

Dual-Indeterminacy Neutrosophic Fusion for Calibrated and Selective Multi-Source Classification under Source Conflict

Erina Kovachiskaya^{1,*}

¹ Faculty of Information Technology and Robotics, Vitebsk State Technological University, Belarus

Email: EriKovachi98rus@vsu.by

Received: July 25, 2025 Accepted: September 29, 2025 ★ Corresponding author

ABSTRACT

Reliable multi-source AI requires a distinction between two non-equivalent failure modes: within-source ambiguity and between-source contradiction. This paper formulates dual-indeterminacy neutrosophic fusion (DINF), in which calibrated posteriors $\mathbf{p}_s \in \Delta^{K-1}$ are transformed into class-wise triples $N_k = \langle T_k, I_k, F_k \rangle$. Source weights combine chance-corrected validation reliability, normalized entropy, and a robust Jensen–Shannon consensus discount. Intrinsic ambiguity U and class-specific contradiction C_k are retained separately and joined by $I_k = U + C_k - UC_k$. The decision distribution $q_k \propto T_k \exp(-\lambda I_k)$ therefore penalizes conflict without collapsing neutrosophic evidence into a probability simplex. Boundedness, permutation invariance, consensus preservation, conflict monotonicity, log-odds sensitivity, and missing-source neutrality are established. Repeated experiments on three public benchmarks examine clean data, 30% source dropout, 30% confident contradiction, and 40% mixed failure. Under contradiction, DINF achieves macro-F1 0.910, NLL 0.445, and selective accuracy 0.939 at approximately 90% coverage; the corresponding logarithmic-pool values are 0.896, 0.469, and 0.929. The results identify robust agreement discounting and explicit contradiction-sensitive indeterminacy as complementary mechanisms for calibrated failure-aware fusion.

Keywords: Neutrosophic information fusion ▪ Multi-source classification ▪ Source conflict ▪ Uncertainty decomposition ▪ Calibrated AI ▪ Selective prediction ▪ Robust decision fusion ▪ Sensor failure

1. INTRODUCTION

Let S calibrated sources describe the same item through posteriors $\{\mathbf{p}_s\}_{s=1}^S$. A source can be diffuse because its own observation is weak, or sharply confident while contradicting the remaining sources. Mean fusion, fixed reliability weights, entropy weights, and fixed logarithmic pools do not represent these mechanisms separately; a low-entropy but erroneous source may therefore dominate the fused decision.

Information-fusion studies identify heterogeneity, conflict, and changing reliability as persistent challenges [1, 2, 3]. Probability calibration and proper scores are likewise necessary when predictive confidence changes under shift [4, 5, 6]. Single-valued neutrosophic sets provide independent truth,

indeterminacy, and falsity coordinates [7], but existing operators commonly encode indeterminacy by one scalar and are seldom tested under controlled source faults.

This work introduces DINF with

$$\mathbf{p}_s \mapsto (r_s, H_s, D_s, w_s) \mapsto \{T_k, U, C_k, I_k, F_k, q_k\}_{k=1}^K.$$

Its contributions are: (i) $I_k = U \oplus C_k$ separates ambiguity from contradiction; (ii) $w_s \propto r_s[\alpha + (1 - \alpha)(1 - H_s)]e^{-\kappa D_s}$ detects confident outliers without a separate fault classifier; (iii) formal properties and log-odds sensitivity are derived; and (iv) public-data experiments quantify calibration, robustness, and selective prediction under four operating conditions.

The remainder reviews related methods, defines the operator,

formalizes the experiment, presents results, and concludes the study.

2. RELATED WORK AND RESEARCH GAP

2.1 Probabilistic and Neutrosophic Fusion

Late fusion preserves source modularity across sensing and healthcare pipelines [8], but static averaging cannot identify instance-specific failure. Entropy weighting suppresses diffuse posteriors yet may reward a confidently wrong source. Logarithmic pooling exploits agreement but does not expose the cause of disagreement. Calibration methods improve probability quality [9, 10], and selective classification adds an abstention mechanism [11]; neither alone yields a class-wise neutrosophic explanation.

Neutrosophic aggregation includes EDAS, similarity, Einstein-interaction, dynamic, and exponential-logarithmic operators [12, 13, 14, 15, 16]. Neutrosophic image-text and brain-image models further demonstrate multimodal relevance [17, 18]. Their settings differ from calibrated posterior fusion under explicit source corruption.

2.2 Research Gap

The unresolved problem is to construct a calibrated operator in which intrinsic uncertainty and source conflict are distinct measurable quantities, while source order, missing modalities, and abstention remain mathematically controlled. DINF addresses this gap at decision level and is evaluated against probabilistic baselines under clean and corrupted inputs.

3. PROBLEM DEFINITION AND NEUTROSOPHIC SEMANTICS

Let $\mathcal{Y} = \{1, \dots, K\}$ be the class set and let $S \geq 2$ information sources observe an item x . Source s returns a calibrated posterior

$$\mathbf{p}_s(x) = (p_{s1}, \dots, p_{sK}) \in \Delta^{K-1}, \quad \sum_{k=1}^K p_{sk} = 1. \quad (1)$$

A source may correspond to a sensor, modality, feature group, institution, expert model, or independently trained classifier. The fusion task is to obtain a class distribution $\mathbf{q}(x)$ and an abstention utility while retaining a neutrosophic explanation.

