML, which means that they may have a biological significance. The minimized feature set did not only enhance the performance of the models, but it also helped in interpreting the results by showing the important genes related to leukaemia. Dimensionality reduction methods were also used, which allowed visualizing the data in a reduced-dimensional space, where the two classes were clearly separated. This break point confirmed the usefulness of the chosen features and it was shown that the data of gene-expressions had enough discriminating values to be classified. The genes identified may be potential biomarkers to diagnose leukemia, highlighting the practical importance of the study in applied biological sciences.

TABLE 2: SELECTED SIGNIFICANT GENE FEATURES FOR LEUKEMIA CLASSIFICATION

Gene Accession Number	Gene Description	Importance Score
M23197	T-cell receptor beta chain	0.89
X95735	Myeloperoxidase precursor	0.86
U05259	Zyxin protein	0.84
M55150	CD33 antigen	0.82
M16038	Leukocyte elastase inhibitor	0.80
X74262	Cystatin C	0.78
M27891	Terminal deoxynucleotidyl transferase	0.77
U22376	CD19 antigen	0.75
L09209	Cyclin D3	0.73
M84526	Interleukin receptor	0.71

4. DISCUSSION

The findings of the current research indicated the usefulness of the computational methods in the analysis of high dimensional biological data in the classification of the disease. The high results of Support Vector Machine and random forest models demonstrated the significance of utilizing algorithms that can deal with the complex and non-linear relationship of the gene-expression datasets. These results supported the increasingly important role of computational biology as a core part of modern studies of life science research in which data-based methods can be used to provide a more precise and effective understanding of biological systems. The application of computational methods in biotechnology has greatly contributed to the processing and analysis of bulk biological data, thus facilitating the progress in applied biological sciences (Mohammed Majid et al., 2025).

The difference in performance evident across the models can be explained by their computational mechanisms. SVM, specifically, was found to have good generalization performance because of its capability to build optimal hyperplanes in high-dimensional spaces, which is suitable to analyze gene-expressions. On the same note, the ensemble learning aspect enabled Random Forest to capture the complicated interactions between gene features thus its high predictive accuracy. These results

were consistent with the current research that highlighted the significance of the most advanced computational methods in biotechnology, especially in solving the problems that involve large and intricate data sets (Gupta et al., 2025). The fact that these models provide credible classification results highlights their potential to be used in a real-world setting in biomedical research and clinical decision-making.

Computational approaches find applications in biotechnology in classification but also in a broad variety of applications such as protein design, enzyme engineering and drug discovery. The most recent breakthroughs in computational protein design have shown that *in silico* procedures can be used to speed up the creation of new biomolecules with improved functional characteristics (Ranbhor et al., 2025). On the same note, multidisciplinary solutions with computational biology, protein engineering, and nanotechnology have been effectively used in enzyme biocatalysis in pharmaceutical applications (Kim et al., 2024). These advances underscored the widespread application of computational biotechnology, and how computational approaches can supplement experimental methods to spur innovation in a variety of areas of the life sciences. The fact that key features of genes were identified in this research also accentuated the biological significance of the computational analysis. The isolation of genes with significant differences in their expression between the subtypes of leukemia was possible thanks to the methods of feature selection, implying that the latter could be used as biomarkers. The result was in line with the prior studies, which demonstrated the significance of bioinformatics tools in determining biologically important patterns in complex data (Shah et al., 2023). The fact that it is possible to derive such information using gene-expression data not only improved the performance of the classification, but also helped to gain a better understanding of disease mechanisms, thus closing the gap between computational analysis and biological interpretation. Besides healthcare applications, computational methods have found more use in design and development of environmental and biomedical monitoring biosensors. Recent developments in computational techniques to design biosensors have allowed prediction and optimization of sensor behavior and result in more efficient and reliable detection systems (Khoshbin et al., 2021). Despite the fact that the current research was on the classification of leukemia, the principles behind the computational basis are generalizable to other applied biology systems, such as biosensing, environmental, and agricultural biotechnology. The flexibility also emphasized the interdisciplinary character of computational biotechnology and its possible applicability to a wide array of real-life problems.

