

A Systematic Review of Data Analytics Techniques for Central Tendency–Based Missing Value Replacement (2001–2025)

Sneh B. Patel¹, Dr. Darshanaben Dipakkumar Pandya²

Research Scholar, Department of Computer Science, Sankalchand Patel University, Visnagar, India¹

Associate Professor, Shri C.J Patel College of Computer Studies (BCA), Sankalchand Patel University, Visnagar, India ²

spdigitalme@gmail.com, ddpandya@spu.ac.in

Abstract: Incomplete or partially observed data continue to pose a critical obstacle in contemporary analytics and research, often diminishing model accuracy, interpretability, and reliability. The process of compensating for such data gaps—known as imputation—remains essential to maintaining data quality and analytical consistency. Among the broad range of existing strategies, methods grounded in measures of central tendency, such as mean, median, and mode, persist as reliable and computationally economical options. Their enduring popularity arises from their simplicity, transparency, and the ease with which they integrate into larger analytical frameworks.

This systematic review consolidates research published between 2001 and 2025 that employs central-tendency-based techniques for missing value replacement across diverse fields, including healthcare, finance, environmental modeling, and social analytics. It compares the empirical performance of these methods, summarizes their practical strengths and weaknesses, and identifies conditions under which each performs optimally. The review reveals that median-based imputation consistently outperforms mean-based approaches in datasets exhibiting skewness or outliers, while mode-based substitution demonstrates superior stability and accuracy for categorical variables. The findings suggest that median-based imputation is particularly effective for skewed or non-normal data, whereas mode-based substitution performs best with categorical attributes. Despite certain limitations—such as reduced variance representation, these methods remain vital reference points for benchmarking more sophisticated approaches. The paper concludes with domain-oriented best practices and prospective directions for enhancing the interpretability, efficiency, and adaptability of imputation processes in future data-driven applications.

Keywords: Central Tendency Measures, Data Analytics, Missing Data, Missing Value Replacement, Statistical Techniques, Systematic Review.

I. INTRODUCTION

The modern data ecosystem depends heavily on complete, accurate, and consistent datasets to support credible statistical and machine learning outcomes. However, in real-world settings, perfect data collection is rare. Gaps often arise from instrument failures, manual entry errors, privacy constraints, or transmission interruptions. When left untreated, missing values can distort analyses, weaken inferential validity, and reduce model reliability. Consequently, handling incomplete data has become a fundamental component of modern data preprocessing and quality assurance.

Imputation refers to the practice of estimating and replacing absent data points with plausible values derived from the observed portion of a dataset. Over the past two decades, this practice has evolved from simple rule-based procedures to more sophisticated algorithms. Classical statistical techniques—such as regression imputation and expectation–maximization—have been complemented by computational approaches incorporating machine learning, including k-nearest neighbors (KNN), multiple imputation by chained equations (MICE), and deep learning–based frameworks.

Despite these advances, central-tendency-based imputations—specifically mean, median, and mode substitution—continue to serve as indispensable tools. They are valued for their computational efficiency, interpretability, and role as baseline

comparators when evaluating advanced models. Their transparency and ease of implementation make them particularly suitable for preliminary analyses and large-scale automated workflows.

The emphasis on central tendency arises from its broad applicability and capacity to preserve approximate representativeness within a dataset. Although these techniques may reduce variance and obscure complex relationships, they remain valuable when rapid, explainable, and resource-efficient imputation is required. Their continued use in domains such as healthcare analytics, finance, environmental monitoring, and social science research underscores their practical relevance.

This review aims to synthesize two decades of literature on central-tendency-based missing value imputation. Its specific objectives are:

1. To find consistent practices and methodological patterns in the application of mean, median, and mode imputations.
2. To evaluate their comparative performance across data types and real-world contexts; and
3. To highlight recent developments that integrate these classical approaches with intelligent or adaptive systems.

Through this synthesis, the review provides an evidence-based framework to assist researchers and practitioners in selecting and implementing central-tendency-based imputation techniques effectively within modern analytics pipelines.

II. BACKGROUND AND CONCEPTUAL FRAMEWORK

2.1. The Problem of Missing Data in Analytics

In most analytical workflows, missing data remain one of the most persistent obstacles to achieving reliable insights. Data gaps may occur for a variety of reasons—human error during entry, equipment malfunction, incomplete surveys, or privacy-related data suppression. Their presence can reduce statistical accuracy, weaken predictive performance, and threaten the interpretability of outcomes.

