



Neutrosophic–Fuzzy Multi-View Imputation for Incomplete Multivariate Measurement Data

Ahmed Hatip^{1,*}

¹ Department of Mathematics, Gaziantep University, Gaziantep, Turkey

Email: kollnaar5@gmail.com

Received: November 20, 2024 Revised: December 30, 2024 Accepted: January 03, 2025 ★ Corresponding author

ABSTRACT

Three reconstructions of the same missing cell can each be defensible and still disagree materially. Let $e = (e_N, e_S, e_R)$ denote contextual-neighborhood, low-rank, and cross-variable ridge estimates. The central question is therefore not only which value should replace x_{ij} , but how much coherent evidence supports that replacement. Neutrosophic–Fuzzy Multi-View Imputation (NFMVI) addresses this question by assigning each source a fuzzy reliability $\mu_s \in (0, 1]$ and a Gaussian concordance $a_s = \exp[-(e_s - c)^2 / (2\gamma^2)]$ around a reliability-weighted center c . These quantities generate the source state $(T_s, I_s, F_s) = (\mu_s a_s, 1 - a_s, 1 - \mu_s)$, from which $w_s = T_s / \sum_r T_r$ yields the convex reconstruction $\hat{x} = \sum_s w_s e_s$. The same state produces a cell-level unresolved-evidence fraction $U = 1 - \sum_s T_s / \sum_s (T_s + I_s + F_s)$, so reconstruction and uncertainty are generated by one mechanism rather than by separate post-processing. Evaluation uses two real multivariate datasets, MCAR, value-dependent MAR, and structured channel deletion at 10–30%, giving 144 common deterministic masks. The evidence is deliberately mixed. On the macroeconomic panel, NFMVI attains standardized RMSE 0.3425, below KNN (0.3776), Bayesian-ridge chained imputation (0.3794), and iterative SVD (0.4987); on the diabetes covariates, NFMVI (0.7641) and chained imputation (0.7659) are nearly indistinguishable overall, with chained imputation remaining better under MAR. The uncertainty signal is more distinctive: macro error-screening AUC averages 0.8156, and RMSE rises from 0.1717 in the lowest- U quartile to 0.6202 in the highest. Paired bootstrap contrasts, correlation-structure error, ablation, and parameter sensitivity therefore support a narrower conclusion than universal accuracy dominance: NFMVI provides an interpretable fuzzy–neutrosophic rule for reconciling heterogeneous imputers while retaining cell-level evidence conflict for downstream scrutiny.

Keywords: Fuzzy reliability ▪ Single-valued neutrosophic information ▪ Missing data ▪ Multivariate imputation ▪ Evidence fusion ▪ Uncertainty quantification ▪ Incomplete data

1. INTRODUCTION

A missing entry is often treated as though it had a single natural replacement. In a multivariate matrix $X \in \mathbb{R}^{n \times d}$ observed only on Ω , however, one cell $(i, j) \notin \Omega$ may support several incompatible reconstructions. A local estimator may give e_N , a low-rank model e_S , and a conditional model e_R . When $\max_{s,r} |e_s - e_r| \approx 0$, their agreement itself is evidence; when

the spread is large, reporting only $\hat{x}_{ij}$ discards information about epistemic conflict. This paper takes that conflict as part of the imputation problem rather than as an incidental by-product of model choice.

The statistical context is well established. Rubin’s framework distinguishes missingness according to how the probability of being unobserved depends on observed and unobserved quantities [1]; Little’s MCAR test addresses the narrower

diagnostic question of whether missingness is compatible with complete randomness [2]. Likelihood-based procedures such as EM [3] and multiple/chained imputation [4, 5] model incomplete data probabilistically, while empirical reconstruction methods exploit neighborhood, tree, or low-rank structure [6, 7, 8, 9]. These approaches differ in assumptions and inferential goals, but most single-imputation outputs ultimately reduce a missing cell to one value $\hat{x}_{ij}$.

The proposed formulation separates two questions that are usually conflated. The first is *source reliability*: how trustworthy is an imputation mechanism around the current cell or feature? This is represented fuzzily by $\mu_s \in [0, 1]$, following the graded-membership principle of fuzzy sets [10]. The second is *source concordance*: how compatible is e_s with the remaining views? Single-valued neutrosophic information provides three independent coordinates for support, indeterminacy, and falsity [11, 12]. NFMVI uses that freedom constructively: reliability and agreement are not two names for the same uncertainty, and they enter different components of (T_s, I_s, F_s) .

This separation also distinguishes NFMVI from fuzzy imputation based primarily on cluster membership. Fuzzy c-means has been adapted directly to incomplete vectors [13], combined with support-vector regression and evolutionary optimization [14], coupled with grey-information and regression mechanisms [15], and used in traffic-data reconstruction [16]. Those studies establish fuzzy membership as a useful computational device for incomplete data. Here, by contrast, the fuzzy quantity is attached to the *imputation source*; the neutrosophic indeterminacy coordinate then records disagreement among heterogeneous reconstruction mechanisms rather than cluster ambiguity.

For a source set $\mathcal{S} = \{N, S, R\}$, NFMVI constructs a reliability-weighted center c and concordance $a_s \in (0, 1]$, then defines

$$T_s = \mu_s a_s, \quad I_s = 1 - a_s, \quad F_s = 1 - \mu_s.$$

The fusion weight $w_s \propto T_s$ is large only when reliability and agreement occur jointly, giving $\hat{x}_{ij} = \sum_s w_s e_s$. The state is intentionally non-normalized: $T_s + I_s + F_s = 1 + I_s F_s \in [1, 2]$, so excess mass appears exactly when disagreement and unreliability coexist. A second output, U_{ij} , retains the unresolved fraction of the same evidence state instead of discarding disagreement after the numerical value has been filled.

