

**MINISTRY OF EDUCATION AND TRAINING  
UNIVERSITY OF TRANSPORT AND COMMUNICATIONS**

**NHU VAN KIEN**

**RESEARCH AND DEVELOPMENT OF  
MULTI-CRITERIA GROUP DECISION-MAKING MODELS  
IN A LINGUISTIC INFORMATION ENVIRONMENT  
BASED ON HEDGE ALGEBRAS**

**MAJOR: INFORMATION TECHNOLOGY  
CODE NO: 9480201**

**DOCTORAL THESIS SUMMARY**

**HA NOI - 2026**

**This research was completed at the University of Transport and Communications.**

**Supervisors:**

**1. Dr. Hoang Van Thong**

**University of Transport and Communications**

**2. Assoc. Prof. Dr. Nguyen Cat Ho**

**Duy Tan University**

**Reviewer 1:**

**Reviewer 2:**

**Reviewer 3:**

The thesis was defended before the Doctoral-Level Evaluation Council.  
at the University of Transport and Communications  
at... on 2026

The thesis can be referred to:

**- National Library of Vietnam**

**- Library of University of Transport and Communications**

## INTRODUCTION

### 1. Reason for the research topic

Decision-making in engineering and information systems often involves selecting or ranking alternatives based on multiple criteria, such as cost, quality, integration capability, reliability, and safety. When multiple experts, managers, or relevant stakeholder groups are involved, the problem becomes a MCGDM problem. In practice, many criteria are qualitative in nature, including assessments of technological suitability, service quality, readiness for digital transformation, and risk levels. Experts often evaluate these criteria using linguistic terms rather than crisp numerical values. However, linguistic information is characterised by vagueness, uncertainty, subjectivity, and context dependence; therefore, a model capable of preserving semantic structure throughout the decision-making process is required.

Approaches based on fuzzy sets, FAHP, FTOPSIS, 2-tuple, and 4-tuple models, have made important contributions. However, many existing models still depend on membership functions, fuzzy numbers, defuzzification, approximate mappings, or different representational layers, which may reduce transparency and semantic consistency. HA provides a mathematical structure for modelling domains of linguistic terms according to semantic order. The fuzzyness measure and the SQM enable the quantification of linguistic terms while contributing to semantic preservation.

A review of related studies reveals three research gaps as follows: TOPSIS/FTOPSIS have not fully exploited the semantic structure of HA in determining PIS/NIS and ranking alternatives. HA-based MCGDM has not adequately integrated criterion weight determination under consistency control with alternative ranking. HA-based linguistic MCGDM has mainly considered static contexts, whereas practical decision-making problems evolve over time and involve historical and missing data. Therefore, it is necessary to develop MCGDM models in a HA-based linguistic information environment, which provides the foundation for the HA-TOPSIS, EFAHP-4-tuple-HA, and dynamic HA-based L-FAHP-L-FTOPSIS models.

### 2. Research Objectives

The general objective of the thesis is to develop MCGDM models in a HA-based linguistic information environment, including scale construction, semantic quantification, criterion-weight determination, consistency control, aggregation of group evaluations, ranking of alternatives, and extension to dynamic contexts. The specific objectives are to construct an HA-TOPSIS model; develop an integrated EFAHP - 4-tuple - HA model; construct a dynamic HA-based L-FAHP - L-FTOPSIS model; and validate the proposed models through two application problems: software supplier selection and assessment of digital transformation readiness of small and medium-sized enterprises (SMEs).

### 3. Research object, scope, and methodology

The research object of the thesis is the linguistic MCGDM problem under an HA-based approach.

The scope of the research focusses on semantic representation and quantification, determination under consistency control, aggregation of expert opinions, ranking of alternatives, and handling of dynamic contexts. The thesis does not aim to cover the entire field of MCDM/MCGDM, validate the models across all application domains, or develop a complete decision support system.

Research methodology includes literature analysis, mathematical modelling, algorithm development, property analysis, experimentation, sensitivity analysis, and comparative analysis of results.

### 4. Structure of the thesis

The thesis consists of an introduction, four core chapters, a Conclusion, a list of published works, references, and appendices. The Introduction presents the rationale, objectives, research content, research object, scope, methodology, new contributions, and structure of the thesis.

Chapter 1, reviews the MCGDM problem in a linguistic information environment and synthesises the theoretical foundations of fuzzy sets, FAHP, FTOPSIS, the 2-tuple, 4-tuple models, and HA, thus identifying the research gaps.

Chapter 2, develops the HA-TOPSIS model for linguistic MCGDM, including the construction of an HA scale, semantic quantification, determination of PIS/NIS, measurement of semantic distances, calculation of the closeness coefficient, and classification of alternatives [CT1].

Chapter 3, presents the integrated EFAHP–4-tuple–HA model for determining the weights under consistency control in an HA environment, aggregating expert opinions, and ranking alternatives in the semantic domain [CT2].

Chapter 4, extends the research to dynamic MCGDM through the dynamic HA-based L-FAHP–L-FTOPSIS model, addressing changes in the set of alternatives, the set of criteria, the set of experts, missing data, and integration of historical data through a semantic decay mechanism [CT3].

The Conclusion summarises the obtained results and identifies the contributions, limitations, and future research directions. The appendices provide data tables, computational results, and experimental evidence to validate the proposed models.

## CHAPTER 1. RESEARCH OVERVIEW OF THE LINGUISTIC MCGDM PROBLEM AND THEORETICAL FOUNDATIONS BASED ON HEDGE ALGEBRAS

### 1.1. The linguistic MCGDM problem

The linguistic MCGDM problem is formulated as follows:

(i) *Input data of the problem:*

Consider three sets: the set of alternatives  $\bar{A} = \{A_1, A_2, \dots, A_m\}$ ,  $m \geq 2$ ; the set of criteria  $\bar{C} = \{C_1, C_2, \dots, C_n\}$ ,  $n \geq 2$ ; and the set of experts  $\bar{D} = \{D_1, D_2, \dots, D_h\}$ ,  $h \geq 2$ .

For each expert  $D_e \in \bar{D}$ , the evaluation of alternative  $A_i \in \bar{A}$  with respect to criterion  $C_j \in \bar{C}$  is represented in the form of a decision evaluation matrix:

$$A_{D_e} = (a_{ij}^e)_{m \times n}, \quad (1.1)$$

where  $a_{ij}^e$  denotes the linguistic evaluation provided by expert  $D^e \in \bar{D}$  for alternative  $A_i \in \bar{A}$  with respect to criterion  $C_j \in \bar{C}$ ;  $e = 1, \dots, h$ ;  $i = 1, \dots, m$ ; and  $j = 1, \dots, n$ .

When the problem takes into account the relative importance of the criteria, the weight vector is given by:

$$\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_n)^T, \text{ with } \tilde{w}_j \geq 0 \text{ and } \sum_{j=1}^n \tilde{w}_j = 1. \quad (1.2)$$

Within the scope of this thesis, both the alternatives' evaluations and the relative importance of the criteria are represented by linguistic terms. Therefore, the model must handle the vagueness, uncertainty, subjectivity and semantic structure of the input information while maintaining semantic order during aggregation and ranking.

(ii) *Output of the problem:* Ranking of alternatives  $A_i \in \bar{A}$

This thesis focusses on selection and ranking problems with two core stages: criterion weight determination and ranking of alternatives in both static and dynamic contexts. In a dynamic context, the set of alternatives, the set of criteria, the set of experts, and the evaluation data may change over time, together with historical data and missing data. Therefore, HA are used as the foundation for representing and aggregating linguistic information to develop a methodological framework for the linguistic MCGDM problem in a sequential, cumulative and extensible direction: starting from the HA-TOPSIS model for ranking alternatives in the semantic domain, then developing the integrated EFAHP–4-tuple–HA model, and finally extending the framework to a dynamic context through the construction of the dynamic HA-based L-FAHP–L-FTOPSIS model.

## 1.2. Fundamental theoretical background

### 1.2.1. Fuzzy Set Theory and linguistic variables

Zadeh's fuzzy set theory represents the degree of membership of an element by means of a membership function, thereby providing a foundation for modelling vague information in MCDM and MCGDM. Linguistic variables allow linguistic terms to serve as values of variables and thus support computing with words (CWW). However, the semantics of linguistic terms depend on membership functions and design parameters, which motivates the use of CWW models to handle semantics in a more rigorous manner.

### 1.2.2. Models of computing with words in a linguistic environment

CWW models provide a basis for processing linguistic evaluations in MCDM/MCGDM. The extension-principle-based approach operates on membership functions but requires an approximation back to linguistic labels. The symbolic model uses indices, which is computationally convenient but may lead to information loss during rounding. The 2-tuple model often assumes that the linguistic scale is uniformly and symmetrically distributed; its computation still relies on real-valued indices and inverse mapping to labels and, therefore, lacks an algebraic foundation for handling semantic relations. These limitations provide the basis for selecting HA as the fundamental theoretical foundation of this thesis as its fundamental theoretical foundation.

### 1.2.3. Theory of Hedge Algebras

#### 1.2.3.1. Definition of HA

**Definition 1.1.** [53] *An HA  $\mathcal{AX}$  of a linguistic variable  $X$  is a four-component structure:  $\mathcal{AX} = (X, G, H, \leq)$ . In this structure,  $X$  is the set of linguistic terms of the variable  $X$ , with  $X \subseteq \text{Dom}(X)$ ;  $G = \{g^-, g^+\}$  is the set of generators, where  $g^-$  and  $g^+$  are the negative and positive primary generators, respectively, satisfying the semantic relation  $g^- \leq g^+$ .  $H$  is the set of hedges of the linguistic variable  $X$ .  $\leq$  is a semantic ordering relation defined on  $X$ . It is assumed that the value domain contains constant elements  $(0, W, 1)$ , denoting the smallest element, the neutral element, and the greatest element in  $X$ , respectively, such that  $0 \leq g^- \leq W \leq g^+ \leq 1$ .*

#### 1.2.3.2. Fuzziness measure of linguistic values

**Definition 1.2.** [52] *Let  $\mathcal{AX}^* = (X, G, H, \sigma, \phi, \leq)$  be a complete linear HA (Com HA) of the linguistic variable  $X$ . A mapping  $f_{\mathfrak{M}}: X \rightarrow [0, 1]$  is called a fuzziness measure of the linguistic terms in  $X$  if the following conditions are satisfied:*

- i)  $f_{\mathfrak{M}}$  is a complete measure on  $X$ , that is:  $f_{\mathfrak{M}}(g^-) + f_{\mathfrak{M}}(g^+) = 1$  and  $\sum_{h \in H} f_{\mathfrak{M}}(hu) = f_{\mathfrak{M}}(u)$ ,  $\forall u \in X$ ;
- ii)  $f_{\mathfrak{M}}(x) = 0$ ,  $\forall x$  satisfying  $H(x) = \{x\}$ ; in particular,  $f_{\mathfrak{M}}(0) = f_{\mathfrak{M}}(W) = f_{\mathfrak{M}}(1) = 0$ ;
- iii)  $\forall x, y \in X$ ,  $\forall h \in H$ ,  $\mu(h) = \frac{f_{\mathfrak{M}}(hx)}{f_{\mathfrak{M}}(x)} = \frac{f_{\mathfrak{M}}(hy)}{f_{\mathfrak{M}}(y)}$ , This ratio is independent of  $x$  and  $y$ , and is referred to as the fuzziness measure of the hedges, denoted by  $\mu(h)$ .

From (i) and (iii), the recursive formula to calculate the fuzziness measure  $x = h_m \dots h_1 g$ , with  $g \in \{g^-, g^+\}$  is obtained as follows:  $f_{\mathfrak{M}}(x) = \mu(h_m) \dots \mu(h_1) f_{\mathfrak{M}}(g)$ , where  $\sum_{h \in H} \mu(h) = 1$ .

#### 1.2.3.3. Semantic Quantification of linguistic values

**Definition 1.3.** [52] *Let  $\mathcal{AX}^* = (X, G, H, \sigma, \phi, \leq)$  be a linear ComHA. A mapping  $\vartheta: X \rightarrow [0, 1]$  is called a SQM of  $\mathcal{AX}^*$  if it satisfies the following conditions:*

- (i)  $\vartheta$  is a one-to-one mapping from  $X$  into  $[0, 1]$  such that  $\vartheta(0) = 0$ ,  $\vartheta(1) = 1$  and it preserves the semantic order on  $X$ ; that is  $\forall x, y \in X$ ,  $x < y \Rightarrow \vartheta(x) < \vartheta(y)$ ;
- (ii)  $\vartheta$  is continuous in the sense that:  $\forall x \in X$ ,  $\vartheta(\sigma x) = \infimum \vartheta(H(x))$  và  $\vartheta(\phi x) = \supremum \vartheta(H(x))$ .

