

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**

DỰ THẢO LUẬN ÁN TIẾN SĨ

HÀ NỘI - 2026

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

DỰ THẢO LUẬN ÁN TIẾN SĨ

CÁN BỘ HƯỚNG DẪN KHOA HỌC:

1. TS. HOÀNG VĂN THÔNG
2. PGS. TSKH. NGUYỄN CÁT HỒ

HÀ NỘI - 2026

LỜI CAM ĐOAN

Tác giả xin cam đoan đây là công trình nghiên cứu của riêng tác giả. Các kết quả được viết chung với các tác giả khác đều được sự đồng ý của đồng tác giả trước khi đưa vào luận án. Các kết quả trong luận án là trung thực và chưa từng được công bố trong bất kỳ công trình nào khác.

Tác giả

Nhữ Văn Kiên

LỜI CẢM ƠN

Tác giả xin chân thành cảm ơn TS. Hoàng Văn Thông, PGS. TSKH. Nguyễn Cát Hồ đã luôn nhiệt tình hướng dẫn và tạo điều kiện thuận lợi để tác giả hoàn thành luận án này. Tác giả chân thành cảm ơn PGS. TS. Lưu Quốc Đạt đã có nhiều hỗ trợ, đóng góp quý báu trong quá trình nghiên cứu và thời gian hoàn thành luận án này.

Tác giả xin chân thành gửi lời cảm ơn đến Ban Giám hiệu Trường Đại học Giao thông Vận tải, Khoa Công nghệ thông tin, Phòng Đào tạo Sau đại học và các Phòng chức năng trong Trường đã quan tâm giúp đỡ, tạo điều kiện thuận lợi trong suốt thời gian học tập, nghiên cứu và hoàn thành luận án.

Xin trân trọng cảm ơn Ban Giám đốc Đại học Công nghiệp Hà Nội, Ban Giám hiệu Trường Công nghệ thông tin và Truyền thông, Khoa Khoa học máy tính và các Phòng chức năng trong Trường đã quan tâm giúp đỡ, tạo điều kiện để tác giả có thể thực hiện kế hoạch nghiên cứu đảm bảo tiến độ.

Tác giả xin cảm ơn người thân, bạn bè và đồng nghiệp đã tạo điều kiện, động viên và hỗ trợ tác giả trong quá trình thực hiện luận án.

Cuối cùng, tác giả xin được bày tỏ tình cảm và lòng biết ơn vô hạn tới bố, mẹ và những người thân trong gia đình, đặc biệt là vợ tác giả - Nguyễn Thị Thu, người đã luôn động viên, khích lệ, chia sẻ và gánh vác công việc để tác giả có thời gian nghiên cứu và hoàn thành luận án. Luận án cũng là món quà tinh thần mà tác giả trân trọng gửi tặng đến các thành viên trong gia đình.

Xin trân trọng cảm ơn!

MỤC LỤC

LỜI CAM ĐOAN	i
LỜI CẢM ƠN	ii
MỤC LỤC.....	iii
DANH MỤC CÁC TỪ VIẾT TẮT	vi
DANH MỤC CÁC BẢNG BIỂU	vii
DANH MỤC CÁC HÌNH VẼ	x
DANH MỤC CÁC KÝ HIỆU	xi
MỞ ĐẦU	1
1. Tính cấp thiết của đề tài	1
2. Mục tiêu nghiên cứu.....	4
2.1. Mục tiêu tổng quát	4
2.2. Mục tiêu cụ thể	4
3. Nội dung nghiên cứu	5
4. Đối tượng và phạm vi nghiên cứu.....	6
4.1. Đối tượng nghiên cứu.....	6
4.2. Phạm vi nghiên cứu.....	6
5. Phương pháp nghiên cứu	7
6. Đóng góp mới của đề tài	7
7. Kết cấu của luận án	9
CHƯƠNG 1. TỔNG QUAN NGHIÊN CỨU VỀ BÀI TOÁN MCGDM TRONG	
MÔI TRƯỜNG THÔNG TIN NGÔN NGỮ VÀ CƠ SỞ LÝ THUYẾT DỰA	
TRÊN ĐẠI SỐ GIA TỬ	11
1.1. Bài toán ra quyết định nhóm đa tiêu chí	11
1.1.1. Bối cảnh bài toán MCDM và MCGDM.....	11
1.1.2. Phát biểu bài toán MCGDM trong môi trường thông tin ngôn ngữ	12
1.1.3. Đặc điểm của bài toán MCGDM trong môi trường ngôn ngữ	13
1.1.4. Phạm vi của bài toán MCGDM	14
1.1.5. Vai trò của bài toán trong các ứng dụng	15
1.2. Cơ sở lý thuyết nền tảng.....	16
1.2.1. Lý thuyết tập mờ và biến ngôn ngữ	16
1.2.2. Các mô hình tính toán với từ trong môi trường ngôn ngữ	19
1.2.3. Lý thuyết Đại số gia tử.....	22

1.3. Các phương pháp nền mà luận án kế thừa và phát triển	32
1.3.1. Phương pháp AHP và vai trò trong xác định trọng số của tiêu chí	32
1.3.2. Phương pháp TOPSIS và vai trò trong xếp hạng phương án	36
1.3.3. Mối liên hệ giữa AHP và TOPSIS trong các mô hình MCGDM	38
1.4. Tổng quan nghiên cứu về MCDM và MCGDM ngôn ngữ	40
1.4.1. Các nghiên cứu về MCDM truyền thống và tiếp cận mở rộng dựa trên lý thuyết tập mờ	40
1.4.2. Các nghiên cứu dựa trên các mô hình ngôn ngữ 2-tuple và 4-tuple	41
1.4.3. Các nghiên cứu dựa trên HA cho bài toán MCDM và MCGDM	42
1.4.4. Các nghiên cứu về MCGDM động trong môi trường ngôn ngữ	44
1.5. Khoảng trống nghiên cứu và định hướng luận án	45
1.6. Kết luận chương 1	49
CHƯƠNG 2. PHÁT TRIỂN MÔ HÌNH TOPSIS NGÔN NGỮ DỰA TRÊN ĐẠI SỐ GIA TỬ	50
2.1. Giới thiệu	50
2.2. Phát triển mô hình ra quyết định HA-TOPSIS	52
2.3. Ví dụ thực nghiệm	60
2.4. Phân tích độ nhạy và so sánh kết quả	65
2.4.1. Phân tích độ nhạy	65
2.4.2. So sánh kết quả và thảo luận	67
2.5. Kết luận chương 2	70
CHƯƠNG 3. PHÁT TRIỂN MÔ HÌNH MCGDM TÍCH HỢP FAHP CẢI TIẾN VỚI 4-TUPLE NGỮ NGHĨA THEO HƯỚNG TIẾP CẬN ĐẠI SỐ GIA TỬ ..	71
3.1. Giới thiệu	71
3.2. Kiểm soát tính nhất quán của ma trận so sánh cặp trong môi trường HA ..	73
3.3. Phát triển mô hình FAHP cải tiến với 4-tuple ngữ nghĩa dựa trên HA	79
3.3.1. Xác định trọng số của các tiêu chí	82
3.3.2. Đánh giá xếp hạng các phương án	87
3.4. Ứng dụng mô hình đề xuất vào đánh giá mức độ sẵn sàng chuyển đổi số của SMEs trong lĩnh vực bán lẻ tại Việt Nam	92
3.5. Phân tích độ nhạy và so sánh kết quả	98
3.5.1. Phân tích độ nhạy	98
3.5.2. So sánh kết quả và thảo luận	101
3.6. Kết luận chương 3	105

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	106
4.1. Giới thiệu	106
4.2. Ngữ nghĩa tập mờ của từ ngôn ngữ trong HA	108
4.3. Các phép kết nhập ngôn ngữ	111
4.3.1. Phép kết nhập ngôn ngữ nhị phân trên thang đo 4-tuple ngữ nghĩa	111
4.3.2. Phép kết nhập với dữ liệu khuyết thiếu	120
4.4. Xây dựng mô hình L-FAHP - L-FTOPSIS động dựa trên HA	124
4.4.1. Xác định trọng số của các tiêu chí	127
4.4.2. Đánh giá xếp hạng các phương án	131
4.5. Ví dụ thực nghiệm	139
4.6. Phân tích độ nhạy và so sánh kết quả	145
4.6.1. Phân tích độ nhạy	145
4.6.2. So sánh kết quả và thảo luận	149
4.7. Kết luận chương 4	153
KẾT LUẬN	154
CÁC CÔNG TRÌNH KHOA HỌC CỦA TÁC GIẢ ĐÃ CÔNG BỐ CÓ LIÊN QUAN ĐẾN LUẬN ÁN	157
TÀI LIỆU THAM KHẢO	158
Tiếng Việt	158
Tiếng Anh	158
PHỤ LỤC A. MỘT SỐ KẾT QUẢ	i
A1. Một số kết quả trong Chương 3	i
A2. Một số kết quả trong Chương 4	vii

DANH MỤC CÁC TỪ VIẾT TẮT

TT	Từ viết tắt	Từ gốc	Diễn giải/Tạm dịch
1	AHP	Analytic Hierarchy Process	Phương pháp phân tích thứ bậc
2	ANP	Analytic Network Process	Phương pháp phân tích mạng
3	CWW	Computing with Words	Tính toán với từ
4	CI	Consistency Index	Chỉ số nhất quán
5	CR	Consistency Ratio	Tỷ lệ nhất quán
6	DM	Decision Maker/Expert	Người ra quyết định/chuyên gia
7	DMCGDM	Dynamic Multi-Criteria Group Decision Making	Ra quyết định nhóm đa tiêu chí động
8	FAHP	Fuzzy Analytic Hierarchy Process	Phương pháp phân tích thứ bậc mờ
9	FTOPSIS	Fuzzy Technique for Order Preference by Similarity to Ideal Solution	Kỹ thuật sắp thứ tự ưu tiên dựa trên sự tương đồng với giải pháp lý tưởng mờ
10	HA	Hedge Algebras	Đại số gia tử
11	L-FAHP	Linguistic Fuzzy Analytic Hierarchy Process	Phương pháp phân tích thứ bậc mờ ngôn ngữ
12	L-FTOPSIS	Linguistic Fuzzy Technique for Order Preference by Similarity to Ideal Solution	Kỹ thuật sắp thứ tự ưu tiên dựa trên sự tương đồng với giải pháp lý tưởng mờ ngôn ngữ
13	GTFN	Generalized Triangular Fuzzy Number	Số mờ tam giác tổng quát
14	MCDM	Multi-Criteria Decision Making	Ra quyết định đa tiêu chí
15	MCGDM	Multi-Criteria Group Decision Making	Ra quyết định nhóm đa tiêu chí
16	NIS	Negative ideal solution	Giải pháp lý tưởng tiêu cực
17	PIS	Positive ideal solution	Giải pháp lý tưởng tích cực
18	TOPSIS	Technique for Order Preference by Similarity to Ideal Solution	Kỹ thuật sắp thứ tự ưu tiên dựa trên sự tương đồng với giải pháp lý tưởng
19	SMEs	Small-and Medium-sized Enterprises	Các doanh nghiệp nhỏ và vừa
20	SQM	Semantically Quantifying Mapping	Hàm ánh xạ định lượng ngữ nghĩa
21	RI	Random Index	Chỉ số ngẫu nhiên