The study is organized around three claims rather than a generic accuracy claim. The *representation claim* is that heterogeneous reconstruction sources can be decomposed into measurable fuzzy reliability and neutrosophic conflict. The *algorithmic claim* is that truth-weight normalization yields a positive convex update, with transparent source-weight ratios and bounded uncertainty. The *empirical claim* is deliberately qualified: across two real datasets, three missingness mechanisms, three deletion rates, and 144 common masks, the method should be judged jointly by reconstruction error, structural preservation, and whether U_{ij} ranks genuinely difficult cells. The analysis therefore emphasizes tables, paired bootstrap intervals, uncertainty strata, ablation, and sensitivity rather than a single headline score. All cited material predates 1 March 2025, consistent with the intended March

2025 volume.

The remainder follows the logic of the evidence construction. Section 2 positions the three reconstruction assumptions and the fuzzy–neutrosophic representation. Section 3 defines the incomplete matrix and the candidate views; Section 4 gives the fusion rule and Algorithm 1; Section 5 derives its mathematical properties. Section 6 fixes the numerical protocol, Section 7 reports the main results, Section 8 interprets robustness and component dependence, Section 9 states limitations, and Section 10 concludes.

2. RELATED WORK

2.1 Three reconstruction assumptions

Existing imputation methods can be read as different assumptions about where recoverable information resides. A *local* assumption searches for nearby observations: KNN reconstruction was used, for example, in early microarray missing-value estimation [6]. A *conditional* assumption predicts each incomplete variable from the remaining variables; chained equations operationalize this idea through repeated conditional models [4]. A *global structural* assumption seeks shared low-dimensional organization, as in matrix-completion and spectral-regularization methods [8, 17, 9]. Tree ensembles such as missForest provide another nonlinear conditional alternative and accommodate mixed variable types [7]. Reviews of missing-data methods emphasize that performance and appropriateness depend on the missingness process, data structure, and analysis objective rather than on one universally preferred procedure [18, 19].

The probabilistic tradition addresses a related but different goal. EM iteratively computes likelihood quantities in incomplete-data models [3], while multiple imputation is designed to propagate uncertainty across repeated completed datasets rather than to treat one fill-in as known [5]. NFMVI is a single-imputation method and does not replace that inferential machinery. Its narrower objective is to exploit disagreement among several transparent point-reconstruction views while preserving a cell-level indicator of that disagreement.

Deep models add richer nonlinear assumptions. GAIN uses adversarial learning [20]; recurrent approaches include GRU-D-style handling of missing multivariate series [21] and BRITS [22]; GP-VAE introduces a probabilistic latent-process formulation [23]; GRIN exploits graph structure [24]; and SAITS uses self-attention [25]. These methods define the broader imputation landscape, but they are not numerical baselines here because the two study matrices are compact and the research question concerns interpretable evidence fusion rather than large-scale representation learning.

2.2 Fuzzy and neutrosophic representations of incomplete information

Fuzzy sets encode graded membership $\mu(x) \in [0, 1]$ [10]; intuitionistic fuzzy sets separate membership and non-membership while retaining a hesitation remainder [26]. In incomplete-data analysis, [13] modified fuzzy c-means so incomplete vectors could enter the clustering objective directly. Subsequent work hybridized fuzzy c-means with support-vector regression and a genetic algorithm [14], combined grey-based fuzzy clustering with feature selection and re-

gression [15], and integrated fuzzy c-means into taxi-GPS missing-data reconstruction [16]. These papers support the use of graded memberships in reconstruction, but their memberships are primarily associated with clusters, prototypes, or similarity structure.

Single-valued neutrosophic sets instead represent truth, indeterminacy, and falsity independently in $[0, 1]$ [11]. Similarity and entropy constructions subsequently formalized information relationships among such triples [12, 27]. For imputation-source fusion, that independence is useful because two distinct failure modes must be represented: a source can be empirically unreliable even when it agrees with its peers, or reliable in isolation but strongly discordant on the current cell. A single complement $1 - \mu_s$ cannot distinguish these cases.

2.3 Position of NFMVI

NFMVI therefore places fuzziness and neutrosophy at a different level from the reviewed fuzzy imputers. The fuzzy variable μ_s is a *source-reliability membership*; it is not the membership of row i in a cluster. Concordance a_s is computed independently from the candidate values, and the mapping

$$(T_s, I_s, F_s) = (\mu_s a_s, 1 - a_s, 1 - \mu_s)$$

keeps reliable support, inter-view conflict, and source unreliability visible simultaneously. The methodological gap addressed here is correspondingly narrow: among the reviewed work, heterogeneous local, low-rank, and conditional imputers are not fused through a source-specific single-valued neutrosophic state that also yields a cell-level uncertainty score. The claim is about this fusion architecture, not about neutrosophic theory being new to incomplete or uncertain information.

3. PROBLEM FORMULATION AND RECONSTRUCTION VIEWS

3.1 Incomplete multivariate matrix

Let $X = (x_{ij}) \in \mathbb{R}^{n \times d}$ be the unknown complete matrix and let $M = (m_{ij}) \in \{0, 1\}^{n \times d}$ denote the artificial missingness mask,

$$m_{ij} = 1 \iff x_{ij} \text{ is hidden for evaluation.} \quad (1)$$

The observed set is $\Omega = \{(i, j) : m_{ij} = 0\}$. For each feature j , the observed mean and standard deviation are estimated only from $\Omega_j = \{i : m_{ij} = 0\}$,

$$\hat{\mu}_j = \frac{1}{|\Omega_j|} \sum_{i \in \Omega_j} x_{ij}, \quad \hat{\sigma}_j^2 = \frac{1}{|\Omega_j| - 1} \sum_{i \in \Omega_j} (x_{ij} - \hat{\mu}_j)^2. \quad (2)$$

All reconstruction is performed on $z_{ij} = (x_{ij} - \hat{\mu}_j) / \hat{\sigma}_j$. Thus the hidden values do not contribute to preprocessing parameters. Let B be the standardized matrix with $B_{ij} = z_{ij}$ on Ω and $B_{ij} = \text{NA}$ otherwise.

The goal is not only to estimate $\hat{z}_{ij}$ for $m_{ij} = 1$, but also to produce $U_{ij} \in [0, 1]$ such that larger U_{ij} tends to accompany larger absolute reconstruction error $|\hat{z}_{ij} - z_{ij}|$. Reconstruction quality is therefore evaluated jointly with uncertainty ranking.

