



Neutrosophic–Fuzzy Evidence Fusion for Uncertainty-Aware Cultivar Identification from Viticultural Chemical Profiles

Amine Saddik¹ Ika Hesti Agustin^{2,*}

¹ Faculty of Applied Sciences, Ibn Zohr University, Ait Melloul, Morocco

² Department of Mathematics, University of Jember, Jember, East Java, Indonesia

Emails: EL.cherrat@gmail.com; ikahesti.fmipa@unej.ac.id

Received: Aeptember 14, 2025 Aevised: Aovember 20, 2025 Accepted: Aanuary 15, 2026 ★ Corresponding author

ABSTRACT

Cultivar identification from viticultural chemical profiles is a multiclass recognition problem in which a hard label alone does not reveal whether global chemometric evidence agrees with the local structure of previously observed samples. This paper proposes Neutrosophic–Fuzzy Viticultural Evidence Fusion (NFVEF), an uncertainty-aware classifier that combines a global discriminant probability vector $G(x) \in \Delta^{K-1}$ with a Gaussian fuzzy-neighborhood vector $L(x) \in \Delta^{K-1}$. For every cultivar k , the two evidence views are converted into

$$T_k = \sqrt{G_k L_k}, \quad F_k = \sqrt{(1 - G_k)(1 - L_k)}, \\ I_k = 1 - T_k - F_k.$$

where I_k is exactly the squared Hellinger disagreement between the Bernoulli support views G_k and L_k . A logarithmic fuzzy opinion pool $H_k \propto G_k^\eta L_k^{1-\eta}$ is then attenuated by neutrosophic disagreement, $R_k \propto H_k \exp(-\kappa I_k)$, before classification. The winning class is accompanied by an uncertainty score $U = 1 - T_{\hat{k}}(1 - I_{\hat{k}})(1 - F_{\hat{k}})$, enabling uncertain chemical profiles to be flagged rather than reported with unqualified confidence. The method is evaluated on the UCI Wine cultivar dataset using 60 repeated stratified splits at four synthetic analytical-perturbation levels $\delta \in \{0, 0.1, 0.2, 0.3\}$ measured relative to training-feature standard deviations. NFVEF obtains mean accuracies of 0.9858, 0.9836, 0.9744, and 0.9728, respectively. At $\delta = 0.3$, its paired accuracy advantage over linear discriminant analysis is 0.00278 with a 95% bootstrap interval [0.00123, 0.00463], while RBF-SVM remains slightly better in raw accuracy. The uncertainty score detects NFVEF errors with mean AUC 0.9520 at the strongest perturbation, and retaining the lowest-uncertainty 90% of cases yields 0.9922 accuracy. The contribution is therefore not universal classifier dominance, but a mathematically interpretable fuzzy–neutrosophic evidence layer for cultivar identification from ambiguous chemical measurements.

Keywords: Neutrosophic sets ▪ Fuzzy evidence ▪ Viticultural analytics ▪ Cultivar identification ▪ Agricultural decision support ▪ Uncertainty quantification ▪ Chemical profiling

1. INTRODUCTION

Agricultural analytics increasingly relies on measurements whose value is not only predictive but also evidential: a chem-

ical or sensor profile should support an agronomic decision while indicating when the supporting information is internally ambiguous. Fuzzy methods are well suited to this setting because a sample need not belong to a linguistic or numerical

class with membership only 0 or 1; instead, $\mu_k(x) \in [0, 1]$ can represent graded support [1]. Recent smart-farming studies span monitoring, irrigation, automation, data analysis, yield prediction, and decision support, precisely because agricultural variables are heterogeneous and imperfectly measured [2, 3]. In viticulture, fuzzy inference has already been used to aggregate grape attributes into quality assessments [4] and to classify wine quality from physicochemical information [5]. These applications motivate a related but distinct question: when a chemical profile is assigned to a cultivar, can the classification mechanism also quantify *agreement* and *contradiction* between global and local evidence?

The distinction matters because classification confidence and evidence agreement are not identical quantities. Let $G_k(x)$ denote support for cultivar k from a global parametric model and $L_k(x)$ denote support from a local fuzzy neighborhood. Two profiles can produce the same final score while having very different evidence geometry: $(G_k, L_k) = (0.80, 0.80)$ represents agreement, whereas $(0.98, 0.62)$ contains substantially more view disagreement. A conventional convex average $\eta G_k + (1 - \eta)L_k$ compresses both views immediately. In contrast, neutrosophic representation retains truth, indeterminacy, and falsity components and therefore offers a language for separating support from contradiction [6–8]. This separation is especially attractive in agricultural decision support, where an uncertain sample may reflect cultivar overlap, atypical chemistry, analytical perturbation, or a profile lying far from the local training structure.

Fuzzy and neutrosophic methods have already entered agricultural modeling through several routes. Fuzzy cognitive maps have represented expert knowledge for cotton-yield prediction [9]; interpretable fuzzy inference systems have been developed to combine expert and data information [10]; fuzzy-rough modeling has supported crop identification and suitability prediction [11, 12]; and hybrid fuzzy-logic and genetic-algorithm modeling has been used for maize-yield prediction under maize–legume farming systems [13]. Neutrosophic formulations have addressed uncertain crop-production optimization [14] and crop-selection risk analysis [15]. Those studies demonstrate the usefulness of soft-computing uncertainty in agriculture, yet they do not directly construct a class-wise neutrosophic state from the disagreement of two probabilistic/fuzzy recognition views.

This paper develops such a construction for cultivar identification. The application uses the UCI Wine data, consisting of 178 chemical analyses from wines grown in the same Italian region but derived from three cultivars; each sample contains 13 measured constituents [16, 17]. The dataset is intentionally modest and well studied, so it should not be presented as a difficult modern benchmark. Its value here is methodological: the three-class chemical setting provides a reproducible agricultural test bed in which evidence fusion, analytical perturbation, uncertainty ranking, ablation, and sensitivity can all be examined without confounding the contribution with a large black-box architecture.

