



Paper Type: Original Article

Weighted HyperRough Set and Sequential HyperRough Set

Takaaki Fujita* 

Independent Researcher, Shinjuku, Shinjuku-ku, Tokyo, Japan; t171d603@gunma-u.ac.jp.

Citation:

Received: 16 August 2025

Revised: 10 December 2025

Accepted: 21 February 2026

Fujita, T. (2026). Weighted HyperRough set and sequential HyperRough set . *Uncertainty Discourse and Applications*, 3(2), 136-151.

Abstract

Rough set theory provides a mathematical framework for approximating subsets using lower and upper bounds defined by equivalence relations, effectively capturing uncertainty in classification and data analysis. Building upon these foundational concepts, further generalizations such as Hyperrough Sets and Superhyperrough Sets have been developed. The Weighted Rough Set is an extension of rough set theory that incorporates importance weights for attributes, enabling more precise classification and approximation. In this paper, we investigate the Weighted Hyperrough Set, Weighted Superhyperrough Set, Sequential Hyperrough Set, and Sequential n -Superhyperrough Set, each of which extends the fundamental ideas of Weighted Rough Sets and Sequential Rough Sets to more complex and high-dimensional decision environments.

Keywords: Rough set, Hyperrough set, Weighted rough set, SuperHyperRough set, Sequential rough set.

1|Introduction and Definitions

This section provides an introduction to the foundational concepts and definitions required for the discussions in this paper. Throughout this paper, all sets under consideration are assumed to be finite.

1.1|Rough Set, HyperRough Set, and Superhyperrough Set

A rough set approximates a subset using lower and upper bounds determined by equivalence classes, thereby capturing both certainty and uncertainty in membership [1–3]. The following definitions formalize these concepts.

Definition 0.1 (Set). [4] A *set* is a well-defined collection of distinct elements or objects. If a is an element of a set A , we write $a \in A$; otherwise, we write $a \notin A$.

Definition 0.2 (Subset). [4] Let A and B be sets. A is called a *subset* of B , denoted $A \subseteq B$, if every element of A is also an element of B . If $A \subseteq B$ but $A \neq B$, then A is called a *proper subset* of B , denoted $A \subset B$.

 Corresponding Author: t171d603@gunma-u.ac.jp

 <https://doi.org/10.48313/uda.vi.74>

 Licensee System Analytics. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0>).

Definition 0.3 (Empty Set). [4] The *empty set*, denoted by $\emptyset$, is the unique set containing no elements. Formally, for any set A , $\emptyset \subseteq A$.

Definition 0.4 (Universal Set). A *universal set*, denoted by U , is the set that contains all elements under consideration in a particular context. Every set discussed is assumed to be a subset of U .

Definition 0.5 (Rough Set Approximation). [5] Let X be a nonempty universe of discourse, and let $R \subseteq X \times X$ be an equivalence relation (also called an indiscernibility relation) on X . The relation R partitions X into disjoint equivalence classes, denoted by $[x]_R$ for each $x \in X$, where

$$[x]_R = \{y \in X \mid (x, y) \in R\}.$$

For any subset $U \subseteq X$, the *lower approximation* $\underline{U}$ and the *upper approximation* $\overline{U}$ are defined by:

(1) *Lower Approximation*:

$$\underline{U} = \{x \in X \mid [x]_R \subseteq U\}.$$

This set contains all elements whose entire equivalence class is contained within U ; these elements *definitely* belong to U .

(2) *Upper Approximation*:

$$\overline{U} = \{x \in X \mid [x]_R \cap U \neq \emptyset\}.$$

This set contains all elements whose equivalence class has a nonempty intersection with U ; these elements *possibly* belong to U .

Thus, the pair $(\underline{U}, \overline{U})$ forms the rough set representation of U , satisfying

$$\underline{U} \subseteq U \subseteq \overline{U}.$$

The *HyperRough Set* extends rough set theory by incorporating multiple attributes. Its formal definition is given below [6–8, 8–11].

Definition 0.6 (HyperRough Set). [6, 11–13] Let X be a nonempty finite universe, and let $T_1, T_2, \dots, T_n$ be n distinct attributes with corresponding domains $J_1, J_2, \dots, J_n$. Define the Cartesian product

$$J = J_1 \times J_2 \times \dots \times J_n.$$

Let $R \subseteq X \times X$ be an equivalence relation on X , with $[x]_R$ denoting the equivalence class of x . A *HyperRough Set* over X is a pair (F, J) , where:

- $F : J \rightarrow \mathcal{P}(X)$ is a mapping that assigns to each attribute value combination $a = (a_1, a_2, \dots, a_n) \in J$ a subset $F(a) \subseteq X$.
- For each $a \in J$, the rough set approximations of $F(a)$ are defined as

$$\underline{F(a)} = \{x \in X \mid [x]_R \subseteq F(a)\}, \quad \overline{F(a)} = \{x \in X \mid [x]_R \cap F(a) \neq \emptyset\}.$$

Here, $\underline{F(a)}$ comprises all elements whose equivalence classes are completely contained within $F(a)$, while $\overline{F(a)}$ contains elements whose equivalence classes intersect $F(a)$. Additionally, the following properties hold for all $a \in J$:

- $\underline{F(a)} \subseteq \overline{F(a)}$.
- If $F(a) = \emptyset$, then $\underline{F(a)} = \overline{F(a)} = \emptyset$.
- If $F(a) = X$, then $\underline{F(a)} = \overline{F(a)} = X$.

Example 0.7 (E-commerce Customer Segmentation). Let $X = \{c_1, \dots, c_6\}$ be a set of six customers. We consider three attributes:

$$\begin{aligned} T_1 &= \text{Region}, & J_1 &= \{\text{North}, \text{South}\}, \\ T_2 &= \text{AgeGroup}, & J_2 &= \{\text{Young}, \text{Adult}\}, \\ T_3 &= \text{Device}, & J_3 &= \{\text{Mobile}, \text{Desktop}\}. \end{aligned}$$

Thus the Cartesian product

$$J = J_1 \times J_2 \times J_3$$

has eight possible combinations. Define an equivalence relation R on X by

$$(c_i, c_j) \in R \iff c_i \text{ and } c_j \text{ share the same Region and AgeGroup.}$$

The resulting equivalence classes are

$$[c_1]_R = \{c_1, c_2\}, \quad [c_3]_R = \{c_3\}, \quad [c_4]_R = \{c_4, c_5\}, \quad [c_6]_R = \{c_6\}.$$

Define the mapping $F: J \rightarrow \mathcal{P}(X)$ on two illustrative attribute-value combinations:

$$F(\text{North, Young, Mobile}) = \{c_1, c_2\},$$

$$F(\text{North, Young, Desktop}) = \{c_2\}.$$

Then the rough approximations are computed as follows:

$$\underline{F}(\text{North, Young, Mobile}) = \{x \in X \mid [x]_R \subseteq \{c_1, c_2\}\} = \{c_1, c_2\},$$

$$\overline{F}(\text{North, Young, Mobile}) = \{x \in X \mid [x]_R \cap \{c_1, c_2\} \neq \emptyset\} = \{c_1, c_2\},$$

and for the second combination:

$$\underline{F}(\text{North, Young, Desktop}) = \{x \in X \mid [x]_R \subseteq \{c_2\}\} = \emptyset,$$

$$\overline{F}(\text{North, Young, Desktop}) = \{x \in X \mid [x]_R \cap \{c_2\} \neq \emptyset\} = \{c_1, c_2\}.$$

This example shows how the lower approximation contains exactly those customers whose entire equivalence class lies inside $F(a)$, whereas the upper approximation includes all customers whose class has any overlap with $F(a)$.

An n -SuperHyperRough Set generalizes rough sets by using power sets of attribute values to produce nuanced approximations under uncertainty. The definition of n -SuperHyperRough Sets is described as follows.

