

BỘ GIÁO DỤC VÀ ĐÀO TẠO
TRƯỜNG ĐẠI HỌC GIAO THÔNG VẬN TẢI

NHỮ VĂN KIÊN

NGHIÊN CỨU PHÁT TRIỂN CÁC MÔ HÌNH
RA QUYẾT ĐỊNH NHÓM ĐA TIÊU CHÍ
TRONG MÔI TRƯỜNG THÔNG TIN NGÔN NGỮ
DỰA TRÊN ĐẠI SỐ GIA TỬ

NGÀNH: CÔNG NGHỆ THÔNG TIN
MÃ SỐ: 9480201

TÓM TẮT LUẬN ÁN TIẾN SĨ

HÀ NỘI - 2026

Công trình được hoàn thành tại:
Trường Đại học Giao thông Vận tải

Người hướng dẫn khoa học:

1. TS. Hoàng Văn Thông

Trường Đại học Giao thông Vận tải

2. PGS. TSKH. Nguyễn Cát Hồ

Đại học Duy Tân

Phản biện 1:

Phản biện 2:

Phản biện 3:

Luận án được bảo vệ trước Hội đồng chấm luận án cấp Trường
tại Trường Đại học Giao thông Vận tải
vào hồi giờ phút ngày tháng năm 2026

Có thể tìm hiểu Luận án tại thư viện:

- Thư viện Quốc Gia Việt Nam

- Thư viện Trường Đại học Giao thông Vận tải

MỞ ĐẦU

1. Tính cấp thiết của đề tài

Ra quyết định trong kỹ thuật và hệ thống thông tin thường phải lựa chọn hoặc xếp hạng phương án theo nhiều tiêu chí như chi phí, chất lượng, khả năng tích hợp, độ tin cậy và an toàn. Khi có nhiều chuyên gia, nhà quản lý hoặc nhóm liên quan cùng tham gia, bài toán trở thành MCGDM. Nhiều tiêu chí trong thực tế mang tính định tính như đánh giá mức độ phù hợp công nghệ, chất lượng dịch vụ, mức độ sẵn sàng chuyển đổi số hoặc rủi ro. Chuyên gia thường đánh giá các tiêu chí này bằng từ ngôn ngữ thay vì giá trị số rõ. Tuy nhiên, thông tin ngôn ngữ có tính mơ hồ, không chắc chắn, chủ quan và phụ thuộc ngữ cảnh, nên cần một mô hình có khả năng bảo toàn cấu trúc ngữ nghĩa trong toàn bộ quy trình ra quyết định. Do đó, cần mô hình MCGDM ngôn ngữ có khả năng bảo toàn cấu trúc ngữ nghĩa trong toàn bộ quy trình ra quyết định.

Các tiếp cận dựa trên tập mờ, FAHP, FTOPSIS, 2-tuple và 4-tuple đã có đóng góp quan trọng. Tuy nhiên, nhiều mô hình vẫn phụ thuộc vào hàm thuộc, số mờ, giải mờ, ánh xạ xấp xỉ hoặc lớp biểu diễn khác nhau, nên có thể làm giảm tính minh bạch và tính thống nhất ngữ nghĩa. Đại số gia tử (HA) cung cấp cấu trúc toán học để mô hình hóa miền từ theo trật tự ngữ nghĩa; độ đo tính mờ và hàm định lượng ngữ nghĩa cho phép định lượng từ ngôn ngữ, đồng thời góp phần bảo toàn ngữ nghĩa.

Từ tổng quan nghiên cứu, cho thấy ba khoảng trống cần được làm rõ như sau: TOPSIS/FTOPSIS chưa khai thác đầy đủ cấu trúc ngữ nghĩa của HA khi xác định PIS/NIS và xếp hạng; MCGDM dựa trên HA chưa tích hợp đầy đủ việc xác định trọng số của tiêu chí có kiểm soát nhất quán với xếp hạng; MCGDM ngôn ngữ dựa trên HA còn chủ yếu xét bối cảnh tĩnh, trong khi bài toán thực tiễn biến đổi theo thời gian, có dữ liệu lịch sử và dữ liệu khuyết thiếu. Vì vậy, phát triển các mô hình MCGDM trong môi trường thông tin ngôn ngữ dựa trên HA là cần thiết, tạo cơ sở cho HA-TOPSIS, EFAHP - 4-tuple - HA và L-FAHP - L-FTOPSIS động dựa trên HA.

2. Mục tiêu nghiên cứu

Mục tiêu tổng quát của luận án là phát triển các mô hình MCGDM trong môi trường thông tin ngôn ngữ dựa trên HA, từ xây dựng thang đo, định lượng ngữ nghĩa, xác định trọng số của tiêu chí, kiểm soát tính nhất quán, kết nhập đánh giá nhóm đến xếp hạng phương án và mở rộng sang bối cảnh động. Các mục tiêu cụ thể gồm: xây dựng mô hình HA-TOPSIS; phát triển mô hình tích hợp EFAHP - 4-tuple - HA; xây dựng mô hình L-FAHP - L-FTOPSIS động dựa trên HA; kiểm chứng qua bài toán lựa chọn nhà cung cấp phần mềm và đánh giá mức độ sẵn sàng chuyển đổi số của SMEs.

3. Đối tượng, phạm vi và phương pháp nghiên cứu

Đối tượng nghiên cứu là bài toán MCGDM ngôn ngữ theo hướng tiếp cận HA.

Phạm vi nghiên cứu tập trung vào biểu diễn và định lượng ngữ nghĩa, xác định trọng số của tiêu chí có kiểm soát nhất quán, kết nhập ý kiến chuyên gia, xếp hạng phương án và xử lý bối cảnh động. Luận án không đặt mục tiêu bao quát toàn bộ MCDM/MCGDM, không kiểm chứng trên mọi miền ứng dụng và không phát triển hệ hỗ trợ ra quyết định hoàn chỉnh.

Phương pháp nghiên cứu gồm phân tích tài liệu, mô hình hóa toán học, phát triển thuật toán, phân tích tính chất, thực nghiệm, phân tích độ nhạy và đối sánh kết quả.

4. Cấu trúc của luận án

Luận án gồm phần Mở đầu, bốn chương nội dung, Kết luận, danh mục công trình công bố, tài liệu tham khảo và phụ lục. Phần Mở đầu trình bày tính cấp thiết, mục tiêu, nội dung, đối tượng, phạm vi, phương pháp nghiên cứu, đóng góp mới và bố cục của luận án.

Chương 1, tổng quan nghiên cứu về bài toán MCGDM trong môi trường thông tin ngôn ngữ; tổng hợp cơ sở lý thuyết về tập mờ, FAHP, FTOPSIS, mô hình 2-tuple, 4-tuple và HA, qua đó xác định khoảng trống nghiên cứu.

Chương 2, phát triển mô hình HA-TOPSIS cho MCGDM ngôn ngữ, gồm xây dựng thang đo HA, định lượng ngữ nghĩa, xác định PIS/NIS, đo khoảng cách ngữ nghĩa, tính hệ số gần và xếp hạng phương án; gắn với [CT1].

Chương 3, trình bày mô hình tích hợp EFAHP - 4-tuple - HA nhằm xác định trọng số của tiêu chí có kiểm soát tính nhất quán trong môi trường HA, kết nhập ý kiến chuyên gia và xếp hạng trên miền ngữ nghĩa; gắn với [CT2].

Chương 4 mở rộng sang MCGDM động với mô hình L-FAHP - L-FTOPSIS dựa trên HA, xử lý biến đổi của tập phương án, bộ tiêu chí, tập chuyên gia, dữ liệu khuyết thiếu và tích hợp dữ liệu lịch sử theo cơ chế suy giảm ngữ nghĩa; gắn với [CT3].

Kết luận tổng hợp kết quả đạt được, chỉ ra đóng góp, hạn chế và hướng phát triển. Các phụ lục cung cấp bảng số liệu, kết quả tính toán và minh chứng thực nghiệm phục vụ kiểm chứng mô hình.

CHƯƠNG 1. TỔNG QUAN NGHIÊN CỨU VỀ BÀI TOÁN MCGDM NGÔN NGỮ VÀ CƠ SỞ LÝ THUYẾT DỰA TRÊN ĐẠI SỐ GIA TỬ

1.1. Bài toán MCGDM ngôn ngữ

Bài toán MCGDM ngôn ngữ được phát biểu như sau:

(i) *Đầu vào của bài toán:*

Cho ba tập gồm tập phương án $\bar{A} = \{A_1, A_2, \dots, A_m\}$, ($m \geq 2$); tập tiêu chí $\bar{C} = \{C_1, C_2, \dots, C_n\}$, ($n \geq 2$) và tập chuyên gia $\bar{D} = \{D_1, D_2, \dots, D_h\}$, ($h \geq 2$).

Với mỗi chuyên gia $D_e \in \bar{D}$ đánh giá phương án $A_i \in \bar{A}$ theo tiêu chí $C_j \in \bar{C}$ được biểu diễn dưới dạng ma trận đánh giá quyết định:

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

trong đó, a_{ij}^e là đánh giá ngôn ngữ của chuyên gia $D^e \in \bar{D}$ đối với phương án $A_i \in \bar{A}$ trên tiêu chí $C_j \in \bar{C}$; với $e = 1, \dots, h$; $i = 1, \dots, m$; $j = 1, \dots, n$.

Khi bài toán có xét tầm quan trọng tương đối của các tiêu chí, ta có véc-tơ trọng số:

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

Trong phạm vi luận án, đánh giá các phương án và tầm quan trọng của tiêu chí đều được biểu diễn bằng các từ ngôn ngữ. Do đó, mô hình phải xử lý tính mơ hồ, không chắc chắn, chủ quan và cấu trúc ngữ nghĩa của thông tin đầu vào, đồng thời bảo toàn trật tự ngữ nghĩa khi kết nhập và xếp hạng.

(i) *Đầu ra của bài toán:* Xếp hạng các phương án $A_i \in \bar{A}$

Luận án tập trung vào bài toán lựa chọn và xếp hạng, với hai giai đoạn cốt lõi là xác định trọng số của tiêu chí và xếp hạng phương án trong bối cảnh tĩnh và động. Trong bối cảnh động, tập phương án, bộ tiêu chí, tập chuyên gia và dữ liệu đánh giá có thể thay đổi theo thời gian, kèm dữ liệu lịch sử, khuyết thiếu. Vì vậy, HA được sử dụng làm nền tảng biểu diễn và kết nhập thông tin ngôn ngữ để phát triển khung phương pháp luận cho bài toán MCGDM ngôn ngữ theo hướng nối tiếp, có tính kế thừa và mở rộng: từ mô hình HA-TOPSIS xếp hạng phương án trên miền ngữ nghĩa, đến kế thừa phát triển mô hình tích hợp EFAHP - 4-tuple - HA và mở rộng sang bối cảnh động xây dựng mô hình L-FAHP - L-FTOPSIS động dựa trên HA.

1.2. Cơ sở lý thuyết nền tảng

1.2.1. Lý thuyết tập mờ và biến ngôn ngữ

Tập mờ của Zadeh biểu diễn mức độ thuộc của phần tử qua hàm thuộc, tạo cơ sở mô hình hóa thông tin mơ hồ trong MCDM và MCGDM. Biến ngôn ngữ cho phép từ ngôn ngữ trở thành giá trị của biến, hỗ trợ tính toán với từ. Tuy nhiên, ngữ nghĩa của từ phụ thuộc vào hàm thuộc và tham số thiết kế, tạo động lực sử dụng các mô hình tính toán với từ để xử lý ngữ nghĩa chặt chẽ hơn.

1.2.2. Các mô hình tính toán với từ trong môi trường ngôn ngữ

Các mô hình tính toán với từ là cơ sở xử lý đánh giá ngôn ngữ trong MCDM/MCGDM. Tiếp cận nguyên lý mở rộng thao tác trên hàm thuộc nhưng lấy xấp xỉ về nhân ngôn ngữ. Mô hình ký hiệu dùng chỉ số, thuận tiện tính toán song dễ mất thông tin khi làm tròn. Mô hình 2-tuple thường giả định thang đo ngôn ngữ phân bố đều và đối xứng, việc tính toán vẫn dựa trên chỉ số số thực và ánh xạ ngược về nhân nên thiếu cơ sở đại số để xử lý các quan hệ ngữ nghĩa. Các giới hạn này là cơ sở để luận án lựa chọn HA làm cơ sở lý thuyết nền tảng.

1.2.3. Lý thuyết Đại số gia tử

1.2.3.1. Định nghĩa HA

Định nghĩa 1.1. [53] Một HA $\mathcal{AX}$ của một biến ngôn ngữ X là một bộ gồm bốn thành phần: $\mathcal{AX} = (X, G, H, \leq)$. Trong đó: X là tập các từ ngôn ngữ của biến X , với $X \subseteq \text{Dom}(X)$; $G = \{g^-, g^+\}$ là tập các phần tử sinh, trong đó g^- và g^+ tương ứng là phần tử sinh nguyên thủy âm và dương, thỏa mãn quan hệ ngữ nghĩa $g^- \leq g^+$. H là tập các gia tử của biến ngôn ngữ X . $\leq$ là quan hệ thứ tự được cảm sinh bởi ngữ nghĩa vốn có của các từ ngôn ngữ trên tập X . Giả thiết rằng miền giá trị có chứa các phần tử hằng $(0, W, 1)$ lần lượt biểu thị phần tử bé nhất, phần tử trung hòa (neutral) và phần tử lớn nhất trong X , sao cho $0 \leq g^- \leq W \leq g^+ \leq 1$.