**Definition 1.4.** [52] Let  $AX^* = (X, G, H, \sigma, \phi, \preceq)$  be a linear ComHA, and  $f\mathfrak{m}(x)$  and  $\mu(h)$  be the fuzziness measures of the linguistic value  $x$  and the hedge  $h$ , respectively. Then, a mapping  $\vartheta: X \rightarrow [0, 1]$  induced by the fuzziness measure is defined as follows:

- (i)  $\vartheta(W) = \theta = f\mathfrak{m}(g^-)$ ,  $\vartheta(g^-) = \theta - \alpha f\mathfrak{m}(g^-) = \beta f\mathfrak{m}(g^-)$ ,  $\vartheta(g^+) = \theta + \alpha f\mathfrak{m}(g^+)$ ;
- (ii)  $\vartheta(h_j x) = \vartheta(x) + \zeta g(h_j x) \left\{ \sum_{i=\zeta g(j)}^{j-\zeta g(j)} \mu(h_i) f\mathfrak{m}(x) - \omega(h_j x) \mu(h_j x) f\mathfrak{m}(x) \right\}$  for all  $j$ , where  $-q \leq j \leq p$ ,  $j \neq 0$ , and  $\omega(h_j x) = \frac{1}{2} [1 + \zeta g(h_j x) \zeta g(h_p h_j x) (\beta - \alpha)]$ ;
- (iii)  $\vartheta(\phi g^-) = 0$ ,  $\vartheta(\sigma g^-) = \theta = \vartheta(\phi g^+)$ ,  $\vartheta(\sigma g^+) = 1$ , and for all  $j$ ,  $-q \leq j \leq p$ ,  
 $\vartheta(\phi h_j x) = \vartheta(x) + \zeta g(h_j x) \left\{ \sum_{i=\zeta g(j)}^{j-\zeta g(j)} \mu(h_i) f\mathfrak{m}(x) \right\} - \frac{1}{2} (1 - \zeta g(h_j x)) \mu(h_i) f\mathfrak{m}(x)$ ,  
 $\vartheta(\sigma h_j x) = \vartheta(x) + \zeta g(h_j x) \left\{ \sum_{i=\zeta g(j)}^{j-\zeta g(j)} \mu(h_i) f\mathfrak{m}(x) \right\} + \frac{1}{2} (1 + \zeta g(h_j x)) \mu(h_i) f\mathfrak{m}(x)$ .

#### 1.2.3.4. System of Semantic Similarity Intervals

**Definition 1.5** [52] Given a fuzziness measure  $f\mathfrak{m}$  on  $X$ , a value  $k > 0$ , and a finite set of terms  $X_k = \{x \in X: |x| \leq k\}$ , a set of intervals  $\{S_k(x): x \in X_k\}$  is constructed on the normalized reference domain  $[0, 1]$ . These intervals are called  $k$ -similarity intervals and satisfy the following conditions:

- (i) These intervals form a partition of  $[0, 1]$ ;
- (ii) Each  $S_k(x)$  contains exactly one value  $\vartheta(x)$  belonging to the set  $\{\vartheta(y): y \in X_k\}$ , where  $\vartheta$  is an SQM induced by  $f\mathfrak{m}$ . Accordingly, every value in  $S_k(x)$  is considered similar to  $\vartheta(x)$ , or equivalently, similar to the meaning of the linguistic term  $x$  at level  $k$ .

#### 1.2.3.5. Linguistic scale represented by the semantic 4-tuple model

**Definition 1.6** [54] Given a linguistic variable  $X$  and the normalised numerical reference domain  $[0, 1]$ , the 4-tuple semantic representation of a linguistic term  $x \in \text{Dom}(X)$  is defined by:

$$(x, \vartheta(x), S_k(x), r_x).$$

Where:  $x \in \text{Dom}(X)$ : is a linguistic term;  $S_k(x) \subseteq [0, 1]$ : is the similarity interval of  $x$  for the given term set  $X_k$ ;  $\vartheta(x) \in [0, 1]$ : is the numeric quantitative semantics of  $x$ ;  $r_x \in S_k(x)$ : is a reference value within the similarity interval.

The concepts of HA, including the algebraic structure of linguistic terms, the fuzziness measure, SQM, the system of semantic similarity intervals, and the semantic 4-tuple scale, enable the representation, quantification, comparison, aggregation, and preservation of semantic structure. These concepts provide the direct foundation for developing HA-based linguistic MCGDM models in Chapters 2, 3, and 4.

### 1.3. The foundational methods inherited and extended by the thesis

In MCGDM, this thesis builds on AHP/FAHP for determining criterion weights through pairwise comparisons and consistency control, and TOPSIS/FTOPSIS for ranking alternatives according to their proximity to PIS and distance from NIS. However, these approaches still rely on crisp numerical values, fuzzy numbers, membership functions, or defuzzification, which may lead to insufficient semantic consistency. Therefore, the thesis develops these methods on the basis of HA, thereby establishing the foundation for HA-TOPSIS, EFAHP - 4-tuple - HA, and HA-based dynamic L-FAHP - L-FTOPSIS models.

### 1.4. Overview of research on linguistic MCDM and MCGDM

Traditional MCDM/MCGDM studies have established various frameworks for selection, evaluation, and ranking, including AHP, TOPSIS, VIKOR, and ELECTRE. In linguistic information environments, fuzzy set-based extensions, such as FAHP and FTOPSIS, together with computing-with-words models, such as the 2-tuple linguistic model and the 4-tuple linguistic model, have contributed to handling vagueness and uncertainty and reducing information loss. However, many

models still rely on membership functions, defuzzification, and approximate mapping or lack mechanisms for ensuring semantic consistency across criterion weight determination, aggregation, and ranking. HA-based MCGDM models have demonstrated the ability to model semantic ordering, quantify linguistic terms, and preserve the domain of structure of the linguistic terms. However, criterion weights are often assumed a priori or directly assigned by experts; consistency control mechanisms have not been fully integrated into HA-based procedures; and these models have been mainly been constructed in static contexts.

### **1.5. Research gaps and thesis orientation**

The review shows that MCDM and MCGDM in linguistic information environments have been studied through AHP, TOPSIS, FAHP, FTOPSIS, the 2-tuple linguistic model, the 4-tuple linguistic model, and HA. However, three main research gaps remain. First, TOPSIS/FTOPSIS has not fully exploited the semantic ordering, fuzziness measure, and SQM function of HA to determine PIS/NIS, measure semantic distances, and derive rankings within the same semantic domain. Second, existing HA-based models have not yet fully developed a procedure for determining criterion weight with consistency control that is semantically related to ranking. Third, HA-based MCGDM models remain primarily static and have not been fully handled by historical data, missing data, and semantic degradation. Therefore, the thesis adopts HA as the foundational approach for developing a methodological framework for linguistic MCGDM problems: starting with the HA-TOPSIS model to rank alternatives within a semantic domain; then inherit and supplementing the criterion-weight determination process with consistency control and develop the EFAHP - 4-tuple - HA model; and finally extending the framework to dynamic contexts by constructing a dynamic HA-based dynamic L-FAHP - L-FTOPSIS model by proposing the dynamic L-FAHP–L-FTOPSIS model based on HA.

### **1.6. Chapter 1 Conclusions**

Chapter 1 systematises the MCGDM problem in linguistic information environments and the related theoretical foundations, including fuzzy sets, CWW, the 2-tuple linguistic model, the 4-tuple semantic model, AHP/FAHP, TOPSIS/FTOPSIS, and HA. The chapter clarifies the requirements for preserving semantic structure, semantic connectivity, and consistency control in criterion weight determination, evaluation aggregation, and alternative ranking. These analyses provide the basis for the development of HA-TOPSIS, EFAHP–4-tuple–HA, and dynamic HA-based L-FAHP–L-FTOPSIS in Chapters 2, 3, and 4, respectively.

## **CHAPTER 2. DEVELOPMENT OF A LINGUISTIC TOPSIS MODEL BASED ON HEDGE ALGEBRAS**

### **2.1. Introduction**

This chapter addresses the first research gap: existing TOPSIS/FTOPSIS models still rely on real numbers, fuzzy numbers, and membership functions and have not yet fully exploited the semantic ordering structure of HA for ranking alternatives. In MCGDM, experts often evaluate alternatives using linguistic terms rather than crisp numerical values; direct conversion into real numbers, fuzzy numbers, or membership functions can degrade semantic order and semantic content. TOPSIS ranks alternatives according to the principle of proximity to PIS and distance from NIS. However, traditional TOPSIS and several FTOPSIS models still rely on crisp numerical values, defuzzification, or intermediate quantitative domains. Therefore, this chapter develops HA-TOPSIS to construct an HA-based linguistic scale, perform semantic quantification using the SQM function, determine PIS/NIS, measure semantic distances, compute the closeness coefficient, and rank alternatives within the same semantic domain.

### **2.2. Development of the HA-TOPSIS decision-making model**

(i) *Input data:*

Let  $\bar{A} = \{A_1, A_2, \dots, A_m\}$ ,  $m \geq 2$ ;  $\bar{C} = \{C_1, C_2, \dots, C_n\}$ ,  $n \geq 2$  and  $\bar{D} = \{D_1, D_2, \dots, D_h\}$ ,  $h \geq 2$ . The expert panel  $\bar{D}$  agrees to construct the HA-based linguistic scales  $L_{(j)k}(\bar{A})$  and  $L_k(\bar{C})$  to evaluate alternatives and determine the weight of each criterion  $C_j \in \bar{C}$ , where  $k$  is an integer indicating that the linguistic terms on the scale have a length not exceeding  $k$ , with  $j = 1, \dots, n$ .

For each expert  $D^e \in \bar{D}$ , the following tasks are performed:

(1) Construct the alternative evaluation matrix according to Equation (2.1):

$$A^e = (a_{ij}^e)_{m \times n}. \quad (2.1)$$

Where  $a_{ij}^e \in L_{(j)k}(\bar{A})$  is the linguistic assessment given by expert  $D^e \in \bar{D}$  for alternative  $A_i \in \bar{A}$  with respect to criterion  $C_j \in \bar{C}$ , with  $e = 1, \dots, h$ ;  $i = 1, \dots, m$ ; and  $j = 1, \dots, n$ .

(2) Construct the criterion weight evaluation matrix according to Equations (2.2):

$$P^e = (\tilde{x}_j^e)_{n \times h}. \quad (2.2)$$

where  $\tilde{x}_j^e \in L_k(\bar{C})$  ( $j = 1, \dots, n$ ) is the linguistic assessment provided by expert  $D^e$  regarding the importance level of criterion  $C_j \in \bar{C}$ . The vector of criterion weights is  $\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_n)^T$ , where  $\tilde{w}_j \geq 0$  and  $\sum_{j=1}^n \tilde{w}_j = 1$ .

(ii) *Output data:* The ranking of the alternatives  $A_i \in \bar{A}$ .

The decision-making procedure of the HA-TOPSIS model consists of the following steps.

**Step 1:** Construct HA-based linguistic scales for evaluating alternatives and determining criterion weights.

The expert panel agrees to establish the set of semantic parameters of HA, denoted by  $(HA, fm(g^-), \mu(h), k)$ , to construct the linguistic scales  $L_{(j)k}(\bar{A})$  for evaluating alternatives and  $L_k(\bar{C})$  for assessing criterion weights.

Based on Definitions 1.4 - 1.5 and Proposition 1.1, the procedure for constructing a linguistic scale includes: (i) determining the fuzziness measures of the generators and hedges; and (ii) establishing the set of linguistic terms used for assessment, with the maximum term length not exceeding  $k$ .

**Step 2:** Construct linguistic matrices to evaluate alternatives and criterion weights.

Each expert  $D^e$  uses linguistic terms  $x \in X_{(j)k}$  on the linguistic scale  $L_{(j)k}(\bar{A})$  to evaluate alternatives  $A_i \in \bar{A}$  with respect to each criterion  $C_j \in \bar{C}$  according to Equation (2.1).

Next, each expert  $D^e$  uses linguistic terms  $x \in X_k$  on the linguistic scale  $L_k(\bar{C})$  to assess the importance level of criteria  $C_j \in \bar{C}$  according to Equation (2.2).

**Step 3:** Compute the semantic quantitative matrices for alternatives and criterion weights.

Using Definitions 1.6 - 1.8 and Proposition 1.2, the HA SQM function is constructed to convert the linguistic terms in matrices  $A^e$  and  $P^e$  into the representation of the numerical semantic value of each linguistic term,  $\vartheta(x) \in [0, 1]$ , for computational purposes. Consequently, the semantic quantitative matrix  $\tilde{A}^e$  is determined by Equation (2.3) as follows:

$$\tilde{A}^e = (\bar{a}_{ij}^e)_{m \times n}, i = 1, \dots, m; j = 1, \dots, n; e = 1, \dots, h. \quad (2.3)$$

Similarly, each element in the criterion-weight evaluation matrix  $P^e$  is also converted into the corresponding numerical semantic quantitative value  $\vartheta(\tilde{x}_j^e)$ .

**Step 4:** Construct the group decision matrix for alternatives and criterion weights.

The aggregated group evaluation matrix, denoted by  $A\tilde{A}$ , is constructed by taking the arithmetic mean of the alternative assessments of all expert decision matrices  $\tilde{A}^e$ , assuming equal expert weights, according to Equation (2.4):

$$A\tilde{A} = (u_{ij})_{m \times n} = \frac{1}{h} \sum_{e=1}^h u_{ij}^e. \quad (2.4)$$

Similarly, the aggregated weight of criterion  $C_j \in \bar{C}$  is determined by Equation (2.5):

$$AP = (\tilde{u}_j) = \frac{1}{h} \sum_{e=1}^h \vartheta(\tilde{x}_j^e). \quad (2.5)$$

The vector of criterion weights is  $\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_n)^T$ , with  $\tilde{w}_j = \frac{\tilde{u}_j}{\sum_{k=1}^n \tilde{u}_k}$ ,  $j = 1, \dots, n$ ;  $\tilde{w}_j \geq 0$  and  $\sum_{j=1}^n \tilde{w}_j = 1$ .

**Step 5:** Compute the weighted decision matrix.

$$Y_{ij} = u_{ij} \times \tilde{w}_j; i = 1, \dots, m; j = 1, \dots, n. \quad (2.6)$$

Where  $u_{ij}$  is an element of matrix  $A\tilde{A}$  and  $\tilde{w}_j$  is the weight of criterion  $C_j \in \bar{C}$ .

**Step 6:** Determine PIS/NIS and measure the distance to the ideal solutions.

- Determine (PIS,  $\mathcal{B}^+$ ) and (NIS,  $\mathcal{B}^-$ ) using Equations (2.7) và (2.8) as follows:

$$\mathcal{B}^+ = \{\mathcal{B}_1^+, \mathcal{B}_2^+, \dots, \mathcal{B}_n^+\} \quad (2.7) \quad \text{and} \quad \mathcal{B}^- = \{\mathcal{B}_1^-, \mathcal{B}_2^-, \dots, \mathcal{B}_n^-\} \quad (2.8)$$

+ For benefit criteria:  $\mathcal{B}_j^+ = \sum_{i=1}^m \max \mathcal{B}_{ij}$ , and  $\mathcal{B}_j^- = \sum_{i=1}^m \min \mathcal{B}_{ij}$ .