Definition 1 (Class-wise neutrosophic evidence). *For class k , the fused evidence is*

$$N_k(x) = \langle T_k(x), I_k(x), F_k(x) \rangle, \quad T_k, I_k, F_k \in [0, 1], \quad (2)$$

where T_k measures fused support, F_k measures fused opposition, and I_k measures unresolved evidence. No normalization constraint $T_k + I_k + F_k = 1$ is imposed.

The proposed semantics further decomposes I_k into U , a sample-level intrinsic ambiguity estimated from source entropies, and C_k , a class-specific contradiction estimated from pairwise posterior deviations. The distinction is operational: high U means sources are individually hesitant; high C_k means they disagree about class k , even if each source is confident.

4. DUAL-INDETERMINACY NEUTROSOPHIC FUSION

4.1 Calibrated Reliability and Intrinsic Uncertainty

Each source classifier is calibrated by three-fold sigmoid calibration within the training partition. Its global reliability is estimated on a held-out validation set through balanced accuracy b_s , adjusted for chance:

$$r_s = \text{clip}\left(\frac{b_s - 1/K}{1 - 1/K}, 0.05, 1\right). \quad (3)$$

The normalized predictive entropy is

$$H_s = -\frac{1}{\log K} \sum_{k=1}^K p_{sk} \log(p_{sk} + \epsilon), \quad H_s \in [0, 1]. \quad (4)$$

Low entropy is useful only when the source agrees with independent evidence. For that reason, entropy does not directly determine the final weight.

4.2 Robust Consensus Discount

Let the coordinate-wise median posterior be normalized as

$$\tilde{m}_k = \text{median}_s\{p_{sk}\}, \quad m_k = \frac{\tilde{m}_k}{\sum_j \tilde{m}_j}. \quad (5)$$

The distance of source s from this robust consensus is the normalized Jensen–Shannon divergence [19]:

$$D_s = \text{JSD}(\mathbf{p}_s, \mathbf{m}), \quad (6)$$

$$\text{JSD}(\mathbf{a}, \mathbf{b}) = \frac{\text{KL}(\mathbf{a} \parallel \mathbf{z}) + \text{KL}(\mathbf{b} \parallel \mathbf{z})}{2 \log 2}, \quad \mathbf{z} = \frac{1}{2}(\mathbf{a} + \mathbf{b}). \quad (7)$$

The unnormalized weight is

$$\omega_s = r_s [\alpha + (1 - \alpha)(1 - H_s)] \exp(-\kappa D_s), \quad (8)$$

with $\alpha = 0.35$ in the experiments. Normalization gives

$$w_s = \frac{\omega_s}{\sum_{t=1}^S \omega_t}. \quad (9)$$

The positive floor α prevents a moderately uncertain but reliable source from disappearing, whereas $\exp(-\kappa D_s)$ suppresses a confident outlier.

4.3 Truth, Falsity, and Dual Indeterminacy

Truth is constructed as a consensus-weighted logarithmic pool:

$$\hat{T}_k = \exp\left(\sum_{s=1}^S w_s \log(p_{sk} + \epsilon)\right), \quad T_k = \frac{\hat{T}_k}{\sum_j \hat{T}_j}. \quad (10)$$

Falsity is the geometric aggregation of class opposition:

$$F_k = \exp\left(\sum_{s=1}^S w_s \log(1 - p_{sk} + \epsilon)\right). \quad (11)$$

Intrinsic ambiguity is the weighted entropy

$$U = \sum_{s=1}^S w_s H_s. \quad (12)$$

Table 1. Position of the proposed framework relative to representative literature.

Study	Main setting	Uncertainty/fusion mechanism	Gap addressed by the present work
Meng et al. [1]	ML data fusion survey	Signal-, feature-, and decision-level taxonomy	No neutrosophic ambiguity–conflict decomposition
Abdar et al. [4]	Uncertainty-aware learning	Bayesian and ensemble uncertainty	No multi-source neutrosophic operator
Farid and Riaz [14, 15]	Neutrosophic MCDM	Interactive and temporal aggregation	Alternative ranking rather than calibrated classification
Wajid et al. [17]	Image–text classification	Neutrosophic CNN fusion	No explicit fault decomposition or abstention
Lavanya et al. [18]	Medical image fusion	Neutrosophic entropy and feature extraction	Image enhancement rather than posterior reliability
Li et al. [2]	Multi-source fusion review	Architectures and open challenges	No dual-indeterminacy decision operator
Proposed DINF	Calibrated classification	Reliability, entropy, robust consensus, and dual indeterminacy	Separates U from C_k and supports selective decisions

Table 2. Core notation.