Definition 0.8 (n -SuperHyperRough Set). [11, 14] Let X be a nonempty finite universe, and let $T_1, T_2, \dots, T_n$ be n distinct attributes with respective domains $J_1, J_2, \dots, J_n$. For each attribute T_i , let $\mathcal{P}(J_i)$ denote its power set. Define the set of all possible attribute value combinations as

$$J = \mathcal{P}(J_1) \times \mathcal{P}(J_2) \times \dots \times \mathcal{P}(J_n).$$

Let $R \subseteq X \times X$ be an equivalence relation on X . An n -SuperHyperRough Set over X is a pair (F, J) , where:

- $F: J \rightarrow \mathcal{P}(X)$ is a mapping that assigns to each attribute value combination $A = (A_1, A_2, \dots, A_n) \in J$ (with $A_i \subseteq J_i$ for all i) a subset $F(A) \subseteq X$.
- For each $A \in J$, the lower and upper approximations are defined as

$$\underline{F}(A) = \{x \in X \mid [x]_R \subseteq F(A)\}, \quad \overline{F}(A) = \{x \in X \mid [x]_R \cap F(A) \neq \emptyset\}.$$

Thus, $\underline{F}(A)$ consists of all elements whose equivalence classes are entirely contained in $F(A)$, and $\overline{F}(A)$ includes those elements whose equivalence classes intersect $F(A)$. The following properties hold for all $A \in J$:

- $\underline{F}(A) \subseteq \overline{F}(A)$.
- If $F(A) = \emptyset$, then $\underline{F}(A) = \overline{F}(A) = \emptyset$.
- If $F(A) = X$, then $\underline{F}(A) = \overline{F}(A) = X$.
- For any $A, B \in J$,

$$\underline{F}(A \cap B) \subseteq \underline{F}(A) \cap \underline{F}(B), \quad \overline{F}(A \cup B) \supseteq \overline{F}(A) \cup \overline{F}(B).$$

Example 0.9 (Online Retail Customer Analysis with 2-SuperHyperRough Set). Let $X = \{c_1, c_2, c_3, c_4, c_5, c_6\}$ be six customers of an e-commerce platform. We consider two attributes:

$$T_1 = \text{Category}, \quad J_1 = \{\text{Electronics, Clothing}\}, \quad T_2 = \text{Channel}, \quad J_2 = \{\text{Online, InStore}\}.$$

The super attribute space is

$$J^* = \mathcal{P}(J_1) \times \mathcal{P}(J_2),$$

whose elements are pairs (A_1, A_2) with $A_1 \subseteq J_1$ and $A_2 \subseteq J_2$.

Define an equivalence relation R on X by

$$(c_i, c_j) \in R \iff c_i \text{ and } c_j \text{ share the same membership tier.}$$

Suppose the membership tiers partition X into

$$[c_1]_R = \{c_1, c_2\}, \quad [c_3]_R = \{c_3, c_4\}, \quad [c_5]_R = \{c_5, c_6\}.$$

Let $F^* : J^* \rightarrow \mathcal{P}(X)$ be defined for two illustrative super-attributes:

$$F^*({\text{Electronics}}, {\text{Online}}) = \{c_1, c_2, c_3\}, \quad F^*({\text{Electronics}}, {\text{Clothing}}, {\text{Online}}) = \{c_2, c_3, c_5\}.$$

For the first combination $A = ({\text{Electronics}}, {\text{Online}})$, the lower and upper approximations are

$$\begin{aligned} F^*(A) &= \{x \in X \mid [x]_R \subseteq \{c_1, c_2, c_3\}\} = \{c_1, c_2\}, \\ \overline{F^*(A)} &= \{x \in X \mid [x]_R \cap \{c_1, c_2, c_3\} \neq \emptyset\} = \{c_1, c_2, c_3, c_4\}. \end{aligned}$$

For the second combination $B = ({\text{Electronics}}, {\text{Clothing}}, {\text{Online}})$, we have

$$\begin{aligned} F^*(B) &= \{x \in X \mid [x]_R \subseteq \{c_2, c_3, c_5\}\} = \emptyset, \\ \overline{F^*(B)} &= \{x \in X \mid [x]_R \cap \{c_2, c_3, c_5\} \neq \emptyset\} = \{c_1, c_2, c_3, c_4, c_5, c_6\}. \end{aligned}$$

This example demonstrates how an n -SuperHyperRough Set uses power-set combinations of attributes to build nuanced lower and upper approximations in a real-world customer segmentation task.

1.2|Weighted Rough Sett

We present a detailed formulation of the Weighted Rough Set model (cf. [15–18]). This model extends classical rough set theory by quantifying the importance of each attribute through numerical weights, thereby allowing a precise and nuanced evaluation of decision systems.

Definition 0.10 (Information System). [15] An *information system* is defined as a quadruple

$$S = (U, A, \{V_a\}_{a \in A}, \{f_a\}_{a \in A}),$$

where:

- U is a finite, nonempty set of objects (e.g., patients, loan applicants, or regions).
- A is a finite set of attributes.
- For each attribute $a \in A$, V_a is a nonempty set called the domain of a (e.g., $V_{\text{Fever}} = \{\text{High}, \text{Normal}\}$).
- For each attribute $a \in A$, $f_a : U \rightarrow V_a$ is a function that assigns to each object $x \in U$ a value $f_a(x) \in V_a$.

Often, the set A is partitioned into condition attributes C and decision attributes D , and the information system is then denoted by

$$S = (U, C \cup \{D\}).$$

For example, in a medical diagnosis system, C might include symptoms such as Fever, Cough, and Fatigue, while D represents the diagnosis (e.g., flu or non-flu).

Definition 0.11 (Weighted Rough Set). [15] Let $S = (U, C \cup \{D\})$ be an information system as defined above. The indiscernibility relation R_C on U is given by

$$(x, y) \in R_C \iff \forall a \in C, f_a(x) = f_a(y),$$

which partitions U into equivalence classes $[x]_C$ containing objects with identical values for all condition attributes. For any subset $X \subseteq U$, the classical rough set lower approximation is

$$\underline{X} = \{x \in U \mid [x]_C \subseteq X\}.$$

For any subset $B \subseteq C$, the *dependency degree* of the decision attribute D on B is defined by

$$\gamma_B = \frac{|\text{POS}_B(D)|}{|U|},$$

where the positive region $\text{POS}_B(D)$ is the union of all equivalence classes in U that are completely contained in a single decision class.

The *significance* of an attribute $a \in B$ is measured by

$$\theta(a) = \gamma_B - \gamma_{B \setminus \{a\}},$$

which quantifies the reduction in dependency when a is removed from B . The *weighted importance* (or weight) of a is then computed as

$$w(a) = \frac{\theta(a)}{\sum_{b \in B} \theta(b)},$$

ensuring that $w(a) \geq 0$ and

$$\sum_{a \in B} w(a) = 1.$$

Furthermore, for any subset $X \subseteq U$ and a given threshold α (with $0 \leq \alpha \leq 1$), the *weighted lower approximation* of X with respect to B is defined as

$$\underline{B}_w(X) = \left\{ x \in U \mid \sum_{a \in B} w(a) \cdot I([x]_B \subseteq X) \geq \alpha \right\},$$

where the indicator function $I(\cdot)$ is defined by

$$I(\varphi) = \begin{cases} 1, & \text{if } \varphi \text{ is true,} \\ 0, & \text{if } \varphi \text{ is false.} \end{cases}$$

This weighted approach refines the approximation by considering the relative contribution of each attribute in classifying the objects.

Example 0.12 (Medical Diagnosis with Numerical Details). Medical diagnosis is the process of identifying diseases or conditions in patients by analyzing symptoms, medical history, and diagnostic test results (cf. [19, 20]). Consider a medical diagnosis system with:

- $U = \{p_1, p_2, p_3, p_4\}$, representing 4 patients.
- Condition attributes $C = \{\text{Fever, Cough, Fatigue}\}$, with domains:
 $V_{\text{Fever}} = \{\text{High, Normal}\}$, $V_{\text{Cough}} = \{\text{Present, Absent}\}$, $V_{\text{Fatigue}} = \{\text{Severe, Mild, None}\}$.
- Decision attribute D indicates the diagnosis (Flu or Non-flu).

