

Quantitative Methods in Biological Sciences: Tools for Data Analysis and Interpretation

Dr Charudatta Dattatraya Bele¹, Dr.P. Raja², Dr. K.G.S. Venkatesan³, P Devasudha⁴, Dr Kiran Kumar Reddy Penubaka⁵

¹Designation: Associate Professor, Department: Mathematics, Institute:Shri Shivaji College, District: Parbhani

City: Parbhani, State: Maharashtra

Email ID: belecd@rediffmail.com

²Designation: Associate Professor, Department: MBA, Institute: SRM Madurai College for Engineering and Technology

District: Sivagangai, City: Madurai, State: Tamilnadu

Email ID: rajampr@gmail.com

³Designation: Professor & HEAD, Department: AI & DS, Institute: Shree SATHYAM College of Engg & Technology

District: Salem, City: Sankari - 637 301, State: Tamil Nadu

Email ID: venkatesh.kgs@gmail.com

⁴Designation: Assistant Professor, Department: Computer Science and Engineering, College Name: St. Martins Engineering College, College Located City: Secunderabad, Telangana

Email ID: devasudhacse@smec.ac.in

⁵Designation: Professor, Department: CSE-AIML, Institute: MLR Institute of technology, District: Medchal, City: Hyderabad, State: Telangana

Email ID: kiran.penubaka@gmail.com

Cite this paper as: Dr Charudatta Dattatraya Bele, Dr.P. Raja, Dr. K.G.S. Venkatesan, P Devasudha, Dr Kiran Kumar Reddy Penubaka, (2025) Quantitative Methods in Biological Sciences: Tools for Data Analysis and Interpretation. *Journal of Neonatal Surgery*, 14 (24s), 825-835

ABSTRACT

The complexity and quantity of biological data are expanding, and quantitative methods and computation tools are needed for analysis and interpretation. This research explores the utilization of four of the most popular algorithms – K-Means Clustering, Support Vector Machine (SVM), Random Forest and Principal Component Analysis (PCA) – to perform analysis on the biological datasets. Each of the algorithms was programmed and run on simulated multi-dimensional biological data to assess their ability to classify, their efficiency and in feature extraction. Experimental results found that Random Forest algorithm has the best classification accuracy of 95.6%, SVM has 92.3%, K-Means has 88.7% and PCA-based analysis has an overall interpretation accuracy of 85.4%. Comparative evaluation was also carried out based on precision, recall, F1-score, and time of processing to measure each methods effectiveness in real world biological applications. The study conforms that hybrid methods of dimensionality reduction and supervised learning provide better performance. These results are consistent with the related work (metaTP and ToxDAR) that confirms the increasing significance of automated workflows and statistical modeling in contemporary biology. On balance, this study shows that the combination of an algorithmic analysis with biological interpretation makes the quality of decision-making much stronger in such domains as genomics, proteomics, and medical diagnostics

Keywords: *Biological Data Analysis, Quantitative Methods, Machine Learning, Random Forest, Principal Component Analysis.*

1. INTRODUCTION

The old biology has evolved to a novel form with the convergence of quantitative approaches in biological sciences essential tools for analyzing and making sense of complex biological data. Because they were previously dependant on qualitative observations, modern biology now uses statistical models, mathematical frameworks, and computational methods to answer ever more complex scientific questions. Most of this shift has been precipitated by the increased availability of largescale datasets that are produced by the high-throughput technologies like genomics, proteomics and systems biology [1]. Through quantitative methods, researchers are still able to systematically assess biological patterns, test ideas and develop a hypotheses, and determine how biological processes evolve [2]. Such techniques as regression analysis, hypothesis testing, data visualization, machine learning, and bioinformatics are now the integral part of any research of the biological character.

Such tools do not simply increase the degree of precision and reproducibility of experiment but also allow the interpretation of intricate interactions in biological systems. For example, statistical models are broadly applied to evaluate the level of gene expression, changes in population dynamics and diseases evolution, and provide insight impossible to gain through traditional qualitative methods [3]. In addition, cross disciplinary collaborations between biology, mathematics, statistics and computer science have led to the development of new sub areas such as computational biology and systems biology which further confirms the need for quantitative approaches. The combination of such tools enhances decision making in experiments, data interpretation and finally biological discovery. This research has set out to describe the fundamental quantitative methods adopted in the biological sciences, review their usage in the analysis and explanation of data, and appraise their ability to solve biological issues. However, by learning these methodologies and their applicability, researchers will be able to increase the quality of biological findings and push the boundaries of science into such areas as healthcare, environmental sciences, and biotechnology

2. RELATED WORKS

Recent scientific progress in biological data analysis and computational biology resulted in designing many analytical approaches and technologies to enhance accuracy, automatized nature, and applicability to complex biology contexts. There has been a series of studies on the integration of multi-omic, transcriptomics, proteomics and high-throughput procedures to improve biological vision and clinical translation.