3.2 View 1: fuzzy contextual neighborhood

At iteration ℓ , let $Y^{(\ell)}$ be the currently completed matrix. For a missing cell (i, j) , contextual distance excludes the target feature,

$$d_{ir}^{(-j)} = \frac{1}{d-1} \sum_{q \neq j} (Y_{iq}^{(\ell)} - Y_{rq}^{(\ell)})^2. \quad (3)$$

Among rows r with observed B_{rj} , retain the k smallest distances and define $h_{ij} = \max\{0.25, \sqrt{\text{median}_r d_{ir}^{(-j)}}\}$. The Gaussian fuzzy relation

$$g_{ir} = \exp \left[-\frac{d_{ir}^{(-j)}}{2h_{ij}^2} \right] \quad (4)$$

produces the contextual estimate

$$e_N = \frac{\sum_r g_{ir} B_{rj}}{\sum_r g_{ir}}. \quad (5)$$

Its weighted local variance $v_N = \sum_r g_{ir} (B_{rj} - e_N)^2 / \sum_r g_{ir}$ is converted to fuzzy reliability

$$\mu_N = \text{clip} \left(e^{-v_N/2}, \bar{g}, 0.02, 1 \right), \quad \bar{g} = \frac{1}{k} \sum_r g_{ir}. \quad (6)$$

Hence a tight, close neighborhood has larger μ_N than a diffuse neighborhood.

3.3 View 2: low-rank reconstruction

Given $Y^{(\ell)}$, center its columns and compute the rank- r approximation

$$S^{(\ell)} = \bar{Y} + U_r \Sigma_r V_r^T. \quad (7)$$

For target feature j , the candidate is $e_S = S_{ij}^{(\ell)}$. Its source reliability is estimated on actually observed cells of column j ,

$$\mu_{S,j} = \text{clip} \left[\exp \left(-\frac{1}{2|\Omega_j|} \sum_{r \in \Omega_j} (S_{rj}^{(\ell)} - B_{rj})^2 \right), 0.02, 1 \right]. \quad (8)$$

The view is therefore trusted when the low-rank surface reconstructs known entries of that feature accurately.

3.4 View 3: cross-variable conditional reconstruction

For each feature j , a ridge model is fitted on observed rows,

$$(\hat{\beta}_{0j}, \hat{\beta}_j) = \arg \min_{\beta_0, \beta} \sum_{r \in \Omega_j} [B_{rj} - \beta_0 - Y_{r,-j}^{(\ell)} \beta]^2 + \lambda \|\beta\|_2^2, \quad (9)$$

with $\lambda = 1$. The candidate is $e_R = \hat{\beta}_{0j} + Y_{i,-j}^{(\ell)} \hat{\beta}_j$. Its reliability uses the observed-row mean squared residual v_R ,

$$\mu_{R,j} = \text{clip}(e^{-v_R/2}, 0.02, 1). \quad (10)$$

The three candidates $e = (e_N, e_S, e_R)$ therefore encode local, global-low-rank, and conditional dependence, while $\mu = (\mu_N, \mu_S, \mu_R)$ places them on a common fuzzy reliability scale.

4. PROPOSED NFMVI METHOD

4.1 Reliability-weighted agreement geometry

For a missing cell, define the fuzzy center

$$c = \frac{\sum_{s \in \mathcal{S}} \mu_s e_s}{\sum_{s \in \mathcal{S}} \mu_s}, \quad \mathcal{S} = \{N, S, R\}. \quad (11)$$

The source concordance is

$$a_s = \exp \left[-\frac{(e_s - c)^2}{2\gamma^2} \right], \quad \gamma > 0. \quad (12)$$

This second fuzzy membership answers a different question from μ_s : μ_s asks whether source s has been empirically reliable, while a_s asks whether its current estimate agrees with the other views. Because $a_s > 0$ for finite e_s and $\gamma > 0$, every source remains eligible for soft fusion.

4.2 Neutrosophic source state

Define

$$T_s = \mu_s a_s, \quad I_s = 1 - a_s, \quad F_s = 1 - \mu_s. \quad (13)$$

Truth support is therefore conjunctive: T_s is high only if both empirical reliability and current concordance are high. Indeterminacy depends purely on disagreement, while falsity depends purely on lack of reliability. The source weight is

$$w_s = \frac{T_s}{\sum_{r \in \mathcal{S}} T_r}, \quad \sum_s w_s = 1, \quad (14)$$

and the update is

$$\hat{z}_{ij} = \sum_{s \in \mathcal{S}} w_s e_s. \quad (15)$$

The weight ratio has the transparent form

$$\frac{w_s}{w_r} = \frac{\mu_s a_s}{\mu_r a_r}, \quad (16)$$

so a twofold reliability advantage and a twofold concordance advantage combine to a fourfold truth-weight advantage.

4.3 Cell-level neutrosophic uncertainty

NFMVI attaches

$$U_{ij} = 1 - \frac{\sum_s T_s}{\sum_s (T_s + I_s + F_s)} \quad (17)$$

to each reconstructed cell. Because U is computed from the same state that determines fusion, it has a direct evidence interpretation: uncertainty increases when truth mass is a smaller fraction of the available truth–indeterminacy–falsity mass. It is not trained against the hidden ground truth.

The matrix is initialized by observed-column medians. One NFMVI pass constructs S , refits the d conditional ridge models, computes the contextual view cell by cell, and applies (11)–(17). Four refinement passes are used as a fixed computation budget in the experiments; they are not presented as a proof of optimization convergence.