Suppose the recorded data is as follows:

Patient	Fever	Cough	Fatigue	Diagnosis
p_1	High	Present	Severe	Flu
p_2	High	Absent	Mild	Flu
p_3	Normal	Present	None	Non-flu
p_4	High	Present	Mild	Flu

Assume that the dependency degree using all condition attributes is $\gamma_C = 0.75$ and that removing Fever gives $\gamma_{C \setminus \{\text{Fever}\}} = 0.50$. Then, the significance of Fever is

$$\theta(\text{Fever}) = 0.75 - 0.50 = 0.25.$$

Similarly, suppose $\theta(\text{Cough}) = 0.15$ and $\theta(\text{Fatigue}) = 0.10$. The weights are then computed as:

$$w(\text{Fever}) = \frac{0.25}{0.25 + 0.15 + 0.10} = 0.50, \quad w(\text{Cough}) = \frac{0.15}{0.50} = 0.30, \quad w(\text{Fatigue}) = \frac{0.10}{0.50} = 0.20.$$

For the set X of patients diagnosed with Flu, say $X = \{p_1, p_2, p_4\}$, and a threshold $\alpha = 0.5$, the weighted lower approximation $\underline{B}_w(X)$ includes any patient p for whom the weighted sum

$$w(\text{Fever}) \cdot I([p]_C \subseteq X) + w(\text{Cough}) \cdot I([p]_C \subseteq X) + w(\text{Fatigue}) \cdot I([p]_C \subseteq X)$$

is at least 0.5. This detailed calculation clearly demonstrates how the weighted rough set model uses attribute significance to filter and classify patients.

Example 0.13 (Credit Scoring with Detailed Computation). Credit scoring is the process of evaluating an individual's creditworthiness based on financial attributes and repayment history using statistical models (cf. [21,22]). Consider a credit scoring system where:

- $U = \{a_1, a_2, a_3, a_4, a_5\}$ represents 5 loan applicants.
- Condition attributes $C = \{\text{Income, Debt Level, Employment Status}\}$ with domains:
 $V_{\text{Income}} = \{\text{Low, Medium, High}\}$, $V_{\text{Debt Level}} = \{\text{Low, High}\}$, $V_{\text{Employment Status}} = \{\text{Unemployed, Employed}\}$.
- D indicates whether an applicant is creditworthy.

Assume that using all condition attributes yields $\gamma_C = 0.80$ and that removing Income results in $\gamma_{C \setminus \{\text{Income}\}} = 0.60$, so that $\theta(\text{Income}) = 0.20$. Further, let $\theta(\text{Debt Level}) = 0.10$ and $\theta(\text{Employment Status}) = 0.10$. The weights are then:

$$w(\text{Income}) = \frac{0.20}{0.40} = 0.50, \quad w(\text{Debt Level}) = \frac{0.10}{0.40} = 0.25, \quad w(\text{Employment Status}) = \frac{0.10}{0.40} = 0.25.$$

For a subset X of creditworthy applicants and a threshold $\alpha = 0.6$, the model computes the weighted lower approximation by verifying that, for each applicant a , the combined weighted evidence from all attributes is strong enough (i.e., at least 60%) to support inclusion in X . This concrete computation ensures that the model effectively distinguishes creditworthy applicants based on weighted financial indicators.

Example 0.14 (Environmental Monitoring with Specific Details). Environmental Monitoring is the systematic collection and analysis of data to assess environmental conditions, detect changes, and inform policy decisions [23–25]. Suppose an environmental monitoring system is designed with:

- $U = \{r_1, r_2, r_3, r_4, r_5, r_6\}$ representing 6 geographical regions.
- Condition attributes $C = \{\text{Air Quality Index (AQI), Water Quality, Vegetation Cover}\}$ with domains:
 $V_{\text{AQI}} = \{\text{Good, Moderate, Poor}\}$, $V_{\text{Water Quality}} = \{\text{Clean, Polluted}\}$, $V_{\text{Vegetation Cover}} = \{\text{Dense, Sparse}\}$.
- D indicates an overall environmental risk level (e.g., low, moderate, or high risk).

Suppose the dependency analysis yields $\theta(\text{AQI}) = 0.30$, $\theta(\text{Water Quality}) = 0.20$, and $\theta(\text{Vegetation Cover}) = 0.10$. The corresponding weights are:

$$w(\text{AQI}) = \frac{0.30}{0.60} = 0.50, \quad w(\text{Water Quality}) = \frac{0.20}{0.60} \approx 0.33, \quad w(\text{Vegetation Cover}) = \frac{0.10}{0.60} \approx 0.17.$$

For a set X of regions classified as low risk, the weighted lower approximation $\underline{B}_w(X)$ will include a region r if the weighted sum of the indicator functions (assessing whether the equivalence class $[r]_C$ is entirely contained in X) is at least a threshold (e.g., $\alpha = 0.5$). This detailed example illustrates how the weighted rough set model can be applied to environmental data for robust risk assessment.

1.3|Sequential Rough Set

We define the Sequential Rough Set (SRS) Model in a detailed and concrete manner. The SRS model refines Pawlak's classical rough set by sequentially partitioning the universe into regions using relevance degrees computed from equivalence classes [26, 27]. This approach improves fault-tolerance and concept approximation without additional assumptions.

Definition 0.15 (Sequential Rough Set Model). [27] Let

$$I = (U, A, \{V_a\}_{a \in A}, r)$$

be an information table where:

- U is a finite, nonempty set of objects.
- A is a finite set of attributes.
- For each $a \in A$, V_a is the domain of a .
- $r : U \rightarrow \prod_{a \in A} V_a$ is the function that assigns values to objects.

Let $B \subseteq A$ be a chosen set of condition attributes, and define the equivalence relation R_B on U by

$$x R_B y \iff \forall a \in B, f_a(x) = f_a(y).$$

Denote by $B(x)$ the equivalence class of x under R_B .

For a target concept $X \subseteq U$ (e.g., the set of objects with a desired property), define the *relevance degree* of x to X as

$$E_B^X(x) = \frac{|B(x) \cap X|}{|X|}, \quad \text{if } X \neq \emptyset,$$

and the relevance degree of x to the complement $\bar{X} = U \setminus X$ as

$$E_B^{\bar{X}}(x) = \frac{|B(x) \cap \bar{X}|}{|\bar{X}|}, \quad \text{if } \bar{X} \neq \emptyset.$$

The Sequential Rough Set Model is constructed by a stepwise (sequential) partitioning process. Initialize

$$X_0 = X \quad \text{and} \quad U_0 = U.$$

For a fixed positive integer n , choose a sequence of threshold pairs

$$(\alpha_1, \beta_1), (\alpha_2, \beta_2), \dots, (\alpha_n, \beta_n) \quad \text{with } 0 \leq \alpha_i, \beta_i \leq 1.$$

For each iteration $i = 1, 2, \dots, n$, define the following subsets on the current universe U_{i-1} :

$$\begin{aligned} P_i(X_{i-1}) &= \left\{ x \in U_{i-1} \mid E_B^{X_{i-1}}(x) > \alpha_i \text{ and } E_B^{\bar{X}_{i-1}}(x) \leq \beta_i \right\}, \\ N_i(X_{i-1}) &= \left\{ x \in U_{i-1} \mid E_B^{X_{i-1}}(x) \leq \alpha_i \text{ and } E_B^{\bar{X}_{i-1}}(x) > \beta_i \right\}, \\ B_{1,i}(X_{i-1}) &= \left\{ x \in U_{i-1} \mid E_B^{X_{i-1}}(x) > \alpha_i \text{ and } E_B^{\bar{X}_{i-1}}(x) > \beta_i \right\}, \\ B_{2,i}(X_{i-1}) &= \left\{ x \in U_{i-1} \mid E_B^{X_{i-1}}(x) \leq \alpha_i \text{ and } E_B^{\bar{X}_{i-1}}(x) \leq \beta_i \right\}. \end{aligned}$$

Then set

$$U_i = B_{2,i}(X_{i-1}) \quad \text{and} \quad X_i = X \cap U_i.$$

The process continues until for some n , we have $B_{2,n}(X_{n-1}) = \emptyset$. Finally, the overall regions are defined by

$$\begin{aligned} \text{POS}_B(X) &= \bigcup_{i=1}^n P_i(X_{i-1}), \quad \text{NEG}_B(X) = \bigcup_{i=1}^n N_i(X_{i-1}), \\ \text{BND}_B(X) &= \bigcup_{i=1}^n B_{1,i}(X_{i-1}). \end{aligned}$$

The model, denoted by

$$\Phi = (I, \text{POS}_B(X), \text{NEG}_B(X), \text{BND}_B(X)),$$

is called the *Sequential Rough Set Model* of X induced by B . Notice that the classical positive region $\{x \in U \mid B(x) \subseteq X\}$ is contained in $\text{POS}_B(X)$, ensuring that this model is a conservative extension of Pawlak's classical rough set.