+ For cost criteria:  $\mathcal{B}_j^+ = \sum_{i=1}^m \min \mathcal{B}_{ij}$ , and  $\mathcal{B}_j^- = \sum_{i=1}^m \max \mathcal{B}_{ij}$ .

- The distances  $d_i^+$  and  $d_i^-$  from alternative  $A_i$  to (PIS,  $\mathcal{B}^+$ ) and (NIS,  $\mathcal{B}^-$ ) are calculated by Equations (2.9) and (2.10):

$$d_i^+ = \left( \sum_{j=1}^n (\mathcal{B}_{ij} - \mathcal{B}_j^+)^2 \right)^{\frac{1}{2}} \quad (2.9); \quad d_i^- = \left( \sum_{j=1}^n (\mathcal{B}_{ij} - \mathcal{B}_j^-)^2 \right)^{\frac{1}{2}} \quad (2.10)$$

where  $i = 1, \dots, m$ ;  $j = 1, \dots, n$ ;  $d_i^+$  and  $d_i^-$  are the distances from  $A_i$  to the PIS and NIS, respectively.

**Step 7:** Rank the alternatives.

Based on the value of the closeness coefficient  $CC_i$ , ( $i = 1, \dots, m$ ) the set of alternatives is ranked in descending order; the alternative with the largest value of  $CC_i$  is the optimal alternative. The closeness coefficient is calculated by Equation (2.11):

$$CC_i = \frac{d_i^-}{(d_i^- + d_i^+)}, i = 1, \dots, m. \quad (2.11)$$

### 2.3. Experimental example

Company A needs to select the most suitable electronic invoice software supplier that meets the required criteria. An expert panel consisting of five decision-makers, denoted by  $D_1, D_2, D_3, D_4$ , and  $D_5$ , was formed to evaluate three suppliers, denoted E1, E2, and E3, based on four criteria: Unit Price (P1), Quality (P2), Technology (P3) and After-sales Support (P4).

**Table 2.1.** Distances from suppliers to PIS/NIS

| Supplier | $d^+$ | $d^-$ |
|----------|-------|-------|
| E1       | 0.053 | 0.100 |
| E2       | 0.104 | 0.060 |
| E3       | 0.076 | 0.071 |

**Table 2.2.** Closeness Coefficients and supplier rankings

| Supplier | Closeness coefficients | Rank |
|----------|------------------------|------|
| E1       | 0.654                  | 1    |
| E2       | 0.364                  | 3    |
| E3       | 0.485                  | 2    |

- The experimental results show that supplier E1 ranks first, E3 ranks second, and E2 ranks third.

- Thus, E1 is closer to the PIS and farther from the NIS than the remaining alternatives.

The weights of the four criteria show that unit price and Technology are the two most influential criteria.

## 2.4. Sensitivity Analysis and comparison of results

### 2.4.1. Sensitivity Analysis

Sensitivity analysis was carried out by changing the fuzziness measure under two negative–positive characteristic scenarios in the semantic domain. Scenario 1 is defined as  $f_{\mathfrak{M}}(U) = 0.4$ ,  $f_{\mathfrak{M}}(I) = 0.6$ ,  $\mu(L) = 0.6$ ,  $\mu(V) = 0$ . Conversely, Scenario 2 is defined as  $f_{\mathfrak{M}}(U) = 0.6$ ,  $f_{\mathfrak{M}}(I) = 0.4$ ,  $\mu(L) = 0.4$ , and  $\mu(V) = 0.6$ .

**Table 2.3.** Closeness Coefficients Obtained under different scenarios

| Supplier | Initial | Scenario 1 | Scenario 2 |
|----------|---------|------------|------------|
| E1       | 0.654   | 0.647      | 0.825      |
| E2       | 0.364   | 0.372      | 0.192      |
| E3       | 0.485   | 0.474      | 0.217      |

**Table 2.4.** Supplier ranking

| Scenario   | Ranking                |
|------------|------------------------|
| Initial    | $E1 \succ E3 \succ E2$ |
| Scenario 1 | $E1 \succ E3 \succ E2$ |
| Scenario 2 | $E1 \succ E3 \succ E2$ |

The results of the sensitivity analysis in Tables 2.3 and 2.4 show that the closeness coefficients of the suppliers vary between scenarios, but not substantially enough to alter the overall ranking is preserved as  $E1 \succ E3 \succ E2$ . This provides important evidence of the stability of the model against the disturbance of semantic parameters.

### 2.4.2. Comparison of results and discussion

**Table 2.5.** Comparison of ranking results among HA-TOPSIS, FTOPSIS, 4-tuple semantic model

| Rank | HA-TOPSIS             |      | FTOPSIS               |      | 4-tuple semantic model |      |
|------|-----------------------|------|-----------------------|------|------------------------|------|
|      | Closeness coefficient | Rank | Closeness coefficient | Rank | Total score            | Rank |
| E1   | 0.65                  | 1    | 0.58                  | 1    | 0.82                   | 1    |
| E2   | 0.36                  | 3    | 0.51                  | 2    | 0.59                   | 3    |
| E3   | 0.49                  | 2    | 0.36                  | 3    | 0.77                   | 2    |

- The comparison results in Table 2.5 show that all three models identify E1 as the best alternative. This confirms the consistency of the experimental data and the main selection result.

- Meanwhile, FTOPSIS produces the ranking  $E1 \succ E2 \succ E3$ , thus reversing the ranks of E2 and E3. This indicates that when alternatives have close evaluation levels, fuzzy-number representation and defuzzification may change the ranking.

- The ranking similarity between HA-TOPSIS and the 4-tuple semantic model reinforces the argument that the HA foundation better preserves the semantic ordering of the data.

## 2.5. Chapter 2 Conclusion

Chapter 2 develops the HA-TOPSIS model for ranking alternatives in the linguistic MCGDM problem in a static context. The model constructs an HA-based linguistic scale, performs semantic quantification using the SQM function, determines PIS/NIS, and computes semantic distances and the closeness coefficient. The experimental example, sensitivity analysis, and comparison of results show that the model maintains semantic consistency, stability, and interpretability. Chapter 3 builds on this model and supplements the criterion-weight determination process with consistency control through the integrated EFAHP-4-tuple-HA model.

## CHAPTER 3. DEVELOPMENT OF AN INTEGRATED MCGDM MODEL COMBINING ENHANCED FAHP WITH THE 4-TUPLE SEMANTIC MODEL UNDER THE HEDGE ALGEBRAS APPROACH

### 3.1. Introduction

Chapter 3 builds on the HA-TOPSIS model developed in Chapter 2 and shifts the focus to criterion weight determination with consistency control in linguistic MCGDM problems. Since criterion weights directly affect evaluation and ranking results, this chapter develops an integrated EFAHP - 4-tuple-HA model to exploit expert knowledge, control judgment consistency, and ensure semantic connectivity among criterion weight determination, evaluation aggregation, and alternative ranking. The model is validated by assessing the level of digital transformation level of retail SMEs in Vietnam.

### 3.2. Consistency Control of pairwise comparison matrices in the HA Environment

In AHP and FAHP, consistency control is an essential condition for assessing the reliability of pairwise comparison matrices. The thesis builds on this rationale and develops it in the HA environment. The consistency ratio (CR) of a linguistic pairwise comparison matrix is determined by Equation (3.1). The pairwise comparison matrix provided by expert  $D^e$  is accepted only if  $CR_{HA}^e < 0.10$ . Otherwise, expert  $D^e$  must revise the pairwise comparisons until the consistency requirement is satisfied.

$$CR_{HA}^e = \frac{CI_{HA}^e}{RI(n)}, \text{ với } e = 1, \dots, h; n \geq 3. \quad (3.1)$$

-  $RI(n)$  denotes the random index of an order  $n$  matrix in Saaty's classical AHP, used as a matrix size normalisation coefficient. The values of  $RI(n)$  are given in Table 3.1 below:

**Table 3.1.** Values of the Random Index ( $RI(n)$ )

| n  | 1    | 2    | 3    | 4    | 5    | 6    | 7    | 8    | 9    | 10   |
|----|------|------|------|------|------|------|------|------|------|------|
| RI | 0.00 | 0.00 | 0.58 | 0.90 | 1.12 | 1.24 | 1.32 | 1.41 | 1.45 | 1.49 |

-  $CI_{HA}^e$  denotes the consistency index of matrix  $P^e$  provided by expert  $D^e$ .

On the basis of the  $CI_{Saaty}$  concept in classical AHP, the thesis develops a formula for computing the consistency index in the HA environment ( $CI_{HA}^e$ ) for the linguistic pairwise comparison matrices  $P^e$  provided by expert  $D^e$  as follows:

$$CI_{HA}^e = \frac{\lambda_{\max}^e - n - (1 - \vartheta(w))}{n - (1 - \vartheta(w))} = \frac{\lambda_{\max}^e - (n + \delta)}{n - \delta}. \quad (3.2)$$

The difference between Equation (3.2) and  $CI_{Saaty}$  in classical AHP lies in replacing the consistency benchmark  $n$  with the soft benchmark  $n + (1 - \vartheta(w))$  and directly embedding the quantitative semantic tolerance of the semantic value of the neutral element  $1 - \vartheta(w)$  into the measurement of consistency deviation in the pairwise comparison matrices, to adapt to the characteristics of the HA-based semantic domain based on HA.

**Proposition 3.1.** For  $n \geq 2$ ,  $\vartheta(w) \in (0, 1)$ , and with  $\mathcal{R}^e$  a positive reciprocal matrix derived from the linguistic matrix  $P^e$ , the index  $CI_{HA}^e$  in Equation (3.2) has the following basic properties:

- (i)  $CI_{HA}^e$  is well-defined on the application domain;
- (ii)  $CI_{HA}^e = 0$  at the HA-consistency benchmark  $\lambda_0^{HA}$ ;
- (iii) For fixed  $n$  and  $\vartheta(w)$ ,  $CI_{HA}^e$  increases monotonically with  $\lambda_{\max}^e$ ;
- (iv) Equation (3.2) directly embeds the quantitative semantic value of the neutral word  $\vartheta(w)$  into the consistency-measurement mechanism;
- (v)  $CI_{HA}^e$  is continuous with respect to  $\lambda_{\max}^e$  and  $\vartheta(w)$ , and is locally stable;
- (iv) (vi) The sign of  $CI_{HA}^e$  allows a direct interpretation of the matrix position relative to the semantic tolerance region (below, equal to, or above the HA-consistency benchmark).

### 3.3. Development of an enhanced FAHP model with 4-tuple semantics based on HA

(i) *Input data:*

Let  $\bar{A} = \{A_1, A_2, \dots, A_m\}$ ,  $m \geq 2$ ;  $\bar{C} = \{C_1, C_2, \dots, C_n\}$ ,  $n \geq 2$  and  $\bar{D} = \{D_1, D_2, \dots, D_h\}$ ,  $h \geq 2$ . The expert panel  $\bar{D}$  constructs two HA-based 4-tuple semantic scales, denoted by  $L_k(\bar{C})$  and  $L_k(\bar{A})$ , for evaluating criterion weights and alternatives of the criteria. For each expert  $D^e \in \bar{D}$ , the following tasks are performed:

(1) Construct a pairwise linguistic comparison matrix to determine criterion weights, according to Equation (3.3):

$$P^e = (x_{ij}^e)_{n \times n}. \quad (3.3)$$

where  $x_{ij}^e \in L_k(\bar{C})$ ,  $x_{ii}^e$  is the neutral element of the scale  $L_k(\bar{C})$ , representing an equally important comparison relation between a criterion and itself,  $e = 1, 2, \dots, h$ ;  $i, j = 1, \dots, n$ ;  $i \neq j$ .

The linguistic pairwise comparison matrix  $P^e$  is accepted only if the consistency ratio satisfies  $CR_{HA}^e < 0.10$ ; otherwise, expert  $D^e$  must revise the pairwise comparisons until the consistency requirement is satisfied.

(2) Construct an alternative evaluation matrix according to Equation (3.4):

$$A^e = (a_{ij}^e)_{m \times n}; a_{ij}^e \in L_k(\bar{A}); e = 1, \dots, h; i = 1, \dots, m; j = 1, \dots, n. \quad (3.4)$$

(ii) *Output data:* The ranking of alternatives  $A_i \in \bar{A}$ .

The EFAHP-4-tuple-HA model is implemented in two stages. In stage 1, EFAHP-HA is developed based on the FAHP method to determine criterion weights through linguistic pairwise comparisons with consistency control in the HA environment. In Stage 2, the 4-tuple -HA model is applied to aggregate, evaluate, and rank the alternatives.

#### 3.3.1. Determination of Criterion Weights

**Step 1.** Designing the Hierarchical Structure for the decision-making problem

First, the MCGDM problem is modelled as a three-level hierarchical structure: the highest level is the goal; the second level consists of the main criteria/sub-criteria; and the lowest level is the set of alternatives.

**Step 2.** Constructing a HA-based 4-tuple semantic scale to evaluate criterion weights

The expert panel establishes the semantic parameter set of HA,  $(HA, f\mathfrak{m}(g^-), \mu(h), k)$  to construct the 4-tuple semantic scale  $L_k(\bar{C})$  used for evaluating criterion weights.