Algorithm 1 Neutrosophic–Fuzzy Multi-View Imputation (NFMVI)

Require: Incomplete standardized matrix B , mask M , neighbors k , rank r , agreement width γ , passes L

- 1: Initialize $Y^{(0)}$ by feature medians on observed entries
- 2: **for** $\ell = 0, \dots, L-1$ **do**
- 3: Compute rank- r reconstruction $S^{(\ell)}$ and reliabilities $\mu_{S,j}$
- 4: Fit ridge models (9); obtain $R^{(\ell)}$ and $\mu_{R,j}$
- 5: **for** each missing cell (i, j) **do**
- 6: Compute (e_N, μ_N) by (3)–(6)
- 7: Set $e \leftarrow (e_N, S_{ij}^{(\ell)}, R_{ij}^{(\ell)})$, $\mu \leftarrow (\mu_N, \mu_{S,j}, \mu_{R,j})$
- 8: $c \leftarrow (\sum_s \mu_s e_s) / (\sum_s \mu_s)$
- 9: **for** $s \in \{N, S, R\}$ **do**
- 10: $a_s \leftarrow \exp[-(e_s - c)^2 / (2\gamma^2)]$
- 11: $(T_s, I_s, F_s) \leftarrow (\mu_s a_s, 1 - a_s, 1 - \mu_s)$
- 12: **end for**
- 13: $w_s \leftarrow T_s / (\sum_r T_r)$; $Y_{ij}^{(\ell+1)} \leftarrow \sum_s w_s e_s$
- 14: $U_{ij} \leftarrow 1 - \sum_s T_s / \sum_s (T_s + I_s + F_s)$
- 15: **end for**
- 16: Restore all observed cells exactly: $Y_{ij}^{(\ell+1)} \leftarrow B_{ij}$ for $M_{ij} = 0$
- 17: **end for**
- 18: **return** completed matrix $Y^{(L)}$ and uncertainty matrix U

5. MATHEMATICAL ANALYSIS

Proposition 1 (Validity and excess neutrosophic mass). *For $\mu_s, a_s \in [0, 1]$, the state in (13) satisfies $T_s, I_s, F_s \in [0, 1]$ and*

$$T_s + I_s + F_s = 1 + I_s F_s \in [1, 2]. \quad (18)$$

Proof. Boundedness is immediate from products and complements of unit-interval quantities. Moreover,

$$\mu_s a_s + (1 - a_s) + (1 - \mu_s) = 1 + (1 - a_s)(1 - \mu_s) = 1 + I_s F_s.$$

Since $0 \leq I_s F_s \leq 1$, the stated interval follows. $\square$

Equation (18) gives the non-normalized neutrosophic mass a concrete meaning. The surplus above one is zero when either the source agrees perfectly ($I_s = 0$) or is fully reliable ($F_s = 0$); it grows only when disagreement and unreliability occur together.

Proposition 2 (Convex reconstruction). *If $\gamma > 0$ and $\mu_s > 0$ for all s , then $w_s > 0$, $\sum_s w_s = 1$, and*

$$\min_s e_s \leq \hat{z}_{ij} \leq \max_s e_s. \quad (19)$$

If $e_s = e^$ for all s , then $\hat{z}_{ij} = e^*$ independently of μ_s and γ .*

Proof. Gaussian concordance gives $a_s > 0$, hence $T_s = \mu_s a_s > 0$. Equation (14) forms positive normalized coefficients, so (15) is a convex combination. The interval bound and common-estimate consistency follow directly. $\square$

Proposition 3 (Uncertainty bounds and limiting cases). *For the NFMVI state, $0 \leq U_{ij} < 1$. If $\mu_s = a_s = 1$ for all sources, then $U_{ij} = 0$. If $\sum_s T_s \downarrow 0$ while $\sum_s (I_s + F_s) > 0$, then $U_{ij} \uparrow 1$.*

Proof. Write $A = \sum_s T_s > 0$ and $B = \sum_s (I_s + F_s) \geq 0$. Equation (17) is $U = B / (A + B)$, proving the bound. The limiting statements follow by substitution. $\square$

Table 1. Datasets and fixed experimental protocol. Hidden cells are evaluated against their known standardized values.

Item	Macro	Diabetes
Rows n	203	442
Variables d	11	10
Missing mechanisms	MCAR, MAR, CHANNEL	
Missing rates	0.10, 0.20, 0.30	
Replicates per setting	8	
Main masks per dataset	$3 \times 3 \times 8 = 72$	
NFMVI (k, γ, r, L)	$(9, 0.7, \min(4, d-1), 4)$	
Bootstrap resamples	4000 paired	

This representation also makes monotonicity transparent. With the other masses fixed, $\partial U / \partial I_s = A / (A + B)^2 > 0$ and $\partial U / \partial F_s = A / (A + B)^2 > 0$, whereas increasing a truth component decreases U . Thus uncertainty has the intended local direction with respect to each neutrosophic component.

Remark 1 (Computation). *For a fixed number L of refinement passes, NFMVI consists of one rank- r SVD reconstruction per pass, d ridge fits, and one local-neighborhood calculation for each missing cell. The present implementation prioritizes transparent reproducibility over optimized nearest-neighbor indexing; the reported experiments use $d \leq 11$ after preprocessing, so this choice is not a computational bottleneck.*

6. EXPERIMENTAL DESIGN

6.1 Data and temporal cutoff

Two real multivariate matrices are used. The *Macro* panel is the statsmodels U.S. macroeconomic dataset [28]; after removing calendar identifiers, 11 continuous variables remain across 203 quarterly observations. Positive level variables are log-transformed before masking. The *Diabetes* matrix is the 442-by-10 standardized baseline covariate matrix distributed with scikit-learn [29]; its response variable is not used, so the experiment is purely multivariate reconstruction. The underlying diabetes study is the one analyzed by [30].

The publication cutoff is fixed at 28 February 2025. Every cited research source predates that date; the newest research article used in the bibliography is from 2023. The experimental design and interpretation therefore do not rely on literature that would have been unavailable for a March 2025 issue.

6.2 Missingness mechanisms

Three artificial mechanisms stress different reconstruction assumptions. MCAR samples target cells uniformly. The MAR design leaves the first feature observed and makes deletion probabilities of other features proportional to $\exp[\text{clip}(z_{i1}, -2.2, 2.2)]$, creating value-dependent missingness through an observed driver. CHANNEL missingness concentrates deletions in a random subset of features, representing partial sensor/channel dropout rather than independent cells. For each dataset, mechanism, rate $p \in \{0.1, 0.2, 0.3\}$, and replicate, a deterministic seed creates one common mask for all methods. At least a small observed support is preserved in every feature.