The proposed Neutrosophic–Fuzzy Viticultural Evidence Fusion (NFVEF) algorithm has four linked stages. First, standardized chemical features $z \in \mathbb{R}^d$ are evaluated by a global linear discriminant model and a local Gaussian fuzzy neighborhood, giving $G, L \in \Delta^{K-1}$. Second, each pair (G_k, L_k) is

transformed into a normalized neutrosophic evidence triple $Z_k = (T_k, I_k, F_k)$. Third, a logarithmic fuzzy opinion pool is attenuated according to I_k before selecting $\hat{y} = \arg \max_k R_k$. Fourth, the winning triple is mapped to an explicit uncertainty $U(x) \in [0, 1]$ that can support screening or selective laboratory review. The pipeline is therefore

$$x \longrightarrow (G, L) \longrightarrow \{(T_k, I_k, F_k)\}_{k=1}^K \longrightarrow R \longrightarrow (\hat{y}, U).$$

The algorithm is stated in pseudocode and its core properties are proved before the numerical study.

The contributions are: (i) a class-wise fuzzy–neutrosophic evidence construction in which indeterminacy has an exact Hellinger interpretation; (ii) a fusion rule that separates fuzzy pooling from neutrosophic disagreement attenuation; (iii) a bounded uncertainty function whose derivatives have the expected signs with respect to T , I , and F ; and (iv) a reproducible cultivar-identification experiment with analytical-noise stress testing, repeated splits, paired bootstrap intervals, class-wise recall, ablation, and parameter sensitivity. Because strong conventional classifiers are included, the analysis also reports where NFVEF does *not* dominate.

The remainder follows an application-centered progression. Section 2 separates agricultural fuzzy systems, neutrosophic decision models, and cultivar-recognition methods. Section 3 defines the agricultural data and the two evidence views. Section 4 presents NFVEF and its pseudocode. Section 5 establishes mathematical properties. Section 6 specifies the numerical protocol, Section 7 reports the main results, Section 8 analyzes robustness and agronomic interpretation, Section 9 states limitations, and Section 10 concludes.

2. RELATED WORK

2.1 Fuzzy reasoning in agriculture and viticulture

The transition from crisp thresholds to graded reasoning is a long-standing feature of fuzzy systems. Zadeh’s membership principle $\mu_A(x) \in [0, 1]$ [1] was followed by extensions for intuitionistic information [18], fuzzy entropy [19], clustering [20], possibilistic membership [21], and fuzzy-rough approximation [22]. In agriculture, these ideas are useful when variables such as quality, suitability, stress, or productivity do not have naturally sharp boundaries. The review by ElBeheiry and Balog [3] identifies sensors, communication, big data, actuators, and data analysis as central smart-farming themes, while Abdullah et al. [2] demonstrate fuzzy monitoring in a broader smart-agriculture architecture.

Application-specific work illustrates several forms of agricultural uncertainty. Tagarakis et al. [4] used fuzzy inference to model grape quality from multiple vineyard attributes, and Petropoulos et al. [5] developed a fuzzy tool for wine-quality classification. Papageorgiou et al. [9] represented causal agronomic knowledge with fuzzy cognitive maps for cotton-yield decision support. For data-driven crop recognition, Acharjya and Rathi [11] combined fuzzy-rough sets with a real-coded genetic algorithm, whereas Anitha and Acharjya [12] combined fuzzy approximation with rough-set and neural-network reasoning for crop suitability. Agricultural variety recognition also has a conventional machine-learning lineage; for example, Koklu and Ozkan [23] classified seven

registered dry-bean varieties using image-derived features. These studies show that fuzziness can operate at rule, feature, neighborhood, or decision levels rather than being restricted to one architecture.

2.2 Neutrosophic information and agricultural decisions

Single-valued neutrosophic representation associates an object with $T, I, F \in [0, 1]$ and does not require the three components to satisfy the probability-simplex constraint [6]. Subsequent work has developed similarity, entropy, and distance constructions for neutrosophic information [7, 8, 24, 25]. The distinction between evidence supporting a proposition, evidence against it, and residual indeterminacy is particularly relevant when inputs come from heterogeneous or partially conflicting sources.

Agricultural neutrosophic studies have so far concentrated largely on planning and multicriteria uncertainty. Kousar et al. [14] formulated multi-objective crop-production optimization in a neutrosophic fuzzy environment, whereas Borah and Dutta [15] developed a quadripartitioned single-valued neutrosophic risk-analysis framework for crop selection. These are decision-analysis contributions rather than pattern classifiers. NFVEF uses neutrosophic information at a different level: for every candidate cultivar k , it derives (T_k, I_k, F_k) directly from the agreement geometry of global and local classification evidence.

2.3 Cultivar recognition and local evidence

The UCI Wine data originate from chemical analyses of three cultivars and were introduced in work comparing multivariate statistical pattern-recognition methods [17]. Linear discriminant analysis has a classical statistical foundation [26], while nearest-neighbor decisions formalize local similarity [27]. Fuzzy-rough nearest-neighbor methods later showed how local neighborhoods can be integrated with graded approximation rather than crisp voting [28]. NFVEF deliberately combines these two scales: global evidence describes the class structure estimated from the complete training set, while local evidence asks whether nearby reference samples support the same conclusion.

The comparison set also includes RBF support-vector machines [29] and random forests [30]. These models are not claimed to be fuzzy or neutrosophic; they serve as strong predictive controls. All implementations use the scikit-learn software ecosystem [31]. This design helps separate two questions: whether NFVEF predicts competitively and whether its explicit (T, I, F) layer adds uncertainty information that ordinary class labels do not provide.