Example 0.16 (Sequential Rough Set Model in Patient Classification). Consider a simplified medical diagnosis scenario with five patients:

$$U = \{u_1, u_2, u_3, u_4, u_5\}.$$

Let there be two symptoms (attributes):

- a_1 : *Fever* with domain $V_{a_1} = \{\text{High}, \text{Normal}\}$,
- a_2 : *Cough* with domain $V_{a_2} = \{\text{Present}, \text{Absent}\}$.

Assume the following data (for simplicity, only symptoms are listed and the decision attribute is given separately):

Patient	Fever	Cough	Diagnosis
u_1	High	Present	Disease
u_2	High	Absent	Disease
u_3	Normal	Present	No Disease
u_4	High	Present	Disease
u_5	Normal	Absent	No Disease

Let the condition attribute set be $B = \{a_1, a_2\}$. The equivalence classes induced by R_B are:

$$B(u_1) = \{u_1, u_4\}, \quad B(u_2) = \{u_2\}, \quad B(u_3) = \{u_3\}, \quad B(u_5) = \{u_5\}.$$

Define the target concept X as the set of patients with disease:

$$X = \{u_1, u_2, u_4\}.$$

Now, compute the relevance degrees for each patient:

$$E_B^X(u_1) = \frac{|B(u_1) \cap X|}{|X|} = \frac{|\{u_1, u_4\}|}{3} = \frac{2}{3} \approx 0.67,$$

$$E_B^X(u_2) = \frac{|\{u_2\}|}{3} = \frac{1}{3} \approx 0.33, \quad E_B^X(u_4) \approx 0.67.$$

For patients not in X (i.e. u_3 and u_5):

$$E_B^{\bar{X}}(u_3) = \frac{|B(u_3) \cap \{u_3, u_5\}|}{2} = \frac{1}{2} = 0.50, \quad E_B^{\bar{X}}(u_5) = 0.50.$$

Also, for u_1 and u_4 , since $B(u_1)$ and $B(u_4)$ contain only patients from X ,

$$E_B^{\bar{X}}(u_1) = E_B^{\bar{X}}(u_4) = 0.$$

Choose thresholds for the first sequential trisection: let $\alpha_1 = 0.5$ and $\beta_1 = 0.2$.

Now, partition the initial universe $U_0 = U$ as follows:

- $P_1(X_0) = \{x \in U \mid E_B^X(x) > 0.5 \text{ and } E_B^{\bar{X}}(x) \leq 0.2\}$. Here, u_1 and u_4 qualify because $E_B^X(u_1) = E_B^X(u_4) = 0.67 > 0.5$ and $E_B^{\bar{X}}(u_1) = E_B^{\bar{X}}(u_4) = 0 \leq 0.2$.
- $N_1(X_0) = \{x \in U \mid E_B^X(x) \leq 0.5 \text{ and } E_B^{\bar{X}}(x) > 0.2\}$. Here, u_3 and u_5 qualify since $E_B^X(u_3) = E_B^X(u_5) = 0$ (as they are not in X) and $E_B^{\bar{X}}(u_3) = E_B^{\bar{X}}(u_5) = 0.50 > 0.2$.
- The remaining patients, u_2 , satisfy $E_B^X(u_2) = 0.33 \leq 0.5$ and $E_B^{\bar{X}}(u_2) = 0 \leq 0.2$, so they belong to the boundary-II region:

$$B_{2,1}(X_0) = \{u_2\}.$$

- In this iteration, assume that $B_{1,1}(X_0) = \emptyset$ (i.e., no object has both high relevance to X and high relevance to $\bar{X}$).

Set $U_1 = B_{2,1}(X_0) = \{u_2\}$ and $X_1 = X \cap U_1 = \{u_2\}$. Since U_1 contains only one object, further partitioning would yield empty regions. Thus, the sequential process terminates with:

$$\text{POS}_B(X) = P_1(X_0) \cup (\text{any additional positive region from further iterations}) = \{u_1, u_4\},$$

$$\text{NEG}_B(X) = N_1(X_0) = \{u_3, u_5\},$$

$$\text{BND}_B(X) = B_{1,1}(X_0) \cup (\text{any additional boundary regions}) = \emptyset.$$

In this example, the SRS model clearly separates the patients: u_1 and u_4 are decisively classified as having the disease, u_3 and u_5 are classified as not having the disease, and u_2 is isolated in a residual (boundary) set that is resolved in the subsequent iteration. This fine-grained process improves the fault-tolerance and precision in classifying patients compared to the classical rough set model.

2|Results of This Paper

This section presents the results obtained in this paper.

2.1|Weighted Hyperrough Set

We now present a detailed and concrete formulation of the Weighted Hyperrough Set, which extends the classical weighted rough set by incorporating multiple attributes with non-uniform weights. This model is designed to capture the varying degrees of importance of each attribute in a multi-attribute decision system.

Definition 0.17 (Weighted Hyperrough Set). Let X be a finite universe of objects and let $R \subseteq X \times X$ be an equivalence relation on X . For each $x \in X$, denote by $[x]_R$ the equivalence class of x under R . Suppose that there are n distinct attributes $T_1, T_2, \dots, T_n$ with corresponding nonempty domains $J_1, J_2, \dots, J_n$. Define the Cartesian product of the domains as

$$J = J_1 \times J_2 \times \dots \times J_n,$$

so that each element $a \in J$ is an n -tuple $a = (a_1, a_2, \dots, a_n)$ where $a_i \in J_i$.

Let $w : \{T_1, T_2, \dots, T_n\} \rightarrow [0, 1]$ be a weight function satisfying

$$\sum_{i=1}^n w(T_i) = 1,$$

which reflects the relative importance of each attribute in decision making.

A *Weighted Hyperrough Set* is defined as a pair (F_w, J) where the mapping

$$F_w : J \rightarrow \mathcal{P}(X)$$

associates to each attribute value combination $a = (a_1, a_2, \dots, a_n)$ a subset $F_w(a) \subseteq X$ representing the objects that satisfy the condition described by a .

Furthermore, for fixed thresholds $0 \leq \alpha, \beta \leq 1$, the *weighted lower approximation* and the *weighted upper approximation* of $F_w(a)$ are defined, respectively, by:

$$\begin{aligned} \underline{F_w}(a) &= \left\{ x \in X \mid \sum_{i=1}^n w(T_i) \cdot I([x]_R \subseteq F_w(a)) \geq \alpha \right\}, \\ \overline{F_w}(a) &= \left\{ x \in X \mid \sum_{i=1}^n w(T_i) \cdot I([x]_R \cap F_w(a) \neq \emptyset) \geq \beta \right\}, \end{aligned}$$

where $I(\cdot)$ is the indicator function defined by

$$I(\varphi) = \begin{cases} 1, & \text{if } \varphi \text{ is true,} \\ 0, & \text{if } \varphi \text{ is false.} \end{cases}$$

This model reduces to the classical Weighted Rough Set when $n = 1$ and to the standard Hyperrough Set when the weight function is uniform (i.e., $w(T_i) = 1/n$ for all i).