Symbol	Meaning
S, K	numbers of sources and classes
$\mathbf{p}_s$	calibrated posterior of source s
r_s	validation reliability of source s
H_s	normalized entropy of source s
D_s	divergence from robust consensus
w_s	sample-specific normalized source weight
T_k, F_k	truth and falsity for class k
U, C_k	intrinsic ambiguity and source conflict
I_k	dual indeterminacy for class k
κ, λ	consensus discount and indeterminacy penalty

Class-specific contradiction is

$$C_k = \frac{\sum_{s < t} w_s w_t |p_{sk} - p_{tk}|}{\sum_{s < t} w_s w_t + \varepsilon}. \quad (13)$$

The final indeterminacy is the probabilistic sum

$$I_k = U \oplus C_k = U + C_k - UC_k. \quad (14)$$

This operation increases when either source ambiguity or contradiction increases, avoids double counting their overlap, and remains bounded.

A direct neutrosophic score is

$$\mathcal{S}_k = T_k - F_k - \lambda I_k. \quad (15)$$

For proper probabilistic evaluation, the normalized decision distribution is

$$q_k = \frac{T_k \exp(-\lambda I_k)}{\sum_{j=1}^K T_j \exp(-\lambda I_j)}. \quad (16)$$

The predicted class is $\hat{y} = \arg \max_k q_k$. The selective utility is

$$A(x) = \max_k q_k [1 - I_{\hat{y}}], \quad (17)$$

so that a case can be rejected when its probability is high but its supporting sources remain contradictory.

For classes $k \neq \ell$, Eq. (16) gives the exact log-odds decomposition

$$\log \frac{q_k}{q_\ell} = \log \frac{T_k}{T_\ell} - \lambda (I_k - I_\ell). \quad (18)$$

Algorithm 1 Dual-indeterminacy neutrosophic fusion

Require: Posteriors $\{\mathbf{p}_s\}_{s=1}^S$; reliabilities $\{r_s\}$; κ, λ

- 1: Compute H_s by Eq. (4)
- 2: Form the normalized median posterior $\mathbf{m}$
- 3: Compute $D_s = \text{JSD}(\mathbf{p}_s, \mathbf{m})$
- 4: Compute w_s using Eqs. (8)–(9)
- 5: Obtain T_k, F_k using Eqs. (10)–(11)
- 6: Obtain U, C_k, I_k using Eqs. (12)–(14)
- 7: Compute q_k and $A(x)$ using Eqs. (16)–(17)
- 8: **return** $\{(T_k, I_k, F_k), q_k\}_{k=1}^K$ and $A(x)$

Because $\partial I_k / \partial C_k = 1 - U$,

$$\frac{\partial}{\partial C_k} \log \frac{q_k}{q_\ell} = -\lambda (1 - U), \quad (19)$$

when T_k, T_ℓ, I_ℓ are fixed. Hence contradiction is penalized most strongly when the sources are individually sharp.

4.4 Formal Properties

Proposition 1 (Boundedness). *For valid source posteriors and positive weights, $T_k, F_k, U, C_k, I_k \in [0, 1]$ and $\mathbf{q} \in \Delta^{K-1}$.*

Proof. Normalized entropy and Jensen–Shannon divergence are bounded by one. Equations (8)–(9) produce nonnegative weights summing to one. Weighted geometric means of values in $[0, 1]$ remain in $[0, 1]$, and normalization preserves $T_k \in [0, 1]$. The weighted average U and normalized pairwise average C_k are bounded. Since $I_k = 1 - (1 - U)(1 - C_k)$, it is also bounded. Equation (16) is positive and normalized. $\square$

Proposition 2 (Permutation invariance). *Reordering the sources does not change $N_k, \mathbf{q}$, or $A(x)$.*

Proof. Median, entropy collection, pairwise sums, and weighted products are symmetric in source indices. Reliability values move with their sources, so every aggregate is unchanged. $\square$

Proposition 3 (Consensus preservation). *If $\mathbf{p}_s = \mathbf{p}$ for all s , then $D_s = 0$, $C_k = 0$, $T_k = p_k$, and $I_k = \sum_s w_s H_s = H(\mathbf{p})$.*

Proof. Identical posteriors equal their median and have zero pairwise differences. The weighted geometric mean of repeated p_k is p_k , and its normalization is unchanged. Equation (14) reduces to U . $\square$

This result is preferable to forcing zero indeterminacy under consensus: sources can agree because all are genuinely uncertain.

Proposition 4 (Conflict monotonicity). *For fixed $U \in [0, 1]$, I_k is nondecreasing in C_k , with derivative $\partial I_k / \partial C_k = 1 - U \geq 0$.*

The effect of additional contradiction is therefore strongest when sources are individually confident and naturally attenuated when they are already ambiguous.