3. APPLICATION SETTING AND EVIDENCE REPRESENTATION

3.1 Chemical cultivar-identification setting

Let

$$\mathcal{D} = \{(x_i, y_i)\}_{i=1}^n, \quad x_i \in \mathbb{R}^d, \quad y_i \in \{1, \dots, K\}, \quad (1)$$

with $n = 178$, $d = 13$, and $K = 3$. The UCI record describes the observations as chemical analyses of wines from three cultivars grown in the same Italian region [16]. The features

are alcohol, malic acid, ash, alkalinity of ash, magnesium, total phenols, flavanoids, nonflavanoid phenols, proanthocyanins, color intensity, hue, OD280/OD315 of diluted wines, and proline. Because cultivar names are not supplied in the repository, the classes are reported neutrally as C1, C2, and C3 rather than assigning unsupported grape-variety names.

Training features are standardized coordinatewise,

$$z_{ij} = \frac{x_{ij} - \hat{\mu}_j}{\hat{\sigma}_j}, \quad j = 1, \dots, d, \quad (2)$$

using training-set statistics only. Thus $z_i \in \mathbb{R}^d$ is dimensionless, which is essential because the chemical variables have heterogeneous physical scales. A perturbed test profile used later in the robustness experiment is generated as

$$x^{(\delta)} = x + \delta \hat{D} \xi, \quad \xi \sim \mathcal{N}(0, I_d), \quad (3)$$

where $\hat{D} = \text{diag}(\hat{\sigma}_1, \dots, \hat{\sigma}_d)$ and $\delta \in \{0, 0.1, 0.2, 0.3\}$. Equation (3) is a controlled analytical perturbation model; it is not asserted to reproduce the noise law of a specific spectrometer or laboratory instrument.

3.2 Global discriminant evidence

Let $G_k(x)$ be the posterior probability returned by linear discriminant analysis after standardization. With class prior π_k , class mean μ_k , and common covariance Σ , the discriminant score has the form

$$g_k(z) = z^\top \Sigma^{-1} \mu_k - \frac{1}{2} \mu_k^\top \Sigma^{-1} \mu_k + \log \pi_k, \quad (4)$$

and normalization yields

$$G_k(z) = \frac{e^{g_k(z)}}{\sum_{\ell=1}^K e^{g_\ell(z)}}, \quad G(z) \in \Delta^{K-1}. \quad (5)$$

The global view is therefore class-wide: every training observation contributes to the fitted class means and pooled covariance, so $G_k(z)$ reflects the global chemometric structure rather than only the nearest reference samples.

3.3 Local Gaussian fuzzy evidence

Global posteriors can be confident even when a query lies in a locally mixed region. To represent neighborhood evidence, let $N_q(z)$ be the q nearest standardized training samples and define

$$w_r(z) = \exp\left[-\frac{\|z - z_r\|_2^2}{2h^2}\right], \quad r \in N_q(z), \quad (6)$$

where h is the median training distance to the q th neighbor. The fuzzy local support for cultivar k is

$$L_k(z) = \frac{\sum_{r \in N_q(z)} w_r(z) \mathbf{1}(y_r = k)}{\sum_{r \in N_q(z)} w_r(z)}, \quad (7)$$

so that $L(z) \in \Delta^{K-1}$ exactly. The denominator is strictly positive because Gaussian weights satisfy $w_r(z) > 0$ for every retrieved neighbor. Unlike crisp q -NN voting, a nearby sample contributes more than a distant member of the same neighborhood. This weighting is consistent with the broader fuzzy-clustering and fuzzy-neighborhood tradition [20, 21, 28].

The two vectors $G, L \in \Delta^{K-1}$ encode complementary statistical scales. Their direct difference $|G_k - L_k|$ is informative but incomplete because it ignores whether both views simultaneously support or reject a class. NFVEF therefore transforms each pair (G_k, L_k) into a three-component evidence state before fusion.

4. PROPOSED NFVEF ALGORITHM

4.1 Class-wise neutrosophic evidence

For each cultivar k , define

$$T_k = \sqrt{G_k L_k}, \quad (8)$$

$$F_k = \sqrt{(1 - G_k)(1 - L_k)}, \quad (9)$$

$$I_k = 1 - T_k - F_k. \quad (10)$$

The product in T_k is large only when both views support the class; analogously, F_k is large only when both views reject it. Indeterminacy I_k is the residual agreement deficit. Although general single-valued neutrosophic sets do not impose $T + I + F = 1$ [6], the construction in (8)–(10) intentionally occupies the normalized subfamily

$$T_k + I_k + F_k = 1, \quad (11)$$

which makes the three components immediately interpretable for two-view evidence fusion.

4.2 Fuzzy opinion pooling and neutrosophic attenuation

Before applying the disagreement term, the two support distributions are fused through a logarithmic fuzzy pool,

$$\tilde{H}_k = \max(G_k, \varepsilon)^\eta \max(L_k, \varepsilon)^{1-\eta}, \quad H_k = \frac{\tilde{H}_k}{\sum_\ell \tilde{H}_\ell}, \quad (12)$$

where $0 \leq \eta \leq 1$ controls the global–local balance. Thus $\eta = 1$ recovers the global evidence and $\eta = 0$ recovers the local evidence in the positive-support limit. The final NFVEF score attenuates a pooled class when its two evidence views disagree:

$$\tilde{R}_k = H_k e^{-\kappa I_k}, \quad R_k = \frac{\tilde{R}_k}{\sum_\ell \tilde{R}_\ell}, \quad \kappa \geq 0. \quad (13)$$

For any classes k and ℓ with positive pooled support, normalization cancels in the odds ratio, giving

$$\log \frac{R_k}{R_\ell} = \log \frac{H_k}{H_\ell} - \kappa(I_k - I_\ell). \quad (14)$$

Hence an excess disagreement $I_k - I_\ell = \Delta I > 0$ multiplies the pooled odds of class k relative to class ℓ by $e^{-\kappa \Delta I}$. The cultivar decision is

$$\hat{y} = \arg \max_{k \in \{1, \dots, K\}} R_k. \quad (15)$$

When $\kappa = 0$, Equation (13) reduces exactly to the fuzzy log-pool $R = H$. Thus the ablation between H and R isolates the contribution of neutrosophic disagreement rather than changing the underlying global or local evidence estimators.