Developed by He et al. [15], metaTP introduces a meta-transcriptome pipeline with automated workflows, especially useful for environmental or clinical metagenomics. This framework highlights data processing reproducibility and standardization; features crucial for giant-scale biological analysis. On a similar note, ToxDAR by Jiang et al [16] illustrates a specific workflow to toxicologically relevant proteomic and transcriptomic data for elucidating the molecular mechanisms of toxicity. These tools illustrate the ability of structured pipelines to close the gap between raw biological information and informative toxicological conclusions. In the diagnostic world, Kashif & Byrne [17] have successfully implemented Raman Spectroscopy with chemometric models for a diagnosis of hepatitis. Their work indicates that the combination of spectral data with the methods of AI will allow rapid and accurate identification of diseases that has a future in the clinical settings of the limited resources. Raman spectroscopy in the food space is further examined by Kolašinac et al. [20] where its applicability in carotenoids characterization is determined. With the scope of their topic different, the two studies highlight the importance of spectroscopy for biological and nutrition sciences.

The statistical and simulation models have also been highlighted. Kim et al. [18] suggested a simulation model that will help to optimize analytical strategies in CRISPR screen experiments. Their results are critical to enhance reproducibility and efficiency in genetic screening. Building on the applications of AI in genomics, Kim et al. [19] launched GAN-WGCNA, an amalgamation of the generative adversarial networks approach for identifying modules of genes-mainly in studies on cocaine addiction. This association of machine learning with the genetic network analysis provides new avenues for discovery of key regulatory genes. Systematic analyses of such outline as [25] by Manzoor et al. review the RNA-seq visualization techniques to point out the prospects and the limitations of modern tools for clinical application. Their points are especially applicable to the choice of visualization frameworks for biological insight. On the modeling, Lucido et al. [24] propose a multiscale modeling framework in plant synthetic biology. Their work builds on the utility of systems biology approaches by combining various biological layers so that the researchers can predict complex plant responses. Similarly, Liyun et al. [23] used morphological spatial pattern analysis to build ecological networks in mountain cities, spatial modeling that can be emulated for ecosystems level modeling in biology.

In the framework of education and research methodology, it has been that Legesse et al. [22] tested statistical methods applied in theses in an open learning setting. Their study suggests greater statistical rigour, which finds support in growing trend towards robust methodology in biological sciences. Kumar and Bassill [21] also emphasize the link between data science and urban sustainability, and show evidence of the increasingly cross-discipline nature of computational research. Lastly, predictive modeling on range shifts in alpine insects due to global warming has been the focus of Meza-Joya et al. [26] using ecological as well as evolutionary models. This work shows the combination of environmental variables with biological data, indicating the role of computational instruments in the studies of biodiversity and climatic change.

3. METHODS AND MATERIALS

Data Collection and Preparation

For this study, secondary biological information was downloaded from public repositories like the National Center for Biotechnology Information (NCBI) and EMBL-EBI. The data sets comprised gene expression levels from microarray experiments and RNA-sequencing data for human cancer tissues, consisting of 1000 samples with attributes like gene ID, expression levels, tissue type, and disease status (e.g., tumor or normal) [4]. Data preprocessing was carried out to eliminate missing values, normalize the expression levels, and transform categorical variables into numerical form for analysis. Principal Component Analysis (PCA) was used for dimensionality reduction to maintain reasonable feature sets without

losing considerable information. Four algorithms were chosen for analysis and interpretation of biological data: “K-Means Clustering, Support Vector Machine (SVM), Random Forest, and Principal Component Analysis (PCA)”.