Example 0.18 (Weighted Hyperrough Set in Medical Diagnosis). Consider a medical diagnosis system in which:

- The universe X is a finite set of patients.
- Two symptoms are considered:
 - $T_1 = \text{Fever}$ with domain $J_1 = \{\text{High}, \text{Normal}\}$,
 - $T_2 = \text{Cough}$ with domain $J_2 = \{\text{Present}, \text{Absent}\}$.
- The Cartesian product $J = J_1 \times J_2$ yields combinations such as (High, Present).
- A weight function is assigned as $w(T_1) = 0.6$ (indicating that fever is more important) and $w(T_2) = 0.4$.

Define the mapping F_w by

$$F_w(\text{High}, \text{Present}) = \{\text{patients diagnosed with severe flu}\}.$$

Then, for a threshold $\alpha = 0.5$, the weighted lower approximation

$$\underline{F}_w(\text{High}, \text{Present})$$

consists of those patients x in X for which the weighted sum

$$0.6 I([x]_R \subseteq F_w(\text{High}, \text{Present})) + 0.4 I([x]_R \subseteq F_w(\text{High}, \text{Present}))$$

meets or exceeds 0.5. In practice, this means that a patient is included in the approximation if the equivalence class $[x]_R$ (i.e., the set of patients with identical symptom profiles) is entirely contained in the set of patients diagnosed with severe flu, and the combined weighted evidence from both symptoms is strong enough (as measured by the threshold).

Theorem 0.19 (Generalization via Weighted Hyperrough Set). *The Weighted Hyperrough Set model generalizes both the classical Weighted Rough Set and the standard Hyperrough Set. More precisely:*

- (1) When $n = 1$, the Cartesian product J reduces to J_1 and the weight function satisfies $w(T_1) = 1$. Consequently, the weighted approximations become

$$\underline{F}_w(a) = \{x \in X \mid I([x]_R \subseteq F_w(a)) \geq \alpha\}, \quad \overline{F}_w(a) = \{x \in X \mid I([x]_R \cap F_w(a) \neq \emptyset) \geq \beta\},$$

which is exactly the definition of the classical Weighted Rough Set.

- (2) When the weight function is uniform (i.e., $w(T_i) = 1/n$ for all i), the weighted sum reduces to the arithmetic mean of the indicator functions, and the thresholds α and β reflect the required proportion of attributes satisfying the conditions. This recovers the standard Hyperrough Set definitions.

Proof: (1) If $n = 1$, then $J = J_1$ and the only attribute T_1 is assigned weight 1. Thus, the lower approximation becomes

$$\underline{F}_w(a) = \{x \in X \mid I([x]_R \subseteq F_w(a)) \geq \alpha\},$$

and similarly for the upper approximation. This is identical to the classical weighted rough set framework.

- (2) If $w(T_i) = 1/n$ for all i , then the weighted sum in the lower approximation is

$$\sum_{i=1}^n \frac{1}{n} I([x]_R \subseteq F_w(a)) = \frac{1}{n} \sum_{i=1}^n I([x]_R \subseteq F_w(a)).$$

The condition for inclusion in the approximation is now that the proportion of attributes for which $[x]_R$ is entirely contained in $F_w(a)$ meets or exceeds the threshold α . This is exactly the criterion used in the standard Hyperrough Set. Thus, the Weighted Hyperrough Set indeed generalizes both models. $\square$

2.2|Weighted SuperHyperrough Set

We now describe the Weighted n -Superhyperrough Set, a further generalization that combines the idea of weighting with the power-set based attribute representation. This model accommodates more nuanced attribute combinations and uncertainties.

Definition 0.20 (Weighted n -Superhyperrough Set). Let X be a finite universe and let $R \subseteq X \times X$ be an equivalence relation on X . Consider n distinct attributes $T_1, T_2, \dots, T_n$ with corresponding domains $J_1, J_2, \dots, J_n$. For each attribute T_i , let $\mathcal{P}(J_i)$ denote the power set of J_i . Define the set of all possible attribute value combinations as

$$J^* = \mathcal{P}(J_1) \times \mathcal{P}(J_2) \times \dots \times \mathcal{P}(J_n).$$

Let $w : \{T_1, T_2, \dots, T_n\} \rightarrow [0, 1]$ be a weight function with

$$\sum_{i=1}^n w(T_i) = 1.$$

A Weighted n -Superhyperrough Set is a pair (F_w^*, J^*) where

$$F_w^* : J^* \rightarrow \mathcal{P}(X)$$

assigns to each attribute combination $A = (A_1, A_2, \dots, A_n) \in J^*$ (with $A_i \subseteq J_i$ for every i) a subset $F_w^*(A) \subseteq X$. For fixed thresholds $0 \leq \alpha, \beta \leq 1$, the weighted lower and upper approximations of $F_w^*(A)$ are given by:

$$\underline{F}_w^*(A) = \left\{ x \in X \mid \sum_{i=1}^n w(T_i) \cdot I([x]_R \subseteq F_w^*(A)) \geq \alpha \right\},$$

$$\overline{F}_w^*(A) = \left\{ x \in X \mid \sum_{i=1}^n w(T_i) \cdot I([x]_R \cap F_w^*(A) \neq \emptyset) \geq \beta \right\}.$$

This construction generalizes the Weighted Hyperrough Set when each J_i is crisp (i.e., when $\mathcal{P}(J_i)$ essentially contains only singleton subsets) and generalizes the classical n -Superhyperrough Set when the weight function is uniform.

Example 0.21 (Weighted n -Superhyperrough Set in Housing Selection). Housing selection is the decision-making process of choosing a residence based on criteria like location, price, style, and personal preferences (cf. [28–31]). Consider a decision-making scenario for housing selection where:

- X is a finite set of houses.
- Two attributes are used:
 - $T_1 = \text{Location}$ with domain $J_1 = \{\text{Urban}, \text{Suburban}\}$,
 - $T_2 = \text{Style}$ with domain $J_2 = \{\text{Modern}, \text{Classic}\}$.
- The power set $\mathcal{P}(J_1)$ includes $\emptyset$, $\{\text{Urban}\}$, $\{\text{Suburban}\}$, and $\{\text{Urban}, \text{Suburban}\}$; similarly for $\mathcal{P}(J_2)$. Thus, the set J^* consists of all combinations such as $(\{\text{Urban}\}, \{\text{Modern}\})$.
- The weight function is chosen uniformly as $w(T_1) = 0.5$ and $w(T_2) = 0.5$.

Define the mapping F_w^* by

$$F_w^*(\{\text{Urban}\}, \{\text{Modern}\}) = \{\text{houses located in urban areas with modern design}\}.$$

For a threshold $\alpha = 0.5$, the weighted lower approximation

$$\underline{F}_w^*(\{\text{Urban}\}, \{\text{Modern}\})$$

comprises those houses $x \in X$ for which the weighted sum

$$0.5 I([x]_R \subseteq F_w^*(\{\text{Urban}\}, \{\text{Modern}\})) + 0.5 I([x]_R \subseteq F_w^*(\{\text{Urban}\}, \{\text{Modern}\}))$$

is at least 0.5. This ensures that only houses with a sufficiently strong overall alignment to the criteria are included, effectively filtering out ambiguous cases.

Theorem 0.22 (Generalization via Weighted n -Superhyperrough Set). *The Weighted n -Superhyperrough Set model generalizes both the Weighted Hyperrough Set and the classical n -Superhyperrough Set. Specifically:*

- (1) *If each attribute domain J_i is crisp (that is, if each $\mathcal{P}(J_i)$ consists only of singleton subsets and the empty set), then J^* is equivalent to the Cartesian product $J = J_1 \times J_2 \times \dots \times J_n$ and the Weighted n -Superhyperrough Set reduces to the Weighted Hyperrough Set.*
- (2) *If the weight function is uniform (i.e., $w(T_i) = \frac{1}{n}$ for every i), then the weighted approximations in the Weighted n -Superhyperrough Set coincide with those of the classical n -Superhyperrough Set.*

Proof: (1) Suppose that each domain J_i is crisp, meaning that every J_i consists solely of distinct, non-overlapping elements. In this case, the power set $\mathcal{P}(J_i)$ effectively contains only the singleton sets (ignoring the trivial empty set). Thus, the Cartesian product

$$J^* = \mathcal{P}(J_1) \times \mathcal{P}(J_2) \times \dots \times \mathcal{P}(J_n)$$

collapses to the standard Cartesian product $J_1 \times J_2 \times \dots \times J_n = J$. As a result, the mapping F_w^* is identical to F_w and the weighted approximations remain unchanged, thereby recovering the Weighted Hyperrough Set.