1.2.3.2. Độ đo tính mờ của giá trị ngôn ngữ

Định nghĩa 1.2. [52] Cho $\mathcal{AX}^* = (X, G, H, \sigma, \phi, \leq)$ là một HA tuyến tính đầy đủ (Com HA) của biến ngôn ngữ X . Một ánh xạ $f_{\mathcal{M}}: X \rightarrow [0, 1]$ được gọi là một độ đo tính mờ của các từ ngôn ngữ trong X nếu thỏa mãn các tính chất sau:

- i) $f_{\mathcal{M}}$ là một độ đo đầy đủ trên X , tức là: $f_{\mathcal{M}}(g^-) + f_{\mathcal{M}}(g^+) = 1$ và $\sum_{h \in H} f_{\mathcal{M}}(hu) = f_{\mathcal{M}}(u)$, $\forall u \in X$;
- ii) $f_{\mathcal{M}}(x) = 0$, $\forall x$ thỏa mãn $H(x) = \{x\}$; đặc biệt, $f_{\mathcal{M}}(0) = f_{\mathcal{M}}(W) = f_{\mathcal{M}}(1) = 0$;
- iii) $\forall x, y \in X$, $\forall h \in H$, $\mu(h) = \frac{f_{\mathcal{M}}(hx)}{f_{\mathcal{M}}(x)} = \frac{f_{\mathcal{M}}(hy)}{f_{\mathcal{M}}(y)}$, tỷ số này không phụ thuộc vào x và y , và nó được gọi là độ đo tính mờ của các gia tử, ký hiệu là $\mu(h)$.

Từ (i) và (iii) ta có công thức tính đệ quy độ đo tính mờ của $x = h_m \dots h_1 g$ với $g \in \{g^-, g^+\}$ như sau: $f_{\mathcal{M}}(x) = \mu(h_m) \dots \mu(h_1) f_{\mathcal{M}}(g)$, trong đó $\sum_{h \in H} \mu(h) = 1$.

1.2.3.3. Định lượng ngữ nghĩa của giá trị ngôn ngữ

Định nghĩa 1.3. [52] Cho $\mathcal{AX}^* = (X, G, H, \sigma, \phi, \leq)$ là một ComHA tuyến tính. Một ánh xạ $\vartheta: X \rightarrow [0, 1]$ được gọi là một SQM của $\mathcal{AX}^*$ nếu thỏa mãn các điều kiện sau:

- (i) ϑ là ánh xạ 1-1 từ tập X vào đoạn $[0, 1]$ sao cho $\vartheta(0) = 0$, $\vartheta(1) = 1$ và bảo toàn thứ tự ngữ nghĩa trên X ; tức là $\forall x, y \in X$, $x < y \Rightarrow \vartheta(x) < \vartheta(y)$;
- (ii) ϑ liên tục theo ngữ nghĩa: $\forall x \in X$, $\vartheta(\sigma x) = \infimum \vartheta(H(x))$ và $\vartheta(\phi x) = \supremum \vartheta(H(x))$.

Định nghĩa 1.4. [52] Cho $\mathcal{AX}^* = (X, G, H, \sigma, \phi, \leq)$ là một ComHA tuyến tính và $f_{\mathcal{M}}(x)$ và $\mu(h)$ lần lượt là một độ đo tính mờ của giá trị ngôn ngữ x và của gia tử h . Khi đó, một ánh xạ $\vartheta: X \rightarrow [0, 1]$ được cảm sinh bởi độ đo tính mờ được xác định như sau:

- (i) $\vartheta(W) = \theta = f_{\mathcal{M}}(g^-)$, $\vartheta(g^-) = \theta - \alpha f_{\mathcal{M}}(g^-) = \beta f_{\mathcal{M}}(g^-)$, $\vartheta(g^+) = \theta + \alpha f_{\mathcal{M}}(g^+)$;
- (ii) $\vartheta(h_j x) = \vartheta(x) + \zeta g(h_j x) \left\{ \sum_{i=\zeta g(j)}^{j-\zeta g(j)} \mu(h_i) f_{\mathcal{M}}(x) - \omega(h_j x) \mu(h_j x) f_{\mathcal{M}}(x) \right\}$ với mọi j , trong đó $-q \leq j \leq p$, $j \neq 0$, và $\omega(h_j x) = \frac{1}{2} [1 + \zeta g(h_j x) \zeta g(h_p h_j x) (\beta - \alpha)]$;
- (iii) $\vartheta(\phi g^-) = 0$, $\vartheta(\sigma g^-) = \theta = \vartheta(\phi g^+)$, $\vartheta(\sigma g^+) = 1$, và với mọi j , $-q \leq j \leq p$,
 $\vartheta(\phi h_j x) = \vartheta(x) + \zeta g(h_j x) \left\{ \sum_{i=\zeta g(j)}^{j-\zeta g(j)} \mu(h_i) f_{\mathcal{M}}(x) \right\} - \frac{1}{2} (1 - \zeta g(h_j x)) \mu(h_i) f_{\mathcal{M}}(x)$,
 $\vartheta(\sigma h_j x) = \vartheta(x) + \zeta g(h_j x) \left\{ \sum_{i=\zeta g(j)}^{j-\zeta g(j)} \mu(h_i) f_{\mathcal{M}}(x) \right\} + \frac{1}{2} (1 + \zeta g(h_j x)) \mu(h_i) f_{\mathcal{M}}(x)$.

1.2.3.4. Hệ khoảng tương tự ngữ nghĩa

Định nghĩa 1.5 [52] Cho một độ đo tính mờ f_m trên X , một giá trị $k > 0$, và tập hữu hạn các từ $X_k = \{x \in X: |x| \leq k\}$, ta xây dựng một tập các khoảng $\{S_k(x): x \in X_k\}$ trên miền tham chiếu đã chuẩn hóa $[0, 1]$, gọi là các khoảng tương tự bậc k (k -similarity intervals), thỏa mãn các điều kiện sau:

(i) Các khoảng này tạo thành một phân hoạch của $[0, 1]$;

(ii) Mỗi $S_k(x)$ chứa đúng một giá trị $\vartheta(x)$ thuộc tập $\{\vartheta(y): y \in X_k\}$, trong đó ϑ là một SQM được cảm sinh bởi f_m . Khi đó, mọi giá trị thuộc $S_k(x)$ được xem là tương tự với $\vartheta(x)$, hay nói cách khác, tương tự với ý nghĩa của từ ngôn ngữ x ở mức bậc k .

1.2.3.5. Thang đo ngôn ngữ biểu diễn theo mô hình 4-tuple ngữ nghĩa

Định nghĩa 1.6 [54] Cho một biến ngôn ngữ X và miền tham chiếu số đã chuẩn hóa $[0, 1]$, một biểu diễn 4-tuple ngữ nghĩa của từ ngôn ngữ $x \in \text{Dom}(X)$ được xác định bởi:

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

Trong đó: $x \in \text{Dom}(X)$: là một từ ngôn ngữ; $S_k(x) \subseteq [0, 1]$: là khoảng tương tự của x đối với tập từ X_k đã cho; $\vartheta(x) \in [0, 1]$: là giá trị định lượng ngữ nghĩa số của từ ngôn ngữ x ; $r_x \in S_k(x)$: là giá trị tham chiếu nằm trong khoảng tương tự.

Các khái niệm của HA như cấu trúc đại số ngôn ngữ, độ đo tính mờ, hàm SQM, hệ khoảng tương tự và thang đo 4-tuple ngữ nghĩa cho phép biểu diễn, định lượng, so sánh, kết nhập và bảo toàn cấu trúc ngữ nghĩa. Đây là nền tảng trực tiếp cho phát triển các mô hình MCGDM ngôn ngữ dựa trên HA trong các Chương 2, 3 và 4.

1.3. Các phương pháp nền mà luận án kế thừa và phát triển

Trong MCGDM, luận án kế thừa AHP/FAHP để xác định trọng số của tiêu chí bằng so sánh cặp và kiểm soát nhất quán; kế thừa TOPSIS/FTOPSIS để xếp hạng phương án theo nguyên lý gần PIS và xa NIS. Tuy nhiên, các tiếp cận này còn phụ thuộc vào số rõ, số mờ, hàm thuộc hoặc giải mờ, nên có thể thiếu thống nhất ngữ nghĩa. Vì vậy, luận án đã phát triển chúng trên nền HA, tạo cơ sở cho các mô hình HA-TOPSIS, EFAHP - 4-tuple - HA và L-FAHP - L-FTOPSIS động dựa trên HA.

1.4. Tổng quan nghiên cứu về MCDM và MCGDM ngôn ngữ

Các nghiên cứu MCDM/MCGDM truyền thống đã hình thành nhiều khuôn khổ lựa chọn, đánh giá và xếp hạng như AHP, TOPSIS, VIKOR, ELECTRE. Khi chuyển sang môi trường thông tin ngôn ngữ, các mở rộng dựa trên tập mờ như FAHP, FTOPSIS; các mô hình tính toán với từ như 2-tuple, 4-tuple góp phần xử lý tính mơ hồ, bất định và giảm mất mát thông tin. Tuy nhiên, nhiều mô hình còn phụ thuộc vào hàm thuộc, giải mờ, ánh xạ xấp xỉ hoặc thiếu thống nhất ngữ nghĩa giữa xác định trọng số, kết nhập và xếp hạng. Các mô hình MCGDM dựa trên HA cho thấy khả năng mô hình hóa trật tự ngữ nghĩa, định lượng và bảo toàn cấu trúc miền từ nhưng trọng số của tiêu chí thường được giả định trước hoặc do chuyên gia gán trực tiếp, cơ chế kiểm soát tính nhất quán chưa được tích hợp đầy đủ trong quy trình dựa trên HA và chủ yếu được xây dựng trong bối cảnh tĩnh.

1.5. Khoảng trống nghiên cứu và định hướng luận án

Tổng quan cho thấy MCDM và MCGDM trong môi trường ngôn ngữ đã được nghiên cứu qua AHP, TOPSIS, FAHP, FTOPSIS, 2-tuple, 4-tuple và HA. Tuy nhiên, còn ba khoảng trống chính: Thứ nhất, TOPSIS/FTOPSIS chưa khai thác đầy đủ trật tự ngữ nghĩa, độ đo tính mờ và hàm SQM của HA để xác định PIS/NIS, khoảng cách ngữ nghĩa và xếp hạng trên cùng miền ngữ nghĩa. Thứ hai, các mô hình HA hiện có chưa hoàn thiện quy trình xác định trọng số tiêu chí có kiểm soát tính nhất quán và liên thông với xếp hạng. Thứ ba, MCGDM dựa trên HA còn chủ yếu mô hình tĩnh, chưa xử lý đầy đủ dữ liệu lịch sử, khuyết thiếu và suy giảm ngữ nghĩa. Do đó, luận án định hướng lựa chọn HA làm nền tảng tiếp cận để phát triển khung phương pháp luận cho bài toán MCGDM ngôn ngữ: từ mô hình HA-TOPSIS xếp hạng phương án trên miền ngữ nghĩa; đến kế thừa và bổ sung quy trình xác định trọng số của tiêu chí có kiểm soát tính nhất quán, phát triển mô hình EFAHP - 4-tuple - HA; sau đó mở rộng sang bối cảnh động, xây dựng mô hình L-FAHP - L-FTOPSIS động dựa trên HA.

1.6. Kết luận Chương 1

Chương 1 đã hệ thống hóa bài toán MCGDM trong môi trường thông tin ngôn ngữ và các cơ sở lý thuyết liên quan: tập mờ, CWW, 2-tuple, 4-tuple, AHP/FAHP, TOPSIS/FTOPSIS và HA. Chương làm rõ yêu cầu bảo toàn cấu trúc ngữ nghĩa, liên thông ngữ nghĩa và kiểm soát tính nhất quán trong xác định trọng số, kết nhập đánh giá, xếp hạng phương án. Đây là cơ sở cho phát triển các mô hình: HA-TOPSIS, EFAHP - 4-tuple - HA và L-FAHP - L-FTOPSIS ở các Chương 2, 3 và 4.

CHƯƠNG 2. PHÁT TRIỂN MÔ HÌNH TOPSIS NGÔN NGỮ DỰA TRÊN ĐẠI SỐ GIA TỬ

2.1. Giới thiệu

Chương 2 được triển khai từ khoảng trống nghiên cứu thứ nhất: các mô hình TOPSIS/ FTOPSIS hiện tại còn phụ thuộc vào số thực và số mờ, hàm thuộc; chưa khai thác cấu trúc trật tự ngữ nghĩa của HA để xếp hạng các phương án. Trong MCGDM, chuyên gia thường đánh giá phương án bằng từ ngôn ngữ thay vì số rõ; việc chuyển đổi trực tiếp sang số thực, số mờ hoặc hàm thuộc có thể làm suy giảm trật tự và nội dung ngữ nghĩa. TOPSIS xếp hạng phương án theo nguyên lý gần PIS và xa NIS, song TOPSIS truyền thống và một số mô hình FTOPSIS còn phụ thuộc vào giá trị số rõ và giải mờ hoặc miền định lượng trung gian. Vì vậy, chương này phát triển HA-TOPSIS nhằm xây dựng thang đo ngôn ngữ dựa trên HA, định lượng ngữ nghĩa bằng hàm SQM, xác định PIS/NIS, đo khoảng cách ngữ nghĩa, tính hệ số gần và xếp hạng trên cùng miền ngữ nghĩa.

2.2. Phát triển mô hình ra quyết định HA-TOPSIS

(i) Dữ liệu đầu vào:

Cho $\bar{A} = \{A_1, A_2, \dots, A_m\}$, $m \geq 2$; $\bar{C} = \{C_1, C_2, \dots, C_n\}$, $n \geq 2$ và $\bar{D} = \{D_1, D_2, \dots, D_h\}$, $h \geq 2$. Hội đồng chuyên gia $\bar{D}$ thống nhất xây dựng các thang đo ngôn ngữ $L_{(j)k}(\bar{A})$ và $L_k(\bar{C})$ dựa trên HA để đánh giá các phương án và xác định trọng số của tiêu chí $C_j \in \bar{C}$ (với k là một số nguyên, biểu thị cho các từ ngôn ngữ trên thang đo có độ dài $\leq k$; $j = 1, \dots, n$). Đối với mỗi chuyên gia $D^e \in \bar{D}$ thực hiện:

(1) Xây dựng ma trận đánh giá các phương án theo công thức (2.1):

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

Trong đó, $a_{ij}^e \in L_{(j)k}(\bar{A})$ là đánh giá ngôn ngữ của chuyên gia $D^e \in \bar{D}$ đối với phương án $A_i \in \bar{A}$ trên tiêu chí $C_j \in \bar{C}$, với $i = 1, \dots, m$; $j = 1, \dots, n$; $e = 1, \dots, h$.

