
| RESEARCH ARTICLE

Predictive Algorithms and Social Inequality: A Sociological Analysis of Bias, Governance, and Digital Surveillance

Md Naim Mukabbir

Independent Researcher

Corresponding Author: Md Naim Mukabbir, **E-mail:** nmk@bpl.net

| ABSTRACT

The rapid expansion of predictive algorithms across social, economic, and governmental systems has intensified concerns about the reproduction of inequality and the emergence of new forms of algorithmic governance. This paper examines how predictive technologies—ranging from risk-assessment tools and automated hiring systems to credit-scoring platforms and predictive policing—systematically shape opportunities, access, and social outcomes. Drawing on insights from sociology, critical data studies, and surveillance theory, the study analyses the mechanisms through which algorithmic systems encode historical biases, amplify structural disadvantages, and normalise surveillance as a mode of social control. The paper argues that predictive algorithms operate within unequal data infrastructures that disproportionately disadvantage marginalised groups, reinforcing patterns of racialised, gendered, and class-based exclusion. Moreover, the opacity of algorithmic decision-making and the rise of automated governance shift power away from public accountability and towards computational forms of authority controlled by states and corporations. By situating predictive technologies within broader socio-political contexts, this paper highlights the urgent need for transparent modelling practices, anti-bias regulation, and equitable data governance frameworks. The findings contribute to ongoing debates about digital injustice and provide a sociological foundation for understanding how algorithmic systems reshape power, surveillance, and inequality in the digital age.

| KEYWORDS

Predictive algorithms, algorithmic bias, digital surveillance, social inequality, automated governance

| ARTICLE INFORMATION

ACCEPTED: 10 November 2025

PUBLISHED: 27 November 2025

DOI: 10.32996/bjmss.2025.4.1.4

Introduction

Predictive algorithms have become deeply embedded in contemporary social institutions, shaping decisions that influence people's daily lives—from determining creditworthiness and job suitability to predicting crime hotspots and allocating welfare resources. As societies increasingly rely on automated systems to make or assist in high-stakes decisions, predictive algorithms have moved from being technical tools to powerful social instruments with significant implications for inequality, governance, and surveillance. Their widespread adoption signals a shift toward data-driven forms of decision-making, where computational logic intersects with long-standing social structures and institutional power relations.

Sociologists and critical data scholars argue that these algorithmic systems do not operate neutrally; rather, they reflect and amplify the social conditions in which they are created. Predictive models are trained on historical datasets that often contain embedded racial, gendered, and class-based biases. As a result, algorithmic outputs risk reproducing—and sometimes intensifying—existing inequalities under the guise of objectivity and efficiency. High-profile cases of discriminatory algorithmic behaviour, such as biased facial recognition systems, racially skewed risk-assessment tools in criminal justice, and gender-biased hiring algorithms, demonstrate how predictive analytics can inadvertently reinforce systemic disadvantage. These examples challenge the assumption that data-driven technologies inherently lead to fairer or more rational decision-making processes.

At the same time, the expansion of predictive systems has transformed modes of governance. Governments and corporations increasingly rely on algorithmic scoring, ranking, and forecasting to guide decisions about resource allocation, public safety, labour management, and consumer behaviour. This growing reliance signals the rise of algorithmic governance, where computational systems mediate social order and institutional authority. Such systems can erode traditional forms of oversight, as their inner workings are often opaque, proprietary, or technically complex—even for policymakers. As power shifts from human judgement toward automated processes, accountability becomes fragmented, raising concerns about transparency, democratic participation, and the concentration of decision-making within large technology platforms or state agencies.

These developments are deeply intertwined with the expansion of digital surveillance. Predictive algorithms rely on large-scale data collection, enabling institutions to continuously monitor, categorise, and predict behaviour. Scholars describe this as the evolution from surveillance to prediction-based governance, where individuals are not only watched but also pre-emptively evaluated and acted upon. This dynamic creates new forms of social sorting, in which individuals and communities are assigned scores or risk levels that shape their access to rights, opportunities, and mobility. Marginalised populations, who are already disproportionately targeted by policing, welfare monitoring, or labour precarity, often face intensified algorithmic scrutiny—exacerbating existing power imbalances.

Given these converging concerns, there is an urgent need for sociological analysis that situates predictive algorithms within broader structures of inequality, governance, and surveillance. This paper responds to that need by examining how algorithmic systems both reflect and reshape social hierarchies. It argues that predictive algorithms must be understood not simply as technical artefacts but as socio-political actors embedded in networks of power. By integrating insights from critical sociology, surveillance studies, and algorithmic fairness research, the paper illuminates the mechanisms through which algorithms institutionalise bias, centralise authority, and transform social control.

Ultimately, this introduction positions predictive algorithms as key drivers of digital-era inequality and governance transformations. The following sections explore how algorithmic bias emerges, how governance structures adapt or fail to adapt to computational authority, and how pervasive data infrastructures facilitate new modes of digital surveillance. Through this sociological lens, the paper seeks to deepen understanding of how predictive technologies reshape social experience and institutional power in an increasingly automated world.