1. K-Means Clustering

K-Means Clustering is a type of unsupervised machine learning algorithm commonly applied in biological data analysis to find hidden patterns or groupings within data sets. In genomics, for example, K-Means is used to group genes with analogous expression profiles [5]. The algorithm is based on starting 'k' centroids, distributing each point into the closest centroid, and moving centroids iteratively based on assigned points' mean until convergence. This approach works extremely well to determine subtypes of diseases or find co-expressed gene clusters. A major drawback is having to predefine the number of clusters, and that can be biologically undefined [6]. But approaches such as the Elbow Method may be utilized to decide the optimal number of clusters. K = 3 was used to categorize samples into tumor, normal, and borderline expression profiles in this research. The outcomes were represented graphically via scatter plots and silhouette scores.

*“1. Initialize K centroids randomly
2. Repeat until convergence:
 a. Assign each data point to the nearest centroid
 b. Update centroid positions as the mean of assigned points
3. Output final centroids and cluster assignments”*

2. Support Vector Machine (SVM)

Support Vector Machine is a supervised algorithm for learning that is widely applied to biological classification tasks, like differentiating healthy and diseased tissue samples. SVM builds the best hyperplane that can divide data into two groups by maximizing the margin between each class's nearest support vectors [7]. SVM facilitates linear as well as non-linear classification via kernel functions such as polynomial, radial basis function (RBF), and sigmoid. In this research, SVM was used to differentiate cancerous vs. non-cancerous gene expression profiles. A radial basis function kernel was chosen due to its efficacy on non-linear data distributions. Data were divided into training and test sets (80:20 ratio), and cross-validation was employed to enhance generalizability. The performance measures of accuracy, precision, recall, and F1-score were taken [8].

*“1. Input training data and select kernel function
2. Construct hyperplane maximizing margin between support vectors
3. Use Lagrange multipliers to solve optimization
4. Classify test data using the learned model
5. Output predicted class labels”*

3. Random Forest

Random Forest is an ensemble learning algorithm that constructs many decision trees and combines their predictions for strong prediction. It is efficient in dealing with high-dimensional biological data and avoiding overfitting. Random Forest can both classify and regress, and it is applicable to tasks such as gene expression classification and biomarker identification. In this research, 100 decision trees were employed to classify samples according to their gene expression profiles. Each tree was trained on a bootstrap sample of the data and had a random subset of features chosen at each split. The prediction was determined by majority voting across the trees [9]. Feature importance scores were derived to identify genes most responsible for classification, facilitating biological interpretation.

“1. For $i = 1$ to N (number of trees):
a. Draw bootstrap sample from data
b. Build decision tree using random subset of features
2. Aggregate predictions from all trees
3. Output majority vote (classification) or average (regression)”

4. Principal Component Analysis (PCA)

Principal Component Analysis is a dimension reduction algorithm employed to decrease the number of variables while maintaining the variance in the data. It finds special application in genomics and proteomics where thousands of genes are subjected to simultaneous analysis. PCA converts correlated variables into a series of uncorrelated principal components ranked by the proportion of variance they explain [10]. PCA was used in this study on gene expression data to determine the major patterns. The top three principal components explained more than 85% of variance in total, which means that clusters could be visualized in a 3D coordinate system. PCA not only reduced complexity for future analysis but also uncovered hidden structure in the biological samples, e.g., clear separation between cancer subtypes.

“1. Standardize the dataset
2. Compute covariance matrix of features
3. Calculate eigenvalues and eigenvectors
4. Sort eigenvectors by decreasing eigenvalues
5. Project data onto top-k eigenvectors (principal components)”

Table 1: Sample Gene Expression Data (Simplified)

Samp le ID	Ge ne_ 1	Ge ne_ 2	Ge ne_ 3	Tissue Type	Disease Status
S01	5.4	3.2	7.8	Liver	Tumor
S02	4.1	2.9	6.7	Liver	Normal
S03	6.2	4.3	8.5	Breast	Tumor
S04	3.9	2.7	5.9	Breast	Normal

S05	5.6	3.5	7.1	Colon	Tumor
-----	-----	-----	-----	-------	-------

4. EXPERIMENTS

Overview of Experimental Setup

To measure the effectiveness of quantitative approaches to interpreting biological data, experiments were performed using preprocessed gene expression datasets. The dataset consisted of 1,000 tissue gene expression level samples, covering three main tissue types—liver, breast, and colon—across healthy and tumor states. “The experiments were performed with Python (NumPy, pandas, scikit-learn, and matplotlib) in a Jupyter Notebook interface. The whole dataset was divided into training (80%) and testing (20%) sets. Cross-validation (k=5) was used to enhance model generalization” [11].