DANH MỤC CÁC BẢNG BIỂU

Bảng 1.1. Mô tả thang đo của phương pháp AHP [95].....	32
Bảng 2.1. Độ phức tạp của thuật toán thực hiện các công thức trong HA-TOPSIS.....	59
Bảng 2.2. Thiết lập cấu trúc HA và các tham số mờ xây dựng thang đo ngôn ngữ cho các tiêu chí và trọng số của tiêu chí	61
Bảng 2.3. Đánh giá các nhà cung cấp theo tiêu chí P1	62
Bảng 2.4. Đánh giá các nhà cung cấp theo tiêu chí P2	62
Bảng 2.5. Đánh giá các nhà cung cấp theo tiêu chí P3	62
Bảng 2.6. Đánh giá các nhà cung cấp theo tiêu chí P4	62
Bảng 2.7. Đánh giá trọng số quan trọng của các tiêu chí.....	63
Bảng 2.8. Trung bình đánh giá của các nhà cung cấp theo các tiêu chí.....	63
Bảng 2.9. Véc-tơ trọng số của các tiêu chí.....	63
Bảng 2.10. Ma trận quyết định có trọng số	64
Bảng 2.11. Khoảng cách từ nhà cung cấp đến lựa chọn tốt nhất và xấu nhất.....	64
Bảng 2.12. Hệ số gần và thứ hạng nhà cung cấp	64
Bảng 2.13. Trọng số của các tiêu chí của các kịch bản.....	66
Bảng 2.14. Hệ số gần thu được trong các kịch bản.....	66
Bảng 2.15. Thứ hạng của nhà cung cấp thu được trong các kịch bản	66
Bảng 2.16. So sánh kết quả xếp hạng giữa HA-TOPSIS, FTOPSIS và 4-tuple ngữ nghĩa	67
Bảng 2.17. Hệ số tương quan thứ hạng và mức độ đồng thuận giữa các mô hình đối sánh.....	68
Bảng 3.1. Giá trị chỉ số ngẫu nhiên (RI(n)).....	75
Bảng 3.2. Độ đo tính mờ của các từ ngôn ngữ $x \in X_2$ trong thang đo $L_2(\bar{C})$	83
Bảng 3.3. Giá trị định lượng ngữ nghĩa của từ ngôn ngữ $x \in X_2$ trong thang đo $L_2(\bar{C})$	83
Bảng 3.4. Bảng giá trị GTFN của các từ ngôn ngữ $x \in X_2$ trên thang đo $L_2(\bar{C})$	86
Bảng 3.5. Độ phức tạp của thuật toán thực hiện các công thức trong mô hình.....	91
EFAHP - 4-tuple - HA.....	91
Bảng 3.6. Thiết lập cấu trúc ComHA và bộ tham số ngữ nghĩa.....	93
Bảng 3.7. Thang đo ngôn ngữ 4-tuple xác định trọng số của tiêu chí $L_2(\bar{C})$	94
Bảng 3.8. Thang đo ngôn ngữ 4-tuple để đánh giá các phương án $L_2(\bar{A})$	95
Bảng 3.9. Tổng hợp điểm số và xếp hạng mức độ chuyển đổi số cho 16 SMEs trong lĩnh vực bán lẻ tại Việt Nam.....	97

Bảng 3.10. So sánh kết quả giữa mô hình EFAHP - 4-tuple - HA với FAHP - FTOPSIS và HA-TOPSIS.....	101
Bảng 3.11. Hệ số tương quan thứ hạng và mức độ đồng thuận giữa các mô hình đối sánh.....	102
Bảng 4.1. Thang đo ngôn ngữ 4-tuple ngữ nghĩa $L_2(\bar{C})$	113
Bảng 4.2. Giá trị chỉ số ngẫu nhiên $(RI(n_\tau))$	127
Bảng 4.3. Các TFSs của các từ ngôn ngữ dùng để so sánh cặp giữa hai tiêu chí trên thang đo $L_2(\bar{C}^\tau)$ dựa trên HA	130
Bảng 4.4. Thang đo ngôn ngữ 4-tuple và TFS của các từ ngôn ngữ dựa trên HA dùng để xếp hạng các phương án $L_2(\bar{A}^\tau)$	132
Bảng 4.5. Độ phức tạp của thuật toán thực hiện các công thức trong mô hình L-FAHP - L-FTOPSIS động dựa trên HA	137
Bảng 4.6. Đánh giá trung bình có trọng số của các doanh nghiệp tại thời điểm τ_1	141
Bảng 4.7. Khoảng cách từ các doanh nghiệp tới $\bar{b}^+$ và $\bar{b}^-$ tại thời điểm τ_1	141
Bảng 4.8. Hệ số gần và thứ hạng của các doanh nghiệp tại thời điểm τ_1	142
Bảng 4.9. Đánh giá trung bình có trọng số của các doanh nghiệp tại thời điểm τ_2	142
Bảng 4.10. Khoảng cách từ các doanh nghiệp tới $\bar{b}^+$ và $\bar{b}^-$ tại thời điểm τ_2	142
Bảng 4.11. Hệ số gần và thứ hạng của các doanh nghiệp tại thời điểm τ_2	143
Bảng 4.12. Đánh giá trung bình có trọng số của các doanh nghiệp tại thời điểm τ_3	143
Bảng 4.13. Khoảng cách từ các doanh nghiệp tới $\bar{b}^+$ và $\bar{b}^-$ tại thời điểm τ_3	144
Bảng 4.14. Hệ số gần và thứ hạng của các doanh nghiệp tại thời điểm τ_3	144
Bảng 4.15. So sánh kết quả phân tích độ nhạy trong các kịch bản tại τ_1, τ_2 và τ_3	145
Bảng 4.16. Bảng so sánh kết quả giữa hai mô hình tại thời điểm τ_1	150
Bảng 4.17. Bảng so sánh kết quả giữa hai mô hình tại thời điểm τ_2	151
Bảng 4.18. Bảng so sánh kết quả giữa hai mô hình tại thời điểm τ_3	152
Bảng A1.1. Bộ tiêu chí đánh giá mức độ sẵn sàng chuyển đổi số của SMEs trong lĩnh vực bán lẻ tại Việt Nam.....	i
Bảng A1.2. Ma trận so sánh mờ trung bình cho bảy tiêu chí chính.....	ii
Bảng A1.3. Các giá trị mở rộng tổng hợp mờ của 7 tiêu chí chính và 34 tiêu chí con	iii
Bảng A1.4. Véc-tơ trọng số của 7 tiêu chí chính và 34 tiêu chí con	iv
Bảng A1.5. Kết quả đánh giá có trọng số cho mức độ chuyển đổi số của 16 SMEs bán lẻ tại Việt Nam.	v

Bảng A1.6. Điểm số và thứ hạng cuối cùng sau khi phân tích độ nhảy ở kịch bản 1	vii
Bảng A2.1. Bộ tiêu chí đánh giá mức độ chuyển đổi số của SMEs bán lẻ tại Việt Nam	vii
Bảng A2.2. Ma trận so sánh mờ cho ba tiêu chí chính tại thời điểm τ_1	viii
Bảng A2.3. Ma trận so sánh mờ cho tiêu chí P1 tại thời điểm τ_1	viii
Bảng A2.4. Ma trận so sánh mờ cho tiêu chí P2 tại thời điểm τ_1	ix
Bảng A2.5. Ma trận so sánh mờ cho tiêu chí P3 tại thời điểm τ_1	ix
Bảng A2.6. Trọng số mờ của các tiêu chí chính và tiêu chí con tại thời điểm τ_1	ix
Bảng A2.7. Tổng hợp đánh giá của các doanh nghiệp tại thời điểm τ_1	x
Bảng A2.8. Trọng số mờ của các tiêu chí chính và tiêu chí con tại thời điểm τ_2	xi
Bảng A2.9. Tổng hợp đánh giá của các doanh nghiệp tại thời điểm τ_2	xii
Bảng A2.10. Trọng số mờ của các tiêu chí chính và tiêu chí con tại thời điểm τ_3 ..	xv
Bảng A2.11. Tổng hợp đánh giá của các doanh nghiệp tại thời điểm τ_3	xvi

DANH MỤC CÁC HÌNH VẼ

Hình 1. Cấu trúc luận án.....	9
Hình 1.1. Khung phương pháp MCGDM trong môi trường ngôn ngữ dựa trên HA.	48
Hình 2.1. Quy trình MCGDM của mô hình HA-TOPSIS.....	53
Hình 2.2. Cấu trúc phân cấp của bài toán lựa chọn nhà cung cấp	60
Hình 2.3. Kết quả phân tích độ nhạy mô hình TOPSIS ngôn ngữ dựa trên HA	65
Hình 2.4. Kết quả so sánh giữa HA-TOPSIS với FTOPSIS và 4-tuple ngữ nghĩa ..	69
Hình 3.1. Quy trình MCGDM của mô hình EFAHP - 4-tuple - HA	81
Hình 3.2. Phân tích độ nhạy theo kịch bản nhiều đơn tiêu chí.....	99
Hình 3.3. Phân tích độ nhạy theo kịch bản nhiều toàn cục	100
Hình 4.1. TFSs của các từ ngôn ngữ $x \in X_2$ của F_{L_2}	110
Hình 4.2. Quy trình MCGDM động trong môi trường ngôn ngữ dựa trên HA	126
Hình 4.3. Kịch bản phân tích độ nhạy với nhiễu loạn toàn cục.....	146
Hình 4.4. Độ nhạy của mô hình trước nhiễu loạn toàn cục và thay đổi độ đo tính mờ của HA.	148