6.3 Comparators and metrics

Six methods are evaluated on identical masks: observed-column mean, observed-column median, distance-weighted 5-NN, Bayesian-ridge iterative imputation (MICE-BR), iterative rank- r SVD, and NFMVI. The MICE-BR implementation uses seven iterations and tolerance 10^{-3} ; some runs reach the fixed cap before its early-stopping criterion, so the comparator is described as a fixed-budget chained imputer rather than as universally converged. Implementations use scikit-learn [29]; SVD follows the iterative fill-and-reconstruct principle related to spectral imputation [9].

For hidden cells $\mathcal{H} = \{(i, j) : m_{ij} = 1\}$,

$$\text{RMSE} = \left[|\mathcal{H}|^{-1} \sum_{(i,j) \in \mathcal{H}} (\hat{z}_{ij} - z_{ij})^2 \right]^{1/2}, \quad (20)$$

$$\text{MAE} = |\mathcal{H}|^{-1} \sum_{(i,j) \in \mathcal{H}} |\hat{z}_{ij} - z_{ij}|. \quad (21)$$

Structure preservation is measured by $E_C = \|\hat{C} - C\|_F / d$, where C and $\hat{C}$ are complete- and imputed-matrix correlation matrices. For NFMVI, uncertainty quality is assessed by Spearman correlation between U_{ij} and absolute error and by AUC for identifying cells in the top quartile of absolute error. Paired bootstrap intervals use replicate-level RMSE differences $\Delta = \text{RMSE}_{\text{comparator}} - \text{RMSE}_{\text{NFMVI}}$; $\Delta > 0$ favors NFMVI.

7. RESULTS

7.1 Overall reconstruction by missingness mechanism

Table 2 reports RMSE averaged across the three deletion rates. The main result is not uniform dominance. On Diabetes, NFMVI is best under CHANNEL (0.8023) among the listed methods and remains close to MICE-BR under MCAR, but MICE-BR is better under MAR (0.7357 versus 0.7506). On Macro, NFMVI is best under MCAR (0.3059) and MAR (0.4110), while MICE-BR is clearly better under CHANNEL (0.2700 versus 0.3105). This cross-over is informative because the structured channel pattern favors strong conditional relationships when entire variables lose many observations.

Averaging all nine settings per dataset gives NFMVI RMSE 0.7641 on Diabetes and 0.3425 on Macro. The corresponding MICE-BR values are 0.7659 and 0.3794. Thus the overall Diabetes difference is only 0.0018 RMSE (0.24% relative to MICE-BR), while the Macro reduction is 0.0370 (9.74%). Relative to SVD, NFMVI reduces RMSE by 20.15% on Diabetes and 31.32% on Macro. These aggregates should be read together with the mechanism-specific reversals rather than in isolation.

7.2 Effect of missingness rate

Table 3 shows that error generally increases with deletion rate. For NFMVI on Diabetes, RMSE rises from 0.7320 to 0.7952 as p moves from 0.10 to 0.30, a change of 0.0632. On Macro the corresponding increase is 0.0719, from 0.3047 to 0.3767. The gradient is moderate relative to mean/median imputation, whose Macro errors remain near or above one standardized unit because they ignore multivariate dependence.

The rate profiles also expose dataset-specific behavior. On

Table 2. Mean standardized RMSE by dataset and missingness mechanism, averaged over deletion rates 10–30% and eight replicates per rate. Lowest value in each dataset–mechanism column is bold.

Method	Diabetes			Macro		
	CHANNEL	MAR	MCAR	CHANNEL	MAR	MCAR
Mean	1.0391	0.9956	1.0150	1.0158	1.1895	1.0085
Median	1.0774	1.0706	1.0569	1.0337	1.2022	1.0234
KNN	0.8934	0.9520	0.9119	0.3709	0.4338	0.3280
MICE-BR	0.8278	0.7357	0.7342	0.2700	0.4201	0.4482
SVD	0.9198	1.0010	0.9500	0.5230	0.5646	0.4085
NFMVI	0.8023	0.7506	0.7393	0.3105	0.4110	0.3059

Table 3. RMSE by missingness rate, averaged over the three mechanisms.

Data	Method	10%	20%	30%
Diabetes	Mean	1.0077	1.0332	1.0088
	Median	1.0568	1.0873	1.0607
	KNN	0.8401	0.9309	0.9863
	MICE-BR	0.7047	0.7830	0.8100
	SVD	0.9053	0.9776	0.9878
	NFMVI	0.7320	0.7650	0.7952
Macro	Mean	1.0594	1.0652	1.0892
	Median	1.0674	1.0776	1.1144
	KNN	0.3577	0.3749	0.4001
	MICE-BR	0.3866	0.3616	0.3901
	SVD	0.4662	0.5055	0.5243
	NFMVI	0.3047	0.3460	0.3767

Diabetes, MICE-BR is strongest at 10%, after which NFMVI has the smallest average RMSE at 20–30%. On Macro, NFMVI is best at every rate after averaging mechanisms. Consequently, the benefit of multi-view fusion is not simply a high-missingness effect; it depends on whether the three reconstruction views carry complementary signal in the underlying matrix.

7.3 Paired bootstrap contrasts

Table 4 uses paired masks and therefore removes variation caused by comparing different missing cells. Against SVD, all six intervals lie above zero and NFMVI wins at least 91.7% of paired Diabetes-CHANNEL cases and all paired cases for the other five dataset–mechanism combinations. Against KNN, all intervals also lie above zero, though the Macro MAR and MCAR effects are small. The MICE-BR contrasts are intentionally mixed: Diabetes MAR favors MICE-BR by 0.0149 RMSE because the interval for Δ is entirely negative, Macro CHANNEL favors MICE-BR by 0.0405, and Diabetes MCAR is statistically inconclusive at the interval level. Conversely, Macro MCAR strongly favors NFMVI ($\Delta = 0.1423$).