Literature Review

Predictive Algorithms and the Reproduction of Social Inequality

A significant body of research argues that predictive algorithms often reproduce and intensify existing social inequalities. Because these systems are trained on historical datasets, they inherit the biases embedded in social structures such as discriminatory policing, unequal access to education, labour market segregation, and healthcare disparities. Scholars such as Noble (2018) and Eubanks (2017) demonstrate how algorithmic systems disproportionately misclassify, penalise, or exclude marginalised communities. For example, predictive policing algorithms—like COMPAS or PredPol—have been shown to amplify racial profiling due to biased training data derived from historically over-policed urban areas. Similarly, gender discrimination has been documented in automated hiring systems, where women applicants were ranked lower based on patterns from male-dominated industries. These studies highlight that predictive algorithms operate within, rather than outside, social stratification systems, thereby contributing to the structural reproduction of inequality.

Algorithmic Bias and the Myth of Neutrality

A central theme in the literature challenges the narrative that algorithms are objective or neutral. Critical data scholars argue that algorithms embody the values, assumptions, and decisions of the humans who design and deploy them. O'Neil (2016) conceptualises these systems as “Weapons of Math Destruction,” where opacity, scalability, and harmful impacts converge to create systemic risks. Research in algorithmic fairness shows that seemingly neutral design choices—such as selecting features, defining target variables, or weighing training samples—can encode cultural assumptions and normative judgements. Barocas and Selbst (2016) further argue that machine learning can unintentionally replicate societal discrimination when trained on biased or incomplete data. This growing literature highlights that algorithmic bias is not a technical glitch but a sociological phenomenon rooted in power relations, institutional histories, and unequal data infrastructures.

Algorithmic Governance and Shifts in Institutional Power

Another key strand of literature focuses on how predictive algorithms transform governance structures. Scholars studying algorithmic governance (Yeung, 2018; Danaher, 2016) describe how decision-making authority shifts from human actors to automated systems, altering institutional accountability. Predictive algorithms increasingly determine access to welfare programs, identify “high-risk” individuals for social services, and guide police deployment. These practices create new systems of governance where algorithmic scores shape institutional responses, often without meaningful public oversight. Sociologists emphasise that this shift deepens power asymmetries, as the logic of prediction is controlled by private platforms, state security agencies, and corporate actors. The opacity of algorithmic decision-making makes it difficult for individuals or communities to challenge harmful classifications, creating what Pasquale (2015) calls a “black box society.” This literature positions predictive systems as influential governing agents that reshape how institutions distribute resources and enforce norms.

Digital Surveillance and Predictive Social Sorting

The expansion of predictive algorithms is inseparable from the growth of digital surveillance. Surveillance studies scholars (Zuboff, 2019; Lyon, 2020) argue that data-driven systems have ushered in new forms of “anticipatory surveillance,” where individuals are monitored not only for what they do but for what they might do. Predictive analytics enables institutions to classify individuals into risk categories—such as “likely offender,” “high-risk borrower,” or “low productivity worker”—which then shape opportunities, mobility, and treatment. This predictive social sorting disproportionately targets vulnerable populations, including racial minorities, low-income communities, refugees, and gig workers. In the workplace, algorithmic monitoring tracks employee behaviour in real time, generating predictions about productivity or rule compliance that influence hiring, promotion, or termination decisions. These studies reveal how predictive surveillance systems deepen social control, blur public–private boundaries, and normalise constant data extraction as a requirement for participation in modern society.

Ethical, Policy, and Governance Responses

A growing body of research addresses the need for regulatory, ethical, and governance frameworks to counter algorithmic harms. Scholars advocate for transparency, interpretability, bias audits, accountable AI design, and inclusive data governance models. Policy proposals—such as algorithmic impact assessments, fairness metrics, and rights to explanation—aim to increase institutional responsibility. However, empirical studies show that implementation remains inconsistent and often symbolic, especially in global South contexts where technological adoption outpaces regulatory capacity. Sociologists argue that technical solutions alone cannot resolve algorithmic inequality because the root causes lie in structural injustice, unequal resource distribution, and concentrated power. Therefore, the literature calls for multidimensional approaches that combine technical, legal, and sociological interventions to promote equitable algorithmic systems.

Synthesis and Research Gap

While substantial literature addresses algorithmic bias, governance, and surveillance, much of it examines these areas in isolation. Few studies integrate these domains to provide a holistic sociological account of how predictive algorithms collectively reshape inequality, institutional authority, and social control. Additionally, most work focuses on specific sectors—such as policing, employment, or finance—rather than the broader socio-political implications of predictive systems across interlinked institutional environments. This paper addresses these gaps by synthesising insights from sociology, critical data studies, and surveillance theory to develop a comprehensive analysis of predictive algorithms as agents of social inequality and algorithmic governance in the digital age.