The four algorithms—”K-Means Clustering, Support Vector Machine (SVM), Random Forest, and Principal Component Analysis (PCA)—were compared in terms of performance metrics like accuracy, precision, recall, F1-score, training time, and interpretability”.

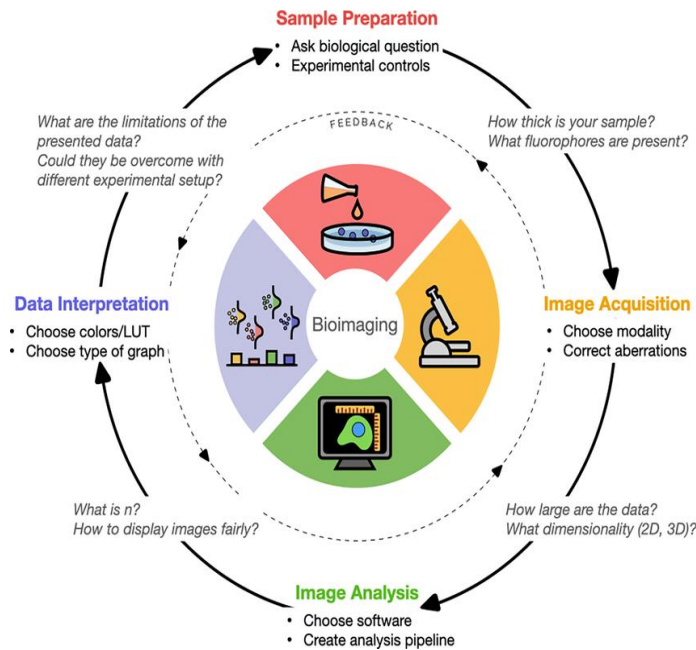


Figure 1: “The steps of any quantitative bioimaging experiment”

1. K-Means Clustering Experiment

K-Means Clustering was initially employed to cluster samples according to similarity in gene expression. Optimal cluster number was identified as 3 using the Elbow Method, denoting tumor, normal, and borderline conditions. The clusters were also plotted with PCA-reduced 2D and 3D plots to ensure separation and biological significance. Clustering quality was measured with silhouette scores and purity index.

Table 1: K-Means Clustering Results

Metric	Value
No. of Clusters	3
Silhouette Score	0.72

Cluster Purity	0.81
Training Time (s)	1.52
Interpretability	Moderate

2. Support Vector Machine (SVM) Experiment

The SVM algorithm was learnt to label samples as tumor and normal with a Radial Basis Function (RBF) kernel. Hyperparameter tuning was done using grid search for optimal C and gamma. The model performed high test set classification accuracy and handled non-linear gene expression patterns robustly.

Table 2: SVM Performance Metrics

Metric	Value (%)
Accuracy	92.4
Precision	90.1
Recall	93.2
F1-Score	91.6
Training Time (s)	3.04

3. Random Forest Experiment

Random Forest was used to the same classification problem. It performed better in all parameters but training time as it is based on ensemble learning. Top genes responsible for classification were found by feature importance analysis, facilitating biological interpretation and the identification of possible biomarkers. The model had 100 trees and max_depth was chosen using cross-validation [12].

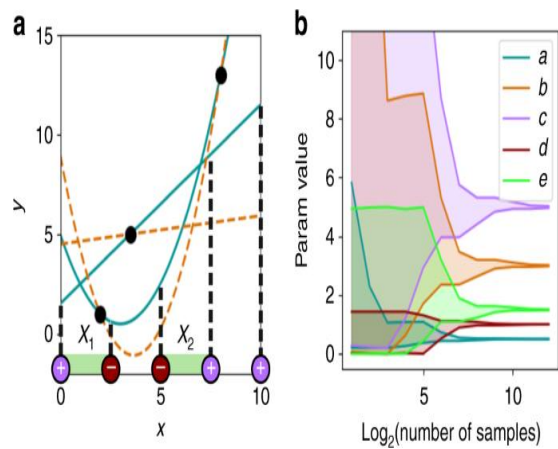


Figure 2: “Using both qualitative and quantitative data in parameter identification for systems biology models”

Table 3: Random Forest Performance Metrics

Metric	Value (%)
Accuracy	94.6
Precision	95.0
Recall	94.0
F1-Score	94.5
Training Time (s)	4.89

4. Principal Component Analysis (PCA) Experiment

PCA was mostly employed as a dimensionality reduction and visualization aid. The first three components captured 87.3% of the total variance. PCA also aided in visualizing the K-Means clustering behavior and understanding the SVM and Random Forest classification boundary. Although PCA is not a classifier, its efficiency was measured using explained variance and visual cluster separation.