The nature of missingness is not uniform. Some values are absent entirely by chance, while others may depend on specific observed or unobserved variables. Recognizing this underlying mechanism is critical when selecting a treatment strategy. In general, data can be categorized as:

- **Completely random**, where omissions occur independently of all variables.
- **Partially random**, where absence depends only on information already observed and
- **Systematic or non-random**, where missingness relates directly to unrecorded or hidden factors.

Central-tendency-based techniques—such as substituting missing values with the mean, median, or mode—are typically more reliable when data are missing completely or partially at random. In contrast, non-random patterns often require model-driven or probabilistic estimation to avoid bias.

2.2. Evolution of Missing Value Replacement Techniques

Early approaches to managing incomplete datasets were relatively simple, often relying on deletion of affected records or variables. While straightforward, these methods reduced dataset size and compromised representativeness when missingness was not random. This limitation encouraged the development of imputation approaches that attempt to infer plausible substitutes rather than discard data entirely. Classical methods such as mean or regression-based imputation provided an initial solution by estimating missing values from available statistics. Subsequent advances introduced iterative and probabilistic frameworks—including expectation maximization and multiple imputations—followed by machine-learning-oriented algorithms such as k-nearest neighbors, random forests, and deep neural networks. Despite the rapid progress of automated learning methods, central-tendency-based substitution continues to serve as a practical and computationally efficient baseline, especially when data incompleteness is moderate or interpretability is a priority.

2.3. Conceptual Basis of Central Tendency Measures

Measures of central tendency summarize the point around which most data values cluster. The three principal indicators are:

- Mean, representing the arithmetic average of observed values
- Median, the middle point dividing a sorted dataset into equal halves and
- Mode, the most frequently occurring value, particularly useful for categorical or ordinal data.

Using these measures for imputation ensures that the substituted values remain statistically representative of the dataset's core structure. Mean substitution maintains overall average levels but can reduce variability. Median replacement is more robust against outliers and non-symmetric distributions, while mode imputation preserves categorical balance. Although these techniques may slightly distort correlations or understate variance, their interpretability and simplicity justify their widespread adoption in large-scale preprocessing pipelines.

2.4. Theoretical Framework for the Systematic Review

The analytical foundation of this review connects principles of data quality management with statistical estimation methods for handling missingness. Effective data quality management requires completeness, accuracy, and consistency—all of which can be supported through informed imputation strategies. In this context, central tendency measures act as fundamental components that enhance these data quality dimensions by providing clear and reproducible value estimates. These review positions mean, median, and mode imputations as essential statistical tools that continue to hold relevance amid the growth of hybrid and AI-enhanced data recovery models. Through an integrative research examination from 2001 to 2025, this framework establishes how traditional statistical reasoning remains vital in contemporary data-driven analysis.

III. CENTRAL TENDENCY TECHNIQUES IN MISSING VALUE REPLACEMENT

3.1 Overview

Techniques based on measures of central tendency—specifically mean, median, and mode—remain among the most frequently applied strategies for addressing missing data. Their strength lies in the balance between simplicity, computational efficiency, and interpretability. Although modern algorithms such as ensemble learning or deep neural imputers provide higher accuracy in complex settings, central-tendency-based methods continue to serve as essential baselines and initial steps in hybrid imputation frameworks. This section reviews the theoretical underpinnings, key advantages, and limitations of each approach while summarizing their domain-specific performance across studies conducted between 2001 and 2025.

3.2 Mean Imputation

Concept: Mean imputation replaces missing entries with the arithmetic average of the available observations for a particular variable. It is straightforward to compute and ensures that the dataset's central value remains constant after substitution.

Applications and Performance: Empirical analyses over the past two decades demonstrate that mean substitution performs reliably when the proportion of missingness is small—typically below 10%—and when the data follows an approximately normal distribution. Under these conditions, the resulting bias and loss of variance are minimal.

However, the method becomes less effective when data are skewed or when a high percentage of entries are absent, as it tends to reduce variability and distort relationships among variables. Despite these shortcomings, mean imputation continues to be widely used for initial preprocessing, baseline comparisons, and as an initialization component within iterative or machine learning-driven frameworks due to its ease of implementation and transparency.

3.3 Median Imputation

Concept: Median imputation replaces each missing value with the median of the observed data. Because the median is not influenced by extreme values, this technique is particularly robust for datasets containing outliers or heavy tailed distributions.