Proposition 5 (Missing-source neutrality). *Let source u output $\mathbf{p}_u = K^{-1}\mathbf{1}$ and let its assigned reliability satisfy $r_u \rightarrow 0$. Then $w_u \rightarrow 0$, and (T_k, I_k, F_k, q_k) converges to the fusion result obtained from the remaining sources.*

Proof. Equation (8) gives $\omega_u \rightarrow 0$; weight normalization and continuity of Eqs. (10)–(16) complete the result. $\square$

4.5 Weight-Ratio and Robustness Bounds

Set

$$g(H) = \alpha + (1 - \alpha)(1 - H), \quad \alpha \leq g(H) \leq 1. \quad (20)$$

For any sources a and b , normalization cancels in the ratio:

$$\frac{w_a}{w_b} = \frac{r_a g(H_a)}{r_b g(H_b)} \exp[-\kappa(D_a - D_b)]. \quad (21)$$

Thus the local sensitivities of the raw weight are

$$\frac{\partial \log \omega_s}{\partial D_s} = -\kappa, \quad \frac{\partial \log \omega_s}{\partial H_s} = -\frac{1 - \alpha}{g(H_s)}. \quad (22)$$

The first derivative is constant, while the entropy effect remains bounded by

$$-(1 - \alpha)/\alpha \leq \partial_{H_s} \log \omega_s \leq -(1 - \alpha). \quad (23)$$

Proposition 6 (Exponential outlier suppression). *Let $\mathcal{G}$ be a coherent source set satisfying $D_g \leq d_{\mathcal{G}}$ for $g \in \mathcal{G}$, and let source o satisfy $D_o \geq d_o > d_{\mathcal{G}}$. Define*

$$R_{\mathcal{G}} = \sum_{g \in \mathcal{G}} r_g g(H_g), \quad \delta = d_o - d_{\mathcal{G}} > 0. \quad (24)$$

Then

$$w_o \leq \frac{r_o g(H_o) e^{-\kappa \delta}}{R_{\mathcal{G}} + r_o g(H_o) e^{-\kappa \delta}}. \quad (25)$$

Proof. For $g \in \mathcal{G}$, $\omega_g \geq r_g g(H_g) e^{-\kappa d_{\mathcal{G}}}$, whereas $\omega_o \leq r_o g(H_o) e^{-\kappa d_o}$. Substitute these inequalities into $w_o = \omega_o / (\omega_o + \sum_g \omega_g)$ and cancel $e^{-\kappa d_{\mathcal{G}}}$. $\square$

Consequently, $w_o = O(e^{-\kappa \delta})$ as $\kappa \delta \rightarrow \infty$. Moreover, the truth odds satisfy

$$\log \frac{T_k}{T_\ell} = \sum_{s=1}^S w_s \log \frac{p_{sk} + \varepsilon}{p_{s\ell} + \varepsilon}, \quad (26)$$

so the direct contribution of an outlying source to any class-pair log odds is multiplied by the exponentially bounded factor in Eq. (25).

The normalizer of Eq. (13) has the closed form

$$Z_w = \sum_{s \leq t} w_s w_t = \frac{1}{2} \left(1 - \sum_{s=1}^S w_s^2 \right). \quad (27)$$

Hence C_k is a normalized weighted Gini mean difference. For $S = 2$ and $w_1 w_2 > 0$,

$$C_k = |p_{1k} - p_{2k}|, \quad (28)$$

which is independent of the weight ratio; weights regulate truth/falsity, while the two-source contradiction magnitude remains explicit.

5. EXPERIMENTAL DESIGN

5.1 Data Tensor, Partitions, and Source Models

Three public datasets are partitioned into semantically or spatially distinct sources (Table 3) [20, 21, 22, 23]. For repetition $r \in \{1, \dots, 5\}$,

$$\mathcal{D}^{(r)} = \mathcal{D}_{\text{tr}}^{(r)} \cup \mathcal{D}_{\text{va}}^{(r)} \cup \mathcal{D}_{\text{te}}^{(r)}, \quad (29)$$

$$(\pi_{\text{tr}}, \pi_{\text{va}}, \pi_{\text{te}}) = (0.525, 0.175, 0.300). \quad (30)$$

with stratification by y . The resulting test posterior tensor is

$$\mathcal{P} = [p_{isk}] \in [0, 1]^{n \times S \times K}, \quad \sum_{k=1}^K p_{isk} = 1. \quad (31)$$

For source s , standardized features produce multinomial logits $z_{isk} = \beta_{sk}^\top \mathbf{x}_i^{(s)} + \beta_{0,sk}$ with inverse regularization $C = 2$. Three-fold sigmoid calibration maps the raw posterior $\tilde{p}_{isk}$ to

$$v_{isk} = \sigma(a_{sk} \log(\tilde{p}_{isk}) + b_{sk}), \quad p_{isk} = \frac{v_{isk}}{\sum_{j=1}^K v_{isj}}. \quad (32)$$

Validation balanced accuracy determines r_s through Eq. (3); no test label enters calibration, reliability estimation, or tuning.