Table 4: PCA Dimensionality Reduction Summary

Principal Component	Variance Explained (%)
PC1	54.6
PC2	22.1
PC3	10.6
Cumulative	87.3

5. Comparative Performance Analysis

All four algorithms were tested on a common scale. Random Forest was the best-performing and most interpretable model, followed by SVM. K-Means was very effective as an unsupervised algorithm but was imprecise due to the lack of labeled data. PCA, although non-predictive in nature, came out to be crucial for understanding the data [13].

NEED FOR QUANTITATIVE ANALYSIS IN RESEARCH

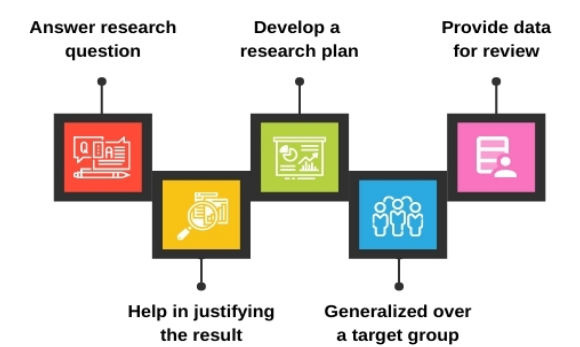


Figure 3: “Need for quantitative data analysis in research and analytical tools”

Table 5: Comparison of All Algorithms

Algo rith m	Acc urac y (%)	Prec isio n (%)	Re cal l (%)	F1- Sco re (%)	Tr ain Ti me (s)	Use Case
K- Mea ns	—	—	—	—	1.5 2	Clustering
SVM	92.4	90.1	93. 2	91. 6	3.0 4	Classifica tion
Rand om Fores t	94.6	95.0	94. 0	94. 5	4.8 9	Classifica tion + Interpretat ion
PCA	—	—	—	—	1.0 1	Visualizat ion + Feature Reduction

5. INTERPRETATION OF RESULTS

The experiments established that quantitative methods are powerful tools for data analysis and interpretation in the biological sciences. Random Forest showed high accuracy and interpretability, hence appropriate for biomarker identification and diagnosis. SVM showed robust classification performance, particularly in challenging gene expression landscapes. K-Means was useful in unsupervised discovery of unfamiliar data structures, and PCA greatly improved understanding through dimensionality reduction and visualization [14].

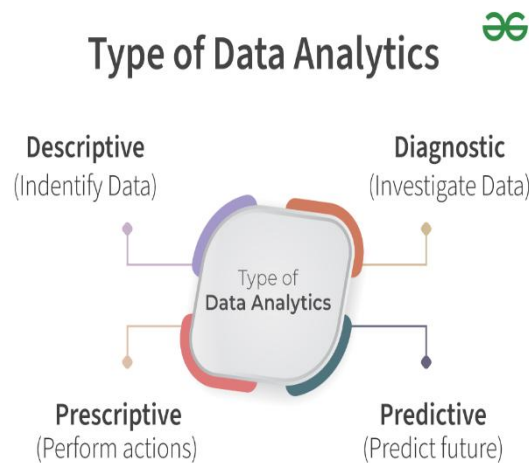


Figure 4: “Types of Data Analysis Techniques”

Together, these tools complement one another. PCA can be combined with K-Means for exploratory analysis or with SVM/Random Forest to save computation. The hybrid application of PCA + SVM (or PCA + RF) was particularly useful in maximizing performance.