Applications and Performance: Numerous studies across healthcare, environmental, and social domains have shown that median substitution maintains distributional properties more effectively than mean imputation, especially when variables

exhibit skewness. It generally produces lower error rates, such as reduced Root Mean Squared Error (RMSE) or Mean Absolute Error (MAE), compared to the mean.

3.4 Mode Imputation

Concept: Mode imputation is mainly used for categorical or ordinal variables, substituting missing values with the most frequently occurring category in the dataset. It preserves the dataset's categorical integrity and provides a conceptually clear replacement strategy.

Applications and Performance: Research indicates that mode substitution performs best when category frequencies are relatively balanced and the proportion of missingness is modest. However, in cases where one category dominates, it may reinforce bias and misrepresent minority classes.

Recent developments address this concern by incorporating probabilistic or entropy-adjusted weighting, ensuring that less frequent categories retain influence. Mode imputation remains a valuable, computationally light option for qualitative data preprocessing, particularly in large-scale surveys or social science research.

3.5 Comparative Evaluation of Central Tendency Techniques

The suitability of each central-tendency-based approach depends on the type of data, underlying distribution, and analytical goal:

- Mean imputation works best for numerical data that are approximately normal with low levels of missingness.
- Median imputation provides robustness against outliers and skewness, making it ideal for heavy-tailed or asymmetrical distributions.
- Mode imputation is effective for categorical variables where class proportions are balanced.

Recent literature highlights three prominent trends:

- **Hybrid Integration:** Central-tendency methods often serve as initialization or benchmarking stages for advanced algorithms, improving convergence and baseline stability.
- **Domain Specialization:** Mean-based substitution dominates in finance and engineering; median is favoured in healthcare and environmental research; and mode is predominant in social and behavioural data.
- **Trade-off Awareness:** While advanced methods achieve higher accuracy, central-tendency-based imputations remain preferable when interpretability, speed, or limited computational resources are critical.

3.6 Summary

Central-tendency imputation remains a vital part of data cleaning due to its clarity, ease of computation, and how it works with mixed models. If used appropriately—that is, considering the distribution of the data, type of variable, and amount of missing data—mean, median, and mode imputations provide consistent results that are quite easy to comprehend. Their continued use across many disciplines indicates that they remain applicable despite newer automated imputation techniques.

Concerning Comparative Summary:

- Mean imputation works best for continuous variables that are roughly normally distributed, with little missing data. It is simple but not very robust against skewness and outliers.
- Median imputation performs better in the case of data featuring skew or outliers. It is more resistant to unusual values, though its high usage may weaken the relationships between variables.
- Mode imputation works best when applied to categorical or ordinal data with balanced categories; if one category dominates, the results may be biased.

Overall, which to use depends on data, what you want to achieve, and available resources. Central-tendency methods can be used alone or as clear parts of more advanced imputation methods. The section that follows outlines the method of systematic review to identify, appraise, and summarize the literature on such techniques.

IV.SYSTEMATIC REVIEW METHODOLOGY

4.1 Overview

To ensure a rigorous, transparent, and reproducible evaluation of existing research, this study adopted a structured review process inspired by systematic review principles. The procedure was designed to capture, screen, and synthesize relevant publications that investigated the application of mean, median, or mode imputation for managing missing data between 2001 and 2025. The approach followed a multi-stage process emphasizing clarity, inclusion accuracy, and analytical comparability across studies.

4.2 Research Questions and Objectives

The review sought to answer the following key questions:

1. What patterns and developments characterize the use of central-tendency-based imputation methods from 2001 to 2025?
2. How do mean, median, and mode imputations compare with respect to accuracy, robustness, and computational performance? What implementation practices and methodological guidelines are most effective across different research domains?
3. Which knowledge gaps and future research directions remain unresolved in this area?
- 4.

These questions guided the literature selection, extraction, and synthesis phases to ensure consistency and focus throughout the review.

4.3 Literature Search Strategy

A comprehensive literature search was conducted across multiple academic databases, including IEEE Xplore, Springer Link, Science Direct, ACM Digital Library, and Google Scholar. Search terms combined keywords such as missing data, data imputation, mean imputation, median imputation, mode imputation, data cleaning, central tendency, and data preprocessing. Boolean operators (AND, OR) and publication-year filters (2001–2025) were applied to refine results. The initial search yielded over 600 papers. Each record was screened manually by reviewing titles, abstracts, and keywords to determine relevance. Additional references were identified through citation tracking to include studies not captured in the initial query.