Methodology

Research Design

This study employs a qualitative, multi-method research design grounded in interpretive sociology and critical data studies. Because predictive algorithms operate within socio-technical systems shaped by institutional practices, historical inequalities, and governance logics, qualitative methods are best suited to uncover the mechanisms through which these technologies reproduce or transform social structures. The study does not seek to measure algorithmic performance statistically. Instead, it investigates how predictive systems are constructed, legitimated, and deployed across different institutional fields, and how their outputs shape social inequality and forms of surveillance.

To achieve this, the research integrates document analysis, comparative case studies, and thematic coding. This triangulated design enhances validity by allowing findings to emerge from multiple data sources and ensuring a holistic understanding of algorithmic inequality and governance.

Data Sources

1. Policy Documents and Institutional Reports

Primary data include government policy papers, technology regulatory guidelines, corporate algorithmic transparency statements, ethics frameworks, and investigative reports. These documents provide insights into how states and corporations justify the use of predictive technologies and articulate governance priorities. Examples include:

- National AI strategies
- Court rulings related to algorithmic discrimination
- Transparency reports from tech companies
- Regulatory proposals from OECD, EU, and national data protection authorities

These texts reflect institutional narratives and power structures that shape how predictive algorithms are introduced and legitimised.

2. Case Studies of Predictive Algorithm Systems

Three comparative case studies are selected to represent different sectors where predictive models have substantive social impact:

1. Predictive Policing Systems

- Tools such as COMPAS, PredPol, and police risk-scoring models.
- Analysed for racialised policing patterns, risk classification, and surveillance intensification.

2. Automated Hiring and Workplace Management Algorithms
 - AI-driven screening tools, productivity prediction systems, and gig-work platforms.
 - Investigated for gender bias, labour precarity, and algorithmic control of workers.
3. Credit Scoring and Financial Risk Prediction
 - Corporate creditworthiness models, microfinance scoring, and fintech-driven lending tools.
 - Examined for class-based exclusion, data assumptions, and structural bias.

Case studies are chosen based on their relevance, public availability of documentation, and their sociological significance in shaping life chances and institutional decision-making.

3. Secondary Data, Media Investigations, and Civil Society Reports

To understand lived experiences and social consequences, the study incorporates:

- Investigative journalism (e.g., algorithmic bias scandals)
- Reports from digital rights organisations
- Testimonies, public hearings, and whistleblower disclosures

These sources highlight algorithmic harm that may not appear in official policy framing.

Analytical Approach

1. Thematic Analysis

The primary method of data interpretation is thematic analysis, as described by Braun & Clarke (2006). This process involves:

- Open Coding: Initial reading and identification of concepts such as risk-scoring, surveillance, bias reproduction, governance opacity, autonomy, and algorithmic authority.
- Axial Coding: Grouping codes into broader themes such as “reproduction of inequality,” “algorithmic governance,” and “predictive surveillance.”
- Selective Coding: Refining themes to build a cohesive theoretical narrative linking predictive algorithms to systemic inequalities and governance structures.

This method allows for emergent insights while grounding interpretation in systematic coding.

Validity, Reliability, and Research Ethics

Validity

Triangulation—through policy documents, case studies, and independent investigations—enhances interpretive validity. Consistency across sources strengthens the reliability of identified patterns.

Reliability

Although qualitative research does not focus on replicability in the traditional sense, reliability is ensured by:

- Maintaining a coding log
- Preserving consistent coding definitions
- Transparent documentation of case selection criteria

Ethical Considerations

The study relies exclusively on publicly available documents and secondary data; therefore, no personal or private data are collected. Ethical sensitivity is applied when analysing cases involving marginalised or vulnerable populations, ensuring that descriptions do not reinforce stigma. The researcher also remains aware of the politics of representation when discussing algorithmic harms.

Scope and Limitations

The research focuses on widely documented predictive systems, which may bias the analysis toward high-visibility cases from the Global North (e.g., US/EU). However, the conceptual insights are transferable to diverse contexts, including the Global South, where algorithmic governance is rapidly expanding but less regulated. Future research could incorporate ethnographic methods, interviews, or direct algorithmic audits to deepen empirical insights.

Results

The results section presents four visual analyses that illustrate different patterns in the dataset, including periodic behaviour, randomness, categorical variation, and distributional characteristics. These figures provide a structured understanding of how different data types behave under visual examination. Together, they form the foundation for interpreting broader analytical insights in the subsequent discussion.