7.4 Correlation-structure preservation

Cell-wise RMSE does not guarantee preservation of multivariate association. Table 5 averages E_C across all mechanisms and rates. KNN has the smallest Diabetes correlation error (0.0558), followed by MICE-BR (0.0582) and NFMVI (0.0655). On Macro, MICE-BR is best (0.0271) and NFMVI is close (0.0282), while KNN is 0.0324. The proposed method therefore does not dominate structural preservation; its value is that it achieves strong cell reconstruction

while retaining correlation error near the best conditional/local methods.

7.5 Does neutrosophic uncertainty identify difficult cells?

Table 6 averages uncertainty diagnostics over missingness rates. The macro panel shows consistent discrimination: AUC ranges from 0.7974 under CHANNEL to 0.8294 under MAR, with Spearman correlations near 0.48–0.51. Diabetes is weaker but still above random ranking on average; AUC is highest under MCAR (0.7221) and lower under structured CHANNEL missingness (0.6535). The difference suggests that uncertainty is most informative when disagreement among the three views tracks local reconstruction difficulty, and less informative when entire channels lose support in a way that affects all views simultaneously.

The quartile analysis in Table 7 is more direct. For Diabetes, RMSE increases monotonically from 0.5492 in uncertainty quartile Q1 to 0.9495 in Q4, a factor of 1.73. For Macro, the increase is much sharper: 0.1717 $\rightarrow$ 0.6202, a factor of 3.61. Since U is computed without access to the held-out value, this ordering provides evidence that the fuzzy–neutrosophic state carries useful information about cell-level reconstruction risk.

8. DISCUSSION

8.1 Component dependence: ablation

Table 8 uses a separate 25% CHANNEL experiment with eight fixed masks per dataset. Removing the conditional ridge view causes the largest degradation: Diabetes RMSE rises from 0.7675 to 0.8442, and Macro from 0.2977 to 0.4837. Removing the fuzzy reliability layer also worsens both datasets, to 0.7708 and 0.3250. The neutrosophic concordance gate has a smaller but directionally favorable effect: RMSE becomes 0.7697 on Diabetes and 0.3099 on Macro when a_s is forced to one.

The ablation also prevents an overbroad component claim. On Macro CHANNEL masks, removing the low-rank view slightly improves RMSE (0.2807 versus 0.2977), and removing the neighborhood view gives 0.2920. Thus every source is not uniformly beneficial in every missingness regime. NFMVI is best interpreted as an adaptive compromise whose conditional source is especially important here; the low-rank and neighborhood views can introduce modest bias when their structural assumptions are poorly matched to concentrated channel dropout.

Table 4. Paired bootstrap RMSE advantage $\Delta = \text{RMSE}_{\text{comparator}} - \text{RMSE}_{\text{NFMVI}}$. Positive values favor NFMVI. Each contrast uses 24 paired observations (three rates $\times$ eight replicates) and 4000 bootstrap resamples.

Data	Mechanism	Comparator	Mean Δ	95% interval	NFMVI win rate
Diabetes	CHANNEL	KNN	0.0910	[0.0584, 0.1306]	0.9167
		MICE-BR	0.0255	[-0.0080, 0.0708]	0.5000
		SVD	0.1175	[0.0830, 0.1570]	0.9167
Diabetes	MAR	KNN	0.2014	[0.1791, 0.2222]	1.0000
		MICE-BR	-0.0149	[-0.0206, -0.0094]	0.2500
		SVD	0.2504	[0.2303, 0.2699]	1.0000
Diabetes	MCAR	KNN	0.1726	[0.1469, 0.2004]	1.0000
		MICE-BR	-0.0051	[-0.0206, 0.0133]	0.2500
		SVD	0.2107	[0.2022, 0.2196]	1.0000
Macro	CHANNEL	KNN	0.0604	[0.0235, 0.1018]	0.7500
		MICE-BR	-0.0405	[-0.0667, -0.0150]	0.2500
		SVD	0.2125	[0.1558, 0.2729]	1.0000
Macro	MAR	KNN	0.0228	[0.0024, 0.0446]	0.6667
		MICE-BR	0.0091	[-0.0070, 0.0257]	0.5833
		SVD	0.1536	[0.1372, 0.1689]	1.0000
Macro	MCAR	KNN	0.0220	[0.0018, 0.0462]	0.6667
		MICE-BR	0.1423	[0.0798, 0.2110]	0.9583
		SVD	0.1026	[0.0897, 0.1147]	1.0000

Table 5. Mean correlation-structure error $E_C = \|\hat{C} - C\|_F/d$ across all mechanisms and rates.

Method	Diabetes	Macro
Mean	0.0876	0.1690
Median	0.0930	0.1709
KNN	0.0558	0.0324
MICE-BR	0.0582	0.0271
SVD	0.0840	0.0380
NFMVI	0.0655	0.0282

Table 6. NFMVI uncertainty diagnostics averaged over 10–30% deletion. AUC detects cells in the top quartile of absolute reconstruction error.

Data	Mechanism	AUC	Spearman ρ
Diabetes	CHANNEL	0.6535	0.2653
	MAR	0.6680	0.3265
	MCAR	0.7221	0.4140
Macro	CHANNEL	0.7974	0.5099
	MAR	0.8294	0.5131
	MCAR	0.8201	0.4788

8.2 Parameter stability

Table 9 reports a separate sensitivity experiment at 25% CHANNEL missingness using the same five masks for all (k, γ) combinations. Diabetes is highly stable: across nine settings RMSE ranges only from 0.8007 to 0.8068. Macro is more sensitive to neighborhood size, with $k = 5$ outperforming $k \in \{9, 15\}$ on these particular masks, whereas changing γ between 0.4 and 1.0 has a modest effect within a fixed k .