DANH MỤC CÁC KÝ HIỆU

TT	Ký hiệu	Diễn giải
1	$\mathcal{AX}$	Đại số gia tử tuyến tính
2	$\mathcal{AX}^*$	Đại số gia tử tuyến tính đầy đủ
3	$\mu(h)$	Độ đo tính mờ của gia tử h
4	$f_m(x)$	Độ đo tính mờ của từ ngôn ngữ x
5	$\vartheta(x)$	Hàm định lượng ngữ nghĩa của từ ngôn ngữ x
6	$\mu_A(x)$	Hàm xác định độ thuộc của giá trị x vào tập mờ A
7	$l(x)$	Độ dài của từ ngôn ngữ x
8	$\mathfrak{J}_m$	Khoảng tính mờ của giá trị ngôn ngữ
9	X_k	Tập tất cả các từ ngôn ngữ có độ dài $\leq k$
10	$X_{(k)}$	Tập các từ ngôn ngữ có độ dài đúng bằng k
11	$\bar{A}$	Tập hữu hạn các phương án
12	$\bar{C}$	Tập hữu hạn các tiêu chí
13	$\bar{D}$	Tập hữu hạn người ra quyết định (chuyên gia)
14	$\oplus$	Phép kết nhập ngôn ngữ nhị phân trên thang đo 4-tuple ngữ nghĩa
15	$\widehat{\oplus}$	Phép kết nhập với dữ liệu khuyết thiếu

MỞ ĐẦU

1. Tính cấp thiết của đề tài

Ra quyết định là một hoạt động vô cùng quan trọng trong hầu hết các lĩnh vực của đời sống kinh tế - xã hội, quản lý và kỹ thuật, bởi chất lượng của quyết định có thể tác động trực tiếp đến hiệu quả đầu tư, hiệu năng vận hành, mức độ an toàn, khả năng thích ứng và sức cạnh tranh của tổ chức [13, 92, 24]. Trong thực tiễn, chất lượng của một phương án thường không thể được đánh giá đầy đủ chỉ bằng một tiêu chí đơn lẻ [77]. Các bài toán như lựa chọn nhà cung cấp giải pháp số [92, 40, 70, 24], lựa chọn nền tảng công nghệ [26, 6], xác định phương án triển khai hệ thống thông tin [92, 24, 20], đánh giá rủi ro công nghệ [85, 35], xếp hạng các phương án đầu tư hạ tầng số [10, 90], hay đánh giá mức độ sẵn sàng chuyển đổi số của doanh nghiệp [118, 73, 16, 74, 56, 102], thường phải xem xét đồng thời trên nhiều tiêu chí như chi phí, chất lượng, khả năng tích hợp, độ tin cậy, tính linh hoạt, mức độ an toàn, hiệu quả vận hành. Giữa các tiêu chí này có thể tồn tại quan hệ hỗ trợ, cạnh tranh hoặc xung đột, làm cho quá trình ra quyết định trở nên phức tạp cả về dữ liệu lẫn phương pháp [12, 76, 24].

Trong bối cảnh đó, ra quyết định đa tiêu chí (MCDM) và ra quyết định nhóm đa tiêu chí (MCGDM) là các khuôn khổ phương pháp luận quan trọng để hỗ trợ phân tích, lựa chọn, đánh giá và xếp hạng phương án. MCDM tập trung vào xử lý các phương án theo nhiều tiêu chí có bản chất và mức độ quan trọng khác nhau [13, 76], trong khi MCGDM mở rộng bài toán theo hướng kết hợp ý kiến của nhiều chuyên gia, nhà quản lý hoặc nhóm liên quan có chuyên môn, kinh nghiệm, ưu tiên và mức độ chấp nhận rủi ro khác nhau [13, 14]. Do đó, MCGDM không chỉ phải xử lý xung đột giữa các tiêu chí, mà còn phải xử lý sự khác biệt trong nhận thức, kinh nghiệm và quan điểm của các chủ thể tham gia ra quyết định [93, 13].

Một đặc điểm nổi bật của các bài toán MCGDM trong thực tiễn là thông tin đánh giá không phải lúc nào cũng tồn tại dưới dạng số rõ [75]. Đối với các tiêu chí định tính như mức độ phù hợp công nghệ, chất lượng dịch vụ, năng lực đổi mới, mức độ sẵn sàng chuyển đổi số, khả năng thích ứng hay mức độ rủi ro, chuyên gia thường sử dụng các từ ngôn ngữ thay vì các giá trị số chính xác [12, 38]. Cách biểu đạt này gần với cơ chế nhận thức tự nhiên của con người, nhưng đồng thời đặt ra yêu cầu phương pháp luận phức tạp hơn: mô hình ra quyết định phải xử lý được tính mơ hồ, tính không chắc chắn, tính chủ quan, tính phụ thuộc ngữ cảnh và cấu trúc ngữ nghĩa của các giá trị ngôn ngữ trong suốt các giai đoạn xác định trọng số của tiêu chí, kết nạp ý kiến nhóm chuyên gia và xếp hạng phương án.

Trên nền tảng lý thuyết tập mờ của Zadeh [115], nhiều phương pháp MCDM và MCGDM đã được phát triển nhằm xử lý thông tin bất định. Các phương pháp như AHP [95], TOPSIS [59, 37], và các mở rộng dựa trên lý thuyết tập mờ, tiêu biểu là FAHP [19, 58, 97, 71], FTOPSIS [21, 98, 114, 64, 45, 60, 11], các mô hình tích hợp FAHP - FTOPSIS [10, 57, 64, 87, 91], ANP - FTOPSIS [66, 69], ANP-DEMATEL [82], cùng các mở rộng dựa trên số mờ [19, 108, 25] và các dạng tập mờ khác [45, 113, 8, 105, 22] đã góp phần quan trọng vào việc biểu diễn và xử lý đánh giá bất định. Bên cạnh đó, các mô hình 2-tuple [49, 61, 36, 7, 5, 39], AHP - 2-tuple [28, 99] và 4-tuple ngữ nghĩa [54, 3] đã làm rõ hơn yêu cầu bảo toàn thông tin ngôn ngữ trong quá trình tính toán và diễn giải kết quả.

Tuy nhiên, khi đặt trong yêu cầu xử lý trực tiếp thông tin ngôn ngữ của bài toán MCGDM, các hướng tiếp cận nêu trên vẫn còn một số giới hạn cần tiếp tục nghiên cứu. Trong nhiều mô hình dựa trên tập mờ, mối liên hệ giữa từ ngôn ngữ và tập mờ biểu diễn còn phụ thuộc đáng kể vào lựa chọn hàm thuộc, dạng số mờ, tham số và kinh nghiệm của người thiết kế [46, 52, 54]. Các bước tính toán mờ, giải mờ hoặc ánh xạ xấp xỉ trở lại nhãn ngôn ngữ có thể làm suy giảm một phần khả năng diễn giải ngữ nghĩa của đánh giá ban đầu [49, 54, 48]. Đối với TOPSIS [59] và FTOPSIS [21], việc xác định giải pháp lý tưởng tích cực (PIS) và giải pháp lý tưởng tiêu cực (NIS), đo khoảng cách và suy ra thứ hạng thường diễn ra trên các biểu diễn số hoặc số mờ trung gian; vì vậy, khoảng cách số học không phải lúc nào cũng tương thích hoặc phản ánh đầy đủ mối quan hệ ngữ nghĩa giữa các từ đánh giá [63]. Đối với bài toán xác định trọng số của tiêu chí, FAHP [106] và các biến thể của nó [19, 58, 97, 71] đã cho thấy khả năng xử lý tính bất định trong so sánh cặp, nhưng các thao tác kỹ thuật như tổng hợp mờ, tính trọng tâm và giải mờ có thể làm suy giảm tính minh bạch và khả năng diễn giải ngữ nghĩa của dữ liệu đầu vào [63, 99]. Trong các mô hình tích hợp, việc xác định trọng số của tiêu chí và xếp hạng phương án trên các lớp biểu diễn khác nhau, chẳng hạn kết hợp số mờ với giá trị số rõ cũng có thể làm giảm tính thống nhất ngữ nghĩa xuyên suốt quy trình MCGDM [63, 117]. Đây là một thách thức đáng chú ý đối với các bài toán mà cả tầm quan trọng của tiêu chí và đánh giá phương án đều được chuyên gia diễn đạt bằng ngôn ngữ tự nhiên [63, 99].

Từ yêu cầu đó, lý thuyết Đại số gia tử (HA) do Nguyễn Cát Hồ và cộng sự đề xuất, cung cấp một cấu trúc toán học dựa trên tiên đề để mô hình hóa miền từ của biến ngôn ngữ như một đại số có cấu trúc trật tự ngữ nghĩa [53, 52]. Trong HA, các từ ngôn ngữ được sinh từ các phần tử sinh dưới tác động của các gia tử; đồng thời, độ đo tính mờ và hàm định lượng ngữ nghĩa cho phép ánh xạ mỗi từ ngôn ngữ tới một giá trị số theo cách vẫn bảo toàn trật tự và nội dung ngữ nghĩa của từ [53, 52, 54]. Nhờ đó, HA tạo ra một nền tảng phù hợp để xây dựng các mô hình MCGDM ngôn ngữ có khả năng kiểm soát cấu trúc ngữ nghĩa rõ hơn, từ xây dựng thang đo, định lượng ngữ nghĩa, chuẩn hóa đánh giá, kết nhập nhóm, đo khoảng cách đến giải pháp lý tưởng cho đến diễn giải kết quả xếp hạng. Các nghiên cứu dựa trên HA đã đạt được một số kết quả trong điều khiển mờ

[86, 72, 103], hồi quy mờ [104], tóm tắt dữ liệu ngôn ngữ [88], dự báo chuỗi thời gian ngôn ngữ [50, 51], hệ hỗ trợ ra quyết định [46, 47], và MCGDM [54, 3, 33, 89].