(2) Xây dựng ma trận đánh giá trọng số của các tiêu chí theo công thức (2.2):

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

Trong đó, $\tilde{x}_j^e \in L_k(\bar{C})$ ($j = 1, \dots, n$) là đánh giá ngôn ngữ của chuyên gia D^e về mức độ quan trọng của tiêu chí $C_j \in \bar{C}$. Véc-tơ trọng số của các tiêu chí $\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_n)^T$, với $\tilde{w}_j \geq 0$ và $\sum_{j=1}^n \tilde{w}_j = 1$.

(ii) Dữ liệu đầu ra: Thứ hạng của các phương án $A_i \in \bar{A}$.

Quy trình ra quyết định của mô hình HA-TOPSIS gồm có các bước sau:

Bước 1: Xây dựng thang đo ngôn ngữ để đánh giá các phương án và xác định trọng số của tiêu chí dựa trên HA.

Hội đồng chuyên gia thống nhất thiết lập các bộ tham số ngữ nghĩa của HA $(HA, fm(g^-), \mu(h), k)$ để xây dựng các thang đo ngôn ngữ $L_{(j)k}(\bar{A})$ dùng để đánh giá các phương án và $L_k(\bar{C})$ dùng để đánh giá trọng số của tiêu chí.

Dựa trên các Định nghĩa 1.4 - 1.5 và Mệnh đề 1.1, qui trình xây dựng thang đo ngôn ngữ bao gồm:

(i) xác định độ đo tính mờ của các phần tử sinh và gia tử; (ii) thiết lập tập các từ ngôn ngữ dùng trong đánh giá với độ dài từ tối đa không vượt quá k .

Bước 2: Xây dựng ma trận ngôn ngữ đánh giá các phương án và trọng số của tiêu chí

Với mỗi chuyên gia D^e sử dụng các từ ngôn ngữ $x \in X_{(j)k}$ trên thang đo ngôn ngữ $L_{(j)k}(\bar{A})$ để đánh giá các phương án $A_i \in \bar{A}$ theo từng tiêu chí $C_j \in \bar{C}$ bởi công thức (2.1).

Tiếp theo, mỗi chuyên gia D^e sử dụng từ ngôn ngữ $x \in X_k$ trên thang đo ngôn ngữ $L_k(\bar{C})$ để đánh giá mức độ quan trọng của các tiêu chí $C_j \in \bar{C}$ theo công thức (2.2).

Bước 3: Tính ma trận định lượng ngữ nghĩa cho các phương án và trọng số của các tiêu chí

Sử dụng các Định nghĩa 1.6 - 1.8 và Mệnh đề 1.2 xây dựng hàm SQM của HA để chuyển đổi các từ ngôn ngữ trong các ma trận A^e và P^e sang biểu diễn giá trị định lượng ngữ nghĩa số của từ ngôn ngữ $\vartheta(x) \in [0, 1]$ phục vụ tính toán. Từ đó ma trận định lượng ngữ nghĩa $\tilde{A}^e$ được xác định theo công thức (2.3) sau:

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

Tương tự, với mỗi phần tử trong ma trận đánh giá trọng số của các tiêu chí P^e cũng được chuyển đổi sang giá trị định lượng ngữ nghĩa số $\vartheta(\tilde{x}_j^e)$ tương ứng.

Bước 4. Xây dựng ma trận quyết định nhóm của các phương án và trọng số của các tiêu chí.

Ma trận đánh giá tổng hợp nhóm, ký hiệu $A\tilde{A}$ được xây dựng bằng cách tính trung bình cộng đánh các phương án giá từ tất cả các ma trận quyết định của chuyên gia $\tilde{A}^e$ nhau (trọng số chuyên gia bằng nhau), theo công thức (2.4) sau đây:

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

Tương tự, trọng số tổng hợp của tiêu chí $C_j \in \bar{C}$ cũng được xác định theo công thức 2.5 dưới đây:

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

Véc-tơ trọng số của các tiêu chí $\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_n)^T$, với $\tilde{w}_j = \frac{\tilde{u}_j}{\sum_{k=1}^n \tilde{u}_k}, j = 1, \dots, n; \tilde{w}_j \geq 0$ và $\sum_{j=1}^n \tilde{w}_j = 1$.

Bước 5: Tính toán ma trận quyết định có trọng số

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

Trong đó, u_{ij} là phần tử của ma trận $A\tilde{A}$ và $\tilde{w}_j$ là véc-tơ trọng số của tiêu chí $C_j \in \bar{C}$.

Bước 6. Xác định PIS/NIS và đo khoảng cách đến giải pháp lý tưởng

- Xác định (PIS, $\mathcal{L}^+$) và (NIS, $\mathcal{L}^-$) bởi các công thức (2.7) và (2.8) như sau:

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

+ Đối với tiêu chí lợi ích: $\mathcal{L}_j^+ = \sum_{i=1}^m \max \mathcal{L}_{ij}$, và $\mathcal{L}_j^- = \sum_{i=1}^m \min \mathcal{L}_{ij}$.

+ Đối với tiêu chí chi phí: $\mathcal{L}_j^+ = \sum_{i=1}^m \min \mathcal{L}_{ij}$, và $\mathcal{L}_j^- = \sum_{i=1}^m \max \mathcal{L}_{ij}$.

- Khoảng cách d_i^+ và d_i^- từ phương án A_i đến (PIS, $\mathcal{L}^+$) và (NIS, $\mathcal{L}^-$) được tính bởi các công thức (2.9) và (2.10) dưới đây:

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

với $i = 1, \dots, m; j = 1, \dots, n; d_i^+$ và d_i^- là khoảng cách ngắn nhất và dài nhất của A_i .

Bước 7. Xếp hạng các phương án

Dựa trên giá trị của hệ số gần CC_i , ($i = 1, \dots, m$) tập các phương án được xếp hạng theo thứ tự và phương án có giá trị CC_i lớn nhất là phương án tối ưu nhất, được tính theo công thức (2.11):

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

2.3. Ví dụ thực nghiệm

Công ty A có nhu cầu lựa chọn nhà cung cấp phần mềm hóa đơn điện tử phù hợp nhất, đáp ứng các tiêu chí mong muốn. Một hội đồng chuyên gia gồm năm người ra quyết định, ký hiệu lần lượt là D_1, D_2, D_3, D_4 và D_5 , được thành lập để tiến hành phỏng vấn ba nhà cung cấp, ký hiệu E1, E2, E3, dựa trên bốn tiêu chí: Đơn giá (P1), Chất lượng (P2), Công nghệ (P3) và Hỗ trợ sau bán hàng (P4).

Bảng 2.1. Khoảng cách từ nhà cung cấp đến lựa chọn tốt nhất và xấu nhất

Nhà cung cấp	d^+	d^-
E1	0.053	0.100
E2	0.104	0.060
E3	0.076	0.071

Bảng 2.2. Hệ số gần và thứ hạng nhà cung cấp

Nhà cung cấp	Hệ số gần	Thứ hạng
E1	0.654	1
E2	0.364	3
E3	0.485	2

Kết quả thực nghiệm cho thấy nhà cung cấp E1 đứng thứ nhất, E3 đứng thứ hai. E2 đứng thứ ba.

- Như vậy, E1 vừa gần PIS, vừa xa NIS hơn so với các phương án còn lại.
- Các trọng số của bốn tiêu chí cho thấy Đơn giá và Công nghệ là hai tiêu chí có ảnh hưởng lớn nhất.

2.4. Phân tích độ nhạy và so sánh kết quả

2.4.1. Phân tích độ nhạy

Phân tích độ nhạy được thực hiện thông qua việc thay đổi độ đo tính mờ theo hai kịch bản đặc trưng âm - dương trên miền ngữ nghĩa của HA như sau: Kịch bản 1: $f_m(U) = 0.4, f_m(I) = 0.6, \mu(L) = 0.6, \mu(V) = 0$. Ngược lại, ở kịch bản 2: Với $f_m(U) = 0.6, f_m(I) = 0.4, \mu(L) = 0.4, \mu(V) = 0.6$.

Bảng 2.3. Hệ số gần thu được trong các kịch bản

Nhà cung cấp	Ban đầu	Kịch bản 1	Kịch bản 2
E1	0.654	0.647	0.825
E2	0.364	0.372	0.192
E3	0.485	0.474	0.217

Bảng 2.4. Thứ hạng của nhà cung cấp

Kịch bản	Xếp hạng
Ban đầu	$E1 > E3 > E2$
Kịch bản 1	$E1 > E3 > E2$
Kịch bản 2	$E1 > E3 > E2$

Kết quả phân tích độ nhạy trong các Bảng 2.3 và 2.4 cho thấy hệ số gần của các nhà cung cấp có thay đổi theo từng kịch bản nhưng không lớn, và thứ hạng tổng thể vẫn được bảo toàn $E1 > E3 > E2$; đây là bằng chứng quan trọng về độ ổn định của mô hình trước nhiều tham số ngữ nghĩa.

2.4.2. So sánh kết quả và thảo luận

Bảng 2.5. So sánh kết quả xếp hạng giữa HA-TOPSIS, FTOPSIS và 4-tuple ngữ nghĩa

Phương án	HA-TOPSIS		FTOPSIS		4-tuple ngữ nghĩa	
	Hệ số gần	Xếp hạng	Hệ số gần	Xếp hạng	Tổng điểm	Xếp hạng
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

- Kết quả đối sánh trong Bảng 2.5 cho thấy cả ba mô hình đều nhận diện E1 là phương án tốt nhất. Điều này khẳng định tính hợp lý của dữ liệu thực nghiệm và kết quả lựa chọn chính.
- Trong khi FTOPSIS cho thứ hạng $E1 > E2 > E3$ (đảo thứ hạng E2, E3), cho thấy khi các phương án có mức đánh giá gần nhau, biểu diễn số mờ và giải mờ có thể làm thay đổi thứ hạng.
- Sự tương đồng thứ hạng giữa HA-TOPSIS và 4-tuple ngữ nghĩa củng cố luận điểm rằng nền tảng HA giúp bảo toàn trật tự ngữ nghĩa của dữ liệu tốt hơn.

2.5. Kết luận Chương 2

Chương 2 đã phát triển mô hình HA-TOPSIS nhằm xếp hạng phương án cho bài toán ra MCGDM ngôn ngữ trong bối cảnh tĩnh. Mô hình xây dựng thang đo ngôn ngữ dựa trên HA, định lượng ngữ nghĩa bằng hàm SQM, xác định PIS/NIS, tính khoảng cách ngữ nghĩa và hệ số gần. Thực nghiệm, phân tích độ nhạy và so sánh cho thấy mô hình duy trì tính nhất quán ngữ nghĩa, tính ổn định và khả năng diễn giải. Chương 3 tiếp tục kế thừa và bổ sung quy trình xác định trọng số của tiêu chí có kiểm soát tính nhất quán bằng mô hình tích hợp EFAHP - 4-tuple – HA.

CHƯƠNG 3. PHÁT TRIỂN MÔ HÌNH MCGDM TÍCH HỢP FAHP CẢI TIẾN VỚI MÔ HÌNH 4-TUPLE NGỮ NGHĨA THEO HƯỚNG TIẾP CẬN ĐẠI SỐ GIA TỬ

3.1. Giới thiệu

Chương 3 kế thừa mô hình HA-TOPSIS ở Chương 2 nhưng chuyển trọng tâm sang xác định trọng số của tiêu chí có kiểm soát tính nhất quán trong bài toán MCGDM ngôn ngữ. Do trọng số ảnh hưởng trực tiếp đến kết quả đánh giá và xếp hạng, chương này phát triển mô hình tích hợp EFAHP - 4-tuple - HA nhằm khai thác tri thức chuyên gia, kiểm soát tính nhất quán của phán đoán và bảo đảm liên thông ngữ nghĩa giữa xác định trọng số của tiêu chí, kết nhập đánh giá và xếp hạng phương án. Mô hình được kiểm chứng qua bài toán đánh giá mức độ chuyển đổi số của SMEs bán lẻ tại Việt Nam.

3.2. Kiểm soát tính nhất quán của ma trận so sánh cặp trong môi trường HA

Trong AHP và FAHP, kiểm soát tính nhất quán là điều kiện quan trọng để đánh giá độ tin cậy của ma trận so sánh cặp. Luận án kế thừa tinh thần này nhưng phát triển trong môi trường HA.

Tỷ lệ nhất quán (Consistency Ratio – CR) của ma trận so sánh cặp ngôn ngữ được xác định bởi công thức (3.1). Ma trận so sánh cặp của chuyên gia D^e chỉ được chấp nhận nếu $CR_{HA}^e < 0.10$. Ngược lại, chuyên gia D^e phải điều chỉnh lại các so sánh cặp cho đến khi đạt yêu cầu nhất quán:

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

- $RI(n)$: Là chỉ số ngẫu nhiên của ma trận bậc n theo AHP cổ điển của Saaty, được sử dụng như một hệ số chuẩn hóa theo kích thước ma trận. Các giá trị $RI(n)$ được cho trong Bảng 3.1 dưới đây:

Bảng 3.1. Giá trị chỉ số ngẫu nhiên ($RI(n)$)

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

- Trên cơ sở kế thừa CI_{Saaty} của AHP cổ điển, luận án đã phát triển công thức tính chỉ số nhất quán cho các ma trận so sánh cặp ngôn ngữ (P^e) của chuyên gia D^e trong môi trường HA (CI_{HA}^e) như sau:

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

Điểm khác biệt của công thức (3.2) so với CI_{Saaty} của AHP cổ điển là đã thay mốc nhất quán n bằng mốc mềm $n + (1 - \vartheta(w))$ và nhúng trực tiếp dung sai ngữ nghĩa của giá trị định lượng ngữ nghĩa phần tử trung hòa $1 - \vartheta(w)$ vào đo độ lệch nhất quán của ma trận so sánh cặp ngôn ngữ để thích nghi với đặc trưng của miền ngữ nghĩa dựa trên HA.