Algorithm 1 Neutrosophic–Fuzzy Viticultural Evidence Fusion (NFVEF)

Require: Training data $\mathcal{D}$, query x , neighbors q , balance η , attenuation κ

- 1: Fit training scaler and transform $x \mapsto z$
- 2: Fit LDA and compute $G = (G_1, \dots, G_K)$ using (5)
- 3: Find $N_q(z)$ and compute Gaussian weights $w_r(z)$ using (6)
- 4: Compute local fuzzy vector $L = (L_1, \dots, L_K)$ using (7)
- 5: **for** $k = 1, \dots, K$ **do**
- 6: $T_k \leftarrow \sqrt{G_k L_k}$
- 7: $F_k \leftarrow \sqrt{(1 - G_k)(1 - L_k)}$
- 8: $I_k \leftarrow 1 - T_k - F_k$
- 9: $\tilde{H}_k \leftarrow \max(G_k, \varepsilon)^\eta \max(L_k, \varepsilon)^{1-\eta}$
- 10: **end for**
- 11: Normalize $H_k \leftarrow \tilde{H}_k / (\sum_\ell \tilde{H}_\ell)$
- 12: $\tilde{R}_k \leftarrow H_k e^{-\kappa I_k}$ for all k ; normalize to R
- 13: $\hat{k} \leftarrow \arg \max_k R_k$
- 14: $U \leftarrow 1 - T_{\hat{k}}(1 - I_{\hat{k}})(1 - F_{\hat{k}})$
- 15: **return** cultivar $\hat{k}$, uncertainty U , and triples $\{(T_k, I_k, F_k)\}_{k=1}^K$

4.3 Decision-level uncertainty

For the winning class $\hat{k}$, define reliability and uncertainty by

$$\mathcal{R}(x) = T_{\hat{k}}(1 - I_{\hat{k}})(1 - F_{\hat{k}}), \quad U(x) = 1 - \mathcal{R}(x). \quad (16)$$

A confident sample should therefore exhibit high truth support, low indeterminacy, and low falsity simultaneously. This differs from ordinary $1 - \max_k R_k$: U is explicitly tied to the neutrosophic components used to explain the decision.

Algorithm 1 is deliberately modular. The global learner could be replaced by another calibrated model and the local evidence could be generated by another fuzzy relation, but the present study fixes both components before testing so that the numerical experiment evaluates one reproducible method rather than a family selected after observing results.

5. MATHEMATICAL PROPERTIES

Proposition 1 (Validity and Hellinger interpretation)

For $G_k, L_k \in [0, 1]$, the quantities in (8)–(10) satisfy $T_k, F_k, I_k \in [0, 1]$ and (11). Moreover,

$$I_k = H^2(\text{Bern}(G_k), \text{Bern}(L_k)), \quad (17)$$

where squared Hellinger distance is defined as one minus the Bhattacharyya coefficient.

Proof. Nonnegativity of T_k and F_k is immediate. By Cauchy–Schwarz,

$$\begin{aligned} & \sqrt{G_k L_k} + \sqrt{(1 - G_k)(1 - L_k)} \\ & \leq \sqrt{[G_k + (1 - G_k)][L_k + (1 - L_k)]} = 1. \end{aligned}$$

so $I_k \geq 0$; the upper bound follows because $T_k + F_k \geq 0$. Equation (11) holds by definition. The Bernoulli Bhattacharyya coefficient is precisely $\sqrt{G_k L_k} + \sqrt{(1 - G_k)(1 - L_k)}$, yielding (17). $\square$

A consequence is

$$I_k = 0 \iff G_k = L_k, \quad (18)$$

so the indeterminacy coordinate is not an arbitrary penalty. It is the square of the Hellinger metric between two Bernoulli class-support views; equivalently, $\sqrt{I_k}$ is a proper metric distance while I_k is its squared disagreement. For small $\Delta_k = L_k - G_k$ with $0 < G_k < 1$, a Taylor expansion gives

$$I_k = \frac{\Delta_k^2}{8G_k(1-G_k)} + O(|\Delta_k|^3), \quad (19)$$

showing that equal absolute disagreement is more consequential when global support lies close to 0 or 1 than near $1/2$.

Proposition 2 (Fusion limits) *Assume strictly positive G_k, L_k . The pool in (12) satisfies $H = G$ at $\eta = 1$ and $H = L$ at $\eta = 0$. The final score satisfies $R = H$ when $\kappa = 0$. If $G_k = L_k$ for every k , then $I_k = 0$ for every class and therefore $R = H$ for any $\kappa \geq 0$.*

Proof. Substituting $\eta \in \{0, 1\}$ into (12) gives the first two identities after normalization. Substituting $\kappa = 0$ into (13) gives $e^0 = 1$. If $G_k = L_k$, Proposition 1 and (18) give $I_k = 0$, so the attenuation is again unity. $\square$

Proposition 3 (Uncertainty bounds and monotonicity)
For $T, I, F \in [0, 1]$, the uncertainty function

$$U(T, I, F) = 1 - T(1 - I)(1 - F) \quad (20)$$

satisfies $0 \leq U \leq 1$ and

$$\begin{aligned} \frac{\partial U}{\partial T} &= -(1 - I)(1 - F) \leq 0, \\ \frac{\partial U}{\partial I} &= T(1 - F) \geq 0, \quad \frac{\partial U}{\partial F} = T(1 - I) \geq 0. \end{aligned} \quad (21)$$

Proof. The product $T(1 - I)(1 - F)$ lies in $[0, 1]$, giving the bounds. Direct differentiation yields (21). $\square$

Thus, on the ambient cube $[0, 1]^3$, increasing truth cannot increase uncertainty, whereas increasing indeterminacy or falsity cannot decrease it. NFVEF triples additionally satisfy $T + I + F = 1$, so admissible changes of the three coordinates are coupled; the partial derivatives in (21) should therefore be read as coordinatewise sensitivity statements for the uncertainty mapping rather than as three independently varying agronomic quantities.