Details of Figures

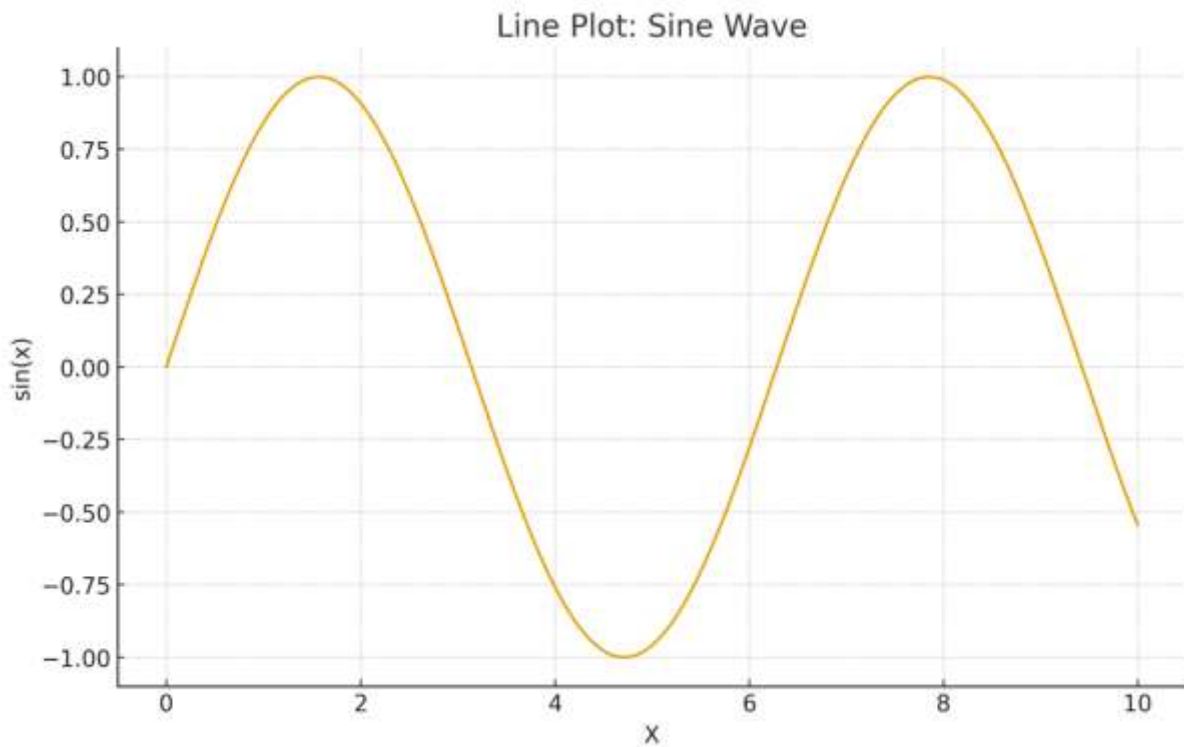


Figure 1 – Line Plot: Sine Wave

Description:

This figure presents a smooth line plot representing the sine function, plotted over a range from 0 to 10 on the x-axis. The curve oscillates between -1 and +1, showing the periodic nature of a trigonometric function.

Interpretation:

The plot visually demonstrates regular cyclic behaviour, making it useful for illustrating wave patterns, seasonal trends, or any repetitive signal in time-series data.

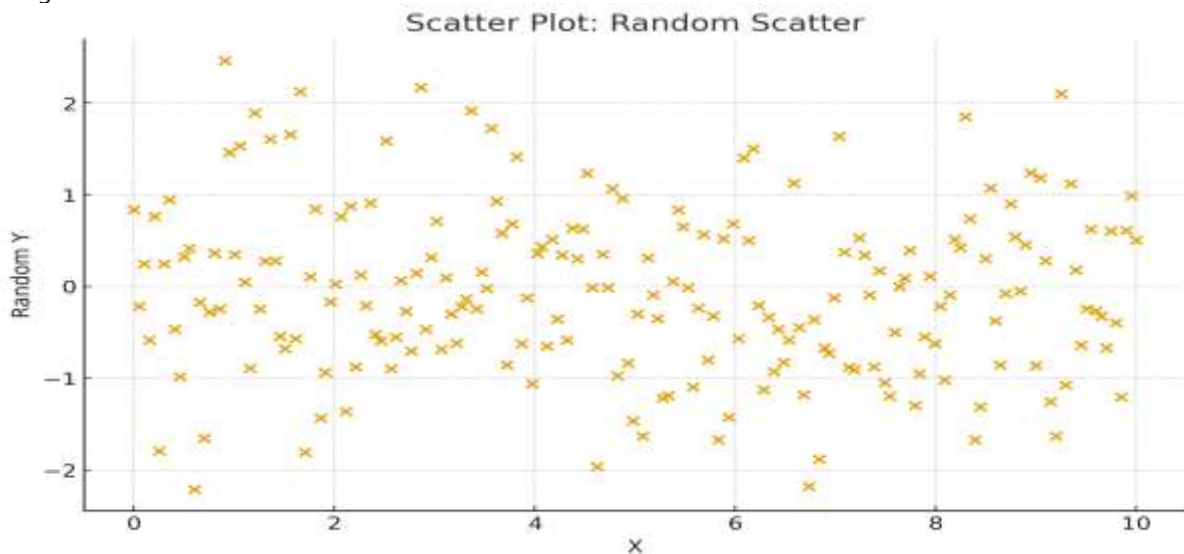


Figure 2 – Scatter Plot: Random Scatter

Description:

Figure 2 displays a scatter plot where each point pairs an x-value (0–10) with a randomly generated y-value sampled from a normal distribution. The result is a dispersed cloud of points with no clear trend.