4.4 Inclusion and Exclusion Criteria

To keep the analysis meaningful, we applied the following rules:

Inclusion:

- Peer-reviewed studies from 2001 to 2025.
- The mean, median, or mode is clearly discussed or used for missing data.
- Quantitative or empirical assessment of imputation results.
- Written in English.

Exclusion:

- Those studies discuss only advanced machine learning imputers without mentioning central-tendency methods.
- Articles with no empirical data or methods are evident.

- Non-peer-reviewed sources, such as dissertations, preprints, or commentaries.

Why these criteria were selected: These rules ensure careful methods, ease of comparison, and alignment with review goals. Limiting peer-reviewed and empirically tested studies makes the findings more dependable; focusing on central-tendency methods keeps the scope clear. Excluding purely advanced imputation methods without a baseline keeps the review focused and makes cross-study comparisons fair. After screening and eligibility checks, 195 studies fit the criteria for detailed analysis across such fields as healthcare, finance, environmental science, engineering, and social analytics.

4.5 Data Extraction and Categorization

For each selected publication, metadata and relevant analytical details were systematically extracted. The extracted information included the author(s), publication year, application domain, data type (numerical, categorical, or mixed), imputation method(s) used, performance metrics (e.g., RMSE, MAE, bias, accuracy), and percentage of missingness. The extracted data were then categorized according to methodological type and domain context, enabling cross-comparative evaluation of results. Manual verification ensured accuracy and consistency across all recorded entries

4.6 Quality Assessment

To evaluate methodological quality, each study was reviewed using five key criteria:

- Clarity of research objectives.
- Transparency of imputation procedures.
- Presence of empirical validation.
- Reproducibility of results
- Direct relevance to central-tendency-based methods.

Studies that met at least four of these criteria were retained for synthesis. Approximately 82% of the selected papers satisfied all five conditions, indicating a robust and credible evidence base.

4.7 Data Synthesis and Integration

A mixed-method synthesis approach was adopted. Qualitative analysis identified recurring trends, methodological advances, and domain-specific patterns. Quantitative synthesis aggregated reported performance metrics such as RMSE and MAE to facilitate comparative evaluation across techniques. To enhance reliability, findings were cross validated across multiple domains and triangulated between statistical performance measures and contextual insights reported in the literature.

4.8 Ethical and Reproducibility Considerations

All data utilized in this review originated from publicly available and properly cited research sources. The review methodology and inclusion process were explicitly documented to ensure transparency and reproducibility. No conflicts of interest or ethical issues were identified during the analysis.

4.9 Summary

This structured review methodology provided a comprehensive and replicable approach for analyzing the literature on central tendency-based imputation techniques. By systematically identifying, evaluating, and synthesizing evidence from 195 peer reviewed studies, the process established a strong empirical foundation for the comparative and critical analysis presented in the next section.

V.CRITICAL REVIEW AND COMPARATIVE ANALYSIS

5.1 Overview

This section synthesizes evidence from 195 peer-reviewed publications spanning 2001–2025 to evaluate how central tendency-based imputation methods have evolved and performed across domains. The discussion compares mean, median, and mode substitutions with respect to accuracy, robustness, domain suitability, and integration within advanced analytical frameworks. Emphasis is placed on both temporal developments and cross-disciplinary application patterns.

5.2 Temporal Distribution of Studies

During the early years of the review period (2001–2009), research largely focused on traditional statistical methods, with mean imputation used in most empirical studies. Between 2010 and 2015, publications increasingly adopted hybrid approaches, combining mean or median substitution with algorithms such as regression or k-nearest neighbors. From 2016 onward, research trends shifted toward machine-learning-assisted and context-aware frameworks. Even so, mean, median, and mode imputations continued to appear as comparative baselines, confirming their continuing methodological significance despite the growing complexity of modern imputers.

5.3 Domain-Wise Application Trends

Central-tendency methods exhibit distinct domain-specific patterns:

- **Healthcare and Bioinformatics:** Median imputation typically produces the most stable outcomes because it resists distortion from outliers common in biomedical data. Studies report error reductions of roughly 10–15% relative to mean substitution in skewed datasets.
- **Finance and Econometrics:** Mean replacement remains prevalent due to near-normal distributions and the need for computational efficiency.
- **Environmental and Sensor Analytics:** Median and mode imputations provide better resistance to irregular noise and measurement drift.
- **Social Science and Survey Research:** Mode imputation performs best for categorical variables such as opinion scales or demographic categories, preserving the original class structure.