For n training samples and d features, fitting LDA has the usual data-dependent linear-algebra cost, while query-time global scoring is $O(Kd)$. Brute-force neighbor retrieval is $O(nd)$ per query and local aggregation is $O(qK)$. The neutrosophic transformation and fusion are $O(K)$, so for the present small reference set the dominant query cost is the neighborhood search. Spatial indexing or approximate-neighbor search can replace brute force when n becomes large without altering Equations (8)–(16).

6. EXPERIMENTAL PROTOCOL

6.1 Dataset and repeated evaluation

The UCI Wine dataset contains class counts (59, 71, 48) and no missing values [16]. Each experiment uses a stratified 70/30 train–test split. To reduce dependence on one favorable partition, 60 independent stratified splits are generated

for every perturbation level in (3). Test perturbations use training-feature standard deviations only; the training data remain unperturbed. Consequently the experiment evaluates robustness to analytical variation at inference rather than robustness to mislabeled or corrupted training observations.

The principal NFVEF setting is

$$q = 9, \quad \eta = 0.80, \quad \kappa = 1.0. \quad (22)$$

The choice gives the global chemometric view larger weight while preserving local evidence and a moderate disagreement attenuation. These values are fixed for the main experiment. A separate sensitivity grid later evaluates $q \in \{5, 9, 15\}$, $\eta \in \{0.65, 0.80, 0.95\}$, and $\kappa \in \{0, 1, 2\}$.

6.2 Comparators and metrics

Predictive comparators are LDA, logistic regression, distance-weighted 7-NN, RBF-SVM with $C = 2$, and a 70-tree random forest. These cover linear probabilistic discrimination, local classification, kernel discrimination, and tree ensembles; their classical foundations include Fisher [26], Cover and Hart [27], Cortes and Vapnik [29], and Breiman [30]. All computational models are implemented with scikit-learn [31]. The fuzzy/neutrosophic ablation includes global evidence G , local evidence L , the fuzzy log-pool H , and complete NFVEF R .

For test labels y_i and predictions $\hat{y}_i$, accuracy and macro- F_1 are reported. Let $e_i = \mathbf{1}(\hat{y}_i \neq y_i)$. The diagnostic value of uncertainty is measured by

$$\text{AUC}_{\text{err}} = \Pr\{U(x_{\text{error}}) > U(x_{\text{correct}})\}, \quad (23)$$

whenever both correct and incorrect samples occur in a split. Values near one indicate that errors are ranked as substantially more uncertain than correct decisions.

A second diagnostic retains the lowest-uncertainty 90% of each test split. With test size 54, the implementation keeps $\lceil 0.9 \times 54 \rceil = 49$ cases, giving realized coverage

$$\hat{C}_{0.9} = 49/54 = 0.9074. \quad (24)$$

The corresponding selective accuracy is descriptive rather than a deployed rejection policy because the threshold is computed from the test uncertainty ordering. It answers a narrow question: *does U rank the safer profiles near the bottom?*

Paired contrasts use the same 60 train–test partitions for both methods. If d_r is a replicate-level metric difference, the reported 95% interval is obtained from 4000 paired bootstrap resamples of $\{d_r\}_{r=1}^{60}$. The fixed random seed is 20251217 and replicate-level outputs are retained with the source code.

7. NUMERICAL RESULTS

7.1 Classification under analytical perturbation

Table 1 reports mean accuracy. On clean profiles, LDA is marginally higher than NFVEF (0.9861 versus 0.9858), a difference of only -3.1×10^{-4} . At $\delta = 0.2$, NFVEF reaches 0.9744, slightly above LDA (0.9731), logistic regression (0.9719), RBF-SVM (0.9738), random forest (0.9704), and kNN (0.9565). At the strongest perturbation $\delta = 0.3$, NFVEF

remains at 0.9728 while LDA falls to 0.9701 and random forest to 0.9608. RBF-SVM obtains the best raw accuracy at this level (0.9762). The result supports competitive robustness under the stated perturbation model, not universal superiority.

Table 1. Mean classification accuracy over 60 repeated stratified splits. The perturbation δ in (3) is measured in training-feature standard-deviation units.

Method	$\delta = 0$	$\delta = 0.1$	$\delta = 0.2$	$\delta = 0.3$
NFVEF	0.9858	0.9836	0.9744	0.9728
LDA	0.9861	0.9843	0.9731	0.9701
Logistic regression	0.9818	0.9790	0.9719	0.9719
7-NN, distance weighted	0.9630	0.9630	0.9565	0.9494
RBF-SVM	0.9818	0.9827	0.9738	0.9762
Random forest	0.9731	0.9790	0.9704	0.9608

The decrease in NFVEF accuracy from clean to 0.3σ perturbation is

$$\Delta_{0 \rightarrow 0.3} = 0.97284 - 0.98580 = -0.01296, \quad (25)$$

which corresponds to about 1.30 percentage points. Macro- F_1 follows almost the same trajectory (Table 2), indicating that overall accuracy is not being maintained by sacrificing one cultivar completely.