Using Definition 1.5 and Proposition 1.1, the fuzzyness measures are computed for linguistic terms  $x \in X_k$  on the 4-tuple semantic scale  $L_k(\bar{C})$ . Next, by applying Definitions 1.6 to 1.8 and Proposition 1.2, each linguistic term  $x \in X_k$  is mapped to a semantic quantitative value  $\vartheta(x) \in [0, 1]$ . Subsequently, Definitions 1.9 to 1.12 and Proposition 1.3 are applied to determine the semantic similarity interval  $S_k(x)$  for each linguistic term  $x \in X_k$  on the scale  $L_k(\bar{C})$ . On this basis, the linguistic scale  $L_k(\bar{C})$  for evaluating criterion weights is represented as a 4-tuple semantic representation, specifically:  $L_k(\bar{C}) = \{(x, \vartheta(x), S_k(x), r_x): x \in X_k, r_x \in S_k(x)\}$ .

**Step 3.** Constructing the pairwise comparison matrix and checking consistency.

Each expert  $D^e$  uses the linguistic terms  $x \in X_k$  on the linguistic scale  $L_k(\bar{C})$  to evaluate the relative importance of the criteria, thus constructing the linguistic pairwise comparison matrix  $P^e = (x_{ij}^e)_{n \times n}$  according to Equation (3.3). The elements on the main diagonal represent the comparison relation of a criterion with itself and are therefore assigned the neutral element  $w$  on the linguistic scale  $L_k(\bar{C})$ , that is:  $x_{ii}^e = w$ ,  $(i = 1, \dots, n)$ .

Then, each element  $x_{ij}^e \in P^e$  is mapped to its semantic quantitative value  $\vartheta(x_{ij}^e)$  and further transformed into a generalized triangular fuzzy number (GTFN)  $u_{ij}^e = (L_{ij}^e, C_{ij}^e, R_{ij}^e; \omega_{ij}^e)$ ,

Accordingly, matrix  $P^e$  is transformed into the GTFN matrix  $\mathcal{R}^e$  as shown in Equation (3.5):

$$\mathcal{R}^e = (u_{ij}^e)_{n \times n}. \quad (3.5)$$

where  $u_{ij}^e = (L_{ij}^e, C_{ij}^e, R_{ij}^e; \omega_{ij}^e)$ ,  $(u_{ij}^e)^{-1} = (\frac{1}{R_{ij}^e}, \frac{1}{C_{ij}^e}, \frac{1}{L_{ij}^e}; \omega_{ij}^e)$ , ( $e = 1, \dots, h$ ;  $i, j = 1, \dots, n$  and  $i \neq j$ ); In this thesis, the reciprocal operation is not interpreted as a closed operation on the same representation space; rather, it is used mainly to establish the dual relationship among elements in the fuzzy pairwise comparison matrix. Subsequent computational steps are performed only after the elements have been transferred into a compatible representation domain to ensure the formal consistency of the model. The component  $\omega_{ij}^e$  is the height of the fuzzy number, representing the maximum degree of membership or the confidence level of the evaluation; therefore, it is preserved and not reciprocated.

Equation (3.1) is used to check the consistency ratio  $CR_{HA}^e$  of the linguistic pairwise comparison matrices  $P^e$ . Matrix  $P^e$  is accepted only if  $CR_{HA}^e < 0.10$ ; otherwise, expert  $D^e$  must revise the linguistic pairwise comparisons until the consistency requirement is satisfied.

**Step 4.** Aggregation of the evaluations provided by experts

After the matrices  $P^e$  have satisfied the consistency requirement, they are aggregated into a group fuzzy pairwise comparison matrix  $AP$  using the arithmetic mean of the GTFNs according to Equation (3.6):

$$AP = (\tilde{u}_{ij})_{n \times n} \quad (3.6)$$

$$\text{where } \tilde{u}_{ij} = \left( \frac{\sum_{e=1}^h L_{ij}^e}{h}, \frac{\sum_{e=1}^h C_{ij}^e}{h}, \frac{\sum_{e=1}^h R_{ij}^e}{h}, \frac{\sum_{e=1}^h \omega_{ij}^e}{h} \right), e = 1, \dots, h; i, j = 1, \dots, n \text{ and } i \neq j.$$

**Step 5.** Computing fuzzy synthetic extent values and the criterion weight vector

From the aggregated matrix  $AP$ , the fuzzy synthetic extent value of each criterion, denoted by  $P_i$ , is determined through normalization by Equation (3.7) as follows:

$$\begin{aligned} \mathbb{P}_i &= (\mathbb{D}_i, \mathbb{K}_i, \mathbb{T}_i; \min(\omega_{ij})) \\ &= \left( \frac{\sum_{j=1}^n L_{ij}}{\sum_{i=1}^n L_{ij} + \sum_{k=1, k \neq i}^n \sum_{j=1}^n R_{kj}}, \frac{\sum_{j=1}^n C_{ij}}{\sum_{i=1}^n \sum_{j=1}^n C_{ij}}, \frac{\sum_{j=1}^n R_{ij}}{\sum_{j=1}^n R_{ij} + \sum_{k=1, k \neq i}^n \sum_{j=1}^n L_{kj}}; \min(\omega_{ij}) \right), \text{ with } i, j = 1, \dots, n; i \neq j. \end{aligned} \quad (3.7)$$

Each  $\mathbb{P}_i$  represents the aggregated priority level of criterion  $C_j$  after all pairwise comparison relations in the group matrix. Next, the centroid of each fuzzy number  $P_i$  is determined by  $m_i = (\bar{A}_{\mathbb{P}_i}, \bar{L}_{\mathbb{P}_i})$ , with  $i = 1, 2, \dots, n$ , and is computed by Equation (3.8):

$$\bar{A}_{\mathbb{P}_i} = \frac{(\mathbb{D}_i + \mathbb{K}_i + \mathbb{T}_i)}{3}, \quad \bar{L}_{\mathbb{P}_i} = \frac{\min(\omega_{ij})}{3}. \quad (3.8)$$

The distance from the centroid  $m_i = (\bar{A}_{\mathbb{P}_i}, \bar{L}_{\mathbb{P}_i})$  to the reference point  $\mathbb{C} = (\bar{A}_{min}, \bar{L}_{min})$  is computed by Equation (3.9):

$$D(\mathbb{P}_i, \mathbb{C}) = \sqrt{(\bar{A}_{\mathbb{P}_i} - \bar{A}_{min})^2 + \left(\bar{L}_{\mathbb{P}_i} - \frac{\psi}{3} \bar{L}_{min}\right)^2} \quad (3.9)$$

where  $\bar{A}_{min} = \min(\mathbb{D}_i)$ ,  $\bar{L}_{min} = \min(\omega_{ij})$ , and  $\psi = \min(\omega_{ij})$ ;

Finally, the criterion weight vector  $\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_n)^T$  corresponding to the fuzzy comparison matrix is calculated by Equation (3.10) as follows:

$$\tilde{w}_i = \frac{D(\mathbb{P}_i, \mathbb{C}_i)}{\sum_{i=1}^n D(\mathbb{P}_i, \mathbb{C}_i)} = \frac{\sqrt{(\bar{A}_{\mathbb{P}_i} - \bar{A}_{min})^2 + (\bar{L}_{\mathbb{P}_i} - \frac{\psi}{3} \mathcal{L}_{min})^2}}{\sum_{i=1}^n \sqrt{(\bar{A}_{\mathbb{P}_i} - \bar{A}_{min})^2 + (\bar{L}_{\mathbb{P}_i} - \frac{\psi}{3} \mathcal{L}_{min})^2}}, i = 1, \dots, n \quad (3.10)$$

### 3.3.2. Evaluation and ranking of alternatives

**Step 6.** Constructing a 4-tuple semantic scale for evaluating alternatives.

To evaluate alternatives according to the criteria, the expert panel establishes the semantic parameter set of HA,  $(HA, f\mathfrak{m}(g^-), \mu(h), k)$  to construct the 4-tuple semantic scale  $L_k(\bar{A})$  for alternative evaluation. The procedure to construct this scale is analogous to that for  $L_k(\bar{C})$  in Step 2, but it is used for evaluating and ranking alternatives. Accordingly, the linguistic scale  $L_k(\bar{A})$  is defined by:  $L_k(\bar{A}) = \{(x, \vartheta(x), S_k(x), r_x) : x \in X_k, r_x \in S_k(x)\}$ . Where each linguistic term  $x \in X_k$  on the scale  $L_k(\bar{A})$  is represented as a 4-tuple semantic representation.

**Step 7.** Collecting and converting alternative evaluations into 4-tuple semantic representations.

Each expert  $D^e$  ( $e = 1, 2, \dots, h$ ), uses the linguistic terms  $x \in X_k$  on the linguistic scale  $L_k(\bar{A})$  to evaluate  $m$  alternatives under  $n$  criteria. These evaluations are organised into an alternative evaluation matrix  $A^e = (a_{ij}^e)_{m \times n}$  according to Equation (3.4).

Then, by applying Definitions 1.9 to 1.12 and Proposition 1.3, each element  $a_{ij}^e$  in matrix  $A^e$  is converted into the corresponding 4-tuple semantic representation  $(a_{ij}^e, \vartheta(a_{ij}^e), S_k(a_{ij}^e), r_{a_{ij}^e})$ .

**Step 8.** Aggregating and ranking alternatives

- *Aggregation of expert evaluations:*

The alternative evaluation matrices of experts are aggregated into an aggregated decision matrix of the expert group, denoted by:  $AA = (\bar{a}_{ij})_{m \times n}$ . Here,  $\bar{a}_{ij}$  is the aggregated evaluation of alternative  $A_i$  with respect to criterion  $C_j$ . To aggregate expert opinions, the thesis employs the arithmetic mean aggregation operator on the representative component  $r_x$  of the 4-tuple semantic representations. Specifically, given a set  $T_s = \{\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_p\}$ , where each element  $\tilde{x}_j$  has the form:  $\tilde{x}_j = (x_j, \vartheta(x_j), S_k(x_j), r_j), j = 1, \dots, p$ , the aggregation operator  $g$  is defined by:

$$g(\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_p) = \Delta \left( \frac{1}{p} \sum_{j=1}^p r_j \right), \quad (3.11)$$

where,  $\Delta(\cdot)$  is the inverse mapping operator from the representative value in the normalized domain to a corresponding 4-tuple semantic representation, and  $\frac{1}{p} \sum_{j=1}^p r_j \in S_k(x)$ . In this problem, for each pair  $(i, j)$ , the aggregated element  $\bar{a}_{ij}$  is determined by applying operator  $g$  to the evaluations of all experts:  $\bar{a}_{ij} = g((a_{ij}^1, \vartheta(a_{ij}^1), S_k(a_{ij}^1), r_{ij}^1), \dots, (a_{ij}^h, \vartheta(a_{ij}^h), S_k(a_{ij}^h), r_{ij}^h))$ .

- *Weighted aggregation according to the criteria:*

After constructing the aggregated decision matrix  $AA$ , the final aggregated score for each alternative  $A_i \in \bar{A}$  is computed by the weighted mean aggregation operator according to the criterion weight vector  $\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_p)^T$ , with  $\sum_{j=1}^p \tilde{w}_j = 1$  obtained in Step 5. The weighted aggregation is then defined by the following.

$$Y(\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_p) = \Delta \left( \sum_{j=1}^p \tilde{w}_j r_j \right) \quad (3.12)$$

The result of the aggregation is a 4-tuple semantic representation:

$$\left( x, \vartheta(x), S_k(x), \sum_{j=1}^p \tilde{w}_j r_j \right), \text{ where } \sum_{j=1}^p \tilde{w}_j r_j \in S_k(x).$$

Finally, based on the total scores, the alternatives  $A_i \in \bar{A}$  are ranked in descending order; the alternative with the highest score is ranked first.

### 3.4. Application of the proposed model to the assessment of the readiness for digital transformation of SMEs in Vietnam's retail sector

In this application, the empirical dataset comprises four experts, 16 SMEs, seven main criteria, and 34 sub-criteria. Specifically, four experts in retail research and management participated in evaluating the importance levels of a criterion system consisting of seven main criteria and 34 sub-criteria, developed on the basis of Decision No. 2158/QĐ-BTTTT dated 07 November 2023 of the Ministry of Information and Communications. Meanwhile, 16 SMEs in Vietnam's retail sector were surveyed to assess their current digital transformation status.

On this basis, the expert panel reached a consensus on the construction of two HA-based 4-tuple linguistic scales, denoted by  $L_k(\bar{C})$  and  $L_k(\bar{A})$ , for assessing criterion importance and the level of fulfilment level, respectively. Accordingly, the four experts used the scale  $L_k(\bar{C})$  to perform pairwise comparisons and determine the relative importance levels of the criteria, whereas the 16 surveyed SMEs reported their fulfillment levels for each evaluation criterion on the scale  $L_k(\bar{A})$ .

The empirical computation process of the EFAHP - 4-tuple - HA model was organised into two stages: (i) determining the weights of the criterion with consistency control using the EFAHP–HA method in the HA environment and (ii) aggregating, evaluating, and ranking the digital transformation readiness of enterprises using the HA-based 4-tuple linguistic model.