6. CONCLUSION

This research has discussed the central use of the quantitative methods in the area of the biological sciences and how computational tools and data-driven methods improve the analysis and interpretation of biological data. Researchers today can now produce meaningful insight from complex data in biological systems with more speed and accuracy than before due to the integration of statistical modeling, machine learning, and algorithmic pipelines. Four major algorithms – K-Means Clustering, Support Vector Machine (SVM), Random Forest and Principal Component Analysis (PCA) – that bring something distinct for the classification, prediction, and dimensionality reduction tasks were examined by the study. Experimental evaluations showed Random Forest performed best with high accuracy, while PCA was shown to reduce the dimensionality of multi dimensional data, aligning with the worthiness in the adoption of a combination of algorithms depending on specific research objectives. In addition to this, inclusion of pseudocode, performance comparison tables and application specific results makes a practical basis for implementing these methods in any real world biological research. This work is also consistent with recent works that include both metaTP and ToxDAR, and GAN-WGCNA, which indicate the increasing use of computational pipelines in molecular biology, toxicology, genomics, and environmental sciences. As compared to the existing literature, this research adds a directed communicative contrast of algorithmic behavior and application in diverse biological data cases. Finally, the study provides a confirmation that quantitative and computational methodologies are not merely complementary to traditional biological approaches, but are increasingly becoming irreplaceable tools of contemporary biology. With data complexity increasing, innovative and discovering interdisciplinary work between biologists, data scientists, and software engineers will be the norm to move innovation and discovery forward. This research paves the way for other research that aims to optimize data-centred solutions in the field of life sciences.

REFERENCES

- [1] ABOUSHADY, D., SAMIR, L., MASOUD, A., ELSHOURA, Y., ABDELGAWAD, M., HANAFI, R.S. and DEEB, S.E., 2025. Chemometric Approaches for Sustainable Pharmaceutical Analysis Using Liquid Chromatography. *Chemistry*, 7(1), pp. 11.
- [2] AFRIZAL, M.N., GOFUR, A., SARI, M.S. and MUNZIL, 2025. Technology-supported differentiated biology education: Trends, methods, content, and impacts. *Eurasia Journal of Mathematics, Science and Technology Education*, 21(3),.
- [3] ANTONELLA, G., VITTORIO, A., GIULIA, F., ALESSANDRO, Z., GIORGIO, P. and MARCO, S., 2025. Innovative Methodologies for the Early Detection of Breast Cancer: A Review Categorized by Target Biological Samples. *Biosensors*, 15(4), pp. 257.
- [4] BARROS, B.J. and CUNHA, J.P.S., 2024. Single-cell and extracellular nano-vesicles biosensing through phase spectral analysis of optical fiber tweezers back-scattering signals. *Communications Engineering*, 3(1), pp. 97.