Mặc dù vậy, trong các nghiên cứu MCGDM dựa trên HA hiện có, nhiều mô hình chủ yếu được phát triển trong bối cảnh tĩnh, trọng số của tiêu chí thường được giả định trước hoặc do chuyên gia gán trực tiếp, cơ chế kiểm soát tính nhất quán trong phán đoán so sánh cặp của chuyên gia chưa được tích hợp đầy đủ trong một quy trình dựa trên HA [54, 3, 33, 89]. Đồng thời, bài toán MCGDM động trong môi trường thông tin ngôn ngữ, nơi tập phương án, bộ tiêu chí, hội đồng chuyên gia hoặc dữ liệu đánh giá có thể thay đổi theo thời gian, vẫn đặt ra yêu cầu cần có cơ chế kết nhập và tích hợp dữ liệu lịch sử với dữ liệu hiện tại một cách hợp lý [17]. Những vấn đề này cho thấy cần tiếp tục phát triển các mô hình MCGDM dựa trên HA theo hướng vừa bảo toàn cấu trúc ngữ nghĩa, vừa đáp ứng các yêu cầu tính toán trong xác định trọng số, kiểm soát nhất quán, kết nhập nhóm và xếp hạng phương án. Trên cơ sở phân tích các nghiên cứu liên quan và định hướng của luận án, có thể thấy ba khoảng trống khoa học cần tiếp tục làm rõ, như sau:

Thứ nhất, các mô hình TOPSIS [59] và FTOPSIS [21] hiện tại còn phụ thuộc vào số thực và số mờ, hàm thuộc, giải mờ hoặc ánh xạ xấp xỉ; chưa khai thác đầy đủ cấu trúc trật tự ngữ nghĩa, độ đo tính mờ và hàm định lượng ngữ nghĩa của HA để xếp hạng các phương án trên miền ngữ nghĩa.

Thứ hai, các mô hình MCGDM dựa trên HA hiện tại [54, 3, 33, 89] trọng số của tiêu chí thường giả định biết trước hoặc gán trực tiếp, chưa có quy trình từ so sánh cặp ngôn ngữ, kiểm soát tính nhất quán đến suy ra véc-tơ trọng số, nên độ tin cậy và liên thông với xếp hạng còn hạn chế. Một số mô hình tích hợp FAHP-FTOPSIS [10, 57, 87] dùng lớp biểu diễn khác nhau, chuyển đổi giữa nhãn ngôn ngữ, số mờ, số rõ hoặc giải mờ, có thể làm giảm tính minh bạch, khả năng diễn giải và thống nhất ngữ nghĩa.

Thứ ba, các mô hình MCGDM ngôn ngữ dựa trên HA [54, 3, 33, 89] hiện chủ yếu tập trung vào bối cảnh tĩnh, còn thiếu cơ chế hình thức để xử lý quá trình ra quyết định động; đồng thời cần tích hợp dữ liệu lịch sử với dữ liệu hiện tại và xử lý dữ liệu khuyết thiếu theo một cơ chế suy giảm ngữ nghĩa nhất quán.

Từ các khoảng trống trên, luận án hướng tới lựa chọn HA làm nền tảng tiếp cận để phát triển khung phương pháp luận cho bài toán MCGDM ngôn ngữ theo hướng nổi bật, có tính kế thừa và mở rộng. Cụ thể, luận án phát triển mô hình HA-TOPSIS nhằm xếp hạng phương án trên miền ngữ nghĩa thông qua thang đo ngôn ngữ, định lượng ngữ nghĩa, xác định PIS/NIS, đo khoảng cách ngữ nghĩa và tính hệ số gần. Tiếp đó, luận án kế thừa và bổ sung quy trình xác định trọng số của tiêu chí có kiểm soát tính nhất quán của các ma trận so sánh cặp ngôn ngữ, đồng thời phát triển mô hình MCGDM tích hợp EFAHP - 4-tuple - HA vừa xác định trọng số của các tiêu chí có kiểm soát tính nhất quán vừa xếp hạng các phương án liên thông trên cùng miền ngữ nghĩa. Sau đó, luận án mở rộng bài toán sang bối cảnh động, phát triển khung phương pháp xây dựng mô hình L-FAHP - L-FTOPSIS động dựa trên HA.

Những phân tích trên cho thấy đề tài có tính cấp thiết cả về lý luận, phương pháp luận và thực tiễn. Về lý luận, đề tài góp phần làm rõ hơn cơ sở biểu diễn, định lượng và xử lý thông tin ngôn ngữ trong MCGDM dựa trên HA. Về phương pháp luận, đề tài hướng tới xây dựng khung phương pháp có khả năng liên kết tương đối thống nhất các thành phần cốt lõi của quy trình MCGDM, bao gồm xây dựng thang đo ngôn ngữ, xác định trọng số của tiêu chí, kiểm soát tính nhất quán, kết nhập ý kiến chuyên gia, đo khoảng cách ngữ nghĩa, xếp hạng phương án và mở rộng sang môi trường động. Về ứng dụng, các mô hình đề xuất có khả năng áp dụng cho các bài toán như lựa chọn nhà cung cấp, đánh giá mức độ sẵn sàng chuyển đổi số của doanh nghiệp, quản trị hệ thống thông tin và đánh giá theo chu kỳ, trong điều kiện dữ liệu vừa mang tính ngôn ngữ vừa biến đổi theo thời gian dưới tác động của công nghệ, thị trường và chính sách.

Vì vậy, việc thực hiện đề tài “**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**” có cơ sở khoa học rõ ràng, gắn với nhu cầu thực tiễn của các bài toán ra quyết định, đồng thời đáp ứng yêu cầu hoàn thiện khung phương pháp luận cho MCGDM trong môi trường thông tin ngôn ngữ. Đây là cơ sở trực tiếp để luận án phát triển các mô hình HA-TOPSIS, EFAHP - 4-tuple - HA và L-FAHP - L-FTOPSIS động dựa trên HA trong các Chương tiếp theo.

2. Mục tiêu nghiên cứu

2.1. Mục tiêu tổng quát

Mục tiêu tổng quát của luận án là phát triển các mô hình MCGDM trong môi trường thông tin ngôn ngữ dựa trên HA, khai thác cấu trúc đại số, độ đo tính mờ và hàm ánh xạ định lượng ngữ nghĩa (SQM) để biểu diễn, định lượng thông tin, bảo toàn trật tự và cấu trúc ngữ nghĩa. Trên nền tảng đó, luận án hướng tới xây dựng khung phương pháp luận liên kết các thành phần cốt lõi của MCGDM: xây dựng thang đo ngôn ngữ, xác định trọng số của tiêu chí có kiểm soát tính nhất quán, kết nhập đánh giá nhóm và xếp hạng phương án trong bối cảnh tĩnh và động; qua đó xem xét tính nhất quán ngữ nghĩa, độ ổn định, độ tin cậy và khả năng diễn giải của kết quả.

2.2. Mục tiêu cụ thể

Để đạt được mục tiêu tổng quát nêu trên, luận án tập trung thực hiện các mục tiêu cụ thể sau:

(i). Hệ thống hóa cơ sở lý thuyết và nghiên cứu về MCDM/MCGDM trong môi trường thông tin ngôn ngữ, tập trung vào các tiếp cận dựa trên lý thuyết tập mờ, FAHP, FTOPSIS và HA; phân tích hạn chế, xác định khoảng trống khoa học làm cơ sở phát triển và mở rộng hướng tiếp cận HA cho các mô hình MCGDM.

(ii). Phát triển mô hình HA-TOPSIS cho MCGDM ngôn ngữ, sử dụng thang đo ngôn ngữ và hàm SQM của HA, xác định PIS/NIS, khoảng cách ngữ nghĩa, hệ số gần và xếp hạng phương án trên miền ngữ nghĩa thống nhất; kiểm chứng bằng thực nghiệm, phân tích độ nhạy và đối sánh với các mô hình liên quan.

(iii). Phát triển mô hình MCGDM tích hợp FAHP cải tiến với biểu diễn 4-tuple ngữ nghĩa theo hướng tiếp cận HA, nhằm xác định trọng số của tiêu chí có kiểm soát tính nhất quán của ma trận so sánh cặp ngôn ngữ, kết nhập đánh giá nhóm và xếp hạng phương án trên nền tảng ngữ nghĩa HA; kiểm chứng khả năng áp dụng, độ ổn định và khả năng diễn giải thông qua thực nghiệm, phân tích độ nhạy và so sánh kết quả.

(iv). Phát triển khung phương pháp xây dựng mô hình MCGDM động trong môi trường thông tin ngôn ngữ dựa trên HA, nhằm xử lý sự thay đổi theo thời gian của tập phương án, bộ tiêu chí, hội đồng chuyên gia và dữ liệu đánh giá; đồng thời xây dựng phép kết nhập ngôn ngữ, xử lý dữ liệu khuyết thiếu và cơ chế suy giảm ngữ nghĩa để tích hợp dữ liệu lịch sử với đánh giá hiện tại; phát triển mô hình L-FAHP - L-FTOPSIS động và kiểm chứng bằng thực nghiệm, phân tích độ nhạy và đối sánh kết quả.

3. Nội dung nghiên cứu

Để thực hiện các mục tiêu đề ra, luận án tập trung vào các nội dung nghiên cứu chính sau:

(i). Tổng quan, hệ thống hóa cơ sở lý thuyết của MCDM/MCGDM trong môi trường thông tin ngôn ngữ, gồm lý thuyết tập mờ, FAHP, FTOPSIS, mô hình 2-tuple, 4-tuple ngữ nghĩa và HA; phân tích các nghiên cứu liên quan, nhận diện hạn chế và xác định khoảng trống làm cơ sở phát triển các mô hình MCGDM dựa trên HA.