The results obtained from applying the EFAHP-4-tuple-HA model to rank the digital transformation readiness of 16 retail SMEs (E1–E16) in Vietnam are presented in Table 3.2 as follows:

**Table 3.2.** Summary of scores and rankings of digital transformation readiness for 16 SMEs

| Enterprise | Total score   | 4-tuple semantic model |                       |                |               | Rank     |
|------------|---------------|------------------------|-----------------------|----------------|---------------|----------|
|            |               | $x$                    | $S_k(x)$              | $\vartheta(x)$ | $r_x$         |          |
| E1         | 0.3522        | L                      | [0.219, 0.450)        | 0.336          | 0.3522        | 2        |
| <b>E2</b>  | <b>0.3930</b> | <b>L</b>               | <b>[0.219, 0.450)</b> | <b>0.336</b>   | <b>0.3930</b> | <b>1</b> |
| E3         | 0.2318        | L                      | [0.219, 0.450)        | 0.336          | 0.2318        | 16       |
| E4         | 0.3380        | L                      | [0.219, 0.450)        | 0.336          | 0.3380        | 5        |
| E5         | 0.3446        | L                      | [0.219, 0.450)        | 0.336          | 0.3446        | 3        |
| E6         | 0.3381        | L                      | [0.219, 0.450)        | 0.336          | 0.3381        | 4        |
| E7         | 0.2876        | L                      | [0.219, 0.450)        | 0.336          | 0.2876        | 10       |
| E8         | 0.3250        | L                      | [0.219, 0.450)        | 0.336          | 0.3250        | 6        |
| E9         | 0.2956        | L                      | [0.219, 0.450)        | 0.336          | 0.2956        | 8        |
| E10        | 0.3010        | L                      | [0.219, 0.450)        | 0.336          | 0.3010        | 7        |
| E11        | 0.2767        | L                      | [0.219, 0.450)        | 0.336          | 0.2767        | 13       |
| E12        | 0.2826        | L                      | [0.219, 0.450)        | 0.336          | 0.2826        | 11       |
| E13        | 0.2770        | L                      | [0.219, 0.450)        | 0.336          | 0.2770        | 12       |
| E14        | 0.2950        | L                      | [0.219, 0.450)        | 0.336          | 0.2950        | 9        |
| E15        | 0.2765        | L                      | [0.219, 0.450)        | 0.336          | 0.2765        | 14       |
| E16        | 0.2761        | L                      | [0.219, 0.450)        | 0.336          | 0.2761        | 15       |

- The of the criteria of the weighting results indicate that pillar P7 - Human and organisational capabilities has the largest weight, equal to 0,28. It is followed by P6 - Risk management and cybersecurity, with a weight of 0.16. The remaining pillars are relatively evenly distributed: P1 = 0.10; P2 = 0.10; P3 = 0.11; P4 = 0.12; P5 = 0.12. Regarding the ranking, all 16 SMEs were assigned to the same linguistic class L, namely the Low level. However, due to the 4-tuple representation, the model can still distinguish different degrees of readiness through the reference value  $r_x$ . This explains why the thesis adopts the 4-tuple representation, which shows that the model not only performs classification, but also enables fine-grained ranking within the same linguistic class.

- The results in Table 3.2 indicate that the ranking order of the 16 SMEs is as follows: E2 > E1 > E5 > E6 > E4 > E8 > E10 > E9 > E14 > E7 > E12 > E13 > E11 > E15 > E16 > E3. In particular, E2 ranks first, E1 ranks second, E5 ranks third, and E3 ranks last.

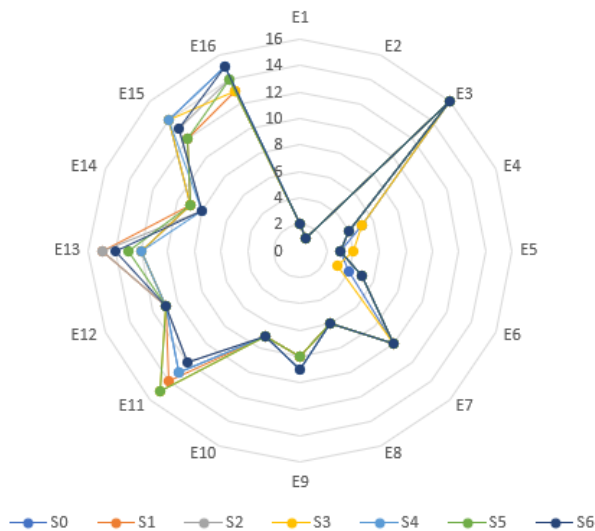
### 3.5. Sensitivity Analysis and result comparison

#### 3.5.1. Sensitivity Analysis

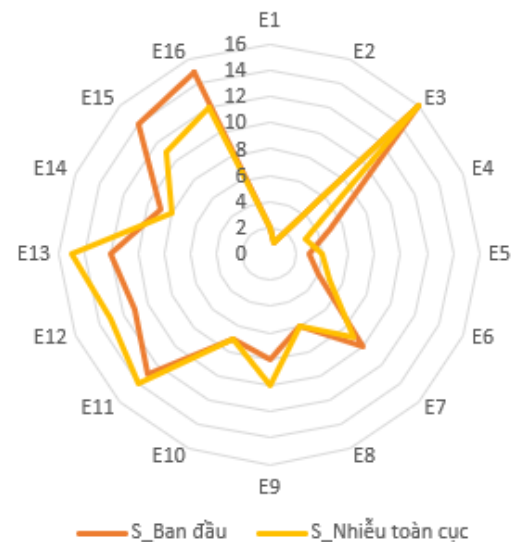
Sensitivity analysis was conducted under two perturbation scenarios: single-criterion perturbation and global perturbation applied to the pairwise comparison matrix of the seven main criteria. Specifically:

In Scenario 1, single-criterion perturbation was generated by successively increasing the priority level of each main criterion in P1 to P7 relative to the remaining six criteria on  $L_2(\tilde{C})$ , in order to separately examine the influence of each criterion on the ranking results. The results shown in Figure 3.1 indicate that the model exhibits high stability for the top and lowest-ranked enterprise groups. Ranking fluctuations occur mainly in the middle-ranked enterprise group, where the total scores are relatively close.

Scenario 2 generated a global perturbation by increasing all pairwise comparison values of the seven main criteria by one level on the same scale  $L_2(\tilde{C})$ . The results shown in Figure 3.2 indicate that the overall ranking structure is maintained, particularly for leading enterprises such as E2 and E1 and for the lowest-ranked enterprise, E3. The fluctuations continue to concentrate in the middle-ranked enterprise group, reflecting the fine-grained discrimination capability of the model when the scores are close.



**Figure 3.1.** Sensitivity analysis under the single-criterion perturbation scenario



**Figure 3.2.** Sensitivity analysis under the global perturbation scenario

In summary, the results of the sensitivity analysis confirm that the EFAHP - 4-tuple - HA model exhibits high stability while maintaining the necessary sensitivity to distinguish alternatives with close evaluation levels.

### 3.5.2. Comparison and discussion of results

To further clarify the suitability, stability, and interpretability of the model, the thesis compares EFAHP - 4-tuple - HA with HA-TOPSIS based on HA and FAHP-FTOPSIS based on fuzzy sets. The results show that the three models consistently identify E2 as the best enterprise and E3 as the lowest-ranked enterprise. This indicates that the main conclusion of the proposed model is consistent with that of the comparison models.

Although the FAHP-FTOPSIS model still identifies E2 as the best alternative, it exhibits substantial fluctuations in the upper- and middle-ranked enterprise groups. These differences indicate that fuzzy mapping and defuzzification may lead to semantic degradation and rank reversal when alternatives have close-evaluation results.

In summary, the EFAHP - 4-tuple - HA model demonstrates good semantic preservation, reduces information loss during computation, and enhances the reliability and interpretability of the ranking results.

### 3.6. Conclusion of Chapter 3

Chapter 3 developed an integrated EFAHP - 4-tuple - HA model that determines criterion weights with consistency control using EFAHP-HA and applies the 4-tuple linguistic model to evaluate and rank alternatives in a connected manner within the same semantic domain. The experiment involving 16 retail SMEs in Vietnam, together with sensitivity analysis and comparisons with HA-TOPSIS and FAHP-FTOPSIS, shows that the model is stable, interpretable, and suitable for linguistic MCGDM. Chapter 4 extends the study to dynamic MCGDM based on HA, considering historical data, missing data, and a semantic degradation mechanism.

## CHAPTER 4. DEVELOPMENT OF A METHODOLOGICAL FRAMEWORK FOR CONSTRUCTING A DYNAMIC MCGDM MODEL IN A LINGUISTIC INFORMATION ENVIRONMENT BASED ON HEDGE ALGEBRAS

### 4.1. Introduction

Chapter 4 builds upon the HA-TOPSIS and EFAHP - 4-tuple - HA models developed in Chapters 2 and 3 and extends the linguistic MCGDM problem to a dynamic context in the HA environment. In this context, the set of alternatives, the set of criteria, the set of experts, the criterion weights, and the evaluation data may vary over time. The thesis develops linguistic aggregation operations on the 4-tuple semantic scale, an aggregation operation with missing data, and a semantic decay mechanism. It also constructs a dynamic L-FAHP - L-FTOPSIS model based on HA and validates the model through the evaluation of SMEs at multiple time points.

### 4.2. Fuzzy-set semantics of linguistic terms in HA

**Definition 4.1.** The mapping  $f_z: L_k \rightarrow F_{L_k}$  is called the triangular fuzzy semantic (TFS) mapping of a linguistic term. For each linguistic term  $x \in X_k$  on the linguistic scale  $L_k$  represented by the 4-tuple semantic representation  $\ddot{x} = (x, \vartheta(x), S_k(x), r_x) \in L_k$ , the mapping  $f_z(\ddot{x})$  is defined by:

$$f_z(\ddot{x}) = (l, m, u). \quad (4.1)$$

where  $(l, m, u)$  is determined according to the TFS partitioning procedure described above.

### 4.3. Linguistic aggregation operations

#### 4.3.1. Binary linguistic aggregation operation on the 4-tuple semantic scale

**Definition 4.2.** Consider two linguistic terms  $\ddot{x}, \ddot{y} \in L_k$  represented by 4-tuple semantic representations:  $\ddot{x} = (x, S_k(x), \vartheta(x), r_x)$ ,  $\ddot{y} = (y, S_k(y), \vartheta(y), r_y)$ , where  $w$  is the neutral element of  $L_k$ . Let  $\ddot{z} = \ddot{x} \oplus \ddot{y} = (z, S_k(z), \vartheta(z), r_z)$  denote the aggregation result of  $\ddot{x}$  and  $\ddot{y}$ . The reference component  $r_z$  is determined as follows:

$$r_z = r_x + r_y = \begin{cases} fm(g^-) \times T\left(\frac{r_x}{fm(g^-)}, \frac{r_y}{fm(g^-)}\right), & \text{if } x, y < w \\ fm(g^-) + fm(g^+) \times S\left(\frac{r_x - fm(g^-)}{fm(g^+)}, \frac{r_y - fm(g^-)}{fm(g^+)}\right), & \text{if } w \leq x, y \\ \frac{fm(y)}{fm(x) + fm(y)} \times r_x + \frac{fm(x)}{fm(x) + fm(y)} \times r_y, & \text{if } x < w < y \text{ and } x \neq 0, y \neq 1 \end{cases} \quad (4.2)$$

where  $T(a, b) = a \times b$  and  $S(a, b) = a + b - a \times b$ .

After  $r_z$  is determined, the remaining components of  $\ddot{z}$  are derived by selecting a linguistic term  $z \in L_k$  such that  $r_z \in S_k(z)$ , while  $\vartheta(z)$  is the corresponding semantic quantitative value of  $z$ . Therefore, the result of the aggregation operation  $\oplus$  is always mapped back into the domain of 4-tuple semantic representations.

From a semantic perspective, Equation (4.2) divides the aggregation operation into three cases based on the positions of the two operands  $x$  and  $y$  with respect to the neutral element  $w$ . Let  $r_x = \vartheta(x)$ ,  $r_y = \vartheta(y)$  and  $r_w = \vartheta(w)$  be the corresponding semantic quantitative values. The three cases are described as follows:

- Case 1 (both operands lie on the negative side of  $w$ ):  $r_x < r_w$  and  $r_y < r_w$ . Then,  $r_z$  is computed according to the first branch of Equation (4.2).
- Case 2 (both operands lie on the non-negative side of  $w$ ):  $r_x \geq r_w$  and  $r_y \geq r_w$ . Then,  $r_z$  is computed according to the second branch of Equation (4.2).
- Case 3 (the mixed case):  $r_x \leq r_w \leq r_y$  or  $r_y \leq r_w \leq r_x$ . Then,  $r_z$  is computed according to the third branch of Equation (4.2) as a convex combination of  $r_x$  and  $r_y$ , with weights depending on the fuzziness measures of the two corresponding linguistic terms.