Table 2. NFVEF performance and uncertainty diagnostics. Values are means over 60 splits; accuracy SD is shown in parentheses.

δ	Accuracy	Macro- F_1	Error AUC	Sel. Acc.
0.0	.9858 (.0142)	.9859	.9718	.9986
0.1	.9836 (.0137)	.9839	.9708	.9980
0.2	.9744 (.0187)	.9748	.9613	.9942
0.3	.9728 (.0165)	.9733	.9520	.9922

Selective accuracy uses realized coverage 0.9074.

7.2 Uncertainty as an error-screening signal

The more distinctive numerical result concerns uncertainty ranking. Error AUC decreases smoothly from 0.9718 on clean profiles to 0.9520 at $\delta = 0.3$, yet remains far above the noninformative value 0.5. Thus, even when perturbation makes classification harder, wrong cultivar assignments tend to receive larger $U(x)$ than correct assignments. At approximately 90% coverage, retained-case accuracy remains 0.9922 at the strongest perturbation, compared with the full-sample 0.9728. This is a diagnostic property of the ranking, not evidence that 9.26% of future samples can be rejected without an externally calibrated operating threshold.

The paired bootstrap results in Table 3 sharpen the predictive comparisons. At $\delta = 0.3$, NFVEF exceeds LDA by 0.00278 with interval $[0.00123, 0.00463]$ and random forest by 0.01204 with interval $[0.00586, 0.01821]$. Against RBF-SVM, however, the mean difference is -0.00340 and the interval crosses zero. At $\delta = 0.2$, the NFVEF-LDA difference 0.00123 also has an interval spanning zero. These intervals discourage overinterpreting very small differences on a compact dataset.

7.3 Cultivar-wise behavior

Table 4 shows that C2 is consistently the most difficult class. At $\delta = 0.2$, recalls are (0.9824, 0.9571, 0.9889) for (C1, C2, C3); at $\delta = 0.3$ they are (0.9852, 0.9532, 0.9856). The spread between the easiest and hardest cultivar at $\delta = 0.3$ is therefore approximately 0.0324. Because the dataset does

Table 3. Paired bootstrap accuracy contrasts, NFVEF minus comparator; 4000 bootstrap resamples of 60 replicate-wise differences.

δ	Comparator	Mean	2.5%	97.5%
0.0	LDA	-.00031	-.00185	.00123
	SVM	.00401	-.00093	.00926
	RF	.01265	.00648	.01914
0.2	LDA	.00123	-.00093	.00370
	SVM	.00062	-.00463	.00586
	RF	.00401	-.00093	.00926
0.3	LDA	.00278	.00123	.00463
	SVM	-.00340	-.00957	.00248
	RF	.01204	.00586	.01821

not provide cultivar names, this pattern should be interpreted statistically rather than attributed to a particular grape variety or chemical mechanism.

Table 4. NFVEF class recall, mean (SD), over repeated splits.

δ	C1	C2	C3
0.0	.9963 (.0140)	.9738 (.0309)	.9900 (.0240)
0.1	.9963 (.0140)	.9643 (.0324)	.9956 (.0168)
0.2	.9824 (.0280)	.9571 (.0408)	.9889 (.0251)
0.3	.9852 (.0268)	.9532 (.0416)	.9856 (.0327)

8. ROBUSTNESS, ABLATION, AND INTERPRETATION

8.1 What is contributed by each evidence layer?

Table 5 isolates the components at $\delta = 0.2$. Global LDA evidence alone obtains 0.97315 accuracy, while local fuzzy evidence alone obtains 0.95988. Their logarithmic pool raises accuracy to 0.97438. Complete NFVEF retains the same mean accuracy but increases error AUC slightly from 0.96053 to 0.96125. The selective accuracy of the log-pool is marginally higher (0.99490 versus 0.99422). Consequently, the neutrosophic attenuation is best interpreted here as an *uncertainty-structuring mechanism*; it is not responsible for a large accuracy gain on this dataset.

Table 5. Ablation at $\delta = 0.2$.

Variant	Accuracy	Error AUC	Sel. Acc.
Global evidence	.97315	.95262	.99116
Local fuzzy evidence	.95988	.94236	.98469
Fuzzy log-pool	.97438	.96053	.99490
NFVEF	.97438	.96125	.99422

This ablation clarifies the division of labor between the two uncertainty mechanisms. The value of a fuzzy-neutrosophic classifier cannot be inferred merely from the presence of T, I, F coordinates. In NFVEF, the fuzzy pool H carries most of the predictive improvement over either isolated evidence source, whereas I_k provides a geometrically defined disagreement quantity and a modest improvement in error ranking. The decomposition is therefore useful even when $\arg \max_k R_k = \arg \max_k H_k$ for many samples.

8.2 Parameter sensitivity

The separate 30-split sensitivity experiment at $\delta = 0.2$ spans all 27 combinations of (q, η, κ) . Mean accuracy ranges only from 0.97346 to 0.97716, a width of

$$0.97716 - 0.97346 = 0.00370, \quad (26)$$

while error AUC ranges from 0.96116 to 0.97614. The principal configuration (9, 0.80, 1) yields accuracy 0.97593 and error AUC 0.96958 in this separate split set. The highest observed accuracy is 0.97716 at (9, 0.65, 2), whereas the highest AUC 0.97614 occurs elsewhere in the grid. This noncoincidence reinforces the distinction between maximizing classification accuracy and maximizing the ability to flag questionable decisions.

Table 6. Selected sensitivity settings at $\delta = 0.2$ (30 independent splits).

q	η	κ	Accuracy	Error AUC
5	.65	0	.97346	.96657
9	.65	2	.97716	.96325
9	.80	1	.97593	.96958
9	.95	0	.97593	.97011
15	.95	1	.97469	.97614

Full 27-setting grid is included in the reproducibility package.