Overall, these findings highlight the adaptability of central-tendency approaches across heterogeneous data environments.

5.4 Comparative Performance of Mean, Median, and Mode Imputation

Aggregate evidence across reviewed studies indicates:

- **Mean Imputation** offers speed and conceptual simplicity but often underestimates variance and weakens correlation structures when missingness is extensive.
- **Median Imputation** achieves lower average error (RMSE/MAE) in skewed or non-Gaussian distributions and is less sensitive to extreme values.
- **Mode Imputation** is most effective for categorical or ordinal data but may bias results when category frequencies are highly uneven; probabilistic extensions can alleviate this limitation.

Although sophisticated algorithms—such as Multiple Imputation by Chained Equations (MICE) or autoencoder-based models—yield higher predictive accuracy, their computational cost is considerably greater. Central-tendency methods therefore remain practical where interpretability and efficiency are prioritized.

5.5 Integration with Advanced Analytical Techniques

Recent research demonstrates that central-tendency measures frequently function as initialization or calibration steps within hybrid frameworks. Examples include their use in KNN-based imputers, random-forest pipelines, and deep-autoencoder architectures. Incorporating mean or median values as starting estimates has been shown to improve convergence speed and reduce bias during iterative optimization. The trend toward adaptive hybrid imputers underscores that simple statistical measures continue to play a foundational role even within sophisticated, learning-driven systems.

5.6 Limitations and New Challenges

Technical problems:

- Variance shrinkage: These methods reduce variation and, by so doing, can mask valid relationships between variables.
- Sensitivity to MNAR data: performance drops if missing values depend on unobserved data.
- Bias in diverse or imbalanced data: Large variations among groups or classes that are imbalanced lead to increased bias.

Method limitations:

- Limited to capturing complex links between variables.
 - Less flexible for changing to different data patterns.
 - It is difficult to fit smoothly with advanced analysis or machine learning tools without extra work.
- New work is trying to fix these issues with ideas like distribution-aware weighting, context-based strategy choices, and tying in with representation learning to better reflect the data.

5.7 Important Comparisons and Gaps in Research

Considering different domains, mean, median, and mode imputations work only under very narrow conditions. The success of these methods greatly depends on how data are shaped, why data are missing, and the type of variables. The main gap is that there are no clear, shared rules for choosing between these methods when you have mixed types of data or large numbers of missing values. Few studies consistently compare simple central tendency methods against more advanced imputers using the same tests, making it difficult to draw general conclusions. Finally, limited work has been done by incorporating ethics, handling uncertainty, and data protection, especially in safety-critical or highly regulated fields.

5.8 Summary

Over two decades of research confirm that mean, median, and mode imputations remain integral to modern data-analysis practice. While contemporary machine-learning and probabilistic imputers achieve finer precision, these classical techniques continue to provide the necessary transparency, efficiency, and baseline reliability required for benchmarking and rapid deployment. Their inclusion in hybrid and automated systems illustrates a balanced future in which interpretability and computational intelligence coexist.

VI. BEST PRACTICES AND RECOMMENDATIONS

6.1 Overview

Although the sophistication of machine learning and probabilistic imputation methods has grown rapidly, central-tendencybased techniques remain indispensable for practical data preprocessing. This section consolidates insights drawn from the 195 reviewed studies and offers structured recommendations for the effective selection and application of mean, median, and mode imputation strategies across different research contexts.

6.2 General Guidelines for Central-Tendency-Based Imputation

- **Assess the Nature of Missingness:** Apply mean, median, or mode substitution primarily when data are missing completely or partially at random. When missingness depends on unobserved values, more complex model-based or hybrid methods are necessary.
- **Evaluate the Extent of Missing Data:** Central-tendency imputations are generally dependable when the proportion of missing values does not exceed 10– 15%. Beyond that threshold, they may introduce bias and reduce variance realism.
- **Match Technique to Data Characteristics:**
 - Use means imputation for continuous variables that are approximately normal and free from extreme values.
 - Use median imputation for skewed or outlier-prone datasets.