(2) Now assume that the weight function is uniform, so that $w(T_i) = \frac{1}{n}$ for all $i = 1, 2, \dots, n$. Then for any attribute combination $A \in J^*$, the weighted sum in the lower approximation becomes

$$\sum_{i=1}^n \frac{1}{n} I([x]_R \subseteq F_w^*(A)) = \frac{1}{n} \sum_{i=1}^n I([x]_R \subseteq F_w^*(A)).$$

This expression is identical to the corresponding term in the classical n -Superhyperrough Set, where the decision to include an object x is based on the proportion of attributes for which $[x]_R$ satisfies the condition. The same argument applies to the upper approximation. Therefore, with a uniform weight function, the Weighted n -Superhyperrough Set model recovers the classical model. $\square$

2.3|Sequential Hyperrough Set

This section presents the Sequential Hyperrough Set (SHRS) Model in a detailed and concrete manner. The SHRS model is designed to generalize both the classical Sequential Rough Set Model and the Hyperrough Set Model by incorporating a hyper attribute space and a sequential partitioning process.

Definition 0.23 (Sequential Hyperrough Set Model). Let

$$I = (U, A, \{V_a\}_{a \in A}, r)$$

be an information system, where U is a finite nonempty set of objects, A is a finite set of attributes, and for each $a \in A$, V_a is its domain with $r : U \rightarrow \prod_{a \in A} V_a$ assigning values to objects. Let $B \subseteq A$ be a set of condition attributes. Define the equivalence relation R_B on U by

$$x R_B y \iff \forall a \in B, f_a(x) = f_a(y),$$

and denote by $B(x)$ the equivalence class of x under R_B .

Assume there are n distinct attribute domains $J_1, J_2, \dots, J_n$ associated with features of interest, and define the hyper attribute space as their Cartesian product

$$J = J_1 \times J_2 \times \dots \times J_n.$$

Let $F : J \rightarrow \mathcal{P}(U)$ be a mapping such that for every $a = (a_1, a_2, \dots, a_n) \in J$, $F(a) \subseteq U$ represents the set of objects exhibiting the attribute combination a .

For a target concept $X \subseteq U$, define for each object $x \in U$ and for each $a \in J$ the hyperrough relevance degrees by

$$E_B^X(x, a) = \frac{|B(x) \cap F(a) \cap X|}{|X|}, \quad E_B^{\bar{X}}(x, a) = \frac{|B(x) \cap F(a) \cap (U \setminus X)|}{|U \setminus X|},$$

with the usual convention that if X (or $U \setminus X$) is empty, the corresponding relevance is defined as zero.

The Sequential Hyperrough Set Model is constructed by a sequential partitioning process. Initialize

$$X_0 = X \quad \text{and} \quad U_0 = U.$$

For a fixed number of iterations m , choose a sequence of threshold pairs

$$(\alpha_1, \beta_1), (\alpha_2, \beta_2), \dots, (\alpha_m, \beta_m)$$

with $0 \leq \alpha_i, \beta_i \leq 1$ for each i . For each iteration $i = 1, 2, \dots, m$, define the following subsets of the current universe U_{i-1} :

$$\begin{aligned} P_i(X_{i-1}) &= \left\{ x \in U_{i-1} \mid \exists a \in J \text{ such that } E_B^{X_{i-1}}(x, a) > \alpha_i \text{ and } E_B^{\bar{X}}(x, a) \leq \beta_i \right\}, \\ N_i(X_{i-1}) &= \left\{ x \in U_{i-1} \mid \exists a \in J \text{ such that } E_B^{X_{i-1}}(x, a) \leq \alpha_i \text{ and } E_B^{\bar{X}}(x, a) > \beta_i \right\}, \\ B_{1,i}(X_{i-1}) &= \left\{ x \in U_{i-1} \mid \exists a \in J \text{ such that } E_B^{X_{i-1}}(x, a) > \alpha_i \text{ and } E_B^{\bar{X}}(x, a) > \beta_i \right\}, \\ B_{2,i}(X_{i-1}) &= \left\{ x \in U_{i-1} \mid \forall a \in J, E_B^{X_{i-1}}(x, a) \leq \alpha_i \text{ and } E_B^{\bar{X}}(x, a) \leq \beta_i \right\}. \end{aligned}$$

Then update

$$U_i = B_{2,i}(X_{i-1}) \quad \text{and} \quad X_i = X \cap U_i.$$

The overall regions for X induced by B are then defined as

$$\begin{aligned} \text{POS}_B(X) &= \bigcup_{i=1}^m P_i(X_{i-1}), \quad \text{NEG}_B(X) = \bigcup_{i=1}^m N_i(X_{i-1}), \\ \text{BND}_B(X) &= \bigcup_{i=1}^m B_{1,i}(X_{i-1}). \end{aligned}$$

Thus, the Sequential Hyperrough Set Model is the tuple

$$\Phi_H = (I, \text{POS}_B(X), \text{NEG}_B(X), \text{BND}_B(X)).$$

This model generalizes the Sequential Rough Set Model (when J is a singleton) and the standard Hyperrough Set Model (when $m = 1$).

Theorem 0.24 (Generalization of Sequential Hyperrough Set). *The Sequential Hyperrough Set Model Φ_H generalizes:*

- (1) *The classical Sequential Rough Set Model, when $n = 1$ (i.e., when $J = J_1$ and the mapping F reduces to the identity on J_1);*
- (2) *The Hyperrough Set Model, when the sequential process is executed in a single iteration ($m = 1$).*

Proof: (1) If $n = 1$, then $J = J_1$ and the mapping $F : J_1 \rightarrow \mathcal{P}(U)$ partitions U based on a single attribute. The relevance degrees $E_B^X(x, a)$ then become equivalent to those defined in the Sequential Rough Set Model. Thus, Φ_H coincides with the classical Sequential Rough Set Model.

(2) If $m = 1$, the sequential process consists of one partitioning step using thresholds (α_1, β_1) over the hyper attribute space J . This single-step partitioning is exactly the standard Hyperrough Set Model. Therefore, Φ_H generalizes the Hyperrough Set Model. $\square$

Example 0.25 (Sequential Hyperrough Set in Housing Selection). Consider a housing selection scenario:

- Let U be a set of houses.
- Two attributes are considered:
 - $T_1 = \text{Location}$ with domain $J_1 = \{\text{Urban, Suburban}\}$,
 - $T_2 = \text{Style}$ with domain $J_2 = \{\text{Modern, Classic}\}$.
- The Cartesian product $J = J_1 \times J_2$ contains combinations such as (Urban, Modern) and (Suburban, Classic).
- Define $F((\text{Urban, Modern}))$ as the set of houses located in urban areas with modern design.
- Let the target concept X be the set of houses that are highly desirable to a buyer.

For each house $x \in U$ and for each attribute combination $a \in J$, compute the relevance degrees $E_B^X(x, a)$ and $E_B^{\bar{X}}(x, a)$. Suppose that for the first iteration, thresholds $\alpha_1 = 0.6$ and $\beta_1 = 0.3$ are chosen. Then, houses for which there exists an $a \in J$ with $E_B^X(x, a) > 0.6$ and $E_B^{\bar{X}}(x, a) \leq 0.3$ are placed in the positive region $P_1(X_0)$; ambiguous cases fall into the boundary region $B_{2,1}(X_0)$ and are further partitioned in subsequent iterations. This sequential refinement ultimately yields a detailed classification that enables a buyer to identify houses matching their preferences with higher accuracy.