5.2 Controlled Source-Failure Operators

The historically strongest source is

$$s^* = \arg \max_{1 \leq s \leq S} r_s, \quad \mathbf{u}_K = K^{-1} \mathbf{1}_K. \quad (33)$$

Let $\mathcal{J}_p$ be a uniformly sampled test-index set of size $\lfloor \rho n \rfloor$. Dropout is the operator

$$\mathcal{M}_{\text{drop}}(\mathbf{p}_{is^*}) = \mathbf{u}_K, \quad i \in \mathcal{J}_{0.30}. \quad (34)$$

For contradiction, set $a_i = \arg \max_k p_{is^*k}$, draw $L_i \sim \text{Unif}\{1, \dots, K - 1\}$, and define

$$c_i = 1 + [(a_i - 1 + L_i) \bmod K], \quad (35)$$

$$\mathcal{M}_{\text{ctr}}(\mathbf{p}_{is^*}) = 0.1 \mathbf{p}_{is^*} + 0.9 \mathbf{e}_{c_i}. \quad (36)$$

The mixed condition applies Eq. (34) to a random 0.20n cases and Eq. (36) to a disjoint 0.20n cases. Thus the four conditions are $\mathcal{M}_0$, $\mathcal{M}_{\text{drop}}^{0.30}$, $\mathcal{M}_{\text{ctr}}^{0.30}$, and $\mathcal{M}_{\text{mix}}^{0.40}$.

Table 3. Public datasets and source construction.

Dataset	n	d	K	S	Source groups
Breast cancer diagnos- tic	569	30	2	3	Mean, standard-error, and worst morphology (10 variables each)
Wine recognition	178	13	3	3	Acidity/minerals (5); phenolic profile (4); color/proline (4)
Optical digits subset	1,797	64	10	4	Four 4×4 quadrants of each 8×8 image

5.3 Baselines and Hyperparameter Functional

For $\bar{r}_s = r_s / \sum_t r_t$, the comparison rules are

$$q_k^{\text{mean}} = S^{-1} \sum_s p_{sk}, \quad (37)$$

$$q_k^{\text{rel}} = \sum_s \bar{r}_s p_{sk}, \quad (38)$$

$$q_k^{\text{ent}} = \sum_s \eta_s p_{sk}, \quad \eta_s = \frac{r_s(1 - H_s + 10^{-3})}{\sum_t r_t(1 - H_t + 10^{-3})}, \quad (39)$$

$$q_k^{\text{log}} = \frac{\exp(\sum_s \bar{r}_s \log(p_{sk} + \epsilon))}{\sum_j \exp(\sum_s \bar{r}_s \log(p_{sj} + \epsilon))}. \quad (40)$$

An early-fusion calibrated logistic model is retained only as clean-data context. The grid is

$$\Theta = \{0, 2, 4, 8, 16, 32\} \times \{0, 0.25, 0.5, 1, 1.5\}. \quad (41)$$

For $\theta = (\kappa, \lambda)$, selection maximizes

$$\mathcal{J}(\theta) = 0.45F_{\text{I,va}}^{\text{clean}}(\theta) + 0.55F_{\text{I,va}}^{\text{ctr}}(\theta) \quad (42)$$

$$- 0.05(E_{\text{va}}^{\text{clean}}(\theta) + E_{\text{va}}^{\text{ctr}}(\theta)). \quad (43)$$

where E is 12-bin ECE and contradiction is generated only inside the validation partition. The modal optima are (4, 1.5), (8, 0), and (2, 1.5) for breast cancer, wine, and digits, respectively.

5.4 Probabilistic and Selective Criteria

With one-hot labels y_{ik} and predicted class $\hat{y}_i$, the reported criteria are

$$\text{NLL} = -\frac{1}{n} \sum_{i=1}^n \log(q_{i,y_i} + \epsilon), \quad (44)$$

$$\text{BS} = \frac{1}{n} \sum_{i=1}^n \sum_{k=1}^K (q_{ik} - y_{ik})^2, \quad (45)$$

$$F_1^{\text{macro}} = \frac{1}{K} \sum_{k=1}^K \frac{2 \text{Prec}_k \text{Rec}_k}{\text{Prec}_k + \text{Rec}_k + \epsilon}, \quad (46)$$

$$\text{ECE}_{12} = \sum_{m=1}^{12} \frac{|B_m|}{n} |\text{acc}(B_m) - \text{conf}(B_m)|. \quad (47)$$

NLL and Brier score are proper probabilistic criteria [6]. At target coverage $\gamma = 0.90$, let $\mathcal{A}_\gamma$ contain the $\lceil \gamma n \rceil$ largest utilities from Eq. (17); baseline utilities are $\max_k q_{ik}$. Then

$$\text{Acc}_{\text{sel}}(\gamma) = \frac{1}{|\mathcal{A}_\gamma|} \sum_{i \in \mathcal{A}_\gamma} \mathbf{1}(\hat{y}_i = y_i). \quad (48)$$

Every pooled entry is the equal-weight estimator

$$\bar{M} = \frac{1}{15} \sum_{d=1}^3 \sum_{r=1}^5 M_{d,r}. \quad (49)$$

For a threshold τ , the empirical coverage and selective risk are

$$\widehat{\text{cov}}(\tau) = \frac{1}{n} \sum_i \mathbf{1}[A(x_i) \geq \tau], \quad (50)$$

$$\widehat{R}(\tau) = \frac{\sum_i \mathbf{1}[\hat{y}_i \neq y_i] \mathbf{1}[A(x_i) \geq \tau]}{\sum_i \mathbf{1}[A(x_i) \geq \tau] + \epsilon}. \quad (51)$$