(ii). Phát triển mô hình TOPSIS ngôn ngữ dựa trên HA, trong đó giải quyết các vấn đề xây dựng thang đo ngôn ngữ, định lượng ngữ nghĩa, xác định PIS/NIS, đo khoảng cách, tính hệ số gần và xếp hạng phương án trên miền ngữ nghĩa thống nhất. Mô hình được thực nghiệm trên bài toán lựa chọn nhà cung cấp phần mềm, kết hợp phân tích độ nhạy và so sánh kết quả.

(iii). Phát triển mô hình MCGDM tích hợp FAHP cải tiến với biểu diễn 4-tuple ngữ nghĩa dựa trên HA. Nội dung nghiên cứu bao gồm kiểm soát tính nhất quán của ma trận so sánh cặp trong môi trường HA, xác định trọng số của tiêu chí, biểu diễn, kết nhập đánh giá nhóm và xếp hạng phương án trên cùng miền ngữ nghĩa. Mô hình được áp dụng đánh giá mức độ sẵn sàng chuyển đổi số của SMEs tại Việt Nam, kèm phân tích độ nhạy và đối sánh với các mô hình liên quan.

(iv). Phát triển khung phương pháp MCGDM động dựa trên HA để xử lý sự thay đổi theo thời gian của tập phương án, bộ tiêu chí, hội đồng chuyên gia và dữ liệu đánh giá; xây dựng phép kết nhập ngôn ngữ, xử lý dữ liệu khuyết thiếu và tích hợp dữ liệu lịch sử với dữ liệu hiện tại theo cơ chế suy giảm ngữ nghĩa; trên cơ sở đó phát triển mô hình L-FAHP - L-FTOPSIS động dựa trên HA. Mô hình được thực nghiệm qua nhiều thời điểm trên bài toán đánh giá mức độ sẵn sàng chuyển đổi số của SMEs, kết hợp phân tích độ nhạy và đối sánh.

4. Đối tượng và phạm vi nghiên cứu

4.1. Đối tượng nghiên cứu

Đối tượng nghiên cứu của luận án là bài toán MCGDM trong môi trường thông tin ngôn ngữ theo hướng tiếp cận HA. Trong đó, luận án tập trung nghiên cứu các mô hình MCGDM có khả năng biểu diễn, định lượng, kết nhập và xếp hạng phương án trên cơ sở đánh giá ngôn ngữ của nhiều chuyên gia theo nhiều tiêu chí. Trên nền tảng đó, đối tượng nghiên cứu trực tiếp của luận án là các mô hình MCGDM dựa trên HA, bao gồm:

- (i) Mô hình TOPSIS ngôn ngữ dựa trên HA cho bài toán xếp hạng phương án trong môi trường thông tin ngôn ngữ;
- (ii) Mô hình MCGDM tích hợp FAHP cải tiến với 4-tuple ngữ nghĩa dựa trên HA để xác định trọng số của tiêu chí có kiểm soát nhất quán và xếp hạng phương án;
- (iii) Phát triển khung phương pháp xây dựng mô hình MCGDM động trong môi trường thông tin ngôn ngữ dựa trên HA nhằm xử lý dữ liệu đánh giá thay đổi theo thời gian, đồng thời có khả năng tích hợp dữ liệu lịch sử theo cơ chế suy giảm ngữ nghĩa.

4.2. Phạm vi nghiên cứu

Về phương pháp luận: Luận án giới hạn nghiên cứu trong khuôn khổ bài toán MCGDM trong môi trường thông tin ngôn ngữ. Trọng tâm nghiên cứu tập trung vào ba vấn đề cốt lõi: biểu diễn và định lượng ngữ nghĩa của thang đo ngôn ngữ, xác định trọng số của tiêu chí có kiểm soát tính nhất quán trong môi trường ngôn ngữ và phát triển các phép toán kết nhập ngôn ngữ nhằm giải quyết bài toán MCGDM ngôn ngữ trong bối cảnh tĩnh và động.

Về cách tiếp cận: Luận án giới hạn trong hướng tiếp cận dựa trên HA tuyến tính và thang đo ngôn ngữ được xây dựng trên cơ sở các tham số ngữ nghĩa của HA. Các mô hình được phát triển nhằm xử lý trực tiếp thông tin ngôn ngữ trên miền từ, bảo toàn trật tự và cấu trúc ngữ nghĩa của các đánh giá, thay vì chỉ dựa vào hàm thuộc hay các bước giải mờ hoặc ánh xạ xấp xỉ theo cách tiếp cận tập mờ truyền thống. Trong phạm vi luận án, bài toán động được hiểu theo ba khía cạnh: sự lặp lại của quá trình ra quyết định qua nhiều thời điểm, sự thay đổi của tập phương án, bộ tiêu chí và hội đồng chuyên gia theo thời gian, và việc ra quyết định có xét đến dữ liệu lịch sử.

Về ứng dụng: Các mô hình đề xuất được kiểm chứng trên hai bài toán lựa chọn nhà cung cấp phần mềm và đánh giá, xếp hạng mức độ sẵn sàng chuyển đổi số của SMEs bán lẻ tại Việt Nam trong bối cảnh tĩnh và động; đồng thời thực hiện phân tích độ nhạy và so sánh với một số phương pháp liên quan. Luận án không đặt mục tiêu bao quát toàn bộ các phương pháp MCDM và MCGDM, không xử lý mọi dạng bất định phức tạp, không kiểm chứng trên tất cả các miền ứng dụng và không phát triển một hệ hỗ trợ ra quyết định hoàn chỉnh. Các mở rộng này được xem là hướng nghiên cứu tiếp theo.

5. Phương pháp nghiên cứu

Để thực hiện các mục tiêu nghiên cứu đã đề ra, luận án sử dụng kết hợp các phương pháp nghiên cứu chủ yếu sau đây:

Thứ nhất, phương pháp phân tích, tổng hợp tài liệu và phỏng vấn chuyên gia để khảo cứu có hệ thống các hướng tiếp cận MCDM và MCGDM trong môi trường thông tin ngôn ngữ theo hướng tiếp cận dựa trên lý thuyết tập mờ, mô hình 2-tuple, mô hình 4-tuple ngữ nghĩa và HA trước đó, từ đó xác định hạn chế của các mô hình hiện có và làm rõ khoảng trống nghiên cứu.

Thứ hai, phương pháp mô hình hóa toán học và tiếp cận hình thức dựa trên HA làm nền tảng cho việc biểu diễn, định lượng và xử lý thông tin ngôn ngữ; phương pháp phát triển mô hình và thuật toán để xây dựng các mô hình HA-TOPSIS, EFAHP - 4-tuple - HA và mô hình L-FAHP - L-FTOPSIS động dựa trên HA.

Thứ ba, phương pháp suy luận, chứng minh và phân tích tính chất nhằm làm rõ cơ sở khoa học, khả năng bảo toàn ngữ nghĩa, kiểm soát tính nhất quán và tích hợp dữ liệu động của các mô hình đề xuất.

Thứ tư, phương pháp thực nghiệm, phân tích độ nhạy và đối sánh được sử dụng để kiểm chứng tính đúng đắn, độ tin cậy, độ ổn định và khả năng ứng dụng của các mô hình trên các bài toán MCGDM trong thực tế.

6. Đóng góp mới của đề tài

Bám sát mục tiêu, nội dung và phạm vi nghiên cứu, đề tài đã đạt được một số đóng góp mới có ý nghĩa về phương pháp luận, mô hình và ứng dụng như sau:

(i) Luận án phát triển mô hình TOPSIS ngôn ngữ dựa trên HA, ký hiệu HA-TOPSIS, cho bài toán xếp hạng phương án trong MCGDM ngôn ngữ. Điểm mới tập trung ở việc xây dựng thang đo ngôn ngữ trên nền HA và sử dụng hàm SQM của HA để xác định PIS/NIS, khoảng cách ngữ nghĩa và hệ số gần trên cùng miền ngữ nghĩa. Mô hình kế thừa nguyên lý so sánh với giải pháp lý tưởng của TOPSIS [59], nhưng phát triển quy trình tính toán theo hướng xử lý trực tiếp thông tin ngôn ngữ dựa trên HA, thay vì phụ thuộc vào biểu diễn số mờ, giải mờ hoặc ánh xạ xấp xỉ như trong nhiều mô hình FTOPSIS [21, 98, 114, 64, 45, 60, 11]. Qua đó, các thành phần cốt lõi của TOPSIS được thực hiện trong một không gian định lượng ngữ nghĩa thống nhất, phù hợp với trật tự của thang đo HA. Mô hình được kiểm chứng trên bài toán lựa chọn nhà cung cấp phần mềm; phân tích độ nhạy và đối sánh được sử dụng để đánh giá tính hợp lý, độ ổn định và khả năng diễn giải trong phạm vi thực nghiệm. Đóng góp này được trình bày tại Chương 2 và công bố trong [CT1].