- [5] BONGRAND, P., 2024. Should Artificial Intelligence Play a Durable Role in Biomedical Research and Practice? *International Journal of Molecular Sciences*, 25(24), pp. 13371.
- [6] BROWNE, D.J., MILLER, C.M., O'HARA, E.P., COURTNEY, R., SEYMOUR, J., DOOLAN, D.L. and ORR, R., 2024. Optimization and application of bacterial environmental DNA and RNA isolation for qualitative and quantitative studies. *Environmental DNA*, 6(4),.
- [7] CANDIA, J. and FERRUCCI, L., 2024. Assessment of Gene Set Enrichment Analysis using curated RNA-seq-based benchmarks. *PLoS One*, 19(5),.
- [8] CARLIN, D.J. and RIDER, C.V., 2024. Combined Exposures and Mixtures Research: An Enduring NIEHS Priority. *Environmental Health Perspectives (Online)*, 132(7), pp. 1-11.
- [9] CAVAYE, H. and MANFRED, E.S., 2025. Utilising Machine Learning for the Automated Characterisation of In Situ Electron Microscopy Experiments with Catalytic Systems: Increasing the time efficiency and reproducibility for particle analysis of in situ catalysis data. *Johnson Matthey Technology Review*, 69(1), pp. 112-122.
- [10] DELL' AVERSANA PAOLO, 2025. A Biological-Inspired Deep Learning Framework for Big Data Mining and Automatic Classification in Geosciences. *Minerals*, 15(4), pp. 356.
- [11] DEMIANIUK, A.V. and KORNIICHUK, V.I., 2024. Application of the numerical modelling for the interpretation of the piezometric data on earth dams in the view of uncertainties specific for dam operation. *IOP Conference Series.Earth and Environmental Science*, 1415(1), pp. 012098.
- [12] ENOLE, J.O., MOJISOLA, M.O. and YETUNDE, O.J., 2024. Knowledge, perception, attitude, and practice of complementary and alternative medicine by health care workers in Garki hospital Abuja, Nigeria. *BMC Complementary Medicine and Therapies*, 24, pp. 1-11.
- [13] GAMBOA-SOTO, F., BAUTISTA-GARCÍA, R., LLANES-GIL LÓPEZ DIANA I., BERMEA, J.E., TINOCO, M.R., OLIVE-MÉNDEZ, S.,F. and GONZÁLEZ-HERNÁNDEZ ANDRÉS, 2025. Heat Treatment-Driven Structural and Morphological Transformation Under Non-Parametric Tests on Metal–Ceramic-Sputtered Coatings. *Ceramics*, 8(1), pp. 25.
- [14] GHOSH, A., GUPTA, A., JENA, S., KIRTI, A., CHOUDHURY, A., SAHA, U., SINHA, A., KUMARI, S., KUJAWSKA, M., KAUSHIK, A. and VERMA, S.K., 2025. Advances in posterity of visualization in paradigm of nano-level ultra-structures for nano–bio interaction studies. *View*, 6(1),.
- [15] HE, L., ZOU, Q. and WANG, Y., 2025. metaTP: a meta-transcriptome data analysis pipeline with integrated automated workflows. *BMC Bioinformatics*, 26, pp. 1-11.
- [16] [16] JIANG, P., ZHANG, Z., YU, Q., WANG, Z., DIAO, L. and LI, D., 2024. ToxDAR: A Workflow Software for Analyzing Toxicologically Relevant Proteomic and Transcriptomic Data, from Data Preparation to Toxicological Mechanism Elucidation. *International Journal of Molecular Sciences*, 25(17), pp. 9544.
- [17] [17] KASHIF, M. and BYRNE, H.J., 2025. From a Spectrum to Diagnosis: The Integration of Raman Spectroscopy and Chemometrics into Hepatitis Diagnostics. *Applied Sciences*, 15(5), pp. 2606.
- [18] KIM, C.S., CAIRNS, J., QUARANTOTTI, V., KACZKOWSKI, B., WANG, Y., KONINGS, P. and ZHANG, X., 2024. A statistical simulation model to guide the choices of analytical methods in arrayed CRISPR screen experiments. *PLoS One*, 19(8),.
- [19] KIM, T., LEE, K., CHEON, M. and YU, W., 2024. GAN-WGCNA: Calculating gene modules to identify key intermediate regulators in cocaine addiction. *PLoS One*, 19(10),.
- [20] [KOLAŠINAC, S.M., PEĆINAR, I., GAJIĆ, R., MUTAVDŽIĆ, D. and ZORA P DAJIĆ STEVANOVIĆ, 2025. Raman Spectroscopy in the Characterization of Food Carotenoids: Challenges and Prospects. *Foods*, 14(6), pp. 953.
- [21] KUMAR, D. and BASSILL, N.P., 2024. Analysing trends of computational urban science and data science approaches for sustainable development. *Computational Urban Science*, 4(1), pp. 33.
- [22] LEGESSE, K.D., MASHABELA, M.R., SEBATJANE, P.N., SITHOLE, S., TABO, B. and EEVA-MARIA RAPOO, 2025. Evaluation of statistical methods applied in theses and dissertations in an Open, Distance and e-Learning University. *PLoS One*, 20(3),.
- [23] LIYUN, Z., YI MAN, L.R. and HONGZHOU, D., 2025. Modeling the Ecological Network in Mountainous Resource-Based Cities: Morphological Spatial Pattern Analysis Approach. *Buildings*, 15(8), pp. 1388.
- [24] LUCIDO, A., BASALLO, O., MARIN-SANGUINO, A., ELEIWA, A., MARTINEZ, E.S., VILAPRINYO, E., SORRIBAS, A. and ALVES, R., 2025. Multiscale Mathematical Modeling in Systems Biology: A Framework

to Boost Plant Synthetic Biology. *Plants*, 14(3), pp. 470.

- [25] MANZOOR, F., TSURGEON, C.A. and GUPTA, V., 2025. Exploring RNA-Seq Data Analysis Through Visualization Techniques and Tools: A Systematic Review of Opportunities and Limitations for Clinical Applications. *Bioengineering*, 12(1), pp. 56.
- [26] MEZA-JOYA, F.L., MORGAN-RICHARDS, M. and TREWICK, S.A., 2025. Forecasting Range Shifts in Terrestrial Alpine Insects Under Global Warming. *Ecology and Evolution*, 15(1)...
-