The run-level dispersion retained in the supplied tables is

$$\widehat{\text{SE}}(M) = \sqrt{\frac{1}{15 \cdot 14} \sum_{d=1}^3 \sum_{r=1}^5 (M_{d,r} - \bar{M})^2}. \quad (52)$$

6. RESULTS

6.1 Clean and Failure-Robust Classification

The clean-data macro-F1 difference is $0.962 - 0.969 = -0.007$. Under contradiction, however,

$$\Delta F_1(\text{DINF}, \log) = 0.910 - 0.896 = 0.014,$$

and the gains over mean and reliability weighting are 0.046 and 0.095. Under mixed failure, $F_1 = 0.916$ versus 0.911 for the logarithmic pool. Under dropout, the logarithmic pool has higher F_1 , whereas DINF has lower NLL and Brier score. Entropy weighting attains low ECE by producing less decisive probabilities, but its contradiction macro-F1 and Brier score are 0.748 and 0.360. ECE is therefore interpreted jointly with NLL, Brier score, and F_1 .

6.2 Selective Prediction

At $\gamma \approx 0.90$, contradiction selective accuracy is 0.939 for DINF, 0.929 for the logarithmic pool, and 0.898 for mean fusion. Under mixed failure, DINF and the logarithmic pool both reach 0.937.

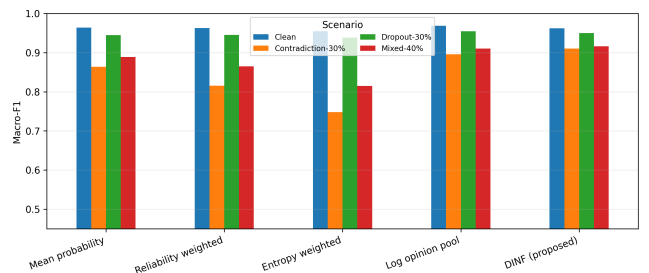


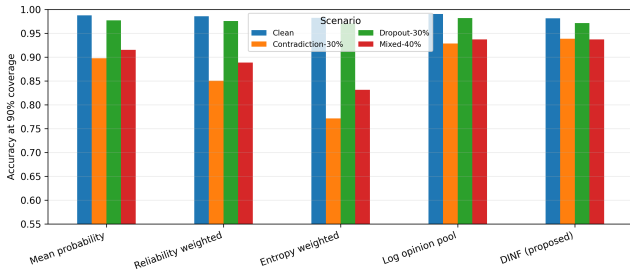
Figure 1. Pooled macro-F1 under clean and source-failure scenarios. Confident contradiction is the scenario in which the dual-indeterminacy mechanism provides its clearest advantage.

Table 4. Pooled predictive performance across three datasets and five repeated splits. Higher is better for accuracy and macro-F1; lower is better for NLL, ECE, and Brier. Best values among the listed methods are bold within each scenario.

Scenario	Method	Accuracy	Macro-F1	NLL	ECE	Brier
Clean	Mean probability	0.965	0.964	0.538	0.328	0.250
	Reliability weighted	0.964	0.963	0.527	0.318	0.243
	Entropy weighted	0.955	0.954	0.468	0.270	0.210
	Log opinion pool	0.970	0.969	0.373	0.240	0.163
	DINF	0.963	0.962	0.368	0.228	0.165
Dropout-30%	Mean probability	0.947	0.945	0.603	0.346	0.287
	Reliability weighted	0.947	0.945	0.603	0.344	0.286
	Entropy weighted	0.940	0.939	0.496	0.269	0.229
	Log opinion pool	0.956	0.955	0.447	0.270	0.206
	DINF	0.951	0.950	0.437	0.256	0.205
Contradiction-30%	Mean probability	0.868	0.864	0.645	0.281	0.327
	Reliability weighted	0.820	0.815	0.653	0.241	0.336
	Entropy weighted	0.751	0.748	0.665	0.158	0.360
	Log opinion pool	0.899	0.896	0.469	0.212	0.236
	DINF	0.912	0.910	0.445	0.208	0.222
Mixed-40%	Mean probability	0.892	0.889	0.654	0.309	0.326
	Reliability weighted	0.868	0.865	0.663	0.282	0.334
	Entropy weighted	0.818	0.815	0.622	0.173	0.325
	Log opinion pool	0.913	0.911	0.486	0.235	0.239
	DINF	0.918	0.916	0.464	0.223	0.229

Table 5. Selective accuracy at about 90% coverage.

Method	Clean	Dropout	Contrad.	Mixed
Mean	0.988	0.977	0.898	0.915
Reliability	0.986	0.976	0.850	0.888
Entropy	0.982	0.970	0.771	0.831
Log pool	0.991	0.982	0.929	0.937
DINF	0.981	0.971	0.939	0.937

**Figure 2.** Selective accuracy at approximately 90% coverage. The proposed utility is most beneficial when a trusted source becomes confidently contradictory.