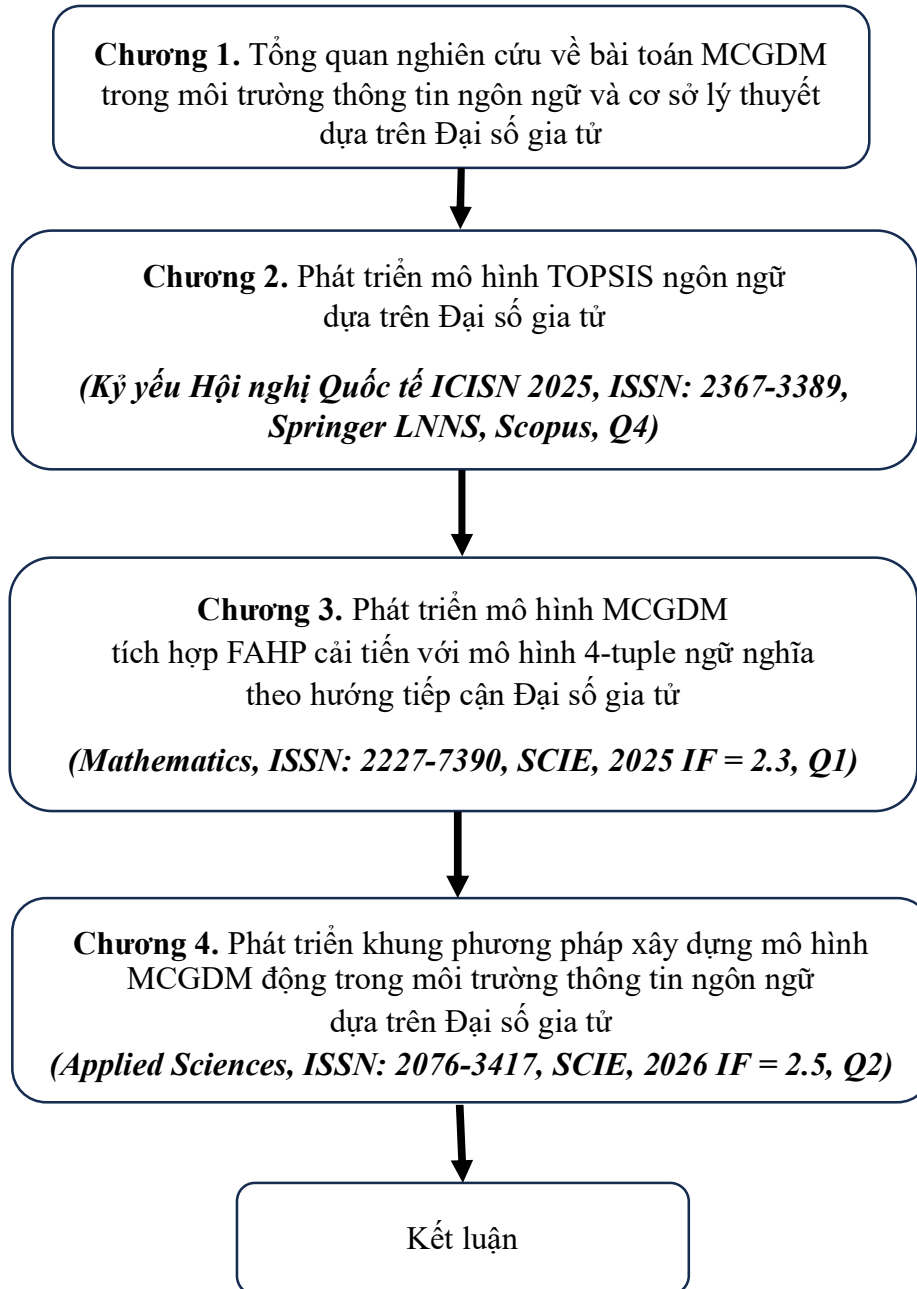
(ii) Luận án phát triển mô hình MCGDM tích hợp FAHP cải tiến với biểu diễn 4-tuple ngữ nghĩa theo hướng tiếp cận HA, ký hiệu EFAHP - 4-tuple - HA, liên kết xác định trọng số của tiêu chí với kết nhập đánh giá nhóm và xếp hạng phương án trên cùng nền tảng ngữ nghĩa. Điểm mới của mô hình thể hiện ở cơ chế kiểm soát tính nhất quán của ma trận so sánh cặp ngôn ngữ trong môi trường HA; xác định trọng số bằng EFAHP theo hướng tiếp cận HA; đồng thời biểu diễn, kết nhập đánh giá chuyên gia và xếp hạng phương án bằng 4-tuple ngữ nghĩa dựa trên HA. Trên cơ sở kế thừa FAHP và mô hình 4-tuple ngữ nghĩa [58, 54], cách tích hợp này tạo sự liên thông ngữ nghĩa giữa giai đoạn xác định mức độ quan trọng của tiêu chí và giai đoạn đánh giá phương án. Mô hình được kiểm chứng trên bài toán đánh giá mức độ sẵn sàng chuyển đổi số của 16 SMEs bán lẻ tại Việt Nam, kết hợp phân tích độ nhạy và đối sánh để xem xét tính khả thi, độ ổn định và khả năng diễn giải trong phạm vi nghiên cứu. Đóng góp này được trình bày tại Chương 3 và công bố trong [CT2].

(iii) Luận án phát triển khung phương pháp xây dựng mô hình MCGDM động trong môi trường thông tin ngôn ngữ dựa trên HA, mở rộng hướng tiếp cận HA từ MCGDM tĩnh sang quá trình ra quyết định qua nhiều thời điểm. Điểm mới gồm phát triển phép kết nhập ngôn ngữ nhị phân trên thang đo 4-tuple ngữ nghĩa với các tính chất đại số cần thiết; xây dựng phép kết nhập cho dữ liệu khuyết thiếu; và thiết lập cơ chế suy giảm ngữ nghĩa để tích hợp dữ liệu lịch sử với đánh giá hiện tại. Trên cơ sở đó, luận án phát triển mô hình L-FAHP - L-FTOPSIS động dựa trên HA để xác định trọng số của tiêu chí và xếp hạng phương án khi tập phương án, bộ tiêu chí, hội đồng chuyên gia hoặc dữ liệu đánh giá thay đổi theo thời gian. Mô hình được kiểm chứng trên bài toán đánh giá mức độ sẵn sàng chuyển đổi số của SMEs bán lẻ tại Việt Nam trong bối cảnh động, kết hợp phân tích độ nhạy và so sánh kết quả. Đóng góp này được hiện thực trong Chương 4 và công bố trong [CT3].

Tổng hợp các kết quả trên, luận án đã hình thành một khung phương pháp luận với ba mô hình đề xuất cho bài toán MCGDM ngôn ngữ dựa trên HA theo hướng kế thừa và mở rộng. HA-TOPSIS tập trung vào xếp hạng trên miền ngữ nghĩa HA; EFAHP - 4-tuple - HA bổ sung xác định trọng số của tiêu chí có kiểm soát tính nhất quán và liên kết với kết nhập, xếp hạng; khung MCGDM động tiếp tục mở rộng quy trình sang nhiều thời điểm, xử lý dữ liệu khuyết thiếu và tích hợp dữ liệu lịch sử. Các đóng góp này phù hợp với phạm vi nghiên cứu của luận án và tạo cơ sở cho các nghiên cứu tiếp theo về MCGDM ngôn ngữ dựa trên HA trong thực tiễn.

7. Kết cấu của luận án

Luận án “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” gồm: Phần mở đầu, bốn chương nội dung (xem Hình 1), phần kết luận chung, danh mục công trình công bố, tài liệu tham khảo và các phụ lục, cụ thể như sau:



Hình 1. Cấu trúc luận án

Phần mở đầu: Trình bày tính cấp thiết của đề tài, mục tiêu, đối tượng và phạm vi nghiên cứu, phương pháp nghiên cứu, khoảng trống khoa học, các đóng góp mới và bố cục tổng thể của luận án.

Chương 1: Tổng quan nghiên cứu về bài toán MCGDM trong môi trường thông tin ngôn ngữ và cơ sở lý thuyết dựa trên HA: Hệ thống hóa các mô hình MCDM/ MCGDM, các hướng tiếp cận dựa trên lý thuyết mờ, các mở rộng, FAHP, FTOPSIS và HA trong môi trường thông tin ngôn ngữ; phân tích những hạn chế khi xử lý thông tin ngôn ngữ, xác định trọng số của tiêu chí và MCGDM động; đồng thời hệ thống hóa cơ sở lý thuyết về HA được sử dụng trong các chương tiếp theo.

- *Chương 2:* Phát triển mô hình TOPSIS ngôn ngữ dựa trên HA [CT1]: tập trung phát triển mô hình HA-TOPSIS cho bài toán MCGDM trong môi trường thông tin ngôn ngữ. Nội dung chương gồm xây dựng thang đo ngôn ngữ dựa trên HA, định lượng ngữ nghĩa các từ ngôn ngữ, thiết lập quy trình HA-TOPSIS và kiểm chứng mô hình qua bài toán lựa chọn nhà cung cấp, kết hợp phân tích độ nhạy và so sánh với các mô hình 4-tuple ngữ nghĩa [54], FTOPSIS [21].

- *Chương 3:* Phát triển mô hình MCGDM tích hợp FAHP cải tiến với mô hình 4-tuple ngữ nghĩa theo hướng tiếp cận HA [CT2]: Chương này tập trung xây dựng thang đo 4-tuple ngữ nghĩa, đề xuất công thức tính chỉ số nhất quán mới trong môi trường HA, xác định trọng số của tiêu chí bằng phương pháp EFAHP-HA và áp dụng mô hình 4-tuple ngữ nghĩa để xếp hạng các phương án trên miền ngôn ngữ; đồng thời kiểm chứng mô hình qua bài toán đánh giá mức độ chuyển đổi số của SMEs tại Việt Nam kèm phân tích độ nhạy, so sánh với các mô hình HA-TOPSIS [CT1], FAHP-FTOPSIS [58, 21].

- *Chương 4:* Phát triển khung phương pháp xây dựng mô hình MCGDM động trong môi trường thông tin ngôn ngữ dựa trên HA [CT3]: Nội dung chương này đề xuất phép kết nhập từ ngôn ngữ 4-tuple ngữ nghĩa nhị phân đảm bảo các tính chất đại số và phép kết nhập với dữ liệu khuyết thiếu có cơ chế suy giảm ngữ nghĩa để tích hợp dữ liệu lịch sử với dữ liệu hiện tại và xây dựng mô hình hai giai đoạn L-FAHP - HA và L-FTOPSIS - HA nhằm giải quyết bài toán MCGDM động trong điều kiện tập phương án, bộ tiêu chí và hội đồng chuyên gia có thể thay đổi theo thời gian. Mô hình được thực nghiệm trên bài toán đánh giá mức độ chuyển đổi số của SMEs qua nhiều thời điểm, kết hợp phân tích độ nhạy và so sánh với phương pháp đối chứng.

- *Kết luận chung và các phụ lục:* Tóm tắt các kết quả đạt được, nhấn mạnh các đóng góp mới của luận án về phương pháp luận và ứng dụng, đồng thời nêu ra một số hạn chế và định hướng nghiên cứu tiếp theo. Các phụ lục cung cấp chi tiết các bảng số liệu và kết quả tính toán nhằm hỗ trợ việc kiểm chứng các mô hình đề xuất.

Các kết quả chính của luận án đã được công bố trong 03 công trình khoa học: 01 báo cáo tại Hội nghị Quốc tế ICISN 2025 (*Scopus Q4*) và 02 bài báo trên các tạp chí quốc tế thuộc danh mục SCIE (*Mathematics* được WoS và Scopus xếp hạng **Q1 top 10%** và *Applied Sciences*, *Q2*), qua đó cho thấy tính mới và giá trị khoa học của các kết quả nghiên cứu.