Mệnh đề 3.1. Với $n \geq 2$, $\vartheta(w) \in (0, 1)$, và $\mathcal{R}^e$ là ma trận dương đối nghịch đảo suy ra từ ma trận ngôn ngữ P^e , chỉ số CI_{HA}^e trong công thức (3.2) có các tính chất cơ bản sau:

- (i) CI_{HA}^e được xác định tốt trên miền áp dụng;
- (ii) $CI_{HA}^e = 0$ tại mốc HA-nhất quán λ_0^{HA} ;
- (iii) Với n và $\vartheta(w)$ cố định, CI_{HA}^e tăng đơn điệu theo λ_{max}^e ;
- (iv) Công thức (3.2) nhúng trực tiếp giá trị định lượng ngữ nghĩa của từ trung hòa $\vartheta(w)$ vào cơ chế đo nhất quán;
- (v) CI_{HA}^e liên tục theo λ_{max}^e , $\vartheta(w)$ và ổn định cục bộ;
- (iv) Dấu của CI_{HA}^e cho phép diễn giải trực tiếp vị trí của ma trận so với vùng dung sai ngữ nghĩa (nhỏ hơn mốc, bằng mốc, hay lớn hơn mốc HA-nhất quán).

3.3. Phát triển mô hình FAHP cải tiến với 4-tuple ngữ nghĩa dựa trên HA

(i) Dữ liệu đầu vào:

Cho $\bar{A} = \{A_1, A_2, \dots, A_m\}$, $m \geq 2$; $\bar{C} = \{C_1, C_2, \dots, C_n\}$, $n \geq 2$ và $\bar{D} = \{D_1, D_2, \dots, D_h\}$, $h \geq 2$. Hội đồng chuyên gia $\bar{D}$ thông nhất xây dựng hai thang đo 4-tuple ngữ nghĩa dựa trên HA, ký hiệu $L_k(\bar{C})$ và $L_k(\bar{A})$ để đánh giá trọng số của tiêu chí và đánh giá các phương án.

Đối với mỗi chuyên gia $D^e \in \bar{D}$, thực hiện:

(1) Xây dựng ma trận so sánh cặp ngôn ngữ để xác định trọng số các tiêu chí, theo công thức (3.3):

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

Trong đó, $x_{ij}^e \in L_k(\bar{C})$, x_{ii}^e là phần tử trung hòa của thang đo $L_k(\bar{C})$, biểu thị quan hệ so sánh quan trọng bằng nhau giữa một tiêu chí với chính nó, với $e = 1, 2, \dots, h$; $i, j = 1, \dots, n$; $i \neq j$.

Ma trận so sánh cặp ngôn ngữ P^e chỉ được chấp nhận nếu đạt tỷ lệ nhất quán $CR_{HA}^e < 0.10$; ngược lại, chuyên gia D^e phải điều chỉnh lại các so sánh cặp cho đến khi đạt yêu cầu nhất quán.

(2) Xây dựng ma trận đánh giá các phương án theo công thức (3.4):

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

(ii) Dữ liệu đầu ra: Thứ hạng của các phương án $A_i \in \bar{A}$.

Mô hình EFAHP - 4-tuple - HA thực hiện theo hai giai đoạn. Giai đoạn 1, phát triển EFAHP-HA dựa trên phương pháp FAHP để xác định trọng số của tiêu chí thông qua so sánh cặp ngôn ngữ có kiểm soát tính nhất quán trong môi trường HA. Giai đoạn 2, áp dụng mô hình 4-tuple - HA để tổng hợp, đánh giá và xếp hạng các phương án.

3.3.1. Xác định trọng số của các tiêu chí

Bước 1. Thiết kế cấu trúc phân cấp cho bài toán ra quyết định

Cấu trúc phân cấp bài toán MCGDM gồm ba mức: Ở cấp cao nhất là mục tiêu tổng thể; cấp thứ hai là bộ tiêu chí chính/ tiêu chí con; và cấp thấp nhất là tập phương án.

Bước 2. Xây dựng thang đo 4-tuple ngữ nghĩa dựa trên HA để đánh giá trọng số của tiêu chí

Hội đồng chuyên gia thông nhất thiết lập bộ tham số ngữ nghĩa của HA ($HA, f_m(g^-), \mu(h), k$) để xây dựng thang đo 4-tuple ngữ nghĩa $L_k(\bar{C})$ dùng để đánh giá trọng số của tiêu chí.

Sử dụng Định nghĩa 1.5 và Mệnh đề 1.1 tính độ đo tính mờ cho các từ ngôn ngữ $x \in X_k$ trên thang đo 4-tuple ngữ nghĩa $L_k(\bar{C})$. Tiếp theo, áp dụng các Định nghĩa 1.6 đến 1.8 và Mệnh đề 1.2 mỗi từ ngôn ngữ $x \in X_k$ được ánh xạ thành một giá trị định lượng ngữ nghĩa $\vartheta(x) \in [0, 1]$. Sau đó, áp dụng các Định nghĩa từ 1.9 đến 1.12 và Mệnh đề 1.3 để xác định khoảng tương tự ngữ nghĩa $S_k(x)$ cho từng từ ngôn ngữ $x \in X_k$ trên thang đo $L_k(\bar{C})$. Trên cơ sở đó, thang đo ngôn ngữ $L_k(\bar{C})$ để đánh giá

trọng số của tiêu chí được biểu diễn dưới dạng 4-tuple ngữ nghĩa, cụ thể: $L_k(\bar{C}) = \{(x, \vartheta(x), S_k(x), r_x): x \in X_k, r_x \in S_k(x)\}$.

Bước 3. Xây dựng ma trận so sánh cặp tiêu chí và kiểm tra tính nhất quán.

Mỗi chuyên gia D^e sử dụng các từ ngôn ngữ $x \in X_k$ trên thang đo ngôn ngữ $L_k(\bar{C})$ để đánh giá mức độ quan trọng tương đối giữa các tiêu chí, từ đó xây dựng ma trận so sánh cặp ngôn ngữ $P^e = (x_{ij}^e)_{n \times n}$ theo công thức (3.3). Các phần tử trên đường chéo chính biểu thị quan hệ so sánh một tiêu chí với chính nó nên được gán bằng phần tử trung hòa w trên thang đo ngôn ngữ $L_k(\bar{C})$, tức là: $x_{ii}^e = w, (i = 1, \dots, n)$.

Sau đó, mỗi phần tử $x_{ij}^e \in P^e$ được ánh xạ sang giá trị định lượng ngữ nghĩa $\vartheta(x_{ij}^e)$ và tiếp tục chuyển đổi thành số mờ tam giác tổng quát (GTFN) $u_{ij}^e = (L_{ij}^e, C_{ij}^e, R_{ij}^e; \omega_{ij}^e)$, Từ đó, ma trận P^e được chuyển thành ma trận GTFN $\mathcal{R}^e$ bởi công thức (3.5) dưới đây:

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

Trong đó, $u_{ij}^e = (L_{ij}^e, C_{ij}^e, R_{ij}^e; \omega_{ij}^e)$, $(u_{ij}^e)^{-1} = \left(\frac{1}{R_{ij}^e}, \frac{1}{C_{ij}^e}, \frac{1}{L_{ij}^e}; \omega_{ij}^e\right)$, $(e = 1, \dots, h; i, j = 1, \dots, n \text{ và } i \neq j)$; trong luận án này phép nghịch đảo không được hiểu như một phép toán đóng trên cùng không gian biểu diễn, mà được sử dụng chủ yếu để thiết lập quan hệ đối ngẫu của các phần tử trong ma trận so sánh cặp mờ. Các bước tính toán tiếp theo chỉ được thực hiện sau khi các phần tử đã được đưa về một miền biểu diễn tương thích, nhằm bảo đảm tính nhất quán hình thức của mô hình. Thành phần ω_{ij}^e là độ cao của số mờ, biểu thị mức độ thuộc cực đại hoặc mức tin cậy của đánh giá, nên được giữ nguyên và không lấy nghịch đảo.

Sử dụng công thức (3.1) để kiểm tra tỷ lệ nhất quán CR_{HA}^e của các ma trận so sánh cặp ngôn ngữ P^e . Ma trận P^e chỉ được chấp nhận nếu $CR_{HA}^e < 0.10$, ngược lại chuyên gia D^e phải điều chỉnh lại các so sánh cặp ngôn ngữ cho đến khi đạt yêu cầu nhất quán.

Bước 4. Tổng hợp các đánh giá do các chuyên gia cung cấp

Sau khi các ma trận P^e đã đạt yêu cầu về tính nhất quán, các ma trận này được tổng hợp thành một ma trận so sánh cặp mờ nhóm AP bằng phép tính trung bình cộng các GTFN theo công thức (3.6):

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

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

Bước 5. Tính giá trị mở rộng tổng hợp mờ và véc-tơ trọng số của các tiêu chí

Từ ma trận tổng hợp AP , giá trị mờ tổng hợp của từng tiêu chí, ký hiệu P_i , được xác định theo phép chuẩn hóa bởi công thức (3.7) như sau:

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

Mỗi $\mathbb{P}_i$ biểu thị mức độ ưu tiên tổng hợp của tiêu chí C_j sau khi xét toàn bộ các quan hệ so sánh cặp trong ma trận nhóm. Tiếp theo, trọng tâm (centroid) của mỗi số mờ P_i được xác định bởi $m_i = (\bar{A}_{\mathbb{P}_i}, \bar{L}_{\mathbb{P}_i})$, với $i = 1, 2, \dots, n$, và được tính theo công thức (3.8):

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

Khoảng cách từ trọng tâm $\mathbf{m}_i = (\bar{A}_{\mathbb{P}_i}, \bar{\mathcal{L}}_{\mathbb{P}_i})$ đến điểm tham chiếu $\mathbf{C} = (A_{min}, \mathcal{L}_{min})$ được tính theo công thức (3.9):

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

trong đó $A_{min} = \min(\mathbb{D}_i)$, $\mathcal{L}_{min} = \min(\omega_{ij})$, $\psi = \min(\omega_{ij})$;

Cuối cùng, véc-tơ trọng số của tiêu chí $\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_n)^T$ tương ứng với ma trận so sánh mờ được tính như công thức (3.10) sau:

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

3.3.2. Đánh giá xếp hạng các phương án

Bước 6. Xây dựng thang đo 4-tuple ngữ nghĩa để đánh giá các phương án

Để đánh giá các phương án theo các tiêu chí, hội đồng chuyên gia thống nhất thiết lập bộ tham số ngữ nghĩa của HA $(HA, f_{\mathfrak{M}}(g^-), \mu(h), k)$ để xây dựng thang đo 4-tuple ngữ nghĩa $L_k(\bar{A})$ dùng để đánh giá các phương án theo các tiêu chí. Quy trình xây dựng thang đo này hoàn toàn tương tự như xây dựng thang đo $L_k(\bar{C})$ trong Bước 2, nhưng được sử dụng cho mục đích đánh giá xếp hạng các phương án. Khi đó, thang đo ngôn ngữ $L_k(\bar{A})$ được xác định bởi:

$$L_k(\bar{A}) = \{(x, \vartheta(x), S_k(x), r_x): x \in X_k, r_x \in S_k(x)\}.$$

Trong đó, mỗi từ ngôn ngữ $x \in X_k$ trên thang đo $L_k(\bar{A})$ được biểu diễn dưới dạng một 4-tuple ngữ nghĩa.

Bước 7. Thu thập và chuyển đổi các đánh giá phương án sang biểu diễn 4-tuple ngữ nghĩa

Với mỗi chuyên gia D^e ($e = 1, 2, \dots, h$) sử dụng các từ ngôn ngữ $x \in X_k$ trên thang đo ngôn ngữ $L_k(\bar{A})$ để đánh giá xếp hạng cho m phương án trên n tiêu chí. Các đánh giá này được tổ chức thành một ma trận quyết định đánh giá các phương án $A^e = (a_{ij}^e)_{m \times n}$ theo công thức (3.4).

Sau đó, áp dụng các Định nghĩa 1.9 đến 1.12 và Mệnh đề 1.3, mỗi phần tử a_{ij}^e trong ma trận A^e được chuyển đổi sang biểu diễn dưới dạng 4-tuple ngữ nghĩa $(a_{ij}^e, \vartheta(a_{ij}^e), S_k(a_{ij}^e), r_{a_{ij}^e})$ tương ứng.

Bước 8. Tổng hợp và xếp hạng các phương án

- *Tổng hợp các đánh giá của chuyên gia:*

Từ các ma trận đánh giá các phương án của các chuyên gia được tổng hợp thành ma trận quyết định tổng hợp của nhóm chuyên gia được ký hiệu là: $AA = (\bar{a}_{ij})_{m \times n}$. Trong đó, $\bar{a}_{ij}$ là đánh giá tổng hợp của phương án A_i theo tiêu chí C_j .

Để tổng hợp ý kiến chuyên gia, luận án sử dụng phép kết nhập trung bình cộng trên thành phần đại diện r_x của các 4-tuple ngữ nghĩa. Cụ thể, với một tập $T_s = \{\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_p\}$, trong đó mỗi phần tử $\tilde{x}_j$ có dạng:

$$\tilde{x}_j = (x_j, \vartheta(x_j), S_k(x_j), r_j), j = 1, \dots, p,$$

toán tử tổng hợp g được xác định bởi:

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

trong đó, $\Delta(\cdot)$ là toán tử ánh xạ ngược từ giá trị đại diện trong miền chuẩn hóa sang một 4-tuple ngữ nghĩa tương ứng, và

$$\frac{1}{p} \sum_{j=1}^p r_j \in S_k(x).$$

Trong bài toán này, với mỗi cặp (i, j) , phần tử tổng hợp $\bar{a}_{ij}$ được xác định bằng cách áp dụng toán tử g cho các đánh giá của tất cả chuyên gia:

$$\bar{a}_{ij} = g((a_{ij}^1, \vartheta(a_{ij}^1), S_k(a_{ij}^1), r_{ij}^1), \dots, (a_{ij}^h, \vartheta(a_{ij}^h), S_k(a_{ij}^h), r_{ij}^h)).$$

- Tổng hợp có trọng số theo các tiêu chí:

Sau khi xây dựng ma trận quyết định tổng hợp AA , điểm tổng hợp cuối cùng cho từng phương án $A_i \in \bar{A}$ được tính bằng phép kết nhập trung bình có trọng số theo véc-tơ trọng số của tiêu chí $\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_p)^T$, với $\sum_{j=1}^p \tilde{w}_j = 1$ thu được trong Bước 5. Khi đó, phép tổng hợp có trọng số được xác định bởi:

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

Kết quả của phép tổng hợp là một 4-tuple ngữ nghĩa có dạng:

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

Cuối cùng, dựa trên tổng điểm các phương án $A_i \in \bar{A}$ được xếp hạng theo thứ tự giảm dần; phương án có điểm cao nhất được xếp hạng thứ nhất.