**Proposition 4.2.** *For all  $\ddot{x}, \ddot{y} \in L_k$ , the aggregation operation  $\oplus$  defined in Definition 4.1 satisfies the following properties:*

- (i) *Closure property on the linguistic scale  $L_k$ , i.e.,  $\ddot{z} = \ddot{x} \oplus \ddot{y} \in L_k$ ;*
- (ii) *Commutativity, i.e.,  $\ddot{x} \oplus \ddot{y} = \ddot{y} \oplus \ddot{x}$ ;*
- (iii) *Monotonicity, if  $\ddot{x}_1 \leq \ddot{x}_2$  then  $\ddot{x}_1 \oplus \ddot{y} \leq \ddot{x}_2 \oplus \ddot{y}$ ;*
- (iv) *There exists a neutral element  $w \in L_k$ , such that  $\ddot{x} \oplus w = w \oplus \ddot{x} = \ddot{x}$ ;*
- (v) *Partial associativity, that is, the aggregation operation is associative only when all operands lie on the same side with respect to  $w$ . Specifically, if  $x, y, z \geq w$  or  $x, y, z \leq w$  then  $(\ddot{x} \oplus \ddot{y}) \oplus \ddot{z} = \ddot{x} \oplus (\ddot{y} \oplus \ddot{z})$ .*

**Definition 4.3.** *Let  $n$  linguistic terms  $\ddot{x}_i \in L_k$ , with  $i = 1, 2, \dots, n$ , be represented by their corresponding 4-tuple semantic representations  $\ddot{x}_i = (x_i, \vartheta(x_i), S_k(x_i), r_{x_i})$  and suppose that these linguistic terms have been sorted in ascending order with respect to  $x_i$ . Let  $\ddot{z} = (z, \vartheta(z), S_k(z), r_z)$  denote the linguistic term obtained by aggregating all  $\ddot{x}_i$ . Then, the multi-ary aggregation operation on the 4-tuple semantic scale  $L_k$  is recursively defined by Equation (4.3):*

$$\ddot{z} = \sum_{i=1}^n \ddot{x}_i = \left( \sum_{i=1}^{n-1} \ddot{x}_i \right) \oplus \ddot{x}_n \quad (4.3)$$

#### 4.3.2. Aggregation operation with missing data

**Definition 4.4.** *Let  $L_k$  be a 4-tuple linguistic scale constructed based on HA. Each linguistic term  $\ddot{x} \in L_k$  is represented as,  $\ddot{x} = (x, S_k(x), \vartheta(x), r_x)$ , where  $w \in L_k$  is the neutral element of the scale  $L_k$ . Let  $\emptyset$  denote the null term representing the case of missing data, with  $\emptyset \notin L_k$ . Suppose that  $\ddot{x}^{(\tau-1)} = (x^{(\tau-1)}, S_k(x^{(\tau-1)}), v(x^{(\tau-1)}), r_{x^{(\tau-1)}})$  and  $\ddot{x}^{(\tau)} = (x^{(\tau)}, S_k(x^{(\tau)}), v(x^{(\tau)}), r_{x^{(\tau)}})$  are respectively the evaluation values of an object at two time points  $\tau - 1$  and  $\tau$ , with  $\tau = 1, 2, \dots, t$ ;*

$t \geq 1$ . Let  $\ddot{x}_{agg}^{(\tau)} = (x_{agg}^{(\tau)}, S_k(x_{agg}^{(\tau)}), v(x_{agg}^{(\tau)}), r_{x_{agg}^{(\tau)}})$  denote the aggregated value at time point  $\tau$  determined from  $\ddot{x}^{(\tau-1)}$  and  $\ddot{x}^{(\tau)}$  in the case where one of the two values is missing, as follows:

(i) Case 1: With  $\tau = 1$  or  $\ddot{x}^{(\tau-1)} = \emptyset$ . When there are no historical data, or when the data at time point  $\tau - 1$  are missing, the aggregated value at time point  $\tau$  is taken as the current evaluation itself:

$$\ddot{x}_{agg}^{(\tau)} = \emptyset \oplus \ddot{x}^{(\tau)} = \ddot{x}^{(\tau)} \quad (4.4)$$

(ii) Case 2:  $\tau > 1$ ,  $\ddot{x}^{(\tau-1)} \neq \emptyset$  and  $\ddot{x}^{(\tau)} = \emptyset$ . When the current evaluation data is missing, the aggregated value at time point  $\tau$  is determined by the semantic decay mechanism from the historical evaluation:

$$\ddot{x}_{agg}^{(\tau)} = \ddot{x}^{(\tau-1)} \oplus \ddot{x}^{(\tau)} = \ddot{x}^{(\tau-1)} \oplus \emptyset \quad (4.5)$$

In this case, the reference component of the aggregated value  $r_{x_{agg}^{(\tau)}}$  is determined by the semantic decay equation (4.6):

$$r_{x_{agg}^{(\tau)}} = \vartheta(w) + \lambda \left( r_{x^{(\tau-1)}} - \vartheta(w) \right). \quad (4.6)$$

Where:

$\vartheta(w)$  is the quantitative semantic value of the neutral element  $w \in L_k$ ;

$\lambda \in [0, 1]$  is the semantic decay coefficient, also referred to as the forgetting rate:

For  $\lambda = 1$ , no semantic decay occurs; therefore  $r_{x_{agg}^{(\tau)}} = r_{x^{(\tau-1)}}$ .

For  $\lambda = 0$ , the historical information is completely forgotten; therefore  $r_{x_{agg}^{(\tau)}} = \vartheta(w)$ .

For  $0 < \lambda < 1$ , the value  $r_{x_{agg}^{(\tau)}}$  lies between  $r_{x^{(\tau-1)}}$  and  $\vartheta(w)$ , thereby modeling the gradual decay of the influence of the historical evaluation toward the neutral element.

After  $r_{x_{agg}^{(\tau)}}$  is determined, the remaining components of the 4-tuple semantic representation  $\ddot{x}_{agg}^{(\tau)}$  are determined by selecting a linguistic term  $x_{agg}^{(\tau)} \in L_k$  such that  $r_{x_{agg}^{(\tau)}} \in S_k(x_{agg}^{(\tau)})$ ,  $v(x_{agg}^{(\tau)})$  is the semantic quantitative value of the term  $x_{agg}^{(\tau)}$ .

Set  $s_{\tau-1} = r_{x^{(\tau-1)}}$ ,  $s_{\tau} = r_{x^{(\tau)}}$ .

On this basis, the following proposition is obtained:

**Proposition 4.2.** Let  $s_{\tau-1} \in [0, 1]$  be the semantic reference value at time point  $\tau - 1$ ,  $\vartheta(w) \in [0, 1]$  be the semantic quantitative value of the neutral element, and  $\lambda \in [0, 1]$  be the semantic decay coefficient. Then, the semantic decay aggregation operation defined by Equation (4.6) satisfies the following properties:

(i). Range preservation:  $\min\{s_{\tau-1}, \vartheta(w)\} \leq s_{\tau} \leq \max\{s_{\tau-1}, \vartheta(w)\}$ , meaning that the decayed value  $s_{\tau}$  always lies within the interval bounded by the historical value  $s_{\tau-1}$  and the neutral value  $\vartheta(w)$ ;

(ii). Convergence to the neutral element: If data are missing for  $n$  consecutive time points and  $0 < \lambda < 1$ , then  $s_{\tau+n}$  converges to  $\vartheta(w)$  as  $n \rightarrow \infty$ ;

(iii). Monotonicity with respect to the decay coefficient: For  $n \geq 1$ , if  $0 \leq \lambda_1 \leq \lambda_2 \leq 1$  then  $|s_{\tau+n}(\lambda_1) - \vartheta(w)| \leq |s_{\tau+n}(\lambda_2) - \vartheta(w)|$ , which means that the smaller  $\lambda$  is, the faster  $s_{\tau+n}$  approaches  $\vartheta(w)$ .

(iv). Stability at the neutral benchmark: if  $s_{\tau-1} = \vartheta(w)$ , then  $s_{\tau} = \vartheta(w)$  for all  $\lambda \in [0, 1]$ , that is, when the historical value is already at the neutral benchmark, the decay operation does not change this value.

#### 4.4. Construction of a dynamic L-FAHP - L-FTOPSIS model based on HA

(i) *Input data:*

(1) Number of evaluation time points:  $T = \{1, 2, \dots, t\}$ , with  $t \geq 1$ .

(2) At each time point  $\tau$  (with  $\tau = 1, 2, \dots, t$ ), perform the following:

- Consider the following three finite sets:  $\bar{A}^\tau = \{A_1^\tau, A_2^\tau, \dots, A_{m_\tau}^\tau\}$ , ( $m_\tau \geq 2$ );  $\bar{C}^\tau = \{C_1^\tau, C_2^\tau, \dots, C_{n_\tau}^\tau\}$ , ( $n_\tau \geq 2$ ) and  $\bar{D}^\tau = \{D_1^\tau, D_2^\tau, \dots, D_{h_\tau}^\tau\}$  ( $h_\tau \geq 2$ ).

- The expert panel  $\bar{D}^\tau$  agrees to construct two 4-tuple semantic scales based on HA, denoted respectively by  $L_k(\bar{C}^\tau)$  and  $L_k(\bar{A}^\tau)$ , to determine criterion weights and evaluate alternatives.

- For each expert  $D_e^\tau \in \bar{D}^\tau$  perform the following:

+ Construct a pairwise linguistic comparison matrix to determine criterion weights, according to Equation (4.7):

$$P_{D_e^\tau}^\tau = (x_{ij}^e(\tau))_{n_\tau \times n_\tau}. \quad (4.7)$$

where each element  $x_{ij}^e(\tau) \in L_k(\bar{C}^\tau)$ . The main-diagonal elements  $x_{ii}^e(\tau)$  are assigned the neutral element of the linguistic scale  $L_k(\bar{C}^\tau)$ , representing the equal-importance comparison relation between a criterion and itself, with  $e = 1, 2, \dots, h_\tau$ ;  $\tau = 1, 2, \dots, t$ ;  $i, j = 1, 2, \dots, n_\tau$ ;  $i \neq j$ .

+ Construct the alternative evaluation matrix according to Equation (4.8):

$$A_{D_e^\tau}^\tau = (a_{ij}^e(\tau))_{m_\tau \times n_\tau}. \quad (4.8)$$

where  $a_{ij}^e(\tau) \in L_k(\bar{A}^\tau)$ ,  $e = 1, \dots, h_\tau$ ;  $i = 1, \dots, m_\tau$ ;  $j = 1, \dots, n_\tau$ ;  $\tau = 1, 2, \dots, t$ .

(3) From the time point  $\tau - 1$  onwards, historical data are taken into account.

(ii) *Output data:* Ranking of the alternatives  $A_i^\tau \in \bar{A}^\tau$ , with historical data taken into account

This decision-making procedure consists of two stages. Stage 1 develops the linguistic FAHP model based on HA, denoted by L-FAHP - HA, to determine the weights based on FAHP, whereas Stage 2 develops the Linguistic FTOPSIS model based on HA, denoted by L-FTOPSIS - HA, to evaluate, aggregate and rank alternatives across multiple time points based on FTOPSIS. In both stages, expert evaluation information is represented on the 4-tuple semantic linguistic scale constructed from HA, allowing aggregation operations to be performed directly in the linguistic domain while preserving the order and inherent semantic structure of linguistic terms. At the same time, the model establishes a semantic decay mechanism to integrate historical data with current data by gradually shifting them toward the neutral element of the scale.

##### 4.4.1. Determination of criterion weights

**Step 1.1.** Construct the linguistic pairwise comparison matrix and check consistency at time point  $\tau$ .

At time point  $\tau$ , each expert  $D_e^\tau$  uses linguistic terms  $x \in X_k$  on the linguistic scale  $L_k(\bar{C}^\tau)$  to construct the linguistic pairwise comparison matrix, denoted by  $P_{D_e^\tau}^\tau$ , according to Equation (4.7).

To ensure the reliability of judgments, the matrix  $P_{D_e^\tau}^\tau$  is checked for consistency through the consistency ratio  $CR_{HA}^e(\tau)$ , determined by Equation (4.9):

$$CR_{HA}^e(\tau) = \frac{CI_{HA}^e(\tau)}{RI(n_\tau)}, e = 1, \dots, h; n_\tau \geq 3. \quad (4.9)$$

where,  $RI(n_\tau)$  is the random index of a matrix of order  $n_\tau$  in classical AHP, as given in Table 4.1. Meanwhile  $CI_{HA}^e(\tau)$  is the consistency index of the linguistic pairwise comparison matrix  $P_{D_e^\tau}^\tau$  provided by expert  $D_e^\tau$  in the HA environment, as shown in Equation (4.10):

$$CI_{HA}^e(\tau) = \frac{\lambda_{\max}^e(\tau) - n_\tau - (1 - \vartheta(w)^\tau)}{n_\tau - (1 - \vartheta(w)^\tau)} \quad (4.10)$$

To compute  $CR_{HA}^e(\tau)$ , the linguistic matrix  $P_{D_e^\tau}^\tau$  is first converted into the fuzzy semantic matrix  $FP_{D_e^\tau}^\tau$  by mapping each element  $x_{ij}^e(\tau)$  to the corresponding fuzzy semantic set through the mapping  $fz(x_{ij}^e(\tau))$  according to Equation (4.1). The matrix  $P_{D_e^\tau}^\tau$  is accepted only if it achieves the consistency ratio  $CR_{HA}^e(\tau) < 0.10$ . Otherwise, expert  $D_e^\tau$  must revise the pairwise comparisons until the consistency requirement is satisfied.

**Step 1.2.** Aggregate the linguistic pairwise comparison matrices of experts at time point  $\tau$ .

After all matrices  $P_{D_e^\tau}^\tau$  have satisfied the consistency ratio, the linguistic evaluations of the entire expert panel at time point  $\tau$  are aggregated to obtain the group pairwise comparison matrix, denoted by  $AP^\tau$ . The aggregation is performed in the 4-tuple linguistic domain by applying the binary linguistic aggregation operation constructed in Equations (4.2) - (4.3), according to Equation (4.11):

$$AP^\tau = (\tilde{x}_{ij}(\tau))_{n_\tau \times n_\tau}. \quad (4.11)$$

where each element  $\tilde{x}_{ij}(\tau)$  of the matrix  $AP^\tau$  is determined by:

$$\tilde{x}_{ij}(\tau) = x_{ij}^1(\tau) \oplus x_{ij}^2(\tau) \oplus \dots \oplus x_{ij}^{n_\tau}(\tau), \quad i, j = 1, \dots, n_\tau; i \neq j. \quad (4.12)$$

The matrix  $AP^\tau$  represents the collective evaluation of the expert panel regarding the relative importance of the criteria at time point  $\tau$ .

**Step 1.3.** Construct the integrated pairwise comparison matrix that incorporates historical data.

This step is applied when  $\tau > 1$ , when the decision-making process at the current time point needs to incorporate information from previous times. To integrate historical data, let:  $\bar{C}_u^\tau = \bar{C}^{\tau-1} \cup \bar{C}^\tau$ , denote the set of criteria appearing in two consecutive time points  $\tau - 1$  and  $\tau$ , and  $n_\tau^u$  be the cardinality of this set. Then, the integrated pairwise comparison matrix incorporating historical data is denoted by  $\overline{AP}^\tau$ , whose elements are determined by Equation (4.13) as follows:

$$\tilde{p}_{ij}^*(\tau) = \begin{cases} \emptyset \oplus \tilde{x}_{ij}(\tau) & \text{if } C_i^{\tau-1} \notin \bar{C}^{\tau-1} \text{ and } C_i^\tau \in \bar{C}^\tau, \\ \tilde{x}_{ij}(\tau-1) \oplus \tilde{x}_{ij}(\tau) & \text{if } C_i^{\tau-1}, C_i^\tau \in \bar{C}^{\tau-1} \cap \bar{C}^\tau, \\ \tilde{x}_{ij}(\tau-1) \oplus \emptyset & C_i^{\tau-1} \in \bar{C}^{\tau-1} \text{ and } C_i^{\tau-1} \notin \bar{C}^\tau. \end{cases} \quad (4.13)$$

where  $\oplus$ : is the aggregation operation with missing data;  $\oplus$ : is the binary linguistic aggregation operation;  $i, j = 1, \dots, n_\tau^u; i \neq j$ .