8.3 Interpretation for viticultural chemical analytics

The agricultural use case is best understood as a decision-support layer around laboratory-derived chemical profiles. Existing fuzzy viticulture work shows how multiple grape attributes can be combined into an interpretable quality index [4], while fuzzy wine analysis has demonstrated the usefulness of graded rules for quality classification [5]. NFVEF addresses a different stage: given a profile to be assigned to a known cultivar class, it asks whether global chemometric structure and nearby reference profiles tell the same story.

For a winning class $\hat{k}$, three qualitatively different situations are possible. If $T_{\hat{k}} \approx 1$ and $(I_{\hat{k}}, F_{\hat{k}}) \approx (0, 0)$, the global and local views agree on strong support. If $I_{\hat{k}}$ is elevated, the views conflict even when the final score still selects that cultivar. If $F_{\hat{k}}$ is elevated, both views jointly reject the candidate and another class should normally dominate after normalization. This separation can support a laboratory workflow in which a high- U profile is remeasured, checked for sample handling, or routed to a more specialized identification model rather than being accepted as an ordinary high-confidence prediction.

The formulation also aligns with a broader agricultural trend toward hybrid uncertainty-aware systems. Fuzzy-rough crop identification uses graded approximations to manage heterogeneous agricultural data [11]; neutrosophic optimization and risk analysis retain indeterminacy in crop-planning decisions [14, 15]; and hybrid fuzzy-logic and genetic-algorithm models combine interpretable soft computing with data-driven maize-yield prediction [13]. NFVEF contributes a smaller, mathematically explicit component to that landscape: a two-view evidence transform that can sit between a predictive model and an agronomic decision.

9. LIMITATIONS

The first limitation is the dataset. With only 178 observations and three well-separated classes, the UCI Wine benchmark cannot establish field-level generalizability. The original statistical study already used it for comparative pattern recognition [17], and modern classifiers can achieve very high accuracy. The present experiment therefore validates the mechanics of the proposed evidence layer rather than demonstrating readiness for field deployment or commercial cultivar certification. Future validation should use larger

multi-vintage, multi-region, and multi-instrument collections in which cultivar, terroir, season, and measurement platform vary simultaneously.

Second, the perturbation model in (3) is synthetic and independent across features after scaling. Real analytical errors may be heteroscedastic, correlated, batch dependent, or affected by calibration drift. A stronger study would estimate a covariance Ω from repeated laboratory measurements and replace $\delta\hat{D}\xi$ with a domain-specific error model $\xi \sim \mathcal{N}(0, \Omega)$ or an empirical bootstrap of instrument replicates.

Third, NFVEF currently uses LDA as the global source and Euclidean Gaussian neighborhoods as the local source. This pairing was selected for transparency, not because it is universally optimal. Stronger nonlinear models can outperform NFVEF in raw accuracy, as RBF-SVM does at $\delta = 0.3$ in Table 1. The uncertainty layer should therefore be studied with calibrated nonlinear global probabilities, metric learning, and fuzzy-rough neighborhoods. The deeper issue is not which classifier has the largest third decimal place, but whether (T, I, F) remains stable and useful when evidence sources become more complex.

Finally, the 90% selective analysis is retrospective because the threshold is taken from the test uncertainty ordering. Deployment requires a validation-set threshold θ satisfying an operating constraint such as $\Pr(U \leq \theta) \approx c$ on held-out historical data. Such calibration is essential before attaching economic or agronomic consequences to a rejection decision.

10. CONCLUSION

This paper formulated cultivar identification as an evidence-fusion problem rather than a hard multiclass decision alone. Given global support G_k and local fuzzy support L_k , NFVEF constructs

$$\begin{aligned} T_k &= \sqrt{G_k L_k}, \\ I_k &= 1 - \sqrt{G_k L_k} - \sqrt{(1 - G_k)(1 - L_k)}, \\ F_k &= \sqrt{(1 - G_k)(1 - L_k)}. \end{aligned}$$

where I_k is exactly a squared Hellinger disagreement. The fuzzy log-pool $H_k \propto G_k^\eta L_k^{1-\eta}$ provides the predictive fusion, while $e^{-\kappa I_k}$ explicitly penalizes disagreement before normalization. The associated uncertainty $U = 1 - T(1 - I)(1 - F)$ is bounded and monotone in the intended directions.

On the UCI cultivar data, NFVEF maintains mean accuracy from 0.9858 without perturbation to 0.9728 under 0.3σ synthetic analytical noise. Its paired NFVEF-LDA bootstrap interval is entirely positive at the strongest tested perturbation, yet RBF-SVM remains marginally higher in raw accuracy, so no universal dominance claim is warranted. More importantly for the proposed contribution, NFVEF uncertainty ranks errors effectively: mean error-detection AUC remains 0.9520 at $\delta = 0.3$, and low-uncertainty cases show substantially higher retained-case accuracy. Ablation demonstrates that fuzzy pooling supplies most of the classification gain, whereas the neutrosophic layer primarily structures disagreement and diagnostic uncertainty.

The resulting framework is therefore best viewed as an interpretable fuzzy-neutrosophic evidence layer for viticultural chemical identification decisions. Future work should eval-

uate it on multi-site cultivar datasets, repeated-instrument measurements, spectral and metabolomic profiles, and field-deployed validation thresholds, where explicit indeterminacy may have greater operational value than on a compact benchmark.