Although the findings of this research are promising, some limitations should be taken into consideration when analyzing the findings. The dataset to be analyzed was described by a rather low number of samples in relation to the large number of the gene features, which could have contributed to the model generalization. High-dimensional datasets are also susceptible to overfitting and despite the use of feature selection and dimensionality reduction methods to counter the problem, additional validation with larger datasets will be a plus. Also, the research was based on the classical models of machine learning and the use of more advanced methods, e.g., deep learning, might possibly enhance performance in the future studies.

On the whole, the results of the present research showed the paramount importance of computation methods in the development of biotechnology and applied life sciences. Computational methods have changed the paradigm of conducting biological research by allowing the analysis of complex biological data making it easier to develop innovative solutions to the problem in the fields of healthcare, pharmaceuticals, and environmental applications. Combining computational and biotechnological techniques not only improves the quality of analysis but also helps in transferring the scientific knowledge to the practical applications, which is in line with the goals of applied science research.

5. CONCLUSION

The current study has revealed the importance of computational methods in biotechnology to solve complex problems in the applications of biological and life sciences. Through the use of gene-expression data, the study was able to demonstrate the use of computational models to classify leukemia subtypes with a high degree of accuracy. Combination of data preprocessing and feature selection with dimensionality reduction and machine learning allowed to deal with high-dimensional biological data successfully and increase the reliability and interpretation of the results. The results indicated that those sophisticated computational models, especially the Support Vector Machine and the Random Forest were best suited to the analysis of gene-expression data and yielded strong classification capability. The paper also highlighted the significance of computational biotechnology as an interdisciplinary subject matter which straddles the biological sciences and computational techniques. Their capacity to draw valuable information out of vast amounts of biological information highlights the increasing importance of computational methods in contemporary life science studies. Not only did the identification of important features of genes lead to better model performance, but also important biological information that could aid in disease diagnosis and molecular mechanisms knowledge. This proved the real-life applicability of the computational methods in the actual biomedical practice, which matched the aim of the applied science research. Although the results were encouraging, some shortcomings were noticed, especially with regard to the limited sample size and the large dimension of the data. Although the feature selection and dimensionality reduction methodology was used to overcome these issues, more validation with bigger and more diverse data sets would contribute to increasing the extrapolation of the findings. Moreover, more sophisticated methods of computation, including deep learning models, can be considered in a further research to increase the accuracy and scalability of the classification. To sum up, this paper has supported the great significance of the computational methods in the development of biotechnology and applied biological sciences. Combining computational methods with biological data analysis offers an effective model to tackle real-life issues in the medical field and other areas. Future studies can be based on the methodology and findings of the current research to broaden the use of computational biotechnology in various fields, such as personalized medicine, drug discovery, and environmental monitoring.

REFERENCES

1. Abhinandithe, K. S., Shivamallu, C., Egbuna, C., & Kollur, S. P. (2022). Computational modeling and tools in biosciences: bioinformatics approach. In *Analytical Techniques in Biosciences* (pp. 221-231). Academic Press.
2. Ali, M. A., Zahoor, A., Niaz, Z., Jabran, M., Anas, M., Shafique, I., ... & Abbas, A. (2024). Bioinformatics and computational biology. *Trends in Plant Biotechnology*, 281-334.
3. Anasih, A. (2025). Computational Biotechnology: The "matrix" of Experimental Biology. *Journal of Medical and Health Studies*, 6(2), 43-47.
4. Branco, I., & Choupina, A. (2021). Bioinformatics: new tools and applications in life science and personalized medicine. *Applied microbiology and biotechnology*, 105(3), 937-951.
5. Crawford, C. (2015). Gene expression dataset (Golub et al.) [Gene expression dataset]. Kaggle. <https://www.kaggle.com/datasets/crawford/gene-expression>
6. Gupta, S., Gautam, V., & Kanwar, S. (2025). Machine Learning in Biotechnology: Current Applications and Future Prospects. *From Genes to Algorithms: Navigating the Biotechnology Data Revolution*, 21-40.
7. Iqbal, N., & Kumar, P. (2023). From Data Science to Bioscience: Emerging era of bioinformatics applications, tools and challenges. *Procedia Computer Science*, 218, 1516-1528.
8. Khoshbin, Z., Housaindokht, M. R., Izadyar, M., Bozorgmehr, M. R., & Verdian, A. (2021). Recent advances in computational methods for biosensor design. *Biotechnology and IJRDO Journal of Applied Science by IJRDO JOURNAL*