In Equation (4.13), three cases are considered corresponding to three dynamic states of the criterion that set: a criterion newly appeared at the current time point, a criterion that was maintained across two time points, and a criterion that is only in historical data but no longer appearing at the current time point. These cases are computed as follows: Case 1 of Equation (4.13) is implemented according to Equation (4.4), Case 2 according to Equations (4.2) and (4.3), and Case 3 according to Equations (4.5) and (4.6). This organisation allows the model to flexibly adapt to changes in the criterion set over time, while ensuring that historical data are not completely discarded, but are integrated in a controlled manner through the linguistic aggregation and semantic decay mechanisms.

**Step 1.4.** Construct the fuzzy pairwise semantic comparison matrix.

Definition 4.1 is applied to convert the linguistic elements in the matrix  $\overline{AP}^\tau$  into fuzzy semantic representations through the mapping function  $fz(\cdot)$ , whereby each linguistic term is assigned a corresponding TFS set. Then, the fuzzy semantic pairwise comparison matrix is denoted by  $\overline{FAP}^\tau$  and is determined by Equation (4.14) as follows:

$$\overline{FAP}^\tau = fz(\tilde{p}_{ij}^*(\tau))_{n_\tau^u \times n_\tau^u} = \begin{pmatrix} fz(\tilde{p}_{11}^*(\tau)) & fz(\tilde{p}_{12}^*(\tau)) & \dots & fz(\tilde{p}_{1n_\tau^u}^*(\tau)) \\ fz(\tilde{p}_{21}^*(\tau)) & fz(\tilde{p}_{22}^*(\tau)) & \dots & fz(\tilde{p}_{2n_\tau^u}^*(\tau)) \\ \vdots & \vdots & \ddots & \vdots \\ fz(\tilde{p}_{n_\tau^u 1}^*(\tau)) & fz(\tilde{p}_{n_\tau^u 2}^*(\tau)) & \dots & fz(\tilde{p}_{n_\tau^u n_\tau^u}^*(\tau)) \end{pmatrix} \quad (4.14)$$

Where  $fz(\tilde{p}_{ij}^*(\tau)) = (l_{ij}(\tau), m_{ij}(\tau), u_{ij}(\tau)), (fz(\tilde{p}_{ij}^*(\tau)))^{-1} = (\frac{1}{u_{ij}(\tau)}, \frac{1}{m_{ij}(\tau)}, \frac{1}{l_{ij}(\tau)})$ , with  $i, j = 1, \dots, n_\tau^u, i \neq j$ ; the function  $fz(\cdot)$  is TFS set mapping of a linguistic term.

It should be noted that, to avoid a zero denominator when computing inverse elements in the fuzzy pairwise comparison matrix, the semantic values are renormalized from the interval  $[0, 1]$  to the interval  $[0.1, 1]$ . This renormalisation does not change the semantic order relation among the linguistic terms but helps to ensure numerical stability during computation.

**Step1.5.** Determine the fuzzy weights and normalised weights of criteria using Buckley's geometric mean method.

Based on the fuzzy semantic pairwise comparison matrix  $\overline{FAP}^\tau$ , the weights of criteria at time point  $\tau$  are computed using the fuzzy geometric mean proposed by Buckley. First, for each criterion  $C_j^\tau$ , the fuzzy geometric mean of the matrix  $\overline{FAP}^\tau$  is computed according to Equation (4.15):

$$RAP_j^\tau = \left( \prod_{j=1}^{n_\tau^u} fz(\tilde{p}_{ij}^*(\tau)) \right)^{1/n_\tau^u}, j = 1, 2, \dots, n_\tau^u. \quad (4.15)$$

Where  $RAP_j^\tau$  represents the aggregated fuzzy priority degree of the  $j^{th}$  criterion relative to all other criteria, and  $fz(\tilde{p}_{ij}^*(\tau))$  is TFS at position  $(i, j)$  of the matrix  $\overline{FAP}^\tau$ .

Next, the fuzzy weight of the  $j^{th}$  criterion, denoted  $F\tilde{w}_j(\tau)$ , is determined by normalizing  $RAP_j^\tau$  by the sum of the fuzzy geometric means of all criteria according to Equation (4.16):

$$F\tilde{w}_j(\tau) = \overline{RAP}_j^\tau(\tau) \otimes \left( \overline{RAP}_1^\tau(\tau) \oplus \overline{RAP}_2^\tau(\tau) \oplus \dots \oplus \overline{RAP}_{n_\tau^u}^\tau(\tau) \right)^{-1}. \quad (4.16)$$

where  $\oplus$ : denotes the addition of TFS sets;  $\otimes$ : is used to multiply the  $j^{th}$  triangular fuzzy number by the fuzzy inverse of the sum; and  $(\cdot)^{-1}$ : is the multiplicative inverse of a positive fuzzy number.

The result obtained is a TFS of the form:  $F\tilde{w}_j(\tau) = (l_{F\tilde{w}_j}(\tau), m_{F\tilde{w}_j}(\tau), u_{F\tilde{w}_j}(\tau))$

For the subsequent computational step in the L-FTOPSIS-HA model, the fuzzy weights are defuzzified to obtain crisp values. In the thesis, defuzzification is performed using the COA method. Specifically, for each fuzzy weight  $F\tilde{w}_j(\tau)$ , the corresponding crisp value  $M_j(\tau)$  is determined by:

$$M_j(\tau) = \frac{l_{F\tilde{w}_j}(\tau) + m_{F\tilde{w}_j}(\tau) + u_{F\tilde{w}_j}(\tau)}{3} \quad (4.17)$$

The value  $M_j(\tau)$  is the crisp representative of the fuzzy weight  $F\tilde{w}_j(\tau)$ . Finally, the values  $M_j(\tau)$  are normalized to obtain the criterion weight vector  $\tilde{w}_j(\tau)$  at time point  $\tau$ , ensuring that the sum of weights is equal to 1. This normalised weight vector is the output of Stage 1 and is used as the input for Stage 2 of the proposed model, namely, the stage of ranking alternatives using the L-FTOPSIS-HA method.

#### 4.4.2. Evaluation and ranking of alternatives.

**Step 2.1.** Construct the linguistic evaluation matrix for alternatives at time point  $\tau$ .

At time point  $\tau$ , each expert  $D_e^\tau$  uses linguistic terms  $x \in X_k$  on the linguistic scale  $L_k(\bar{A}^\tau)$ , to construct the linguistic evaluation matrix for alternatives, denoted by  $\tilde{A}_{D_e^\tau}^\tau$ , with the structure given in Equation (4.8). In this matrix, each element represents the linguistic evaluation by expert  $D_e^\tau$  of the performance level of alternative  $A_i^\tau \in \bar{A}^\tau$  with respect to each criterion  $C_j^\tau \in \bar{C}^\tau$  at time point  $\tau$ .

**Step 2.2.** Aggregate the experts' linguistic evaluation matrices at time point  $\tau$ .

The linguistic evaluations of the entire expert panel at time point  $\tau$  are aggregated to obtain the group alternative evaluation matrix, denoted by  $\tilde{A}\tilde{A}^\tau$ , according to Equation (4.18):

$$\tilde{A}\tilde{A}^\tau = (\tilde{a}_{ij}(\tau))_{m_\tau \times n_\tau}. \quad (4.18)$$

Where each element  $\tilde{a}_{ij}(\tau)$  of the matrix  $A\tilde{A}^\tau$  is determined through the binary linguistic aggregation operation in the 4-tuple semantic domain according to Equation (4.19):

$$\tilde{a}_{ij}(\tau) = a_{ij}^1(\tau) \oplus a_{ij}^2(\tau) \oplus \dots \oplus a_{ij}^{h_\tau}(\tau), i = 1, \dots, m_\tau; j = 1, \dots, n_\tau \quad (4.19)$$

The matrix  $A\tilde{A}^\tau$  reflects the collective evaluation of the expert panel regarding the performance levels of alternatives at time point  $\tau$ .

**Step 2.3.** Construct the integrated alternative evaluation matrix incorporating historical data.

This step is applied when  $\tau > 1$ . To integrate current evaluations with historical data, let:  $\bar{A}_u^\tau = \bar{A}^{\tau-1} \cup \bar{A}^\tau$ ; and  $\bar{C}_u^\tau = \bar{C}^{\tau-1} \cup \bar{C}^\tau$  denote the sets of alternatives and criteria appearing at two consecutive time points  $\tau - 1$  and  $\tau$ , with cardinalities  $m_\tau^u$  and  $n_\tau^u$ , respectively. Similarly, historical integration is also considered simultaneously for the set of criteria. Then, the elements of the integrated alternative evaluation matrix that incorporates historical data, denoted by  $\overline{A\tilde{A}^\tau} = (\tilde{a}_{ij}^*(\tau))_{m_\tau^u \times n_\tau^u}$ , are determined by Equation (4.20) as follows:

$$\tilde{a}_{ij}^*(\tau) = \begin{cases} \emptyset \oplus \tilde{a}_{ij}(\tau) & \text{if } \begin{cases} A_i^\tau \in \bar{A}^\tau \setminus \bar{A}^{\tau-1} \text{ and } C_j^\tau \in \bar{C}^\tau \setminus \bar{C}^{\tau-1} \\ A_i^\tau \in \bar{A}^{\tau-1} \setminus \bar{A}^\tau \text{ and } C_j^\tau \in \bar{C}^\tau \setminus \bar{C}^{\tau-1} \end{cases} \\ \tilde{a}_{ij}(\tau - 1) \oplus \tilde{a}_{ij}(\tau) & \text{if } A_i^\tau \in \bar{A}^\tau \cap \bar{A}^{\tau-1} \text{ and } C_j^\tau \in \bar{C}^\tau \cap \bar{C}^{\tau-1} \\ \tilde{a}_{ij}(\tau - 1) \oplus \emptyset & \text{if } A_i^\tau \in \bar{A}^{\tau-1} \setminus \bar{A}^\tau \text{ and } C_j^\tau \in \bar{C}^{\tau-1} \setminus \bar{C}^\tau \\ \emptyset \oplus \emptyset := \tilde{w} & \text{if } A_i^\tau \in \bar{A}^\tau \setminus \bar{A}^{\tau-1} \text{ and } C_j^\tau \in \bar{C}^{\tau-1} \setminus \bar{C}^\tau. \end{cases} \quad (4.20)$$

Where,  $\oplus$ : is the aggregation operation with missing data;  $\oplus$ : is the binary linguistic aggregation operation;  $\tilde{w}$  is the 4-tuple semantic representation of the neutral element  $w$  in  $L_k$ ;  $i = 1, \dots, m_\tau^u$ ,  $j = 1, \dots, n_\tau^u$ .

In Equation (4.20), four cases are considered corresponding to four dynamic states of the alternative set: an alternative newly appearing at the current time point, an alternative maintained across two consecutive time points, an alternative remaining only in historical data but no longer appearing at the current time point, and the case in which data are missing at both time points due to the addition of a new alternative and the removal of a criterion at the current time point. Therefore, Case 1 of Equation (4.20) is implemented according to Equation (4.4), Case 2 is computed according to Equation (4.2) and (4.3), Case 3 is computed according to Equations (4.5) and (4.6), and Case 4 is intended to ensure the completeness of the formula when the alternative set and the set change simultaneously. Considering the extended matrix on the union domain  $\bar{A}_u^\tau$  and  $\bar{C}_u^\tau$ , for all  $A_i^\tau \in \bar{A}_u^\tau$  and  $C_j \in \bar{C}_u^\tau$ , if  $A_i^\tau \in \bar{A}^\tau \setminus \bar{A}^{\tau-1}$ , and simultaneously  $C_j^\tau \in \bar{C}^{\tau-1} \setminus \bar{C}^\tau$ , then both the historical value and the current value at cell (i, j) are missing. In that case, the convention is adopted:  $\emptyset \oplus \emptyset := \tilde{w} = (w, S_k(w), \mathfrak{g}(w), r_w)$ . This convention is intended only to ensure the completeness of the matrix and the semantic consistency of the HA-based 4-tuple representation. When ranking at time point  $\tau$ , criteria not belonging to  $\bar{C}^\tau$  do not participate in the current decision matrix; therefore, the value  $\tilde{w}$  serves only as a structural trace and does not affect the ranking results at time point  $\tau$ .