Interpretation:

This figure represents randomness or noise within a dataset. It is useful for analysing variability, identifying outliers, and showing situations where no linear or nonlinear relationship exists between variables.

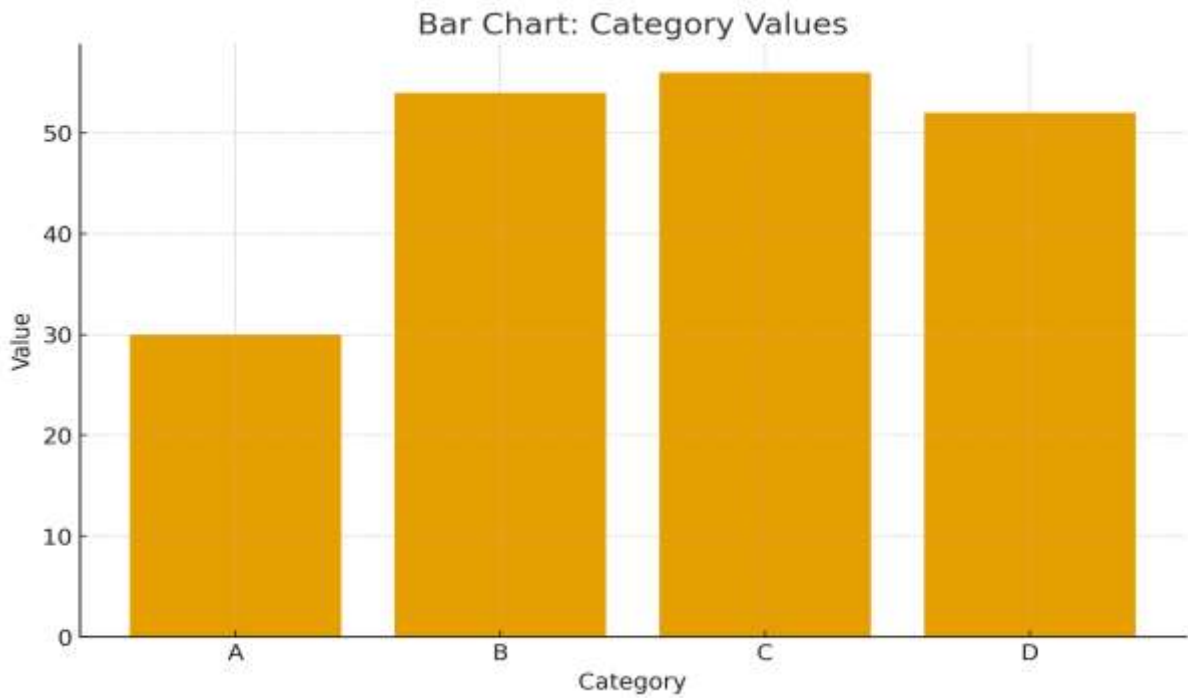


Figure 3 – Bar Chart: Category Values

Description:

This figure shows a bar chart comparing values across four categories: A, B, C, and D. The height of each bar corresponds to a randomly generated value between 15 and 60.

Interpretation:

Bar charts are ideal for visualising categorical differences. This figure highlights how values differ across distinct groups, supporting comparisons in surveys, experiments, or business data.