3.4. Ứng dụng mô hình đề xuất vào đánh giá mức độ sẵn sàng chuyển đổi số của các doanh nghiệp nhỏ và vừa (SMEs) trong lĩnh vực bán lẻ tại Việt Nam

Trong bài toán này, bộ dữ liệu thực nghiệm gồm 4 chuyên gia, 16 SMEs, 7 tiêu chí chính và 34 tiêu chí con. Cụ thể, 4 chuyên gia trong lĩnh vực nghiên cứu và quản trị bán lẻ tham gia đánh giá mức độ quan trọng của bộ tiêu chí gồm 7 tiêu chí chính và 34 tiêu chí con, được xây dựng trên cơ sở Quyết định số 2158/QĐ-BTTTT ngày 07/11/2023 của Bộ Thông tin và Truyền thông; đồng thời, 16 SMEs trong lĩnh vực bán lẻ tại Việt Nam được lựa chọn để khảo sát hiện trạng chuyển đổi số.

Trên cơ sở đó, Hội đồng chuyên gia thống nhất xây dựng hai thang đo 4-tuple ngữ nghĩa dựa trên HA, ký hiệu $L_k(\bar{C})$ và $L_k(\bar{A})$, lần lượt dùng cho đánh giá tầm quan trọng của tiêu chí và đánh giá mức độ đáp ứng của doanh nghiệp. Theo đó, 4 chuyên gia sử dụng thang đo $L_k(\bar{C})$ để thực hiện so sánh cặp và xác định mức độ quan trọng tương đối của các tiêu chí, trong khi 16 SMEs được khảo sát cung cấp mức độ đáp ứng đối với từng tiêu chí đánh giá trên thang đo $L_k(\bar{A})$.

Quá trình tính toán thực nghiệm của mô hình EFAHP - 4-tuple - HA được tổ chức theo hai giai đoạn: (i) xác định trọng số của tiêu chí có kiểm soát tính nhất quán bằng phương pháp FAHP cải tiến trong môi trường HA và (ii) tổng hợp, đánh giá và xếp hạng mức độ sẵn sàng chuyển đổi số của các doanh nghiệp bằng mô hình 4-tuple ngữ nghĩa dựa trên HA.

Kết quả thực hiện mô hình EFAHP - 4-tuple - HA xếp hạng mức độ sẵn sàng chuyển đổi số cho 16 SMEs bán lẻ tại Việt Nam được trình bày trong Bảng 3.2 như sau:

Bảng 3.2. Tổng hợp điểm và xếp hạng mức độ chuyển đổi số cho 16 SMEs tại Việt Nam

Doanh nghiệp	Tổng điểm	Mô hình 4-tuple ngữ nghĩa				Thứ hạng
		x	$S_k(x)$	$\vartheta(x)$	r_x	
E1	0.3522	L	[0.219, 0.450)	0.336	0.3522	2
E2	0.3930	L	[0.219, 0.450)	0.336	0.3930	1
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

- Kết quả xác định trọng số của tiêu chí cho thấy trụ cột P7 - Năng lực con người và tổ chức có trọng số lớn nhất, bằng 0,28. Tiếp theo là P6 - Quản lý rủi ro và an ninh mạng, với trọng số 0,16. Các trụ cột còn lại phân bố tương đối đồng đều: P1 = 0.10; P2 = 0.10; P3 = 0.11; P4 = 0.12; P5 = 0.12.

- Về xếp hạng, tất cả 16 SMEs được gán vào cùng lớp ngôn ngữ L, tức mức Low. Tuy nhiên, nhờ biểu diễn 4-tuple, mô hình vẫn phân biệt được mức độ khác nhau thông qua giá trị tham chiếu r_x ; đây lý do luận án lựa chọn 4-tuple vì nó cho thấy mô hình không chỉ phân lớp mà còn xếp hạng tinh trong cùng một lớp ngôn ngữ.

- Kết quả trong Bảng 3.2 cho thấy thứ hạng của 16 SMEs: E2 > E1 > E5 > E6 > E4 > E8 > E10 > E9 > E14 > E7 > E12 > E13 > E11 > E15 > E16 > E3. Trong đó, E2 đứng thứ nhất, E1 đứng thứ hai, E5 đứng thứ ba và E3 đứng cuối.

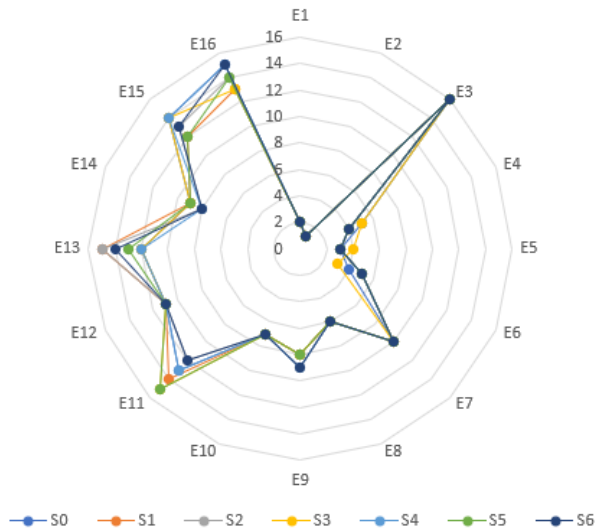
3.5. Phân tích độ nhạy và so sánh kết quả

3.5.1. Phân tích độ nhạy

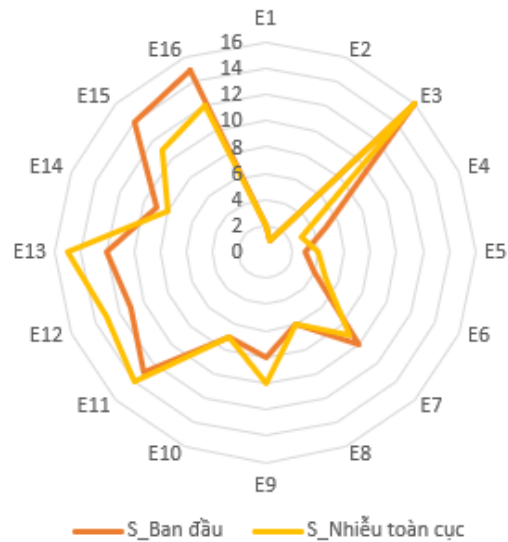
Phân tích độ nhạy theo hai kịch bản nhiều: Nhiều đơn tiêu chí và nhiều toàn cục trên ma trận so sánh cặp giữa bảy tiêu chí chính. Trong đó:

Kịch bản 1, tạo nhiều đơn tiêu chí bằng cách lần lượt tăng một mức ưu tiên của từng tiêu chí chính từ P1 đến P7 so với sáu tiêu chí còn lại trên $L_2(\bar{C})$ nhằm xem xét riêng rẽ mức độ ảnh hưởng của từng tiêu chí đến kết quả xếp hạng. Kết quả trong Hình 3.1 cho thấy mô hình có độ ổn định cao đối với nhóm DoN đứng đầu và xếp cuối. Biến động xếp hạng chủ yếu xuất hiện ở nhóm DoN tầm trung, nơi các tổng điểm tương đối gần nhau.

Kịch bản 2, tạo nhiều toàn cục bằng cách toàn bộ các giá trị so sánh cặp của bảy tiêu chí chính được tăng lên một cấp độ trên cùng thang đo $L_2(\bar{C})$. Kết quả trong Hình 3.2 cho thấy cấu trúc xếp hạng tổng thể vẫn được duy trì, trong đó đối với các doanh nghiệp dẫn đầu như E2, E1 và doanh nghiệp xếp cuối là E3. Sự biến động tiếp tục tập trung ở nhóm doanh nghiệp tầm trung, phản ánh khả năng phân biệt tinh của mô hình trong các trường hợp có điểm số gần nhau.



Hình 3.1. Phân tích độ nhạy theo kịch bản nhiều đơn tiêu chí



Hình 3.2. Phân tích độ nhạy theo kịch bản nhiều toàn cục

Tóm lại, kết quả phân tích độ nhạy khẳng định mô hình EFAHP - 4-tuple - HA có độ ổn định cao, đồng thời vẫn duy trì được độ nhạy cần thiết để phân biệt các phương án có mức đánh giá gần nhau.

3.5.2. So sánh kết quả và thảo luận

Nhằm làm rõ hơn mức độ phù hợp, tính ổn định và khả năng diễn giải của mô hình, luận án so sánh EFAHP - 4-tuple - HA với HA-TOPSIS dựa trên HA và FAHP-FTOPSIS dựa trên tập mờ. Kết quả cho thấy cả ba mô hình đều thống nhất ở doanh nghiệp tốt nhất là E2 và doanh nghiệp thấp nhất là E3. Điều này cho thấy kết luận chính của mô hình đề xuất phù hợp với các mô hình đối sánh.

Trong khi mô hình FAHP-FTOPSIS vẫn nhận diện E2 là phương án tốt nhất nhưng biến động mạnh ở nhóm doanh nghiệp cận trên và tầm trung. Sai khác cho thấy quá trình ánh xạ mờ, khử mờ dễ làm suy giảm ngữ nghĩa khi các phương án có kết quả gần nhau thường bị đảo hạng.

Tóm lại mô hình EFAHP - 4-tuple - HA cho thấy khả năng bảo toàn ngữ nghĩa tốt, giảm mất mát thông tin trong tính toán, đồng thời nâng cao độ tin cậy và khả năng diễn giải của kết quả xếp hạng.

3.6. Kết luận chương 3

Chương 3 đã phát triển mô hình tích hợp EFAHP - 4-tuple – HA, vừa xác định trọng số của tiêu chí có kiểm soát nhất quán bằng EFAHP-HA, vừa áp dụng 4-tuple ngữ nghĩa để đánh giá, xếp hạng các phương án liên thông trên cùng miền ngữ nghĩa. Thực nghiệm trên 16 SMEs bán lẻ tại Việt Nam, phân tích độ nhạy và so sánh với HA-TOPSIS và FAHP-FTOPSIS cho thấy mô hình có tính ổn định, khả năng diễn giải và phù hợp với MCGDM ngôn ngữ. Chương 4 mở rộng sang MCGDM động dựa trên HA, xét dữ liệu lịch sử, dữ liệu khuyết thiếu và cơ chế suy giảm ngữ nghĩa.

CHƯƠNG 4. PHÁT TRIỂN KHUNG PHƯƠNG PHÁP XÂY DỰNG MÔ HÌNH MCGDM ĐỘNG TRONG MÔI TRƯỜNG THÔNG TIN NGÔN NGỮ DỰA TRÊN ĐẠI SỐ GIA TỬ

4.1. Giới thiệu

Chương 4 kế thừa từ các mô hình HA-TOPSIS và EFAHP - 4-tuple – HA trong các Chương 2, 3 và mở rộng bài toán MCGDM ngôn ngữ sang bối cảnh động trong môi trường HA. Trong bối cảnh này, tập phương án, bộ tiêu chí, tập chuyên gia, trọng số và dữ liệu đánh giá có thể thay đổi theo thời gian. Luận án phát triển các phép kết nhập ngôn ngữ trên thang đo 4-tuple ngữ nghĩa, kết nhập với dữ liệu khuyết thiếu và cơ chế suy giảm ngữ nghĩa, đồng thời xây dựng mô hình L-FAHP - L-FTOPSIS động dựa trên HA và kiểm chứng qua đánh giá SMEs tại nhiều thời điểm.

4.2. Ngữ nghĩa tập mờ của từ ngôn ngữ trong HA

Định nghĩa 4.1. Ánh xạ $f_z: L_k \rightarrow F_{L_k}$ được gọi là ánh xạ ngữ nghĩa tập mờ tam giác (triangular fuzzy semantic – TFS) của một từ ngôn ngữ. Với mỗi từ ngôn ngữ $x \in X_k$ trên thang đo L_k được biểu diễn dưới dạng 4-tuple ngữ nghĩa, $\tilde{x} = (x, \vartheta(x), S_k(x), r_x) \in L_k$, ánh xạ $f_z(\tilde{x})$ được xác định bởi,

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

Trong đó, (l, m, u) được xác định theo quy trình phân hoạch ngữ nghĩa tập mờ tam giác nêu trên.

4.3. Các phép kết nhập ngôn ngữ

4.3.1. Phép kết nhập ngôn ngữ nhị phân trên thang đo 4-tuple ngữ nghĩa

Định nghĩa 4.2. Xét hai từ ngôn ngữ $\tilde{x}, \tilde{y} \in L_k$ được biểu diễn dưới dạng 4-tuple ngữ nghĩa: $\tilde{x} = (x, S_k(x), \vartheta(x), r_x)$, $\tilde{y} = (y, S_k(y), \vartheta(y), r_y)$, trong đó w là phần tử trung hoà của L_k . Ký hiệu, $\tilde{z} = \tilde{x} \oplus \tilde{y} = (z, S_k(z), \vartheta(z), r_z)$ là kết quả kết nhập của $\tilde{x}$ và $\tilde{y}$. Thành phần tham chiếu r_z , được xác định như sau:

$$r_z = r_x + r_y$$

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

Trong đó, $T(a, b) = a \times b$ và $S(a, b) = a + b - a \times b$.

Sau khi xác định r_z , các thành phần còn lại của $\tilde{z}$ được suy ra bằng cách chọn từ ngôn ngữ $z \in L_k$ sao cho $r_z \in S_k(z)$, còn $\vartheta(z)$ là giá trị định lượng ngữ nghĩa tương ứng của z . Nhờ đó, kết quả của phép kết nhập $\oplus$ luôn được đưa trở lại miền biểu diễn 4-tuple ngữ nghĩa.