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

Chương này trình bày tổng quan nghiên cứu và các cơ sở lý thuyết nền tảng phục vụ phát triển các mô hình MCGDM trong các chương tiếp theo. Nội dung tập trung vào bài toán MCGDM, cơ sở lý thuyết, các phương pháp nền, tổng quan nghiên cứu liên quan, xác định khoảng trống nghiên cứu và kết luận chương

1.1. Bài toán ra quyết định nhóm đa tiêu chí

1.1.1. Bối cảnh bài toán MCDM và MCGDM

Ra quyết định trong các lĩnh vực kỹ thuật, quản lý và tổ chức hệ thống thường đòi hỏi xem xét đồng thời nhiều tiêu chí, bởi chất lượng của một phương án khó được phản ánh đầy đủ chỉ bằng một tiêu chí đơn lẻ [13, 92, 76]. Chẳng hạn trong việc lựa chọn giải pháp công nghệ, kiến trúc hệ thống, nhà cung cấp, phương án đầu tư hay mức độ sẵn sàng chuyển đổi số thường đòi hỏi xem xét đồng thời nhiều tiêu chí như chi phí, hiệu năng, độ tin cậy, tính bảo mật, khả năng tích hợp, tính linh hoạt và mức độ đáp ứng nghiệp vụ; giữa các tiêu chí này lại có thể tồn tại quan hệ hỗ trợ, cạnh tranh hoặc xung đột, làm cho việc xác định mức độ ưu tiên tổng thể của phương án trở nên phức tạp [92, 12, 24]. Chính đặc trưng đó hình thành bài toán MCDM, trong đó mục tiêu là hỗ trợ lựa chọn, xếp hạng hoặc phân loại các phương án theo nhiều tiêu chí đánh giá đồng thời [14, 76]. Trong MCDM, cấu trúc ưu tiên thường được xem xét từ góc nhìn của một người ra quyết định nhằm xác định phương án hoặc thứ tự ưu tiên phù hợp.

Khi quá trình ra quyết định có sự tham gia của nhiều người ra quyết định/chuyên gia, bài toán được mở rộng thành MCGDM [13]. Nếu trong MCDM trọng tâm chủ yếu nằm ở cơ chế tổng hợp đa tiêu chí dưới góc nhìn của một người ra quyết định, thì trong MCGDM cần đồng thời xử lý thêm sự khác biệt về tri thức, kinh nghiệm, ưu tiên và quan điểm đánh giá giữa nhiều chuyên gia [14, 13, 29]. Do đó, MCGDM không chỉ là mở rộng MCDM về số lượng người tham gia, mà là một khuôn khổ mô hình hóa trong đó quyết định cuối cùng được hình thành từ mức độ quan trọng của các tiêu chí, đánh giá phương án và ý kiến của nhóm chuyên gia [13, 14]. Vì vậy, MCGDM đồng thời bao hàm ba đặc trưng nền tảng: tính đa tiêu chí, tính đa chủ thể trong điều kiện các chuyên gia có thể khác nhau về chuyên môn, kinh nghiệm và tính xung đột tiềm ẩn giữa các quan điểm đánh giá [13].

Trong hầu hết các lĩnh vực, cách tiếp cận MCGDM là phù hợp do phần lớn các quyết định quan trọng đều mang bản chất liên ngành, đòi hỏi dung hòa đồng thời góc nhìn kỹ thuật, quản trị và nghiệp vụ [92, 12]. Chẳng hạn trong bài toán lựa chọn nhà cung cấp

phần mềm hay giải pháp hạ tầng số, quyết định cuối cùng thường phải cân bằng giữa yêu cầu của nhà quản lý, chuyên gia kỹ thuật, chuyên gia nghiệp vụ và người sử dụng cuối [92, 38]; tương tự, trong bài toán đánh giá mức độ sẵn sàng chuyển đổi số, việc xem xét các trụ cột về chiến lược, công nghệ, dữ liệu, nhân lực và quy trình cũng thường đòi hỏi sự tham gia của nhiều nhóm chuyên gia có góc nhìn khác nhau [38].

Trong phạm vi luận án này, MCGDM được xem là khuôn khổ phương pháp luận trọng tâm. Ba chương phương pháp tiếp theo của luận án lần lượt xử lý ba mô hình ra quyết định cốt lõi của bài toán MCGDM trong môi trường thông tin ngôn ngữ: xếp hạng phương án, xác định trọng số của tiêu chí và tích hợp dữ liệu nhiều thời điểm trong bối cảnh động. Vì vậy, việc phân biệt rõ MCDM và MCGDM ở đây không chỉ mang ý nghĩa khái niệm, mà còn có vai trò xác định đúng lớp bài toán mà luận án hướng tới.

1.1.2. Phát biểu bài toán MCGDM trong môi trường thông tin ngôn ngữ

Về hình thức, một bài toán MCGDM có thể được mô tả bởi ba thành phần cơ bản là tập phương án, tập tiêu chí và tập người ra quyết định [79, 14, 13], ký hiệu như sau:

$\bar{A} = \{A_1, A_2, \dots, A_m\}, (m \geq 2)$ là tập hữu hạn các phương án;

$\bar{C} = \{C_1, C_2, \dots, C_n\}, (n \geq 2)$ là tập hữu hạn các tiêu chí đánh giá;

$\bar{D} = \{D_1, D_2, \dots, D_h\}, (h \geq 2)$ là tập hữu hạn người ra quyết định (chuyên gia).

Với mỗi chuyên gia $D_e \in \bar{D}; e = 1, \dots, h$, đưa ra đánh giá đối với phương án $A_i \in \bar{A}; i = 1, \dots, m$, theo tiêu chí $C_j \in \bar{C}; j = 1, \dots, n$, 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:

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

Trong thực tiễn, trọng số có thể được cung cấp trực tiếp dưới dạng số thực, số mờ hoặc giá trị ngôn ngữ; cũng có thể được xác định gián tiếp thông qua ma trận so sánh cặp hoặc phương pháp xác định trọng số phù hợp. Trên cơ sở đó, đầu ra của bài toán là phương án được lựa chọn, thứ hạng các phương án, hoặc kết quả phân lớp các phương án theo những mức xác định trước.

Trong phạm vi luận án này, bài toán MCGDM được xét trong môi trường thông tin ngôn ngữ, tức là các đánh giá a_{ij}^e và trong nhiều trường hợp cả trọng số của tiêu chí được biểu diễn chủ yếu bằng các từ ngôn ngữ thay vì các giá trị số rõ [3]. Điều đó làm thay đổi bản chất của đối tượng xử lý: mô hình không chỉ tổng hợp dữ liệu định lượng, mà phải tổng hợp các đánh giá ngôn ngữ vốn mang tính mơ hồ, không chắc chắn, chủ quan và có cấu trúc ngữ nghĩa như thấp, cao, quan trọng, rất quan trọng, không quan trọng, v.v. Vì vậy, yêu cầu đặt ra là làm thế nào để kết nhập các giá trị ngôn ngữ trong quá trình tính toán. Việc kết nhập này phải bảo toàn được trật tự và nội dung ngữ nghĩa của thông tin đầu vào [3, 54]. Chính từ yêu cầu này, luận án lựa chọn hướng tiếp cận dựa trên HA [53, 52, 54], nhằm phát triển các mô hình và giải quyết bài toán này trong cả môi trường tĩnh và môi trường động trong các Chương tiếp theo.

1.1.3. Đặc điểm của bài toán MCGDM trong môi trường ngôn ngữ

Trong nhiều bài toán MCGDM, dữ liệu đầu vào không phải lúc nào cũng có thể biểu diễn đầy đủ bằng các giá trị số rõ. Khi đánh giá các tiêu chí như mức độ phù hợp công nghệ, khả năng tích hợp, chất lượng dịch vụ, mức độ sẵn sàng chuyển đổi số hay tầm quan trọng của tiêu chí, các chuyên gia thường sử dụng các từ ngôn ngữ thay cho các giá trị số chính xác [38, 24]. Cách biểu đạt này phù hợp với cơ chế nhận thức tự nhiên của con người, nhưng đồng thời tạo ra một lớp khó khăn mới cho mô hình hóa và tính toán [49]. Thông tin ngôn ngữ trong bài toán MCGDM có bốn đặc trưng bản chất. Thứ nhất là tính mơ hồ, bởi ranh giới ngữ nghĩa giữa các từ như quan trọng, rất quan trọng hay cực kỳ quan trọng không tách bạch tuyệt đối [78]. Thứ hai là tính không chắc chắn, do chuyên gia thường đánh giá trong điều kiện thông tin chưa đầy đủ hoặc môi trường ra quyết định đang biến động [75]. Thứ ba là tính chủ quan, bởi cùng một từ có thể được các chuyên gia khác nhau sử dụng với cường độ ngữ nghĩa không hoàn toàn giống nhau, tùy theo chuyên môn, kinh nghiệm và mức độ ưu tiên của họ [78]. Cuối cùng, thông tin ngôn ngữ còn mang tính phụ thuộc ngữ cảnh, bởi ý nghĩa của một từ như cao hay quan trọng có thể thay đổi theo miền ứng dụng và theo tiêu chí đang xét.

Quan trọng hơn là các giá trị ngôn ngữ không chỉ là các nhãn rời rạc có thứ tự, mà còn có cấu trúc ngữ nghĩa nội tại. Chúng thể hiện quan hệ về cường độ, mức độ tăng giảm và sự biến đổi ngữ nghĩa dưới tác động của các gia tử [54]. Vì vậy, đối tượng xử lý trong MCGDM ngôn ngữ không còn chỉ là dữ liệu đánh giá, mà là các đánh giá ngôn ngữ cần được tổng hợp theo cách vẫn duy trì được trật tự và nội dung ngữ nghĩa của thông tin đầu vào.