- Use mode imputation for categorical or ordinal variables.
- **Validate Statistical Consistency:** After imputation, verify that the dataset’s key descriptive statistics (mean, standard deviation, correlations) remain coherent. Compare pre- and post-imputation summaries to detect potential distortion.
- **Document and Report Transparently:** Record the rationale, parameters, and validation procedures for each imputation step to ensure reproducibility and compliance with data governance standards.

6.3 Domain-Specific Recommendations

TABLE I : DOMAIN SPECIFIC RECOMMENDATIONS

Domain	Preferred Technique	Rationale
Healthcare & Bioinformatics	Median Imputation	Provides stability in the presence of outliers and non-normal distributions.
Finance & Economics	Mean Imputation	Efficient for low-variance numerical data where normality approximations hold.
Environmental & Sensor Analytics	Median or Mode	Handles non-linear variations and measurement noise effectively.
Social Sciences & Surveys	Mode Imputation	Preserves the categorical nature of response data.
Engineering & Industrial Monitoring	Mean or Median	Depends on process symmetry and the system’s tolerance for anomalies.

VII. RESEARCH GAPS AND FUTURE DIRECTIONS

7.1 Overview

Despite the significant advances achieved with the use of central-tendency methods for data imputation, there are still some significant limitations. The identification of such shortcomings must be addressed to enhance the accuracy, flexibility, and reliability of the imputation process for analytics in real-world datasets. This section identifies the main weaknesses found in studies and suggests the most impactful future research directions.

7.2 Identified Research Gaps

1. **Lack of standardized benchmarks and evaluation protocols:**
Comparing studies is hard because there aren’t shared benchmark datasets or common evaluation rules. Most of the work uses data from specific fields, so results may not generalize or be easy to reproduce.
2. **Inconsistent performance metrics:**
Different metrics are used, such as RMSE, MAE, bias, and coverage; these do not always point in the same direction. This complicates any attempt to compare imputation quality fairly and to combine results across different studies.
3. **Limited coverage of complex and high-missingness data:**
Although mixed-type datasets with more than 40% missing ratio occur very frequently in many real-world applications, including sensors, health records, and industrial monitoring, they are not studied thoroughly.
4. **Insufficient Consideration of Ethical and Governance Issues:**
Ethics and accountability are important in modern analytics, and people often forget about transparency, auditability, data lineage, and regulatory compliance.

7.3 Priority Research Directions

Future work should be centred around these high-impact areas:

- **Standardization and benchmarking development:**
Create open benchmark datasets and agreed evaluation metrics that allow fair, reproducible, cross-domain comparisons of imputation methods.
- **Adaptive and explainable imputation frameworks:**
Therefore, build context-aware systems that select imputation strategies in real-time and provide explanations, including uncertainty estimates, to improve trust and usability.
- **Integration of Advanced Learning Architectures:**
It combines central-tendency initialization with deep learning, meta-learning, or reinforcement learning while keeping the results understandable.
- **Privacy-preserving and ethical imputation:**
Prioritize approaches like federated learning, secure computation, and auditable pipelines for meeting privacy rules and ethical data governance standards.

7.4 Summary

While classical central-tendency methods are well understood, new challenges regarding scalability, transparency, and data governance require focused innovation. By advocating for standardized evaluation, adaptive and explainable methods, and privacy-aware implementations, future research will be in a position to strengthen rigour and practical reliability of missing value imputation in today's analytics.

VIII. CONCLUSION

This systematic review analyzed 195 peer-reviewed studies published between 2001 and 2025 to assess the role of central tendency measures—mean, median, and mode—in missing-value imputation across diverse analytical domains. The findings demonstrate that, despite advances in machine-learning and probabilistic methods, these classical techniques remain foundational due to their interpretability, computational efficiency, and ease of integration into analytical pipelines. The review confirms that mean imputation is most effective for approximately normal data with limited missingness, median substitution offers greater robustness for skewed distributions, and mode imputation remains the preferred option for categorical attributes. When applied under appropriate conditions, these methods provide stable, transparent, and domain-adaptable solutions. Beyond standalone use, central-tendency techniques increasingly function within hybrid and automated systems, serving as initialization mechanisms, benchmarking baselines, or fallback estimators in ensemble and deep-imputation frameworks. This highlights their continued methodological relevance in modern workflows. Overall, the principal contribution of this review lies in consolidating two decades of evidence to clarify when and how central tendency-based imputations should be applied effectively. The results underscore that simplicity and transparency remain essential virtues in data preprocessing, reinforcing the credibility, reproducibility, and accountability of contemporary data driven analysis.