Về ý nghĩa, công thức (4.2), đã chia phép kết nhập thành ba trường hợp dựa trên vị trí của hai toán hạng x và y đối với phần tử trung hoà w . Gọi $r_x = \vartheta(x)$, $r_y = \vartheta(y)$ và $r_w = \vartheta(w)$ là các giá trị định lượng ngữ nghĩa tương ứng. Ba trường hợp được mô tả như sau:

- Trường hợp 1 (cả hai nằm về phía âm của ω): $r_x < r_w$ và $r_y < r_w$. Khi đó, r_z được tính theo nhánh thứ nhất của công thức (4.2)

- Trường hợp 2 (cả hai nằm về phía không âm của ω): $r_x \geq r_w$ và $r_y \geq r_w$. Khi đó, r_z được tính theo nhánh thứ hai của công thức (4.2)

- Trường hợp 3 (trường hợp hỗn hợp): $r_x \leq r_w \leq r_y$ hoặc $r_y \leq r_w \leq r_x$. Khi đó, r_z được tính theo nhánh thứ ba của công thức (4.2) như một tổ hợp lồi của r_x và r_y , với các trọng số phụ thuộc vào độ đo tính mờ của hai từ tương ứng.

Mệnh đề 4.1. Với mọi $\tilde{x}, \tilde{y} \in L_k$, phép kết nhập $\oplus$ được định nghĩa trong Định nghĩa 4.1 thỏa mãn các tính chất sau:

(i) Tính đóng trên thang đo ngôn ngữ L_k , tức là $\tilde{z} = \tilde{x} \oplus \tilde{y} \in L_k$;

(ii) Tính giao hoán, tức là $\tilde{x} \oplus \tilde{y} = \tilde{y} \oplus \tilde{x}$;

(iii) Tính đơn điệu, nếu $\tilde{x}_1 \leq \tilde{x}_2$ thì $\tilde{x}_1 \oplus \tilde{y} \leq \tilde{x}_2 \oplus \tilde{y}$;

(iv) Tồn tại phần tử trung hòa $w \in L_k$, sao cho $\tilde{x} \oplus w = w \oplus \tilde{x} = \tilde{x}$;

(v) Tính kết hợp bộ phận, tức là phép kết nhập chỉ có tính kết hợp khi tất cả các toán hạng nằm cùng một phía so với phần tử w . Nếu $x, y, z \geq w$ hoặc $x, y, z \leq w$ thì $(\tilde{x} \oplus \tilde{y}) \oplus \tilde{z} = \tilde{x} \oplus (\tilde{y} \oplus \tilde{z})$.

Định nghĩa 4.3. Cho n từ ngôn ngữ $\tilde{x}_i \in L_k$, với $i = 1, 2, \dots, n$, được biểu diễn bởi các 4-tuple ngữ nghĩa tương ứng $\tilde{x}_i = (x_i, \vartheta(x_i), S_k(x_i), r_{x_i})$ và giả sử các từ ngôn ngữ này đã được sắp xếp theo thứ tự tăng dần theo x_i . Ký hiệu $\tilde{z} = (z, \vartheta(z), S_k(z), r_z)$ là từ ngôn ngữ thu được khi tổng hợp toàn bộ các $\tilde{x}_i$. Khi đó, phép kết nhập nhiều ngôi trên thang đo 4-tuple ngữ nghĩa L_k được xác định đệ quy bởi công thức (4.3):

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

4.3.2. Phép kết nhập với dữ liệu khuyết thiếu

Định nghĩa 4.4. Cho thang đo ngôn ngữ 4-tuple L_k được xây dựng dựa trên HA. Mỗi từ ngôn ngữ $\tilde{x} \in L_k$ được biểu diễn dưới dạng, $\tilde{x} = (x, S_k(x), \vartheta(x), r_x)$, trong đó $w \in L_k$ là phần tử trung hoà của thang đo L_k . Ký hiệu $\emptyset$ là từ rỗng biểu thị trường hợp dữ liệu bị khuyết thiếu, với $\emptyset \notin L_k$. Giả sử $\tilde{x}^{(\tau-1)} = (x^{(\tau-1)}, S_k(x^{(\tau-1)}), v(x^{(\tau-1)}), r_{x^{(\tau-1)}})$ và $\tilde{x}^{(\tau)} = (x^{(\tau)}, S_k(x^{(\tau)}), v(x^{(\tau)}), r_{x^{(\tau)}})$ lần lượt là giá trị đánh giá của một đối tượng tại hai thời điểm $\tau - 1$ và τ , với $\tau = 1, 2, \dots, t; t \geq 1$.

Ký hiệu $\tilde{x}_{agg}^{(\tau)} = (x_{agg}^{(\tau)}, S_k(x_{agg}^{(\tau)}), v(x_{agg}^{(\tau)}), r_{x_{agg}^{(\tau)}})$ là giá trị kết nhập tại thời điểm τ được xác định $\tilde{x}^{(\tau-1)}$ và $\tilde{x}^{(\tau)}$ trong trường hợp một trong hai giá trị bị khuyết thiếu, như sau:

(i) Trường hợp 1: Với $\tau = 1$ hoặc $\tilde{x}^{(\tau-1)} = \emptyset$.

Khi chưa có dữ liệu lịch sử, hoặc dữ liệu ở thời điểm $\tau - 1$ bị khuyết thiếu, giá trị kết nhập tại thời điểm τ được lấy bằng chính đánh giá hiện tại:

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

(ii) Trường hợp 2: $\tau > 1$, $\tilde{x}^{(\tau-1)} \neq \emptyset$ và $\tilde{x}^{(\tau)} = \emptyset$.

Khi dữ liệu hiện tại bị khuyết thiếu, giá trị kết nhập tại thời điểm τ được xác định bằng cơ chế suy giảm ngữ nghĩa từ đánh giá lịch sử:

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

Trong trường hợp này, thành phần tham chiếu của giá trị kết nhập $r_{x_{agg}^{(\tau)}}$ được xác định bởi công thức suy giảm ngữ nghĩa (4.6)

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

Trong đó:

$\vartheta(w)$ là giá trị định lượng ngữ nghĩa của phần tử trung hoà $w \in L_k$;

$\lambda \in [0, 1]$ là hệ số suy giảm ngữ nghĩa hay còn gọi là tỷ lệ quên:

Với $\lambda = 1$, không xảy ra suy giảm ngữ nghĩa, do đó $r_{x_{agg}^{(\tau)}} = r_{x^{(\tau-1)}}$.

Với $\lambda = 0$, thông tin lịch sử bị quên hoàn toàn, do đó $r_{x_{agg}^{(\tau)}} = \vartheta(w)$.

Với $0 < \lambda < 1$, giá trị $r_{x_{agg}^{(\tau)}}$ nằm giữa $r_{x^{(\tau-1)}}$ và $\vartheta(w)$, qua đó mô hình hóa sự suy giảm dần ảnh hưởng của đánh giá lịch sử về phía phần tử trung hoà.

Sau khi xác định $r_{x_{agg}^{(\tau)}}$, các thành phần còn lại của 4-tuple ngữ nghĩa $\tilde{x}_{agg}^{(\tau)}$ được xác định bằng cách chọn từ ngôn ngữ $x_{agg}^{(\tau)} \in L_k$ sao cho $r_{x_{agg}^{(\tau)}} \in S_k(x_{agg}^{(\tau)})$, $v(x_{agg}^{(\tau)})$ là giá trị định lượng ngữ nghĩa của từ $x_{agg}^{(\tau)}$.

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

Trên cơ sở đó, ta có mệnh đề sau:

Mệnh đề 4.2. Cho $s_{\tau-1} \in [0, 1]$ là giá trị tham chiếu ngữ nghĩa tại thời điểm $\tau - 1$, $\vartheta(w) \in [0, 1]$ là giá trị định lượng ngữ nghĩa của phần tử trung hòa, và $\lambda \in [0, 1]$ là hệ số suy giảm ngữ nghĩa. Khi đó, phép kết nhập suy giảm ngữ nghĩa theo công thức (4.6) thỏa mãn các tính chất sau:

(i). Tính bảo toàn miền giá trị: $\min\{s_{\tau-1}, \vartheta(w)\} \leq s_{\tau} \leq \max\{s_{\tau-1}, \vartheta(w)\}$, nghĩa là, giá trị suy giảm s_{τ} luôn nằm trong khoảng giới hạn bởi giá trị lịch sử $s_{\tau-1}$ và giá trị trung hòa $\vartheta(w)$;

(ii). Tính hội tụ về phần tử trung hòa: nếu dữ liệu bị khuyết thiếu liên tiếp n thời điểm và $0 < \lambda < 1$ thì $s_{\tau+n}$ sẽ hội tụ về $\vartheta(w)$ khi $n \rightarrow \infty$;

(iii). Tính đơn điệu theo hệ số suy giảm: với $n \geq 1$, nếu $0 \leq \lambda_1 \leq \lambda_2 \leq 1$ thì $|s_{\tau+n}(\lambda_1) - \vartheta(w)| \leq |s_{\tau+n}(\lambda_2) - \vartheta(w)|$, điều đó có nghĩa là λ càng nhỏ thì $s_{\tau+n}$ tiến về $\vartheta(w)$ càng nhanh.

(iv). Tính ổn định tại mốc trung hòa: nếu $s_{\tau-1} = \vartheta(w)$ thì $s_{\tau} = \vartheta(w)$ với mọi $\lambda \in [0, 1]$, tức là, khi giá trị lịch sử đã ở mốc trung hòa, phép toán suy giảm không làm thay đổi giá trị này.

4.4. Xây dựng mô hình L-FAHP - L-FTOPSIS động dựa trên HA

(i) Dữ liệu đầu vào:

(1) Số thời điểm đánh giá: $T = \{1, 2, \dots, t\}$, với $t \geq 1$.

(2) Tại mỗi thời điểm τ (với $\tau = 1, 2, \dots, t$), thực hiện:

- Xét ba tập hữu hạn sau: $\bar{A}^{\tau} = \{A_1^{\tau}, A_2^{\tau}, \dots, A_{m_{\tau}}^{\tau}\}$, ($m_{\tau} \geq 2$); $\bar{C}^{\tau} = \{C_1^{\tau}, C_2^{\tau}, \dots, C_{n_{\tau}}^{\tau}\}$, ($n_{\tau} \geq 2$) và $\bar{D}^{\tau} = \{D_1^{\tau}, D_2^{\tau}, \dots, D_{h_{\tau}}^{\tau}\}$ ($h_{\tau} \geq 2$).

- Hội đồng chuyên gia $\bar{D}^{\tau}$ thống nhất xây dựng hai thang đo 4-tuple ngữ nghĩa dựa trên HA, ký hiệu lần lượt là $L_k(\bar{C}^{\tau})$ và $L_k(\bar{A}^{\tau})$ để xác định trọng số của tiêu chí và đánh giá các phương án.

- Với mỗi chuyên gia $D_e^{\tau} \in \bar{D}^{\tau}$ (với $e = 1, 2, \dots, h_{\tau}$), thực hiện:

+ Xây dựng ma trận so sánh cặp ngôn ngữ để xác định trọng số của tiêu chí, theo công thức (4.7):

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

Trong đó, mỗi phần tử $x_{ij}^e(\tau) \in L_k(\bar{C}^{\tau})$. Các phần tử trên đường chéo chính $x_{ii}^e(\tau)$ là phần tử trung hòa của thang đo ngôn ngữ $L_k(\bar{C}^{\tau})$, biểu thị quan hệ so sánh quan trọng bằng nhau giữa một tiêu chí với chính nó, với $e = 1, 2, \dots, h_{\tau}$; $\tau = 1, 2, \dots, t$; $i, j = 1, 2, \dots, n_{\tau}$; $i \neq j$.

+ Xây dựng ma trận đánh giá các phương án theo công thức (4.8):

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

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

(3) Từ thời điểm $\tau - 1$: có xét đến dữ liệu lịch sử.

(ii) Dữ liệu đầu ra: Thứ hạng của các phương án $A_i^{\tau} \in \bar{A}^{\tau}$ (có xét dữ liệu lịch sử).

Quy trình ra quyết định gồm hai giai đoạn: GD 1 phát triển mô hình Linguistic FAHP dựa trên HA, ký hiệu là L-FAHP - HA để xác định trọng số của tiêu chí dựa trên FAHP, còn GD 2 phát triển mô hình Linguistic FTOPSIS dựa trên HA, ký hiệu là L-FAHP - HA để đánh giá, kết nhập và xếp hạng phương án qua nhiều thời điểm dựa trên FTOPSIS [21]. Trong cả hai giai đoạn, thông tin đánh giá của chuyên gia được biểu diễn trên thang đo ngôn ngữ 4-tuple ngữ nghĩa xây dựng từ HA, cho phép thực hiện trực tiếp các phép kết nhập trên miền ngôn ngữ mà bảo toàn trật tự và cấu trúc ngữ nghĩa nội tại của các từ ngôn ngữ. Đồng thời, mô hình còn thiết lập cơ chế suy giảm ngữ nghĩa để tích hợp dữ liệu lịch sử với dữ liệu hiện tại theo hướng dịch chuyển dần về phần tử trung hòa của thang đo.

4.4.1. Xác định trọng số của các tiêu chí

Bước 1.1. Xây dựng ma trận so sánh cặp ngôn ngữ và kiểm tra tính nhất quán tại thời điểm τ .

Tại thời điểm τ , mỗi chuyên gia D_e^τ sử dụng các từ ngôn ngữ $x \in X_k$ trên thang đo ngôn ngữ $L_k(\bar{C}^\tau)$ để xây dựng ma trận so sánh cặp ngôn ngữ, ký hiệu $P_{D_e^\tau}^\tau$, theo công thức (4.7).

Để bảo đảm độ tin cậy của các phán đoán, ma trận $P_{D_e^\tau}^\tau$ được kiểm tra tính nhất quán thông qua tỷ lệ nhất quán $CR_{HA}^e(\tau)$, xác định theo công thức (4.9):

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

Trong đó, $RI(n_\tau)$ (Random Index) là chỉ số ngẫu nhiên của ma trận bậc n_τ theo AHP cổ điển được tra trong Bảng 4.1. Còn $CI_{HA}^e(\tau)$ là chỉ số nhất quán của ma trận so sánh cặp ngôn ngữ $P_{D_e^\tau}^\tau$ do chuyên gia D_e^τ cung cấp trong môi trường HA như công thức (4.10) dưới đây:

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

Để tính $CR_{HA}^e(\tau)$, trước hết ma trận ngôn ngữ $P_{D_e^\tau}^\tau$ được chuyển sang ma trận mờ ngữ nghĩa $FP_{D_e^\tau}^\tau$ bằng cách ánh xạ mỗi phần tử $x_{ij}^e(\tau)$ sang tập mờ ngữ nghĩa tương ứng thông qua ánh xạ $fz(x_{ij}^e(\tau))$ theo công thức (4.1). Ma trận $P_{D_e^\tau}^\tau$ chỉ được chấp nhận nếu đạt tỷ lệ nhất quán $CR_{HA}^e(\tau) < 0.10$. Ngược lại chuyên gia D_e^τ phải điều chỉnh lại các so sánh cặp cho đến khi đạt yêu cầu nhất quán.