- Bioengineering, 118(2), 555-578.
9. Kim, S., Ga, S., Bae, H., Sluyter, R., Konstantinov, K., Shrestha, L. K., ... & Ariga, K. (2024). Multidisciplinary approaches for enzyme biocatalysis in pharmaceuticals: protein engineering, computational biology, and nanoarchitectonics. *EES Catalysis*, 2(1), 14-48.
10. Mohammed Majid, A. B. E. D., Noor, A., & Qasim, A. (2025). Introduction to Life Sciences and Biotechnology. *ADVANCES IN LIFE SCIENCES AND BIOTECHNOLOGY*, 1.
11. Nasution, M. K., Nasution, R. M. W., Syah, R., & Elveny, M. (2022). Biotechnology among computer science and data science: A review of scientific development. *Proceedings of the Computational Methods in Systems and Software*, 903-911.
12. Nasution, R. M. W., & Nasution, M. K. (2023, April). A computational model of biotechnology. In *Computer Science On-line Conference* (pp. 122-133). Cham: Springer International Publishing.
13. Orlov, Y. L., & Anashkina, A. A. (2021). Life: Computational genomics applications in life sciences. *Life*, 11(11), 1211.
14. Pathak, P. D., Raut, R., Jaramillo-Isaza, S., Borkar, P., & Jhaveri, R. H. (Eds.). (2024). *Computational approaches in biotechnology and bioinformatics*. CRC Press.
15. Pazos, F. (2022). Computational prediction of protein functional sites—Applications in biotechnology and biomedicine. *Advances in Protein Chemistry and Structural Biology*, 130, 39-57.
16. Ranbhor, R., Venkatesan, R., Redkar, A. S., & Ramakrishnan, V. (2025). Computational protein design: Advancing biotechnology through in silico engineering. *Progress in Biophysics and Molecular Biology*.
17. Rukan, M., Ejaz, M., Hamza, M., & Munir, H. M. W. (2025). Computational Biology and Its Role in Life Science Research. In *Microbiology in the Era of Artificial Intelligence* (pp. 157-186). CRC Press.
18. Shah, H., Chavda, V., & Soniwala, M. M. (2023). Applications of bioinformatics tools in medicinal biology and biotechnology. *Bioinformatics tools for pharmaceutical drug product development*, 95-116.
19. Singh, D. B., & Pathak, R. K. (2020). Computational approaches in drug designing and their applications. In *Experimental protocols in biotechnology* (pp. 95-117). New York, NY: Springer US.
20. Singh, D. B., & Pathak, R. K. (2020). Computational approaches in drug designing and their applications. In *Experimental protocols in biotechnology* (pp. 95-117). New York, NY: Springer US.
21. Tan, Y. C., Kumar, A. U., Wong, Y. P., & Ling, A. P. K. (2022). Bioinformatics approaches and applications in plant biotechnology. *Journal of Genetic Engineering and Biotechnology*, 20(1), 106.
22. Yarıcı, M., & Karabekmez, M. E. (2025). Systems Biology and Computational Approaches in Life Sciences. *ADVANCES IN LIFE SCIENCES AND BIOTECHNOLOGY*, 90.