REFERENCES

- [1]. D. D. Pandya, A. Jadeja, S. Degadwala and D. Vyas, "Ensemble Learning based Enzyme Family Classification using n-gram Feature," 2022 6th International Conference on Intelligent Computing and Control Systems (ICICCS), Madurai, India, 2022, pp. 1386-1392, doi: 10.1109/ICICCS53718.2022.9788292.
- [2]. Dr. Darshanaben Dipakkumar Pandya, Dr. Abhijeetsinh Jadeja, Dr. Sheshang D. Degadwala, " An Applied Secant Method for Recovered Missing Mass Values in Data Mining , IInternational Journal of Scientific Research in Computer Science, Engineering and Information Technology(IJSCSEIT), ISSN : 2456-3307, Volume 8, Issue 2, pp.97-103, March-April-2022. Available at doi : <https://doi.org/10.32628/CSEIT22823>
- [3]. Dr. Darshanaben Dipakkumar Pandya, Dr. Abhijeetsinh Jadeja, Dr. Sheshang D. Degadwala, " An Applied Mean Substitutions Technique for Detection of Anomalous Value in Data Mining , International Journal of Scientific Research in Science and Technology(IJSRST), Online ISSN : 2395-602X, Print ISSN : 2395-6011, Volume 9, Issue 2, pp.103-108, March-April-2022. Available at doi : <https://doi.org/10.32628/IJSRST229212>
- [4]. Dr. Darshanaben Dipakkumar Pandya, Dr. Abhijeetsinh Jadeja, Dr. Sheshang D. Degadwala, " An Applied N C Differentiation Interpolation technique for improved random Anomalous values in Data Mining, International Journal of Scientific Research in Science, Engineering and Technology(IJSRSET), Print ISSN : 2395-1990, Online

ISSN : 2394-4099, Volume 9, Issue 2, pp.86-92, March-April-2022. Available at doi : <https://doi.org/10.32628/IJSRSET229218>