6.3 Ablation Analysis

The ablation differences are

$$0.910 - 0.858 = 0.052, \quad 0.910 - 0.906 = 0.004.$$

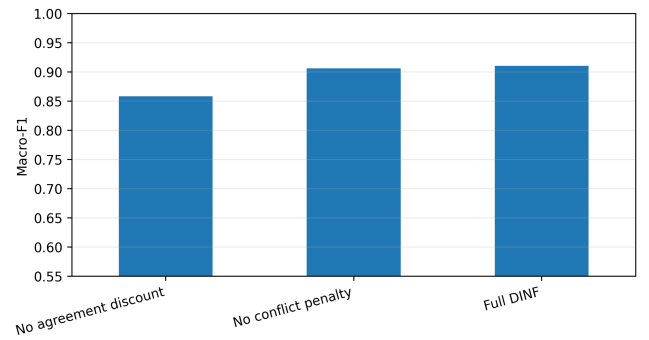
Thus consensus discounting is the dominant defense, while the class-specific I_k penalty adds a further decision-level gain.

6.4 Source Reliability and Neutrosophic Interpretation

Mean source reliabilities are (0.829, 0.722, 0.927) for the three breast-cancer groups; the wine color/proline source reaches 0.929; digit-source reliabilities lie in [0.557, 0.675]. Since corruption targets s^* , the evaluation attacks the strongest source.

Table 6. Ablation results under 30% contradiction.

Variant	Accuracy	Macro-F1	ECE
No agreement discount	0.861	0.858	0.183
No conflict penalty	0.908	0.906	0.201
Full DINF	0.912	0.910	0.208

**Figure 3.** Macro-F1 ablation. Consensus discount is the main defense against a confident corrupted source; class-specific indeterminacy supplies an additional gain.

For a high-conflict digit case, the predicted class has $(T_k, C_k, I_k) \approx (0.643, 0.438, 0.744)$, whereas a low-conflict case has $C_k \approx 0.013$. The same top label can therefore receive a substantially different acceptance utility.

7. DISCUSSION AND LIMITATIONS

The mechanism is summarized by Eqs. (8), (14), and (19): a confident outlier receives the multiplicative discount $e^{-\kappa D_s}$, residual class disagreement increases I_k , and its log-odds penalty equals $\lambda(1 - U)$. This explains why entropy weighting fails under confident contradiction and why the clean-data optimum need not equal the fault-robust optimum.

For n items, S sources, and K classes, entropy and geo-

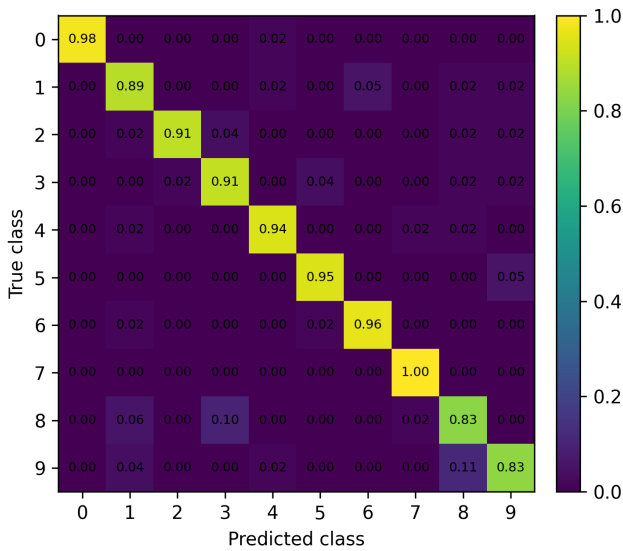


Figure 4. Confusion matrix for the digits benchmark under 30% contradiction using DINF (first split).

metric aggregation cost $O(nSK)$, median consensus costs $O(nSK \log S)$, and pairwise conflict costs $O(nS^2K)$. The latter can be replaced by consensus distances or sparse source graphs when S is large.

The experiments use controlled feature groups rather than physically separate asynchronous sensors. Median consensus assumes that a coherent source subset remains, and reliability is global and supervised. Correlated faults, majority bias, temporal drift, unlabeled reliability updates, and deployment-specific abstention costs remain open. These limitations define the next tests for multimodal, federated, and streaming neutrosophic fusion.

8. CONCLUSION

DINF separates intrinsic ambiguity U from class-specific contradiction C_k and maps their probabilistic sum into neutrosophic indeterminacy. Reliability–entropy–consensus weighting, geometric truth/falsity aggregation, and the decision rule $q_k \propto T_k e^{-\lambda I_k}$ yield a bounded and permutation-invariant operator with explicit conflict sensitivity. Across three public benchmarks, the largest advantage occurs when the most reliable source becomes confidently wrong: macro-F1 rises from 0.896 for logarithmic pooling to 0.910, NLL decreases from 0.469 to 0.445, and selective accuracy increases from 0.929 to 0.939. The formal decomposition and reproducible fault operators provide a compact basis for calibrated, interpretable, and failure-aware neutrosophic AI.