REFERENCES

- [1] L. A. Zadeh, "Fuzzy sets," *Information and Control*, vol. 8, no. 3, pp. 338–353, 1965.
- [2] N. Abdullah, N. A. B. Durani, M. F. B. Shari, K. S. Siong, V. K. W. Hau, W. N. Siong, and I. K. A. Ahmad, "Towards smart agriculture monitoring using fuzzy systems," *IEEE Access*, vol. 9, pp. 4097–4111, 2021.
- [3] N. ElBeheiry and R. S. Balog, "Technologies driving the shift to smart farming: A review," *IEEE Sensors Journal*, vol. 23, no. 3, pp. 1752–1769, 2023.
- [4] A. Tagarakis, S. Koundouras, E. I. Papageorgiou, Z. Dikopoulou, S. Fountas, and T. A. Gemtos, "A fuzzy inference system to model grape quality in vineyards," *Precision Agriculture*, vol. 15, pp. 555–578, 2014.
- [5] S. Petropoulos, C. S. Karavas, A. T. Balafoutis, I. Paraskevopoulos, S. Kallithraka, and Y. Kotseridis, "Fuzzy logic tool for wine quality classification," *Computers and Electronics in Agriculture*, vol. 142, pp. 552–562, 2017.
- [6] H. Wang, F. Smarandache, Y. Zhang, and R. Sunderraman, "Single valued neutrosophic sets," *Multispace and Multistructure*, vol. 4, pp. 410–413, 2010.
- [7] 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.
- [8] 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.
- [9] E. I. Papageorgiou, A. T. Markinos, and T. A. Gemtos, "Fuzzy cognitive map based approach for predicting yield in cotton crop production as a basis for decision support system in precision agriculture application," *Applied Soft Computing*, vol. 11, no. 4, pp. 3643–3657, 2011.
- [10] S. Guillaume and B. Charnomordic, "Learning interpretable fuzzy inference systems with FisPro," *Information Sciences*, vol. 181, no. 20, pp. 4409–4427, 2011.
- [11] D. P. Acharjya and R. Rathi, "An integrated fuzzy rough set and real coded genetic algorithm approach for crop identification in smart agriculture," *Multimedia Tools and Applications*, vol. 81, pp. 35 117–35 142, 2022.
- [12] A. Anitha and D. P. Acharjya, "Crop suitability prediction in vellore district using rough set on fuzzy approximation space and neural network," *Neural Computing and Applications*, vol. 30, no. 12, pp. 3633–3650, 2018.
- [13] K. M. Agboka, H. E. Z. Tonnang, E. M. Abdel-Rahman, J. Odindi, O. Mutanga, and S. Niassy, "Data-driven artificial intelligence (AI) algorithms for modelling potential maize yield under maize–legume farming systems in east africa," *Agronomy*, vol. 12, no. 12, p. 3085, 2022.
- [14] S. Kousar, M. N. Sangi, N. Kausar, D. Pamucar, E. Ozbilge, and T. Cagin, "Multi-objective optimization model for uncertain crop production under neutrosophic fuzzy environment: A case study," *AIMS Mathematics*, vol. 8, no. 3, pp. 7584–7605, 2023.
- [15] G. Borah and P. Dutta, "Fuzzy risk analysis in crop selection using information measures on quadripartitioned single-valued neutrosophic sets," *Expert Systems with Applications*, vol. 255, p. 124750, 2024.
- [16] S. Aeberhard and M. Forina, "Wine," 1992, dataset.
- [17] S. Aeberhard, D. Coomans, and O. Y. de Vel, "Comparative analysis of statistical pattern recognition methods in high dimensional settings," *Pattern Recognition*, vol. 27, no. 8, pp. 1065–1077, 1994.
- [18] K. T. Atanassov, "Intuitionistic fuzzy sets," *Fuzzy Sets and Systems*, vol. 20, no. 1, pp. 87–96, 1986.
- [19] A. De Luca and S. Termini, "A definition of a non-probabilistic entropy in the setting of fuzzy sets theory," *Information and Control*, vol. 20, no. 4, pp. 301–312, 1972.
- [20] J. C. Bezdek, R. Ehrlich, and W. Full, "FCM: The fuzzy c-means clustering algorithm," *Computers & Geosciences*, vol. 10, no. 2–3, pp. 191–203, 1984.
- [21] R. Krishnapuram and J. M. Keller, "A possibilistic approach to clustering," *IEEE Transactions on Fuzzy Systems*, vol. 1, no. 2, pp. 98–110, 1993.
- [22] D. Dubois and H. Prade, "Rough fuzzy sets and fuzzy rough sets," *International Journal of General Systems*, vol. 17, no. 2–3, pp. 191–209, 1990.
- [23] M. Koklu and I. A. Ozkan, "Multiclass classification of dry beans using computer vision and machine learning techniques," *Computers and Electronics in Agriculture*, vol. 174, p. 105507, 2020.
- [24] M. Luo, G. Zhang, and L. Wu, "A novel distance between single valued neutrosophic sets and its application in pattern recognition," *Soft Computing*, vol. 26, pp. 11 129–11 137, 2022.
- [25] J. Ye, "Entropy measures of simplified neutrosophic sets and their decision-making approach with positive and negative arguments," *Journal of Management Analytics*, vol. 8, no. 2, pp. 252–266, 2021.
- [26] R. A. Fisher, "The use of multiple measurements in taxonomic problems," *Annals of Eugenics*, vol. 7, no. 2, pp. 179–188, 1936.
- [27] T. M. Cover and P. E. Hart, "Nearest neighbor pattern classification," *IEEE Transactions on Information Theory*, vol. 13, no. 1, pp. 21–27, 1967.
- [28] C. Cornelis and R. Jensen, "Fuzzy-rough nearest neighbour classification and prediction," *Theoretical Computer Science*, vol. 412, no. 42, pp. 5871–5884, 2011.
- [29] C. Cortes and V. Vapnik, "Support-vector networks," *Machine Learning*, vol. 20, pp. 273–297, 1995.
- [30] L. Breiman, "Random forests," *Machine Learning*, vol. 45, pp. 5–32, 2001.
- [31] 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, pp. 2825–2830, 2011.