Bước 1.2. Kết nhập các ma trận so sánh cặp ngôn ngữ của các chuyên gia tại thời điểm τ .

Sau khi tất cả các ma trận $P_{D_e^\tau}^\tau$, đã đạt tỷ lệ nhất quán, tiến hành kết nhập các đánh giá ngôn ngữ của toàn bộ hội đồng chuyên gia tại thời điểm τ để thu được ma trận so sánh cặp nhóm, ký hiệu AP^τ . Phép kết nhập được thực hiện trên miền ngôn ngữ 4-tuple bằng cách áp dụng phép kết nhập ngôn ngữ nhị phân đã xây dựng trong các công thức (4.2) và (4.3), theo công thức (4.11) sau:

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

Trong đó, mỗi phần tử $\tilde{x}_{ij}(\tau)$ của ma trận AP^τ được xác định bởi:

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

Ma trận AP^τ biểu diễn đánh giá tập thể của hội đồng chuyên gia về mức độ quan trọng tương đối giữa các tiêu chí tại thời điểm τ .

Bước 1.3. Xây dựng ma trận so sánh cặp tổng hợp có tích hợp dữ liệu lịch sử.

Bước này được áp dụng cho trường hợp $\tau > 1$, khi quá trình ra quyết định tại thời điểm hiện tại cần kế thừa thông tin từ các thời điểm trước. Để thực hiện việc tích hợp dữ liệu lịch sử, ký hiệu: $\bar{C}_\tau^\tau = \bar{C}^{\tau-1} \cup \bar{C}^\tau$, là tập hợp các tiêu chí xuất hiện trong hai thời điểm liên tiếp $\tau - 1$ và τ , và n_τ^u là số phần tử của tập này. Khi đó, ma trận so sánh cặp tổng hợp có tích hợp dữ liệu lịch sử được ký hiệu là $\overline{AP}^\tau$, với các phần tử được xác định theo công thức (4.13) như sau:

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

Trong đó, $\oplus$: là phép kết nhập với dữ liệu khuyết thiếu; $\oplus$: là phép kết nhập ngôn ngữ nhị phân; $i, j = 1, \dots, n_\tau^u; i \neq j$.

Trong công thức (4.13), có ba trường được xét tương ứng với ba trạng thái động của tập tiêu chí: tiêu chí mới xuất hiện ở thời điểm hiện tại, tiêu chí được duy trì qua hai thời điểm, và tiêu chí chỉ còn trong dữ liệu lịch sử mà không còn xuất hiện tại thời điểm hiện tại lần lượt được tính toán như sau:

Trường hợp 1 của công thức (4.13) được thực hiện theo công thức (4.4), trường hợp 2 theo các công thức (4.2) và (4.3), còn trường hợp ba được tính theo các công thức (4.5) và (4.6). Cách tổ chức này cho phép mô hình linh hoạt thích ứng với sự thay đổi của tập tiêu chí theo thời gian, đồng thời bảo đảm rằng dữ liệu lịch sử không bị loại bỏ hoàn toàn mà được tích hợp có kiểm soát thông qua cơ chế kết nhập ngôn ngữ và suy giảm ngữ nghĩa.

Bước 1.4. Xây dựng ma trận so sánh cặp ngữ nghĩa mờ

Áp dụng Định nghĩa 4.1 để chuyển đổi các phần tử ngôn ngữ trong ma trận $\overline{AP}^\tau$ sang biểu diễn ngữ nghĩa mờ thông qua hàm ánh xạ $fz(\cdot)$, theo đó mỗi từ ngôn ngữ được gán với một tập mờ tam giác ngữ nghĩa tương ứng. Khi đó, ma trận so sánh cặp ngữ nghĩa mờ được ký hiệu là $\overline{FAP}^\tau$ và được xác định theo công thức (4.14) dưới đây:

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

trong đó: $fz(\tilde{p}_{ij}^*(\tau)) = (l_{ij}(\tau), m_{ij}(\tau), u_{ij}(\tau))$, $(fz(\tilde{p}_{ij}^*(\tau)))^{-1} = (\frac{1}{u_{ij}(\tau)}, \frac{1}{m_{ij}(\tau)}, \frac{1}{l_{ij}(\tau)})$, với $i, j = 1, \dots, n_\tau$, $i \neq j$; hàm $fz(\cdot)$ là phép ánh xạ của tập mờ tam giác của một từ ngôn ngữ.

Cần lưu ý rằng, để tránh xuất hiện mẫu số bằng 0 khi tính các phần tử nghịch đảo trong ma trận so sánh cặp mờ, các giá trị ngữ nghĩa TFS được tái chuẩn hóa từ miền $[0, 1]$ sang miền $[0.1, 1]$. Việc tái chuẩn hóa này không làm thay đổi quan hệ thứ tự ngữ nghĩa giữa các từ ngôn ngữ, nhưng giúp bảo đảm tính ổn định số học trong quá trình tính toán.

Bước 1.5. Xác định trọng số mờ và trọng số chuẩn hóa của các tiêu chí theo phương pháp trung bình hình học của Buckley.

Trên cơ sở ma trận so sánh cặp ngữ nghĩa mờ $\overline{FAP}^\tau$, trọng số của các tiêu chí tại thời điểm τ được tính bởi phép trung bình hình học mờ do Buckley đề xuất. Trước hết, đối với mỗi tiêu chí C_j^τ , trung bình hình học mờ của ma trận $\overline{FAP}^\tau$ được tính theo công thức (4.15):

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

Trong đó, RAP_j^τ biểu diễn mức độ ưu tiên mờ tổng hợp của tiêu chí thứ j so với toàn bộ các tiêu chí còn lại, còn $fz(\tilde{p}_{ij}^*(\tau))$ là phần tử mờ tam giác tại vị trí (i, j) của ma trận $\overline{FAP}^\tau$.

Tiếp theo, trọng số mờ của tiêu chí thứ j , ký hiệu $F\tilde{w}_j(\tau)$ được xác định bằng cách chuẩn hóa giá trị RAP_j^τ theo tổng các trung bình hình học mờ của tất cả các tiêu chí bởi công thức (4.16):

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

Trong đó, $\oplus$: dùng để cộng các số mờ tam giác; $\otimes$: dùng để nhân số mờ tam giác thứ j với nghịch đảo mờ của tổng; và $(\cdot)^{-1}$: là nghịch đảo nhân của số mờ dương.

Kết quả thu được là một số mờ tam giác có dạng: $F\tilde{w}_j(\tau) = (l_{F\tilde{w}_j}(\tau), m_{F\tilde{w}_j}(\tau), u_{F\tilde{w}_j}(\tau))$

Để phục vụ cho bước tính toán tiếp theo trong mô hình L-FTOPSIS-HA, các trọng số mờ được khử mờ để thu được các giá trị số rõ. Trong luận án, quá trình khử mờ được thực hiện bằng phương pháp trọng tâm miền (COA). Cụ thể, với mỗi trọng số mờ $F\tilde{w}_j(\tau)$, giá trị số rõ tương ứng $M_j(\tau)$ được xác định theo công thức:

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

Giá trị $M_j(\tau)$ là đại diện số rõ của trọng số mờ $F\tilde{w}_j(\tau)$. Các giá trị $M_j(\tau)$ được chuẩn hóa để thu được véc-tơ trọng số của tiêu chí $\tilde{w}_j(\tau)$ tại thời điểm τ , bảo đảm thỏa mãn điều kiện tổng các trọng số bằng 1. Véc-tơ trọng số này là kết quả đầu ra của GD 1 và được sử dụng làm đầu vào cho GD 2 của mô hình đề xuất, tức giai đoạn xếp hạng các phương án bằng phương pháp L-FTOPSIS-HA.

4.4.2. Đánh giá xếp hạng các phương án

Bước 2.1. Xây dựng ma trận đánh giá ngôn ngữ cho các phương án tại thời điểm τ .

Tại thời điểm τ , mỗi chuyên gia D_e^τ sử dụng từ ngôn ngữ $x \in X_k$ trên thang đo ngôn ngữ $L_k(\bar{A}^\tau)$, để xây dựng ma trận đánh giá ngôn ngữ cho các phương án, ký hiệu $\tilde{A}_{D_e}^\tau$, theo cấu trúc như trong công thức (4.8). Trong ma trận này, mỗi phần tử biểu diễn đánh giá ngôn ngữ của chuyên gia D_e^τ đối với mức độ đáp ứng của phương án $A_i^\tau \in \bar{A}^\tau$ theo từng tiêu chí $C_j^\tau \in \bar{C}^\tau$ tại thời điểm τ .

Bước 2.2. Kết nhập các ma trận đánh giá ngôn ngữ của chuyên gia tại thời điểm τ .

Kết nhập các đánh giá ngôn ngữ của toàn bộ hội đồng chuyên gia tại thời điểm τ để thu được ma trận đánh giá phương án nhóm, ký hiệu $A\tilde{A}^\tau$, theo công thức (4.18) sau:

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

Trong đó, mỗi phần tử $\tilde{a}_{ij}(\tau)$ của ma trận $A\tilde{A}^\tau$ được xác định thông qua phép kết nhập ngôn ngữ nhị phân trên miền 4-tuple ngữ nghĩa theo công thức (4.19) như sau:

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

Ma trận $A\tilde{A}^\tau$ phản ánh đánh giá tập thể của hội đồng chuyên gia về mức độ đáp ứng của các phương án tại thời điểm τ .

Bước 2.3. Xây dựng ma trận đánh giá phương án tổng hợp có tích hợp dữ liệu lịch sử

Bước này được áp dụng khi $\tau > 1$. Để tích hợp đánh giá hiện tại với dữ liệu lịch sử, ký hiệu: $\bar{A}_u^\tau = \bar{A}^{\tau-1} \cup \bar{A}^\tau$; $\bar{C}_u^\tau = \bar{C}^{\tau-1} \cup \bar{C}^\tau$ là tập hợp các phương án xuất hiện trong hai thời điểm liên tiếp $\tau - 1$ và τ , có số phần tử của hai tập này m_u^τ, n_u^τ . Tương tự, việc tích hợp lịch sử cũng được xét đồng thời đối với tập tiêu chí. Khi đó, các phần tử của ma trận đánh giá tổng hợp có tích hợp dữ liệu lịch sử, ký hiệu $\overline{A\tilde{A}^\tau} = (\tilde{a}_{ij}^*(\tau))_{m_u^\tau \times n_u^\tau}$, được xác định theo công thức (4.20) sau:

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

Trong đó, $\oplus$: là phép kết nhập với dữ liệu khuyết thiếu; $\emptyset$: là phép kết nhập ngôn ngữ nhị phân; $\ddot{w}$ được biểu diễn 4-tuple ngữ nghĩa của phần tử trung hòa w trong L_k ; $i = 1, \dots, m_u^\tau, j = 1, \dots, n_u^\tau$.

Trong công thức (4.20), có bốn trường hợp được xét tương ứng với bốn trạng thái động của tập phương án: phương án mới xuất hiện ở thời điểm hiện tại, phương án được duy trì qua hai thời điểm liên tiếp, phương án chỉ còn trong dữ liệu lịch sử mà không còn xuất hiện tại thời điểm hiện tại và bị khuyết thiếu dữ liệu cả 2 thời điểm bổ sung phương án mới và loại bỏ tiêu chí ở thời điểm hiện tại. Do đó, trường hợp 1 của công thức (4.20) được thực hiện theo công thức (4.4), trường hợp 2 tính theo các công thức (4.2) và (4.3), trường hợp 3 được tính theo các công thức (4.5) và (4.6), còn trường hợp 4 để bảo đảm tính đầy đủ của công thức trong trường hợp tập phương án và tập tiêu chí đồng thời thay

đổi. Xét ma trận mở rộng trên miền hợp $\bar{A}_u^\tau$ và $\bar{C}_u^\tau$, với mọi $A_i^\tau \in \bar{A}_u^\tau$ và $C_j \in \bar{C}_u^\tau$, nếu $A_i^\tau \in \bar{A}^\tau \setminus \bar{A}^{\tau-1}$ đồng thời $C_j^\tau \in \bar{C}^{\tau-1} \setminus \bar{C}^\tau$, thì cả giá trị lịch sử và giá trị hiện tại tại ô (i, j) đều khuyết thiếu; khi đó quy ước: $\emptyset \oplus \emptyset := \ddot{w} = (w, S_k(w), \mathfrak{Y}(w), r_w)$. Quy ước này chỉ nhằm bảo đảm tính đầy đủ của ma trận và tính nhất quán ngữ nghĩa của biểu diễn 4-tuple dựa trên HA. Khi xếp hạng tại thời điểm t, các tiêu chí không thuộc $\bar{C}^\tau$ không tham gia ma trận quyết định hiện tại, nên giá trị $\ddot{w}$ chỉ có ý nghĩa lưu vết cấu trúc và không ảnh hưởng đến kết quả xếp hạng tại thời điểm τ .

Bước 2.4. Chuyển đổi ma trận đánh giá ngôn ngữ tích hợp dữ liệu lịch sử sang ma trận đánh giá mờ.

Tương tự Bước 1.4, áp dụng Định nghĩa 4.1 để chuyển đổi các phần tử ngôn ngữ trong ma trận $\overline{AA}^\tau$ sang biểu diễn ngữ nghĩa mờ bằng hàm ánh xạ $fz(\cdot)$, theo đó mỗi từ ngôn ngữ được gán với một tập mờ tam giác ngữ nghĩa tương ứng. Khi đó, ma trận đánh giá ngữ nghĩa mờ được ký hiệu là $\overline{FAA}^\tau$ và được xác định theo công thức (4.21) sau:

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

Bước 2.5. Tính giá trị đánh giá tổng hợp có trọng số của các phương án.