2.4|Sequential Hyperrough Set

In this section, we define the Sequential n -Superhyperrough Set (SnSHRS) Model. This model extends the Sequential Hyperrough Set Model by allowing each attribute's domain to be represented by its power set. This additional flexibility enables a more granular handling of uncertainty and complex attribute combinations.

Definition 0.26 (Sequential n -Superhyperrough Set Model). Let

$$I = (U, A, \{V_a\}_{a \in A}, r)$$

be an information system and let $B \subseteq A$ be a set of condition attributes. For each attribute T_i (where $i = 1, 2, \dots, n$) with domain J_i , denote its power set by $\mathcal{P}(J_i)$. Define the super attribute space as

$$J^* = \mathcal{P}(J_1) \times \mathcal{P}(J_2) \times \dots \times \mathcal{P}(J_n).$$

Let $F^* : J^* \rightarrow \mathcal{P}(U)$ be a mapping that assigns to each combination $A = (A_1, A_2, \dots, A_n) \in J^*$ (with each $A_i \subseteq J_i$) a subset $F^*(A) \subseteq U$.

For a target concept $X \subseteq U$, define for each $x \in U$ and for each $A \in J^*$ the relevance degrees as

$$E_B^X(x, A) = \frac{|B(x) \cap F^*(A) \cap X|}{|X|}, \quad E_B^{\bar{X}}(x, A) = \frac{|B(x) \cap F^*(A) \cap (U \setminus X)|}{|U \setminus X|}.$$

Select a sequence of threshold pairs

$$(\alpha_1, \beta_1), (\alpha_2, \beta_2), \dots, (\alpha_m, \beta_m)$$

with $0 \leq \alpha_i, \beta_i \leq 1$. Initialize

$$X_0 = X \quad \text{and} \quad U_0 = U.$$

For each iteration $i = 1, \dots, m$, define the following subsets of U_{i-1} :

$$\begin{aligned} P_i^*(X_{i-1}) &= \left\{ x \in U_{i-1} \mid \exists A \in J^* \text{ such that } E_B^{X_{i-1}}(x, A) > \alpha_i \text{ and } E_B^{\bar{X}}(x, A) \leq \beta_i \right\}, \\ N_i^*(X_{i-1}) &= \left\{ x \in U_{i-1} \mid \exists A \in J^* \text{ such that } E_B^{X_{i-1}}(x, A) \leq \alpha_i \text{ and } E_B^{\bar{X}}(x, A) > \beta_i \right\}, \\ B_{1,i}^*(X_{i-1}) &= \left\{ x \in U_{i-1} \mid \exists A \in J^* \text{ such that } E_B^{X_{i-1}}(x, A) > \alpha_i \text{ and } E_B^{\bar{X}}(x, A) > \beta_i \right\}, \\ B_{2,i}^*(X_{i-1}) &= \left\{ x \in U_{i-1} \mid \forall A \in J^*, E_B^{X_{i-1}}(x, A) \leq \alpha_i \text{ and } E_B^{\bar{X}}(x, A) \leq \beta_i \right\}. \end{aligned}$$

Then update

$$U_i = B_{2,i}^*(X_{i-1}) \quad \text{and} \quad X_i = X \cap U_i.$$

The overall regions for X are defined as

$$\begin{aligned} \text{POS}_B(X) &= \bigcup_{i=1}^m P_i^*(X_{i-1}), \quad \text{NEG}_B(X) = \bigcup_{i=1}^m N_i^*(X_{i-1}), \\ \text{BND}_B(X) &= \bigcup_{i=1}^m B_{1,i}^*(X_{i-1}). \end{aligned}$$

Thus, the Sequential n -Superhyperrough Set Model is the tuple

$$\Phi_S = (I, \text{POS}_B(X), \text{NEG}_B(X), \text{BND}_B(X)).$$

This model generalizes the Sequential Hyperrough Set Model by allowing the representation of attribute domains via their power sets, thereby enabling a more detailed treatment of uncertainty and complex attribute interactions.

Theorem 0.27 (Generalization of Sequential n -Superhyperrough Set). *The Sequential n -Superhyperrough Set Model Φ_S generalizes:*

- (1) *The Sequential Hyperrough Set Model, when for each i the power set $\mathcal{P}(J_i)$ contains only singleton subsets (so that J^* is isomorphic to $J_1 \times \dots \times J_n$);*
- (2) *The classical n -Superhyperrough Set Model, when the sequential process consists of a single iteration (i.e., $m = 1$).*

Proof: (1) If every $\mathcal{P}(J_i)$ contains only singleton subsets, then every element $A \in J^*$ is of the form $(\{a_1\}, \{a_2\}, \dots, \{a_n\})$, and J^* becomes naturally isomorphic to the Cartesian product $J_1 \times J_2 \times \dots \times J_n$. In this case, the mapping F^* reduces to the mapping F in the Sequential Hyperrough Set Model, and the computed relevance degrees $E_B^X(x, A)$ coincide with those in the SHRS model. Hence, Φ_S is equivalent to the Sequential Hyperrough Set Model.

(2) If the sequential process is executed in only one iteration (i.e., $m = 1$), then the entire universe U is partitioned in a single step using the thresholds (α_1, β_1) applied over the super attribute space J^* . This procedure exactly corresponds to the classical n -Superhyperrough Set Model. Therefore, Φ_S generalizes the classical n -Superhyperrough Set Model. $\square$

Example 0.28 (Sequential n -Superhyperrough Set in Housing Selection). Consider a housing selection scenario:

- Let U be a set of houses.
- Two attributes are considered:
 - $T_1 = \text{Location}$ with domain $J_1 = \{\text{Urban}, \text{Suburban}\}$,
 - $T_2 = \text{Style}$ with domain $J_2 = \{\text{Modern}, \text{Classic}\}$.
- The power sets $\mathcal{P}(J_1)$ and $\mathcal{P}(J_2)$ provide, for example, $\{\text{Urban}\}$, $\{\text{Suburban}\}$, or $\{\text{Urban}, \text{Suburban}\}$ for location, and similarly for style.
- Define the super attribute space as

$$J^* = \mathcal{P}(J_1) \times \mathcal{P}(J_2).$$

- Let F^* be defined by, for instance,

$$F^*(\{\text{Urban}\}, \{\text{Modern}\}) = \{\text{houses in urban areas with modern design}\}.$$

- Assume the target concept X consists of houses that are highly desirable to a buyer.

For each house $x \in U$ and for each combination $A \in J^*$, compute the relevance degrees

$$E_B^X(x, A) = \frac{|B(x) \cap F^*(A) \cap X|}{|X|} \quad \text{and} \quad E_{\bar{B}}^X(x, A) = \frac{|B(x) \cap F^*(A) \cap (U \setminus X)|}{|U \setminus X|}.$$

Suppose that in the first iteration, we choose thresholds $\alpha_1 = 0.6$ and $\beta_1 = 0.3$, and in a subsequent iteration, $\alpha_2 = 0.5$ and $\beta_2 = 0.2$. The sequential partitioning process then divides the set of houses into positive, negative, and boundary regions based on these thresholds and the complex combinations of location and style. This refined classification enables a buyer to more precisely identify houses that best match their preferences.

Funding

This study did not receive any financial or external support from organizations or individuals.

Acknowledgments

We extend our sincere gratitude to everyone who provided insights, inspiration, and assistance throughout this research. We particularly thank our readers for their interest and acknowledge the authors of the cited works for laying the foundation that made our study possible. We also appreciate the support from individuals and institutions that provided the resources and infrastructure needed to produce and share this paper. Finally, we are grateful to all those who supported us in various ways during this project.

Data Availability

This research is purely theoretical, involving no data collection or analysis. We encourage future researchers to pursue empirical investigations to further develop and validate the concepts introduced here.

Ethical Approval

As this research is entirely theoretical in nature and does not involve human participants or animal subjects, no ethical approval is required.

Conflicts of Interest

The authors confirm that there are no conflicts of interest related to the research or its publication.

Disclaimer

This work presents theoretical concepts that have not yet undergone practical testing or validation. Future researchers are encouraged to apply and assess these ideas in empirical contexts. While every effort has been made to ensure accuracy and appropriate referencing, unintentional errors or omissions may still exist. Readers are advised to verify referenced materials on their own. The views and conclusions expressed here are the authors' own and do not necessarily reflect those of their affiliated organizations.