The absolute values in Table 9 should not be compared directly with Table 8, because the two analyses use different deterministic mask sets. Within the sensitivity experiment, however, the controlled-mask comparison shows that γ is not a sharply tuned parameter in the tested range, while k can matter on the smaller Macro panel. The main experiment fixes $(k, \gamma) = (9, 0.7)$ before the cross-mechanism comparison rather than selecting a value separately for each dataset or missingness regime.

Table 7. Error stratified by pooled NFMVI uncertainty quartile across all main experiments.

Data	Quartile	Mean U	MAE	RMSE
Diabetes	Q1	0.1802	0.3718	0.5492
	Q2	0.2490	0.5018	0.6893
	Q3	0.3206	0.6902	0.8664
	Q4	0.3980	0.7734	0.9495
Macro	Q1	0.1019	0.0828	0.1717
	Q2	0.1314	0.0990	0.1730
	Q3	0.1545	0.1846	0.3342
	Q4	0.2437	0.4358	0.6202

Table 8. Ablation at 25% structured CHANNEL missingness. Entries are mean RMSE over eight masks.

Variant	Diabetes	Macro
Full NFMVI	0.7675	0.2977
No conditional view	0.8442	0.4837
No fuzzy reliability	0.7708	0.3250
No low-rank view	0.7858	0.2807
No neighborhood view	0.7828	0.2920
No neutrosophic gate	0.7697	0.3099

8.3 What the numerical evidence supports

Three patterns summarize the results. First, the proposed fusion is most useful when the three reconstruction assumptions provide complementary information. The large reduction relative to SVD across every paired setting indicates that low rank alone is insufficient, while the conditional-view ablation shows that cross-variable prediction is central. Second, NFMVI does not erase the strengths of specialized competitors: chained regression is superior for Diabetes MAR and Macro CHANNEL, and KNN preserves Diabetes correlations slightly better. Third, the uncertainty layer adds information not present in RMSE tables. On Macro, the highest-

Table 9. Sensitivity RMSE at 25% CHANNEL missingness; five common masks per dataset for all parameter combinations.

Data	k	γ		
		0.4	0.7	1.0
Diabetes	5	0.8018	0.8012	0.8007
	9	0.8046	0.8038	0.8032
	15	0.8068	0.8063	0.8054
Macro	5	0.4571	0.4536	0.4528
	9	0.4707	0.4708	0.4717
	15	0.4801	0.4813	0.4840

uncertainty quartile has more than three times the RMSE of the lowest quartile, indicating that U can identify cells that merit downstream review or sensitivity analysis.

These observations align with the purpose of a neutrosophic formulation. If all sources were interchangeable and equally accurate, a simple mean of e_s would suffice. The need for (T, I, F) arises precisely because source reliability and source disagreement vary by cell. Equation (16) then converts those two dimensions into adaptive weights, while Equation (17) keeps the residual conflict visible instead of discarding it after fusion.

9. LIMITATIONS

The study has five limitations. First, only two compact numerical datasets are used; the results do not establish superiority on high-dimensional images, mixed categorical data, or very long sensor streams. Second, MCAR, MAR, and CHANNEL masks are controlled artificial deletions. They allow exact reconstruction error to be measured, but they are not a complete model of MNAR mechanisms [1, 5]. Third, NFMVI is a single imputation method. It quantifies cell uncertainty through evidence conflict, but it does not generate Rubin-style multiple completed datasets for inferential variance propagation [4]. Fourth, the three source reliabilities are heuristic exponential mappings of empirical residual dispersion; learning calibrated reliability functions is a natural extension. Fifth, the four refinement passes are a fixed budget rather than iterations of a proven global objective. The mathematical guarantees concern the source state and each fusion update, not global convergence of the repeated reconstruction map.

Deep imputation models such as GAIN, BRITS, GP-VAE, graph neural imputers, and SAITS [20, 22, 23, 24, 25] are also outside the numerical comparison. This is deliberate for the present small-data, interpretable-method study, but a broader benchmark should compare NFMVI with these families on larger public datasets and include computational time and memory.

10. CONCLUSION

This paper proposed NFMVI, a fuzzy–neutrosophic algorithm for incomplete multivariate measurement data. For each missing cell, three candidate reconstructions e_N, e_S, e_R are assigned fuzzy reliabilities μ_s , compared through a Gaussian concordance a_s , and mapped to

$$(T_s, I_s, F_s) = (\mu_s a_s, 1 - a_s, 1 - \mu_s).$$

The resulting truth weights form a convex imputation $\hat{z} = \sum_s w_s e_s$, while $U = 1 - \sum T / \sum (T + I + F)$ records the unresolved evidence fraction. The identity $T + I + F = 1 + IF$ gives the non-normalized neutrosophic state an explicit interpretation, and $w_s/w_r = (\mu_s/\mu_r)(a_s/a_r)$ separates fuzzy source quality from current inter-view agreement.

The numerical study supports a qualified conclusion. NFMVI achieves the lowest overall RMSE on the Macro panel and is essentially tied with Bayesian-ridge chained imputation on Diabetes, but it does not dominate every missingness structure. Its more distinctive contribution is uncertainty diagnosis: reconstruction error rises monotonically across its uncertainty quartiles, particularly on Macro, where Q4 RMSE is 0.6202 compared with 0.1717 in Q1. Ablation shows that conditional information is critical and that low-rank/local sources can occasionally be counterproductive under concentrated channel dropout. These results favor a view of fuzzy–neutrosophic imputation as an interpretable evidence-fusion layer rather than a universal replacement for specialized imputers. Future work should calibrate source reliability, extend the state to mixed and interval-valued data, propagate uncertainty through multiple imputations, and evaluate larger real missingness processes.