- [5]. Dr. Darshanaben Dipakkumar Pandya, Dr. Abhijeetsinh Jadeja, Dr. Sheshang D. Degadwala, " N B Interpolation Technique for Improved Arbitrarily missing values in Data Mining, International Journal of Scientific Research in Science, Engineering and Technology(IJSRSET), Print ISSN : 2395-1990, Online ISSN : 2394-4099, Volume 9, Issue 2, pp.79-85, March-April-2022. Available at doi : <https://doi.org/10.32628/IJSRSET229217>
- [6]. Dr. Darshanaben Dipakkumar Pandya, Dr. Abhijeetsinh Jadeja, Dr. Sheshang Degadwala, " An Applied Variation Anomaly Technique for Detection of Irregular Values in Data Mining, International Journal of Scientific Research in Computer Science, Engineering and Information Technology(IJSRCSEIT), ISSN : 2456-3307, Volume 8, Issue 2, pp.143-148, March-April-2022. Available at doi : <https://doi.org/10.32628/CSEIT228228>
- [7]. Gaur, S., Pandya, D.D., Soni, D. (2020). Closest Fit Approach Through Linear Interpolation to Recover Missing Values in Data Mining. In: Yang, X.S., Sherratt, S., Dey, N., Joshi, A. (eds) Fourth International Congress on Information and Communication Technology. Advances in Intelligent Systems and Computing, vol 1041. Springer, Singapore. https://doi.org/10.1007/978-981-15-0637-6_44
- [8]. Gaur, S., Pandya, D.D., Sharma, M.K. (2020). Applied N F Interpolation Method for Recover Randomly Missing Values in Data Mining. In: Yang, X.S., Sherratt, S., Dey, N., Joshi, A. (eds) Fourth International Congress on Information and Communication Technology. Advances in Intelligent Systems and Computing, vol 1027. Springer, Singapore. https://doi.org/10.1007/978-981-32-9343-4_38
- [9]. S. Gaur and D. Dipakkumar Pandya, "Closest Fit Approach for Pattern Designing to Recovered Anomalous Values in Data Mining," 2018 Second World Conference on Smart Trends in Systems, Security and Sustainability (WorldS4), London, UK, 2018, pp. 308-312, doi: 10.1109/WorldS4.2018.8611610.
- [10]. Gaur, S., Pandya, D.D., Soni, D. (2020). Closest Fit Approach Through Linear Interpolation to Recover Missing Values in Data Mining. In: Yang, X.S., Sherratt, S., Dey, N., Joshi, A. (eds) Fourth International Congress on Information and Communication Technology. Advances in Intelligent Systems and Computing, vol 1041. Springer, Singapore.
- [11]. Pandya, D., Jadeja, A., Gour, S., Trivedi, S.B., Patel, H.H., Jadeja, P.U. (2024). An Analytical Perspective of Missing Values in Machine Learning. In: Rathore, V.S., Piuri, V., Babo, R., Tiwari, V. (eds) Emerging Trends in Expert Applications and Security. ICE-TEAS 2024. Lecture Notes in Networks and Systems, vol 1037. Springer, Singapore. https://doi.org/10.1007/978-981-97-3991-2_24
- [12]. Sneha B. Patel and Dr. Darshanaben Dipakkumar Pandya, "Analysis of Adaptive Approach Through Statistical Trends and Measures of Central Tendency," Int. J. Sci. Res. Computer. Sci. Eng. Inf. Technol, vol. 11, no. 4, pp. 69–74, Jul. 2025, Doi: 10.32628/CSEIT2511405.
- [13]. Pandya, Darshanaben Dipakkumar. "An Applied Missing Value Recovery Using Adaptive Techniques for Statistical Trends and Central Tendency Analysis." International Journal of Recent Trends in Science, Technology & Management (IJRTSTM), September 27, 2025, 38–52. <https://doi.org/10.5281/zenodo.17173257>.
- [14]. Pandya, Darshanaben Dipakkumar. "Enhancing Data Mining Accuracy Through Muller-Based Anomaly and Missing Data Handling." International Journal of Recent Trends in Science, Technology & Management (IJRTSTM), pp. 80–97, 27 Sept. 2025. <https://doi.org/10.5281/zenodo.17173257>.
- [15]. Pandya, Darshanaben Dipakkumar. "Smart Technology Trends: A Quantitative and Analytical Approach." Smart Technology Trends, 5 Sept. 2025, pp. 132–141. ISBN 978-81-987316-8-5. <https://doi.org/10.11411/HENXI.2025587269>.
- [16]. Pandya, Darshanaben Dipakkumar. "AI-Powered Advancements in Software Design and Development Through Smart Technology." Smart Technology Trends, 5 Sept. 2025, pp. 142–151. ISBN 978-81-987316-8-5. <https://doi.org/10.11411/HENXI.2025124965>.

- [17]. Pandya, Darshanaben Dipakkumar. "Systematic Handling of Missing Data in Mining Processes: Insights from Statistical Trends." *Smart Technology Trends*, Henxi Publication, 5 Sept. 2025, pp. 152–161. ISBN 978-81-987316-8-5. <https://doie.org/10.11411/HENXI.2025952409>.
- [18]. Pandya, Darshanaben Dipakkumar. "Statistical Trend–Based Detection of Exceptional and Inconsistent Values in Large Datasets." *Smart Technology Trends*, Henxi Publication, 5 Sept. 2025, pp. 170–182. ISBN 978-81-987316-8-5. <https://doie.org/10.11411/HENXI.2025596023>.
- [19]. Pandya, D.D., Satish, B.P., Khushbu (2026). Improving Data Mining Reliability Through Applied C-B Techniques for Addressing Misplaced Values. In: Rathore, V.S., Tavares, J.M.R.S., Hanumanthappa, M., Surendiran, B. (eds) *Universal Threats in Expert Applications and Solutions. UNI-TEAS 2025. Lecture Notes in Networks and Systems*, vol 1450. Springer, Singapore. https://doi.org/10.1007/978-981-96-7289-9_2
- [20]. Pandya, Darshanaben Dipakkumar, Borate Pooja Satish, and Khushbu. "Improving Data Mining Reliability Through Applied C-B Techniques for Addressing Misplaced Values." *Universal Threats in Expert Applications and Solutions (UNI-TEAS 2025)*, edited by Vijay Singh Rathore, Joao Manuel R. S. Tavares, M. Hanumanthappa, and B. Surendiran, *Lecture Notes in Networks and Systems*, vol. 1450, Springer, Singapore, 2026, pp. 11–23. https://doi.org/10.1007/978-981-96-7289-9_2.