Trên cơ sở ma trận đánh giá mờ $\overline{FAA}^\tau$ và trọng số mờ của các tiêu chí $F\tilde{w}_j(\tau)$ thu được trong Bước 1.5, giá trị đánh giá tổng hợp có trọng số của phương án thứ i tại thời điểm τ ký hiệu là Y_i^τ được xác định theo công thức (4.22):

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

Bước 2.6. Xác định khoảng cách từ các phương án đến PIS ($\mathcal{P}^+$) và NIS ($\mathcal{P}^-$) tại thời điểm τ được tính lần lượt theo các công thức (4.23) và (4.24):

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

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

Trong đó, $\mathcal{P}^+ = (1,1,1)$, $\mathcal{P}^- = (0,0,0)$, và $d_i^+(\tau)$ và $d_i^-(\tau)$ lần lượt biểu thị mức độ gần của phương án A_i^τ tới PIS và xa đối với NIS.

Bước 2.7. Tính hệ số gần và thứ hạng cho các phương án.

Hệ số gần của các phương án A_i^τ tại thời điểm τ được tính theo công thức (4.25). Các phương án được xếp hạng dựa trên giá trị CC_i^τ ; trong đó CC_i^τ càng lớn thì phương án càng tối ưu hơn.

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

4.5. Ví dụ thực nghiệm

Bảng 4.1. Kịch bản đánh giá mức độ sẵn sàng chuyển đổi số của SMEs tại Việt Nam

Thời điểm	Tập phương án và bối cảnh	Bộ tiêu chí
τ_1	E1, E2, E3, E4, E5.	03 tiêu chí chính, 11 tiêu chí con.
τ_2	Bổ sung mới: E6, E7, E8.	04 tiêu chí chính, 14 tiêu chí con (bổ sung 1 tiêu chí chính và 3 tiêu chí con).
τ_3	E1 và E2 bị loại bỏ; bổ sung E9, E10.	04 tiêu chí chính, 14 tiêu chí con.

Bảng 4.2. Kết quả xếp hạng mức độ sẵn sàng chuyển đổi số của SMEs ở Việt Nam
tại 3 thời điểm τ_1, τ_2, τ_3

Thời điểm	Kết quả xếp hạng	Nhận xét
τ_1	$E2 > E1 > E3 > E5 > E4$	E2 là phương án nổi trội; thứ hạng phản ánh được sự phân hóa hợp lý giữa nhóm dẫn đầu và nhóm cuối.
τ_2	$E6 > E2 > E1 > E7 > E3 > E8 > E5 > E4$	Khi tập phương án và bộ tiêu chí thay đổi, E6 vươn lên đứng đầu nhờ kết quả hiện tại tốt hơn.
τ_3	$E2 > E9 > E6 > E7 > E10 > E8 > E5 > E1 > E4 > E3$	E2 không còn được đánh giá mới ở τ_3 nhưng vẫn đứng đầu nhờ cơ chế suy giảm ngữ nghĩa cho phép bảo lưu hợp lý thành tích lịch sử tốt của E2.

4.6. Phân tích độ nhạy và so sánh kết quả

4.6.1. Phân tích độ nhạy

Phân tích độ nhạy được thực hiện tại ba thời điểm τ_1, τ_2, τ_3 theo hai kịch bản nhằm kiểm tra độ ổn định và độ nhạy của mô hình đề xuất

Kịch bản 1, nhiều loạn toàn cục bằng cách dịch toàn bộ các phán đoán so sánh cặp ngôn ngữ được dịch tăng một mức trên thang đo $L_2(\bar{C}^t)$.

Kịch bản 2, kết hợp nhiều loạn toàn cục và thay đổi độ đo tính mờ của HA.

Bảng 4.3. So sánh kết quả phân tích độ nhạy trong các kịch bản tại 3 thời điểm τ, τ_2 và τ_3

Thời điểm	Ban đầu	Kết quả phân tích độ nhạy theo hai kịch bản	
		Kịch bản 1	Kịch bản 2
τ_1	$E2 > E1 > E3 > E5 > E4$	$E2 > E1 > E3 > E5 > E4$	$E2 > E1 > E3 > E5 > E4$
τ_2	$E6 > E2 > E1 > E7 > E3 > E8 > E5 > E4$	$E6 > E2 > E7 > E1 > E3 > E8 > E5 > E4$	$E6 > E2 > E7 > E1 > E3 > E8 > E5 > E4$
τ_3	$E2 > E9 > E6 > E7 > E10 > E8 > E5 > E1 > E4 > E3$	$E2 > E9 > E6 > E7 > E10 > E8 > E5 > E1 > E4 > E3$	$E2 > E9 > E6 > E7 > E10 > E1 > E8 > E5 > E4 > E3$

Kết quả tổng quát theo hai kịch bản trong Bảng 4.3 cho thấy thứ hạng SMEs biến động rất nhỏ, nhóm dẫn đầu và xếp cuối hầu như được bảo toàn qua cả ba thời điểm. Trong đó, E2 giữ vị trí dẫn đầu ở thời điểm τ_1 và thời điểm τ_3 , còn E6 giữ ưu thế ở thời điểm τ_2 . Mô hình L-FAHP - L-FTOPSIS động dựa trên HA có độ ổn định cao, độ tin cậy tốt và khả năng diễn giải rõ ràng trong bối cảnh đánh giá mức độ sẵn sàng chuyển đổi số của SMEs bán lẻ tại Việt Nam.

4.6.2. So sánh kết quả

Bảng 4.4. So sánh kết quả mô hình đề xuất với FAHP-FTOPSIS động tại 3 thời điểm τ, τ_2 và τ_3

Thời điểm	FAHP-FTOPSIS động	Mô hình đề xuất
τ_1	$E2 > E1 > E5 > E4 > E3$	$E2 > E1 > E3 > E5 > E4$
τ_2	$E6 > E2 > E8 > E7 > E1 > E5 > E3 > E4$	$E6 > E2 > E7 > E1 > E3 > E8 > E5 > E4$
τ_3	$E9 > E10 > E6 > E2 > E8 > E7 > E1 > E5 > E4 > E3$	$E2 > E9 > E6 > E7 > E10 > E8 > E5 > E1 > E4 > E3$

Kết quả trong Bảng 4.4 cho thấy cả hai mô hình thống nhất ở một số phương án nổi trội, chẳng hạn E2 ở τ_1 và E6 ở τ_2 . Tuy nhiên, tại τ_3 mô hình đối sánh dựa trên lý thuyết tập mờ sử dụng phép kết nhập trung bình cộng của 3 thời điểm và ưu tiên phương án có kết quả tốt ở thời điểm hiện tại là E9 đứng đầu.

Trong khi mô hình đề xuất xếp E2 đứng đầu dựa trên thành tích lịch sử tốt (τ_1 và τ_2). Sự khác biệt này cho thấy mô hình đề xuất có khả năng tích hợp dữ liệu lịch sử với dữ liệu hiện tại theo một cơ chế cho phép ảnh hưởng dữ liệu lịch sử suy giảm hợp lý theo thời gian. Qua đó, mô hình đề xuất không chỉ bảo đảm độ tin cậy và tính ổn định của kết quả xếp hạng, mà còn phản ánh đúng hơn bản chất trưởng thành của quá trình chuyển đổi số trong môi trường ra quyết định động.

4.7. Kết luận chương 4

Trong chương này, luận án đã trình bày một số lý thuyết đại số gia tử mở rộng cho bài toán MCGDM ngôn ngữ động nhằm giải quyết sự thay đổi của bộ tiêu chí, tập phương án và người ra quyết định theo thời gian. Cụ thể trong chương này đã trình bày hai phép kết nhập đó là phép kết nhập ngôn ngữ nhị phân trên thang đo 4-tuple ngữ nghĩa và phép kết nhập với dữ liệu khuyết thiếu, trong đó có cơ chế suy giảm ngữ nghĩa để tích hợp dữ liệu lịch sử trong môi trường HA. Tiếp theo mô hình ra quyết định L-FAHP - L-FTOPSIS động dựa trên HA đã được phát triển. Cuối cùng, mô hình đề xuất đã được ứng dụng trong việc đánh giá mức độ sẵn sàng chuyển đổi số của SMEs trong lĩnh vực bán lẻ tại Việt Nam để minh họa cho những lợi thế và khả năng ứng dụng thực tiễn của lý thuyết và mô hình ra quyết định được đề xuất.

KẾT LUẬN

Luận án đã đạt mục tiêu phát triển các mô hình MCGDM trong môi trường thông tin ngôn ngữ dựa trên HA, theo hướng sử dụng HA làm nền tảng biểu diễn, định lượng ngữ nghĩa, bảo toàn trật tự ngữ nghĩa, kết nhập đánh giá và xếp hạng phương án. Trên cơ sở hệ thống hóa lý thuyết và xác định các khoảng trống nghiên cứu, luận án hình thành một khung phương pháp luận có quan hệ kế thừa, bổ sung và mở rộng qua ba kết quả chính như sau:

Thứ nhất, mô hình HA-TOPSIS được phát triển cho bài toán xếp hạng phương án trên miền ngữ nghĩa, trong đó thang đo ngôn ngữ dựa trên HA và hàm ánh xạ định lượng ngữ nghĩa (SQM) được sử dụng để xác định PIS/NIS, đo khoảng cách ngữ nghĩa và tính hệ số gần. Kết quả được kiểm chứng trên bài toán lựa chọn nhà cung cấp, kết hợp với phân tích độ nhạy và so sánh với các mô hình 4-tuple ngữ nghĩa và FTOPSIS cho thấy mô hình đề xuất bảo toàn trật tự ngữ nghĩa, góp phần nâng cao tính ổn định, độ tin cậy và khả năng diễn giải của kết quả ra quyết định (Chương 2).

Thứ hai, luận án đã phát triển mô hình MCGDM tích hợp EFAHP - 4-tuple - HA nhằm đồng thời xác định trọng số của tiêu chí và xếp hạng phương án trong môi trường ngôn ngữ. Kết quả trọng tâm của mô hình là đề xuất cơ chế kiểm soát tính nhất quán của ma trận so sánh cặp ngôn ngữ trong môi trường HA, trên cơ sở đó xác định trọng số của tiêu chí bằng EFAHP-HA. Tiếp theo, áp dụng mô hình 4-tuple ngữ nghĩa để biểu diễn, kết nhập đánh giá của chuyên gia và xếp hạng phương án. Việc biểu diễn và xử lý dữ liệu theo hướng ngữ nghĩa trên cùng một nền tảng HA góp phần bảo đảm tính liên thông ngữ nghĩa giữa hai giai đoạn xác định trọng số của tiêu chí và xếp hạng phương án. Mô hình

được ứng dụng vào bài toán đánh giá mức độ sẵn sàng chuyển đổi số của 16 SMEs trong lĩnh vực bán lẻ cho thấy mô hình không chỉ bảo toàn tốt hơn về cấu trúc ngữ nghĩa của dữ liệu đầu vào, mà còn duy trì được độ ổn định của thứ hạng khi phân tích độ nhạy và đối sánh với các mô hình HA-TOPSIS và FAHP-FTOPSIS (Chương 3).

Thứ ba, luận án đã phát triển khung phương pháp xây dựng mô hình MCGDM động trong môi trường ngôn ngữ dựa trên HA. Trên nền tảng biểu diễn 4-tuple ngữ nghĩa, luận án đã phát triển phép kết nhập ngôn ngữ nhị phân, phép kết nhập với dữ liệu khuyết thiếu và cơ chế suy giảm ngữ nghĩa nhằm tích hợp dữ liệu lịch sử với dữ liệu hiện tại. Từ khung phương pháp này, luận án xây dựng mô hình L-FAHP - L-FTOPSIS động dựa trên HA để xác định trọng số của tiêu chí và xếp hạng phương án nhất quán qua nhiều thời điểm. Mô hình được kiểm chứng thông qua bài toán đánh giá mức độ sẵn sàng chuyển đổi số của SMEs trong bối cảnh động, kết hợp phân tích độ nhạy và so sánh với phương pháp liên quan. Kết quả thực nghiệm cho thấy mô hình có khả năng phản ánh sự biến đổi của dữ liệu đánh giá theo thời gian, đồng thời duy trì khả năng diễn giải ngữ nghĩa trong phạm vi các kịch bản đã xây dựng (Chương 4).

Bên cạnh những kết quả đạt được, nghiên cứu vẫn còn một số hạn chế cần tiếp tục được nghiên cứu. Thứ nhất, các mô hình được xây dựng chủ yếu trong phạm vi HA tuyến tính với các bộ tham số ngữ nghĩa được xác định theo giả thiết hoặc theo cấu hình thực nghiệm. Thứ hai, phạm vi kiểm chứng thực nghiệm còn giới hạn về số lượng miền ứng dụng, số lượng bộ dữ liệu và quy mô bài toán. Thứ ba, luận án chưa phát triển thành một hệ hỗ trợ ra quyết định hoàn chỉnh để kiểm chứng khả năng triển khai trong môi trường vận hành thực tế.

Từ những hạn chế đó, một số hướng nghiên cứu tiếp theo có thể được triển khai. Trước hết, cần nghiên cứu cơ chế tối ưu và hiệu chỉnh bộ tham số ngữ nghĩa của HA dựa trên dữ liệu thực nghiệm hoặc phản hồi của chuyên gia, nhằm giảm sự phụ thuộc vào các phương pháp khảo sát truyền thống. Tiếp theo, có thể mở rộng các mô hình đề xuất cho bài toán MCGDM quy mô lớn, bài toán có dữ liệu không đầy đủ, dữ liệu không đồng nhất hoặc dữ liệu cập nhật theo thời gian thực. Cuối cùng, cần phát triển hệ hỗ trợ ra quyết định dựa trên các mô hình đã đề xuất và mở rộng ứng dụng sang các lĩnh vực như đánh giá rủi ro, quản lý dự án, y tế, giáo dục, tài chính và môi trường.

**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 Fifth International Conference on Intelligent Systems and Networks*. ICISN, 22–23 March, Hanoi, Vietnam. 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 xếp hạng **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*).