References

- [1] Pawlak, Z. (1982). Rough sets. *International journal of computer & information sciences*, 11, 341–356.
- [2] Pawlak, Z., Wong, S. K. M., & Ziarko, W. (1988). Rough sets: Probabilistic versus deterministic approach. *International journal of man-machine studies*, 29(1), 81–95.
- [3] Pawlak, Z., Grzymala-Busse, J., Slowinski, R., & Ziarko, W. (1995). Rough sets. *Communications of the ACM*, 38(11), 88–95. <https://doi.org/10.1145/219717.219791>
- [4] Jech, T. (2003). *Set theory. The third millennium edition, revised and expanded*. Springer-Verlag, Berlin. <https://doi.org/10.1007/3-540-44761-X>
- [5] Pawlak, Z. (1998). Rough set theory and its applications to data analysis. *Cybernetics & Systems*, 29(7), 661–688. <https://doi.org/10.1080/019697298125470>
- [6] Fujita, T. (2025). Short introduction to rough, hyperrough, superhyperrough, treerough. *Advancing uncertain combinatorics through graphization, hyperization, and uncertainization: Fuzzy, neutrosophic, soft, rough, and beyond: Fifth volume: Various Super-HyperConcepts (Collected Papers)*. <https://www.researchgate.net/publication/388754283>
- [7] Fujita, T. (2025). A study on hyperfuzzy hyperrough sets, hyperneutrosophic hyperrough sets, and hypersoft hyperrough sets with applications in cybersecurity. *Artificial intelligence in cybersecurity*, 2, 14–36. <https://doi.org/10.61356/j.aics.2025.2501>
- [8] Fujita, T. (2024). Hyperrough cubic set and superhyperrough cubic set. *Prospects for Applied Mathematics and Data Analysis*, 4(1), 28–35. <https://doi.org/10.54216/PAMDA.040103>
- [9] Fujita, T., & Smarandache, F. (2025). A concise introduction to hyperfuzzy, hyperneutrosophic, hyperplithogenic, hypersoft, and hyperrough sets with practical examples. *Neutrosophic sets and systems*, 80, 610–615.
- [10] Fujita, T. (2024). Note for intuitionistic hyperrough set, one-directional s-hyperrough set, tolerance hyperrough set, and dynamic hyperrough set. *Soft computing fusion with applications*, 2(3), 157–183. <https://doi.org/10.22105/scfa.vi.51>
- [11] Fujita, T. (2025). *Advancing uncertain combinatorics through graphization, hyperization, and uncertainization: Fuzzy, neutrosophic, soft, rough, and Beyond*. Biblio Publishing. <https://www.researchgate.net/publication/386082978>
- [12] Fujita, T. (2025). Neighborhood hyperrough set and neighborhood superhyperrough set. *Pure mathematics for theoretical computer science*, 5(1), 34–47. <https://www.americaspg.com/article/3834/pdf/stream>
- [13] Fujita, T. (2025). Expanding the horizons of graph theory: Rough tree-width, hyperrough structures, and superhyperrough generalizations. *Applied mathematics on science and engineering*, 2(1), 21–35. <https://doi.org/10.65036/amse212>
- [14] Fujita, T. (In Press). Hyperrough topsis method and superhyperrough topsis method. *Uncertainty discourse and applications*. <https://doi.org/10.48313/uda.vi.53>
- [15] Liu, J., Hu, Q., & Yu, D. (2008). A weighted rough set based method developed for class imbalance learning. *Information sciences*, 178(4), 1235–1256. <https://doi.org/10.1016/j.ins.2007.10.002>
- [16] Liu, J., Hu, Q., & Yu, D. (2008). A comparative study on rough set based class imbalance learning. *Knowledge-Based Systems*, 21(8), 753–763. <https://doi.org/10.1016/j.knosys.2008.03.031>
- [17] Alam, S. I., Hassan, S. I., & Uddin, M. (2018). Fuzzy models and business intelligence in web-based applications. In *Applications of soft computing for the web* (pp. 193–221). Singapore: Springer Singapore. https://doi.org/10.1007/978-981-10-7098-3_12
- [18] Sateesh, N., Srinivasa Rao, P., & Rajya Lakshmi, D. (2023). Optimized ensemble learning-based student's performance prediction with weighted rough set theory enabled feature mining. *Concurrency and computation: Practice and experience*, 35(7), e7601. <https://doi.org/10.1002/cpe.7601>
- [19] Kumar, S. S., & Inbarani, H. H. (2015). Optimistic multi-granulation rough set based classification for medical diagnosis. *Procedia computer science*, 47, 374–382. <https://doi.org/10.1016/j.procs.2015.03.219>
- [20] Adlassnig, K. P. (1986). Fuzzy set theory in medical diagnosis. *IEEE Transactions on systems, man, and cybernetics*, 16(2), 260–265. <https://doi.org/10.1109/TSMC.1986.4308946>
- [21] Thomas, L., Crook, J., & Edelman, D. (2017). *Credit scoring and its applications*. Society for industrial and Applied Mathematics. <https://doi.org/10.1137/1.9781611974560>
- [22] Abdou, H. A., & Pointon, J. (2011). Credit scoring, statistical techniques and evaluation criteria: A review of the literature. *Intelligent systems in accounting, finance and management*, 18(2–3), 59–88. <https://doi.org/10.1002/isaf.325>
- [23] Lovett, G. M., Burns, D. A., Driscoll, C. T., Jenkins, J. C., Mitchell, M. J., Rustad, L., ... & Haeuber, R. (2007). Who needs environmental monitoring? *Frontiers in ecology and the environment*, 5(5), 253–260. [https://doi.org/10.1890/1540-9295\(2007\)5\[253:WNEM\]2.0.CO;2](https://doi.org/10.1890/1540-9295(2007)5[253:WNEM]2.0.CO;2)
- [24] Artiola, J. F., Pepper, I. L., & Brusseau, M. L. (2004). *Environmental monitoring and characterization*. Academic Press.
- [25] Biber, E. (2013). The problem of environmental monitoring. *Envtl. L. Rep. News & Analysis*, 43, 10695. <https://access.heinonline.com/HOL/LandingPage?handle=hein.journals/elrna43&div=104&id=&page=>
- [26] Bi, X., Sun, Y., Yang, Y., Lee, A., Hou, D., Liu, B., ... & Hu, N. (2026). Research on piston aeroengine fault diagnosis method with acoustic emission signal and rough set theory synergy. *Structural Health Monitoring*, 25(2), 1150–1164. <https://doi.org/10.1177/14759217241305708>
- [27] Xu, W., Yan, Y., & Li, X. (2024). Sequential rough set: A conservative extension of Pawlak's classical rough set. *W. Xu et al. Artificial intelligence review*, 58(1), 9. <https://doi.org/10.1007/s10462-024-10976-z>
- [28] Asadi, R., Modiri, A., Hosseinali, F., & Gholizadeh, A. A. (2023). A Patterns of housing and neighborhood selection among special groups; A review of housing selection models in Tehran city. *Journal of Iranian architecture & urbanism (JIAU)*, 14(2), 285–299.
- [29] Lindstrom, B. (1997). A sense of place: Housing selection on Chicago's North Shore. *The sociological quarterly*, 38(1), 19–39. <https://doi.org/10.1111/j.1533-8525.1997.tb02338.x>
- [30] Lim, S. M., & Ha, K. S. (2015). The relations of the life style and housing selection attributes of the middle-aged people. *Journal of the Korea academia-industrial cooperation society*, 16(11), 8074–8088.
- [31] Kizielewicz, B., & Bączkiewicz, A. (2021). Comparison of Fuzzy TOPSIS, Fuzzy VIKOR, Fuzzy WASPAS and Fuzzy MMOORA methods in the housing selection problem. *Procedia computer science*, 192, 4578–4591. <https://doi.org/10.1016/j.procs.2021.09.236>