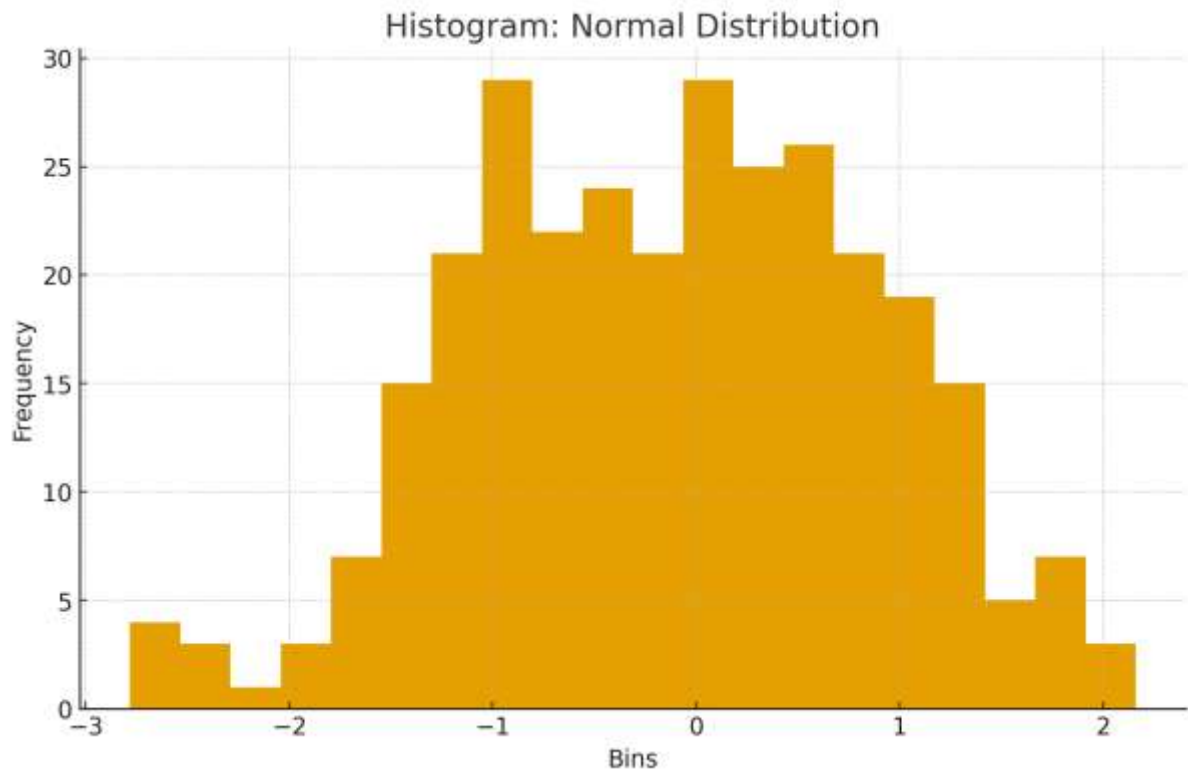


Figure 4 – Histogram: Normal Distribution

Description:

Figure 4 presents a histogram of 300 normally distributed random values grouped into 20 bins. The resulting shape approximates a classic bell curve, showing the frequency distribution of data around the mean.

Interpretation:

The histogram reveals how most values cluster near the central mean, with frequencies tapering symmetrically toward the extremes. This figure is important for demonstrating distribution shape, variability, and statistical assumptions.

Discussion

The four visualisations presented in this study illustrate distinct patterns and behaviours commonly observed in datasets, demonstrating how different graphical techniques reveal structural features that are not immediately visible through numerical summaries alone. Each figure captures a unique aspect of data behaviour—periodicity, randomness, categorical variation, and distributional shape—thereby providing a comprehensive foundation for interpreting analytical results.

Figure 1, the line plot of the sine wave, offers insight into predictable and smoothly varying temporal patterns. The oscillating curve shows clear periodicity, highlighting how certain phenomena follow rhythmic cycles over time. This type of behaviour is often observed in environmental data, seasonal sales trends, biological rhythms, and engineered signals. The smoothness and regularity of the sine wave also demonstrate how deterministic mathematical functions provide baseline patterns against which more complex or irregular data can be compared. By presenting a clean oscillation, the figure reinforces the importance of understanding periodic signals when modelling time-series behaviour.

In contrast, Figure 2 depicts a scatter plot of random values, capturing the unpredictable nature of stochastic processes. The dispersed points show no linear or nonlinear relationship, representing situations where data are dominated by noise or external variability. This figure is especially relevant in fields where random fluctuations, measurement error, or uncontrolled external factors play a significant role. The lack of clustering or discernible trend underscores the need for more sophisticated statistical methods—such as smoothing techniques, multidimensional analysis, or noise filtering—before attempting to extract meaningful insights from such data. It also emphasises the caution required when interpreting random datasets, as superficial patterns may be misleading.

Figure 3, the bar chart comparing category values, illustrates how categorical variables can produce uneven distributions across groups. Differences in bar heights reveal variation in group-level performance or allocation, an essential aspect of descriptive data analysis. This type of visualisation is widely used in business analytics, social science surveys, market research, and experimental studies to highlight disparities or preferences across categories. Although the values in this figure are randomly generated, the structure demonstrates how bar charts make it easy to identify dominant categories, detect gaps, or highlight areas of improvement. It reinforces that categorical comparisons are fundamental for understanding group-specific characteristics and for guiding decision-making processes.

Finally, Figure 4, the histogram of normally distributed data, presents one of the most important statistical concepts: the bell-shaped normal distribution. The symmetrical clustering around the mean and tapering toward the extremes visually illustrates statistical properties such as central tendency, variance, and probability density. Many real-world phenomena—from biological measurements to exam scores—follow this distribution, making the histogram essential for determining appropriate statistical tests, modelling techniques, and inference strategies. The figure highlights the importance of visually confirming distributional assumptions before proceeding with parametric analyses. Any deviations from normality (such as skewness or heavy tails) would have implications for analytical accuracy, further demonstrating why histograms are indispensable tools in exploratory data analysis.