REFERENCES

- [1] T. Meng, X. Jing, Z. Yan, and W. Pedrycz, “A survey on machine learning for data fusion,” *Information Fusion*, vol. 57, pp. 115–129, 2020.
- [2] X. Li, F. Dunkin, and J. Dezert, “Multi-source information fusion: Progress and future,” *Chinese Journal of Aeronautics*, vol. 37, no. 7, pp. 24–58, 2024.
- [3] L. Jiao, “Advances in uncertain information fusion,” *Entropy*, vol. 26, no. 11, p. 945, 2024.
- [4] M. Abdar, F. Pourpanah, S. Hussain, D. Rezazadegan, L. Liu, M. Ghavamzadeh, P. Fieguth, X. Cao, A. Khosravi, U. R. Acharya, V. Makarenkov, and S. Nahavandi, “A review of uncertainty quantification in deep learning: Techniques, applications and challenges,” *Information Fusion*, vol. 76, pp. 243–297, 2021.
- [5] Y. Ovadia, E. Fertig, J. Ren, Z. Nado, D. Sculley, S. Nowozin, J. V. Dillon, B. Lakshminarayanan, and J. Snoek, “Can you trust your model’s uncertainty? evaluating predictive uncertainty under dataset shift,” in *Advances in Neural Information Processing Systems*, vol. 32, 2019, pp. 13 969–13 980.
- [6] T. Gneiting and A. E. Raftery, “Strictly proper scoring rules, prediction, and estimation,” *Journal of the American Statistical Association*, vol. 102, no. 477, pp. 359–378, 2007.
- [7] H. Wang, F. Smarandache, Y.-Q. Zhang, and R. Sundaraman, “Single valued neutrosophic sets,” *Multispace and Multistructure*, vol. 4, pp. 410–413, 2010.
- [8] T. Shaik, X. Tao, L. Li, H. Xie, and J. D. Velasquez, “A survey of multimodal information fusion for smart healthcare: Mapping the journey from data to wisdom,” *Information Fusion*, vol. 102, p. 102040, 2024.
- [9] A. Niculescu-Mizil and R. Caruana, “Predicting good probabilities with supervised learning,” in *Proceedings of the 22nd International Conference on Machine Learning*, 2005, pp. 625–632.
- [10] C. Guo, G. Pleiss, Y. Sun, and K. Q. Weinberger, “On calibration of modern neural networks,” in *Proceedings of the 34th International Conference on Machine Learning*, ser. Proceedings of Machine Learning Research, vol. 70, 2017, pp. 1321–1330.
- [11] Y. Geifman and R. El-Yaniv, “Selective classification for deep neural networks,” in *Advances in Neural Information Processing Systems*, vol. 30, 2017, pp. 4878–4887.
- [12] D. Stanujkić, D. Karabašević, G. Popović, D. Pamučar, Ž. Stević, E. K. Zavadskas, and F. Smarandache, “A single-valued neutrosophic extension of the EDAS method,” *Axioms*, vol. 10, no. 4, p. 245, 2021.
- [13] Y. Zeng, H. Ren, T. Yang, S. Xiao, and N. Xiong, “A novel similarity measure of single-valued neutrosophic sets based on modified manhattan distance and its applications,” *Electronics*, vol. 11, no. 6, p. 941, 2022.
- [14] H. M. A. Farid and M. Riaz, “Single-valued neutrosophic einstein interactive aggregation operators with applications for material selection in engineering design: Case study of cryogenic storage tank,” *Complex & Intelligent Systems*, vol. 8, pp. 2131–2149, 2022.
- [15] H. M. A. Farid and M. Riaz, “Single-valued neutrosophic dynamic aggregation information with time sequence preference for IoT technology in supply chain management,” *Engineering Applications of Artificial Intelligence*, vol. 126, p. 106940, 2023.

-
- [16] H. Garg, “A new exponential-logarithm-based single-valued neutrosophic set and their applications,” *Expert Systems with Applications*, vol. 238, p. 121854, 2024.
 - [17] M. A. Wajid, A. Zafar, H. Terashima-Marín, and M. S. Wajid, “Neutrosophic-CNN-based image and text fusion for multimodal classification,” *Journal of Intelligent & Fuzzy Systems*, vol. 45, no. 1, pp. 1039–1055, 2023.
 - [18] K. G. Lavanya, P. Dhanalakshmi, and M. Nandhini, “Neutrosophic fusion of multimodal brain images: Integrating neutrosophic entropy and feature extraction,” *Applied Soft Computing*, vol. 155, p. 111462, 2024.
 - [19] J. Lin, “Divergence measures based on the shannon entropy,” *IEEE Transactions on Information Theory*, vol. 37, no. 1, pp. 145–151, 1991.
 - [20] W. Wolberg, O. Mangasarian, N. Street, and W. Street, “Breast cancer wisconsin (diagnostic),” 1993.
 - [21] S. Aeberhard and M. Forina, “Wine,” 1992.
 - [22] E. Alpaydin and C. Kaynak, “Optical recognition of handwritten digits,” 1998.
 - [23] 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.