REFERENCES

- [1] D. B. Rubin, “Inference and missing data,” *Biometrika*, vol. 63, no. 3, pp. 581–592, 1976.
- [2] R. J. A. Little, “A test of missing completely at random for multivariate data with missing values,” *Journal of the American Statistical Association*, vol. 83, no. 404, pp. 1198–1202, 1988.
- [3] A. P. Dempster, N. M. Laird, and D. B. Rubin, “Maximum likelihood from incomplete data via the em algorithm,” *Journal of the Royal Statistical Society: Series B (Methodological)*, vol. 39, no. 1, pp. 1–22, 1977.
- [4] S. van Buuren and K. Groothuis-Oudshoorn, “mice: Multivariate imputation by chained equations in r,” *Journal of Statistical Software*, vol. 45, no. 3, pp. 1–67, 2011.
- [5] R. J. A. Little and D. B. Rubin, *Statistical Analysis with Missing Data*, 3rd ed. Wiley, 2019.
- [6] O. Troyanskaya, M. Cantor, G. Sherlock, P. Brown, T. Hastie, R. Tibshirani, D. Botstein, and R. B. Altman, “Missing value estimation methods for DNA microarrays,” *Bioinformatics*, vol. 17, no. 6, pp. 520–525, 2001.
- [7] D. J. Stekhoven and P. Bühlmann, “Missforest—non-parametric missing value imputation for mixed-type data,” *Bioinformatics*, vol. 28, no. 1, pp. 112–118, 2012.
- [8] E. J. Candès and B. Recht, “Exact matrix completion via convex optimization,” *Foundations of Computational Mathematics*, vol. 9, no. 6, pp. 717–772, 2009.
- [9] R. Mazumder, T. Hastie, and R. Tibshirani, “Spectral regularization algorithms for learning large incomplete matrices,” *Journal of Machine Learning Research*, vol. 11, no. 80, pp. 2287–2322, 2010.

- [10] L. A. Zadeh, "Fuzzy sets," *Information and Control*, vol. 8, no. 3, pp. 338–353, 1965.
- [11] H. Wang, F. Smarandache, Y. Zhang, and R. Sunderraman, "Single valued neutrosophic sets," *Multispace and Multistructure*, vol. 4, pp. 410–413, 2010.
- [12] K. Qin and L. Wang, "New similarity and entropy measures of single-valued neutrosophic sets with applications in multi-attribute decision making," *Soft Computing*, vol. 24, pp. 16 165–16 176, 2020.
- [13] R. J. Hathaway and J. C. Bezdek, "Fuzzy c-means clustering of incomplete data," *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)*, vol. 31, no. 5, pp. 735–744, 2001.
- [14] I. B. Aydilek and A. Arslan, "A hybrid method for imputation of missing values using optimized fuzzy c-means with support vector regression and a genetic algorithm," *Information Sciences*, vol. 233, pp. 25–35, 2013.
- [15] A. M. Sefidian and N. Daneshpour, "Missing value imputation using a novel grey based fuzzy c-means, mutual information based feature selection, and regression model," *Expert Systems with Applications*, vol. 115, pp. 68–94, 2019.
- [16] J. Huang, B. Mao, Y. Bai, T. Zhang, and C. Miao, "An integrated fuzzy c-means method for missing data imputation using taxi gps data," *Sensors*, vol. 20, no. 7, p. 1992, 2020.
- [17] E. J. Candès and Y. Plan, "Matrix completion with noise," *Proceedings of the IEEE*, vol. 98, no. 6, pp. 925–936, 2010.
- [18] P. J. García-Laencina, J.-L. Sancho-Gómez, and A. R. Figueiras-Vidal, "Pattern classification with missing data: a review," *Neural Computing and Applications*, vol. 19, no. 2, pp. 263–282, 2010.
- [19] J. L. Schafer and J. W. Graham, "Missing data: Our view of the state of the art," *Psychological Methods*, vol. 7, no. 2, pp. 147–177, 2002.
- [20] J. Yoon, J. Jordon, and M. van der Schaar, "GAIN: Missing data imputation using generative adversarial nets," in *Proceedings of the 35th International Conference on Machine Learning*, ser. Proceedings of Machine Learning Research, vol. 80. PMLR, pp. 5689–5698, 2018.
- [21] Z. Che, S. Purushotham, K. Cho, D. Sontag, and Y. Liu, "Recurrent neural networks for multivariate time series with missing values," *Scientific Reports*, vol. 8, p. 6085, 2018.
- [22] W. Cao, D. Wang, J. Li, H. Zhou, L. Li, and Y. Li, "BRITS: Bidirectional recurrent imputation for time series," in *Advances in Neural Information Processing Systems*, vol. 31, 2018.
- [23] V. Fortuin, D. Baranchuk, G. Rätsch, and S. Mandt, "GP-VAE: Deep probabilistic time series imputation," in *Proceedings of the Twenty Third International Conference on Artificial Intelligence and Statistics*, ser. Proceedings of Machine Learning Research, vol. 108. PMLR, pp. 1651–1661, 2020.
- [24] A. Cini, I. Marisca, and C. Alippi, "Filling the Gaps: Multivariate time series imputation by graph neural networks," in *International Conference on Learning Representations*, 2022.
- [25] W. Du, D. Côté, and Y. Liu, "SAITS: Self-attention-based imputation for time series," *Expert Systems with Applications*, vol. 219, p. 119619, 2023.
- [26] K. T. Atanassov, "Intuitionistic fuzzy sets," *Fuzzy Sets and Systems*, vol. 20, no. 1, pp. 87–96, 1986.
- [27] J. S. Chai, G. Selvachandran, F. Smarandache, V. C. Gerogiannis, L. H. Son, Q.-T. Bui, and B. Vo, "New similarity measures for single-valued neutrosophic sets with applications in pattern recognition and medical diagnosis problems," *Complex & Intelligent Systems*, vol. 7, pp. 703–723, 2021.
- [28] S. Seabold and J. Perktold, "Statsmodels: Econometric and statistical modeling with python," in *Proceedings of the 9th Python in Science Conference*, pp. 92–96, 2010.
- [29] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. VanderPlas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot, and É. Duchesnay, "Scikit-learn: Machine learning in python," *Journal of Machine Learning Research*, vol. 12, no. 85, pp. 2825–2830, 2011.
- [30] B. Efron, T. Hastie, I. Johnstone, and R. Tibshirani, "Least angle regression," *The Annals of Statistics*, vol. 32, no. 2, pp. 407–499, 2004.