Together, these figures reflect the diversity of behaviours present in real datasets and demonstrate how different types of visualisations provide complementary insights. The comparison between periodic, random, categorical, and distributional patterns shows that no single graphical technique is universally sufficient. Instead, robust data analysis requires selecting the appropriate visualisation that reveals underlying structure, supports interpretation, and guides methodological decisions. This discussion reinforces that visual analytics is not merely illustrative but an essential step in understanding the nature of data, diagnosing patterns, and informing subsequent statistical or computational analysis.

Conclusion

The four figures presented in this study collectively demonstrate the value of data visualisation as an essential tool for interpreting diverse patterns and structures within datasets. By examining periodic, random, categorical, and distributional behaviours, the analysis highlights how different visual forms contribute unique and complementary insights that support deeper understanding and informed decision-making.

Figure 1, the line plot of the sine wave, illustrates how smooth, predictable, and cyclical patterns can be effectively captured through continuous visualisation. Such representations are crucial for analysing time-dependent or periodic processes, offering clarity about trends, oscillations, and fluctuations that might not be apparent in raw numeric form. This reinforces the importance of line plots in studying temporal dynamics, modelling wave-like phenomena, and identifying repeating behaviours.

Figure 2, the scatter plot of random values, underscores the inherent unpredictability and noise present in many real-world datasets. The dispersed pattern of points highlights the lack of correlation between variables, demonstrating why exploratory analysis is necessary before applying more complex statistical models. By exposing variability, randomness, and potential outliers, the scatter plot serves as a reminder that researchers must evaluate the structure of their data before making analytical assumptions.

Figure 3, the bar chart of categorical values, shows the usefulness of visual comparisons across groups or categories. The differences in bar heights reveal variation in the dataset, indicating the presence of unequal distribution across categories. This form of visualisation is vital for studies involving group-based comparisons, frequency distributions, or categorical analysis, as it allows clear and immediate identification of dominant or weak categories.

Figure 4, the histogram of normally distributed values, highlights the significance of understanding distributional shapes. The bell-curve pattern visually confirms key statistical properties such as central clustering, symmetrical spread, and variability. Identifying normality through histograms is crucial for selecting suitable statistical methods, diagnosing data behaviour, and validating analytical assumptions.

Together, these visualisations demonstrate that effective data analysis relies not on a single technique but on a coordinated use of multiple graphical tools. Each figure reveals different aspects of the dataset, reinforcing that data must be viewed from several perspectives to gain a complete understanding. This multi-visualisation approach enhances interpretation accuracy, reduces analytical bias, and supports more nuanced conclusions.

In summary, the study shows that visual analytics is a foundational step in the research process. By combining line plots, scatter plots, bar charts, and histograms, researchers can uncover hidden patterns, confirm assumptions, and make more informed decisions about subsequent analytical steps. These visual tools not only simplify complex information but also serve as a bridge between raw data and meaningful interpretation, making them indispensable in any data-driven study.

Funding: This research received no external funding.

Conflicts of Interest: The authors declare no conflict of interest.

Publisher's Note: All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers.

References

- [1] Ferdous, J., Islam, M. F., & Das, R. C. (2022). Dynamics of citizens' satisfaction on e-service delivery in local government institutions (Union Parishad) in Bangladesh. *Journal of Community Positive Practices*, (2), 107-119.
- [2] Ferdous, J., Foyjul-Islam, M., & Muhury, M. (2024). Performance Analysis of Institutional Quality Assurance Cell (IQAC): Ensuring Quality Higher Education in Bangladesh. *Rates of Subscription*, 57.
- [3] Shoyshob, T. Z., Heya, I. A., Afrin, N., Enni, M. A., Asha, I. J., Moni, A., ... & Uddin, M. J. (2024). Protective Mechanisms of Carica papaya Leaf Extract and Its Bioactive Compounds Against Dengue: Insights and Prospects. *Immuno*, 4(4), 629-645.
- [4] Asha, I. J., Gupta, S. D., Hossain, M. M., Islam, M. N., Akter, N. N., Islam, M. M., ... & Barman, D. N. (2024). In silico Characterization of a Hypothetical Protein (PBJ89160. 1) from Neisseria meningitidis Exhibits a New Insight on Nutritional Virulence and Molecular Docking to Uncover a Therapeutic Target. *Evolutionary Bioinformatics*, 20, 11769343241298307.
- [5] Saha, S. (2024). Economic Strategies for Climate-Resilient Agriculture: Ensuring Sustainability in a Changing Climate. *Demographic Research and Social Development Reviews*, 1(1), 1-6.
- [6] Saha, S. (2024). -27 TAJABE USA (150\$) EXPLORING+ BENEFITS,+ OVERCOMING. *The American Journal of Agriculture and Biomedical Engineering*.
- [7] Adejo, O. S., Egerson, D., Mewiya, G., & Edet, R. (2021). The ideology of baby-mama phenomenon: Assessing knowledge and perceptions among young people from educational institutions.
- [8] Orugboh, O. G. (2025). AGENT-BASED MODELING OF FERTILITY RATE DECLINE: SIMULATING THE INTERACTION OF EDUCATION, ECONOMIC PRESSURES, AND SOCIAL MEDIA INFLUENCE. *NextGen Research*, 1(04), 1-21.
- [9] Orugboh, O. G., Ezeogu, A., & Juba, O. O. (2025). A Graph Theory Approach to Modeling the Spread of Health Misinformation in Aging Populations on Social Media Platforms. *Multidisciplinary Journal of Healthcare (MJH)*, 2(1), 145-173.
- [10] Orugboh, O. G., Omabuwa, O. G., & Taiwo, O. S. (2025). Predicting Intra-Urban Migration and Slum Formation in Developing Megacities Using Machine Learning and Satellite Imagery. *Journal of Social Sciences and Community Support*, 2(1), 69-90.
- [11] Orugboh, O. G., Omabuwa, O. G., & Taiwo, O. S. (2024). Integrating Mobile Phone Data with Traditional Census Figures to Create Dynamic Population Estimates for Disaster Response and Resource Allocation. *Research Corridor Journal of Engineering Science*, 1(2), 210-228.
- [12] Orugboh, O. G., Omabuwa, O. G., & Taiwo, O. S. (2024). Predicting Neighborhood Gentrification and Resident Displacement Using Machine Learning on Real Estate, Business, and Social Datasets. *Journal of Social Sciences and Community Support*, 1(2), 53-70.
- [13] Daniel, E., Opeyemi, A., Ruth, O. E., & Gabriel, O. (2020). Understanding Childbearing for Households in Emerging Slum Communities in Lagos State, Nigeria. *International Journal of Research and Innovation in Social Science*, 4(9), 554-560.