**Step 2.4.** Convert the linguistic evaluation matrix integrated with historical data into a fuzzy evaluation matrix.

Similar to Step 1.4, Definition 4.1 is applied to convert the linguistic elements in the matrix  $\overline{A\tilde{A}^\tau}$  into fuzzy semantic representations using the mapping function  $fz(\cdot)$ , whereby each linguistic term is assigned a corresponding TFS set. Then, the fuzzy semantic evaluation matrix is denoted by  $\overline{FA\tilde{A}^\tau}$  and is determined by Equation (4.21) as follows:

$$\overline{FA\tilde{A}^\tau} = fz(\tilde{a}_{ij}^*(\tau))_{m_\tau^u \times n_\tau^u}, i = 1, \dots, m_\tau^u, j = 1, \dots, n_\tau^u. \quad (4.21)$$

**Step 2.5.** Compute the weighted aggregated evaluation values of the alternatives.

Based on the fuzzy evaluation matrix  $\overline{FA\tilde{A}^\tau}$  and the fuzzy weights of criteria  $F\tilde{w}_j(\tau)$  obtained in Step 1.5, the weighted aggregated evaluation value of the  $i^{th}$  at time point  $\tau$  denoted by  $\Upsilon_i^\tau$  is determined according to Equation (4.22):

$$Y_i^\tau = \frac{1}{n_\tau^u} \sum_{j=1}^{n_\tau^u} F \tilde{w}_j(\tau) \times fz(\tilde{a}_{ij}^*(\tau)), i = 1, \dots, m_\tau^u; j = 1, \dots, n_\tau^u. \quad (4.22)$$

**Step 2.6.** Determine the distances from alternatives to PIS ( $\mathcal{B}^+$ ) and NIS ( $\mathcal{B}^-$ ) at time point  $\tau$ , which are computed respectively by Equations (4.23) and (4.24):

$$d_i^+(\tau) = \sqrt{(Y_i^\tau - \mathcal{B}^+)^2}, i = 1, \dots, m_\tau^u, \quad (4.23)$$

$$d_i^-(\tau) = \sqrt{(Y_i^\tau - \mathcal{B}^-)^2}, i = 1, \dots, m_\tau^u. \quad (4.24)$$

where,  $\mathcal{B}^+ = (1, 1, 1)$ ,  $\mathcal{B}^- = (0, 0, 0)$ ;  $d_i^+(\tau)$  and  $d_i^-(\tau)$  represents the shortest and farthest distances of  $A_i^\tau$  at time point  $\tau$ .

**Step 2.7.** Compute the closeness coefficient and ranking of the alternatives.

The closeness coefficient of alternatives  $A_i^\tau$  at time point  $\tau$  is computed by Equation (4.25). The alternatives are ranked based on the value of  $CC_i^\tau$ ; the larger  $CC_i^\tau$  is, the higher the priority of the alternative.

$$CC_i^\tau = \frac{d_i^-(\tau)}{d_i^+(\tau) + d_i^-(\tau)} \quad (4.25)$$

#### 4.5. Experimental example

**Table 4.1.** Scenario for evaluating the digital transformation readiness of SMEs in Vietnam

| Time point | Set of alternatives and context              | Set of criteria                                                                                                 |
|------------|----------------------------------------------|-----------------------------------------------------------------------------------------------------------------|
| $\tau_1$   | E1, E2, E3, E4, and E5.                      | 3 main criteria and 11 sub-criteria.                                                                            |
| $\tau_2$   | Newly added alternatives: E6, E7, and E8.    | 4 main criteria and 14 sub-criteria, including one additional main criterion and three additional sub-criteria. |
| $\tau_3$   | E1 and E2 are removed; E9 and E10 are added. | 4 main criteria and 14 sub-criteria.                                                                            |

**Table 4.2.** Ranking results of the readiness for digital transformation of SMEs in Vietnam at three time points  $\tau_1, \tau_2, \tau_3$

| Time Point | Ranking result                                     | Comment                                                                                                                                                                        |
|------------|----------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| $\tau_1$   | $E2 > E1 > E3 > E5 > E4$                           | E2 is the best-performing alternative; the ranking reasonably reflects the differentiation between the leading group and the bottom group.                                     |
| $\tau_2$   | $E6 > E2 > E1 > E7 > E3 > E8 > E5 > E4$            | When the sets of alternatives and criteria change, E6 rises to the first position due to its stronger current evaluation result.                                               |
| $\tau_3$   | $E2 > E9 > E6 > E7 > E10 > E8 > E5 > E1 > E4 > E3$ | Although E2 is not newly evaluated at time point $\tau_3$ , it still ranks first because the semantic decay mechanism reasonably preserves E2's strong historical performance. |

#### 4.6. Sensitivity Analysis and result comparison

##### 4.6.1. Sensitivity Analysis

Sensitivity analysis was performed at three time points  $\tau_1, \tau_2, \tau_3$  under two scenarios to examine the stability and sensitivity of the proposed model. *Scenario 1:* A global perturbation is introduced by shifting all linguistic pairwise comparison judgments upward by one level on the scale  $L_2(\bar{C}^\tau)$ .

*Scenario 2:* The global perturbation is combined with a change in the fuzziness measure of HA.

**Table 4.3.** Comparison of sensitivity analysis results under different scenarios at three time points  $\tau$ ,  $\tau_2$  and  $\tau_3$ 

| Time point | Initial result                                                                         | Sensitivity analysis results under the two scenarios                                   |                                                                                                                   |
|------------|----------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------|
|            |                                                                                        | Scenario 1                                                                             | Scenario 2                                                                                                        |
| $\tau_1$   | $E2 \succ E1 \succ E3 \succ E5 \succ E4$                                               | $E2 \succ E1 \succ E3 \succ E5 \succ E4$                                               | $E2 \succ E1 \succ E3 \succ E5 \succ E4$                                                                          |
| $\tau_2$   | $E6 \succ E2 \succ \mathbf{E1} \succ \mathbf{E7} \succ E3 \succ E8 \succ E5 \succ E4$  | $E6 \succ E2 \succ \mathbf{E7} \succ \mathbf{E1} \succ E3 \succ E8 \succ E5 \succ E4$  | $E6 \succ E2 \succ \mathbf{E7} \succ \mathbf{E1} \succ E3 \succ E8 \succ E5 \succ E4$                             |
| $\tau_3$   | $E2 \succ E9 \succ E6 \succ E7 \succ E10 \succ E8 \succ E5 \succ E1 \succ E4 \succ E3$ | $E2 \succ E9 \succ E6 \succ E7 \succ E10 \succ E8 \succ E5 \succ E1 \succ E4 \succ E3$ | $E2 \succ E9 \succ E6 \succ E7 \succ E10 \succ \mathbf{E1} \succ \mathbf{E8} \succ \mathbf{E5} \succ E4 \succ E3$ |

The general results under the two scenarios in Table 4.3 show that the SME rankings vary only slightly, while the leading and bottom groups remain almost unchanged at all three time points.

In summary, the sensitivity analysis under the two scenarios shows that the dynamic L-FAHP - L-FTOPSIS model based on HA exhibits high stability, good reliability, and clear interpretability in evaluating the readiness for digital transformation of retail SMEs in Vietnam.

#### 4.6.2. Result comparison

**Table 4.4.** Comparison between the proposed model and the dynamic FAHP-FTOPSIS model at three time points  $\tau_1$ ,  $\tau_2$  và  $\tau_3$ 

| Time point | Dynamic FAHP-FTOPSIS                                                                   | Proposed model                                                                         |
|------------|----------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------|
| $\tau_1$   | $E2 \succ E1 \succ E5 \succ E4 \succ E3$                                               | $E2 \succ E1 \succ E3 \succ E5 \succ E4$                                               |
| $\tau_2$   | $E6 \succ E2 \succ E8 \succ E7 \succ E1 \succ E5 \succ E3 \succ E4$                    | $E6 \succ E2 \succ E7 \succ E1 \succ E3 \succ E8 \succ E5 \succ E4$                    |
| $\tau_3$   | $E9 \succ E10 \succ E6 \succ E2 \succ E8 \succ E7 \succ E1 \succ E5 \succ E4 \succ E3$ | $E2 \succ E9 \succ E6 \succ E7 \succ E10 \succ E8 \succ E5 \succ E1 \succ E4 \succ E3$ |

The results in Table 4.4 show that the two models produce consistent results for several outstanding alternatives, such as  $E2$  at time point  $\tau_1$  and  $E6$  at time point  $\tau_2$ . However, at time point  $\tau_3$  the comparison model based on fuzzy set theory uses the aggregation of arithmetic mean over the three time points and prioritises the alternative with a strong current evaluation result; therefore,  $E9$  ranks first.

In contrast, the proposed model ranks  $E2$  first based on its strong historical performance at  $\tau_1$  and  $\tau_2$ . This difference shows that the proposed model is capable of integrating historical data with current data through a mechanism that allows the influence of historical data to decay reasonably over time. Accordingly, the proposed model not only ensures the reliability and stability of the ranking results, but also better reflects the maturity-orientated nature of the digital transformation process in a dynamic decision-making environment.

#### 4.7. Conclusion of Chapter 4

In this chapter, the thesis has presented several extended theoretical components of HA for the dynamic linguistic MCGDM problem in order to address changes in the set of criteria, the set of alternatives, and decision makers over time. Specifically, this chapter has presented two aggregation operations: the binary linguistic aggregation operation on the 4-tuple semantic scale and the aggregation operation with missing data, which incorporates a semantic decay mechanism to integrate historical data in the HA environment. Next, the dynamic L-FAHP - L-FTOPSIS decision-making model based on HA has been developed. Finally, the proposed model has been applied to assess the digital transformation readiness of SMEs in Vietnam's retail sector to illustrate the practical applicability of the proposed theoretical framework and the decision-making model.

## CONCLUSION

The thesis has achieved its objective of developing MCGDM models in HA-based linguistic information environments, with HA serving as the foundation for representation, semantic quantification, preservation of semantic order, aggregation of assessments, and ranking of alternatives. By systematising the theoretical foundations and the identification of research gaps, the thesis establishes a methodological framework in which the three main results progressively inherit from and extend one another as follows.

First, the HA-TOPSIS model was developed to rank alternatives within a semantic domain. In this model, an HA-based linguistic scale and the SQM function are used to determine PIS/NIS, measure semantic distance, and compute the closeness coefficient. The model was validated using a supplier selection problem. Together with sensitivity analysis and comparisons with the 4-tuple linguistic model and FTOPSIS, the validation results indicate that the proposed model preserves semantic order, thereby contributing to the stability, reliability, and interpretability of DM results (Chapter 2).

Second, the thesis developed an integrated MCGDM model, namely EFAHP-4-tuple-HA, to determine criterion weights and rank alternatives in a linguistic environment. The core result of this model is a consistency control mechanism for linguistic pairwise comparison matrices in the HA environment, on the basis of which the criterion weights are determined using EFAHP-HA. Subsequently, the 4-tuple linguistic model is applied to represent and aggregate expert assessments and rank alternatives. The representation and processing of data on the same HA-based foundation help ensure semantic connectivity between the two stages of criterion weight determination and alternative ranking. The application of the model to the assessment of the readiness for digital transformation of 16 SMEs in Vietnam's retail sector shows that the model preserves the semantic structure of the input data and maintains the stability of the ranking under sensitivity analysis and comparative analysis with HA-TOPSIS and FAHP-FTOPSIS (Chapter 3)

Third, the thesis developed a methodological framework for constructing dynamic MCGDM models in HA-based linguistic information environments. Based on 4-tuple linguistic representation, the thesis developed a binary linguistic aggregation operator, an aggregation operator for missing data, and a semantic decay mechanism to integrate historical data with current data. From this methodological framework, the thesis constructed an HA-based dynamic L-FAHP-L-FTOPSIS model to determine the weights and rank alternatives consistently at multiple time points. The model was validated through the assessment of the readiness for digital transformation of SMEs in a dynamic context, together with sensitivity analysis and comparison with a related method. The experimental results show that the model can reflect temporal changes in the assessment data while maintaining semantic interpretability within the constructed scenarios (Chapter 4).

In addition to these results achieved, the research still has several limitations that require further investigation. First, the models were mainly constructed within the scope of linear HA, with semantic parameter sets determined by assumptions or experimental configurations. Second, the scope of experimental validation remains limited in terms of the number of application domains, datasets, and problem scales. Third, the proposed models have not yet been implemented as a complete decision support system to validate their deployability in a real-world operational environment.

Given these limitations, several future research directions can be pursued. First, mechanisms to optimise and calibrating the semantic parameters of HA should be investigated based on empirical data or expert feedback to reduce dependence on traditional survey-based methods. Next, the proposed models can be extended to large-scale MCGDM problems, as well as to problems involving incomplete, heterogeneous, or real-time updated data. Finally, a decision support system based on the proposed models should be developed and its applications should be extended to domains such as risk assessment, project management, healthcare, education, finance and environmental management.

**CÁC CÔNG TRÌNH KHOA HỌC CỦA TÁC GIẢ ĐÃ CÔNG BỐ  
CÓ LIÊN QUAN ĐẾN LUẬN ÁN**

- [CT1]. Kien, N.V.; Thong, H.V.; Ho, C.N. (2025). “A Multi-Criteria Decision-Making Method Combining Hedge Algebras and the TOPSIS Method”. *Proceedings of the Fif<sup>th</sup> International Conference on Intelligent Systems and Networks*. ICISN, 22–23 March, Hanoi, Vietnam. [https://doi.org/10.1007/978-981-95-1746-6\\_34](https://doi.org/10.1007/978-981-95-1746-6_34) (Kỷ yếu Hội nghị Quốc tế ICISN 2025, ISSN: 2367-3389, Springer, **Scopus, Q4**).
- [CT2]. Kien, N.V.; Thong, H.V.; Ho, N.C.; Dat, L.Q. (2025). “A Novel Integrated Fuzzy Analytic Hierarchy Process with a 4-Tuple Hedge Algebra Semantics for Assessing the Level of Digital Transformation of Enterprises”. *Mathematics*, 13(21), 3539. <https://doi.org/10.3390/math13213539> (ISSN: 2227-7390, SCIE, 2025 IF = 2.3, WoS & Scopus **Q1 (Top 10%)**).
- [CT3]. Thong, H.V.; Dat, L.Q.; Ho, N.C.; Kien, N.V. (2026). “A Novel Linguistic Framework for Dynamic Multi-Criteria Group Decision-Making Using Hedge Algebras”. *Appl. Sci.*, 16, 30. <https://doi.org/10.3390/app16010030> (ISSN: 2076-3417, SCIE, 2026 IF = 2.5, **Q2**; *Correspondence: kiennv@hau.edu.vn*).