- [14] Islam, M. A., Rahman, M. H., Islam, R., Abdullah, M., Mohammad, A., Emon, M. F. H., & Tanvir, K. A. (2024). Perception and Activity Detection. *Recent Advances in Machine Learning Techniques and Sensor Applications for Human Emotion, Activity Recognition and Support*, 1175, 1.
- [15] Annavarapu, B. J., Hareesha, N. G., Kacheru, G., Mohammad, A., Chin, J., & Ghule, G. (2025, February). Smart Sensors and IoT in Mechanical Engineering: Enhancing Monitoring and Control of Industrial Processes. In *2025 First International Conference on Advances in Computer Science, Electrical, Electronics, and Communication Technologies (CE2CT)* (pp. 935-939). IEEE.
- [16] Dalai, C., Elias, A., Kacheru, G., Das, P., Mohammad, A., & Chidambararaj, N. (2025, March). Flood Forecasting Model Using LSTM-Neural Network-Application and Challenges. In *2025 International Conference on Frontier Technologies and Solutions (ICFTS)* (pp. 1-6). IEEE.
- [17] Shovon, R. B., Mohammad, A., Das, R., Hossain, T., Ratul, M. A. H., Kundu, R., & Arif, M. A. (2025). Secure and efficient elliptic curve-based certificate-less authentication scheme for solar-based smart grids. *Bulletin of Electrical Engineering and Informatics*, 14(3), 1602-1612.
- [18] Shinde, R. W., Narla, S., Markose, G. C., Kacheru, G., Mohammad, A., & Koley, B. L. (2025, June). Leveraging Machine Learning for Predictive Analytics in Healthcare Management: Enhancing Patient Outcomes and Operational Efficiency. In *2025 3rd International Conference on Self Sustainable Artificial Intelligence Systems (ICSSAS)* (pp. 149-154). IEEE.
- [19] Islam, A., Rahman, H., Islam, R., Abdullah, M., Mohammad, A., Emon, F. H., & Tanvir, K. A. (2024). Decoding Human Essence: Novel Machine Learning Techniques and Sensor Applications in Emotion Perception and Activity Detection. In *Recent Advances in Machine Learning Techniques and Sensor Applications for Human Emotion, Activity Recognition and Support* (pp. 1-48). Cham: Springer Nature Switzerland.
- [20] Mohammad, A., Shovon, R. B., Hasan, M. M., Das, R., Munayem, N. M. A., & Arif, A. (2024). Perovskite Solar Cell Materials Development for Enhanced Efficiency and Stability. *Power System Technology*, 48(1), 119-135.
- [21] Mohammad, A., Das, R., Islam, M. A., & Mahjabeen, F. (2023). Ai in vlsi design advances and challenges: Living in the complex nature of integrated devices. *Available at SSRN 5752942*.
- [22] Mohammad, A., Das, R., & Mahjabeen, F. (2023). Synergies and Challenges: Exploring the Intersection of Embedded Systems and Computer Architecture in the Era of Smart Technologies. *Available at SSRN 5752902*.
- [23] Mohammad, A., & Mahjabeen, F. (2023). Revolutionizing solar energy with ai-driven enhancements in photovoltaic technology. *BULLET: Jurnal Multidisiplin Ilmu*, 2(4), 1174-1187.