Từ góc độ phương pháp luận, điều đó dẫn tới ba yêu cầu cơ bản đối với mô hình giải quyết bài toán MCGDM ngôn ngữ. Thứ nhất, mô hình phải có khả năng biểu diễn thông tin ngôn ngữ theo cách vừa tính toán được vừa phản ánh được cấu trúc ngữ nghĩa của từ ngôn ngữ. Thứ hai, mô hình phải hỗ trợ tổng hợp ý kiến của nhiều chuyên gia mà không làm suy giảm đáng kể ý nghĩa ban đầu của thông tin. Thứ ba, mô hình phải cho phép xác định trọng số, đo khoảng cách, xếp hạng hoặc lựa chọn phương án trên cơ sở một cơ chế định lượng có ý nghĩa ngữ nghĩa. Đây chính là cơ sở để lựa chọn lý thuyết HA như một nền tảng phù hợp cho việc phát triển các mô hình MCGDM ngôn ngữ trong luận án.

1.1.4. Phạm vi của bài toán MCGDM

Xét theo mục tiêu xử lý tập phương án, các bài toán MCDM và MCGDM thường được phân thành ba dạng cơ bản là bài toán lựa chọn, bài toán xếp hạng và bài toán phân loại [14, 43]. Trong đó, bài toán lựa chọn hướng tới xác định một hoặc một số ít phương án ưu tiên; bài toán xếp hạng nhằm sắp xếp toàn bộ tập phương án theo thứ tự ưu tiên; còn bài toán phân loại tập trung vào việc gán phương án vào các lớp hoặc mức đã xác định trước [14]. Luận án này tập trung chủ yếu vào bài toán lựa chọn và xếp hạng phương án của bài toán MCGDM có dữ liệu đầu vào là thông tin ngôn ngữ. Sự lựa chọn này phù hợp với các bối cảnh ứng dụng như lựa chọn nhà cung cấp giải pháp số và đánh giá mức độ sẵn sàng chuyển đổi số của doanh nghiệp [92, 38, 24].

Xét theo quy trình ra quyết định, luận án không theo đuổi đồng thời mọi hướng giải quyết của MCGDM ngôn ngữ, mà tập trung vào hai giai đoạn cốt lõi có quan hệ trực tiếp tới độ tin cậy của quyết định: xác định trọng số của tiêu chí và xếp hạng phương án. Giai đoạn 1 liên quan đến việc xác định tầm quan trọng tương đối của các tiêu chí trong điều kiện đánh giá ngôn ngữ và nhiều chuyên gia. Giai đoạn 2 liên quan đến việc tổng hợp đánh giá và xác định thứ tự ưu tiên của các phương án trên một không gian biểu diễn ngữ nghĩa nhất quán. Trên cơ sở đó, luận án lựa chọn TOPSIS [59] làm phương pháp nền để phát triển thành mô hình HA-TOPSIS nhằm xếp hạng các phương án, bởi đây là một trong những hướng tiếp cận có logic rõ ràng, được sử dụng rộng rãi và phù hợp để phát triển thành mô hình xếp hạng trên miền ngôn ngữ; đồng thời lựa chọn FAHP [15, 58] làm phương pháp nền để phát triển mô hình cho bài toán xác định trọng số của tiêu chí, vì đây là nhóm phương pháp có vai trò quan trọng trong xử lý so sánh cặp và xác định véc-tơ trọng số của tiêu chí trong MCGDM.

Xét theo diễn biến của môi trường ra quyết định, luận án tiếp cận bài toán MCGDM theo hai dạng: bài toán tĩnh và bài toán động. Trước hết, Chương 2 và Chương 3 tập trung vào bối cảnh tĩnh, trong đó tập phương án, tập tiêu chí, tập chuyên gia và dữ liệu đánh giá được xác định tại một thời điểm ra quyết định. Trên nền tảng đó, Chương 4 mở rộng sang bối cảnh động, nơi quá trình ra quyết định diễn ra qua nhiều thời điểm, đồng thời có thể xuất hiện sự thay đổi của tập phương án, bộ tiêu chí, hội đồng chuyên gia và dữ liệu đánh giá. Trong trường hợp này, bài toán không chỉ đòi hỏi tổng hợp thông tin tại một thời điểm, mà còn đòi hỏi cơ chế tích hợp dữ liệu lịch sử với dữ liệu hiện tại theo một nguyên tắc nhất quán về ngữ nghĩa. Vì vậy, phạm vi nghiên cứu của luận án có thể được xem như một tiến trình phát triển có chủ đích: từ xếp hạng phương án, sang xác định trọng số và tích hợp đánh giá trong cùng một khuôn khổ ngữ nghĩa, rồi tiếp tục mở rộng sang bài toán MCGDM động có xét đến dữ liệu lịch sử. Đây cũng chính là cơ sở logic trực tiếp cho sự hình thành của Chương 2 với mô hình HA-TOPSIS, Chương 3 với mô hình EFAHP - 4-tuple - HA, và Chương 4 với mô hình L-FAHP - L-FTOPSIS động dựa trên HA.

1.1.5. Vai trò của bài toán trong các ứng dụng

Ý nghĩa của bài toán MCGDM ngôn ngữ trở nên rõ hơn trong lĩnh vực CNTT, nơi nhiều quyết định quan trọng phải được đưa ra trên cơ sở các tiêu chí vừa định lượng vừa định tính, trong khi dữ liệu đánh giá thường không đầy đủ và chịu ảnh hưởng mạnh của tri thức chuyên gia [92]. Trong các bài toán như lựa chọn nhà cung cấp phần mềm [92, 70], lựa chọn nền tảng công nghệ [26], đánh giá mức độ sẵn sàng chuyển đổi số của SMEs [75, 56], người ra quyết định không thể chỉ dựa vào một chỉ số kỹ thuật đơn lẻ, mà phải đồng thời cân nhắc nhiều tiêu chí có bản chất khác nhau như chi phí, khả năng tích hợp, độ tin cậy, mức độ an toàn, khả năng mở rộng, chất lượng dịch vụ và mức độ phù hợp với yêu cầu nghiệp vụ [12].

Điểm đáng chú ý là trong các bối cảnh này, nhiều đánh giá đầu vào không tồn tại dưới dạng số rõ mà chủ yếu được biểu diễn bằng các giá trị ngôn ngữ [93], chẳng hạn như mức độ phù hợp công nghệ, mức độ sẵn sàng, mức độ rủi ro, chất lượng dịch vụ hay tầm quan trọng của tiêu chí [62, 93, 78]. Do đó, việc xử lý bài toán không đơn thuần là tổng hợp dữ liệu, mà là tổng hợp các đánh giá ngôn ngữ có cấu trúc ngữ nghĩa, trong đó yêu cầu bảo toàn ý nghĩa của từ ngôn ngữ, tổng hợp nhất quán ý kiến của nhiều chuyên gia, xác định trọng số của tiêu chí có độ tin cậy và xếp hạng phương án một cách có khả năng diễn giải là những vấn đề trọng tâm. Chính những đặc điểm này khiến MCGDM ngôn ngữ trở thành một khuôn khổ phù hợp hơn so với các mô hình ra quyết định chỉ dựa trên dữ liệu số rõ.

Từ góc độ đó, lựa chọn hướng tiếp cận dựa trên HA trong luận án không chỉ có ý nghĩa lý thuyết, mà còn có ý nghĩa phương pháp luận đối với các bài toán ứng dụng. HA cho phép luận án tiếp cận một nền tảng biểu diễn và xử lý thống nhất hơn với thông tin ngôn ngữ, từ xây dựng thang đo, định lượng ngữ nghĩa và đo khoảng cách, đến kết nhập đánh giá, xác định trọng số của tiêu chí và mở rộng sang bối cảnh động [53, 52, 54]. Vì vậy, mục 1.1 không chỉ xác định đối tượng bài toán mà luận án tập trung giải quyết, mà còn tạo tiền đề trực tiếp cho mục 1.2 về cơ sở lý thuyết nền tảng, mục 1.3 về các phương pháp nền mà luận án kế thừa và phát triển, mục 1.4 về tổng quan nghiên cứu và mục 1.5 về khoảng trống nghiên cứu và định hướng phát triển của luận án.

1.2. Cơ sở lý thuyết nền tảng

1.2.1. Lý thuyết tập mờ và biến ngôn ngữ

1.2.1.1. Tập mờ

Khái niệm tập mờ do Zadeh đề xuất vào năm 1965, là một nền tảng quan trọng cho việc mô hình hóa thông tin mơ hồ và không chắc chắn [115]. Khác với tập hợp cổ điển, trong đó quan hệ thuộc chỉ nhận hai giá trị 0 hoặc 1, tập mờ cho phép một phần tử thuộc vào tập với các mức độ khác nhau thông qua hàm thuộc nhận giá trị trong đoạn $[0, 1]$. Nhờ đó, các đánh giá định tính bằng từ ngôn ngữ có thể được chuyển thành các thực thể toán học đủ linh hoạt để phục vụ tính toán trong MCDM và MCGDM.

Định nghĩa 1.1. [115] Cho Z là không gian nền. Một tập mờ F trên Z được xác định là tập hợp mà mỗi phần tử được biểu diễn bởi cặp $(x, \mu_F(x))$, trong đó $x \in Z$ và μ_F là ánh xạ:

$$\mu_F: Z \rightarrow [0, 1], \quad (1.3)$$

trong đó, μ_F được gọi là hàm thuộc của tập F . Tập Z gọi là cơ sở của tập mờ F , ký hiệu $F = \{ \langle x, \mu_F(x) \rangle \mid x \in Z \}$, $\mu_F(x)$ biểu thị mức độ thuộc của x vào F . Giá trị $\mu_F(x)$ càng gần 1 thể hiện mức độ thuộc càng cao.

Nếu $\mu_F(x) = 1$, biểu thị x thuộc hoàn toàn vào F , còn $\mu_F(x) = 0$ biểu thị x không thuộc F . Khác với tập hợp kinh điển chỉ nhận hai giá trị thuộc hoặc không thuộc, tập mờ cho phép biểu diễn mức độ thuộc liên tục trong khoảng $[0, 1]$.

Trường hợp đặc biệt, khi $\mu_F(x)$ chỉ nhận hai giá trị 0 hoặc 1 với mọi $x \in Z$, tập F trở thành một tập kinh điển thông thường.

Một số hàm thuộc phổ biến trong ứng dụng tập mờ:

Tập mờ hình tam giác: $\mu_F(x) = \max\left(\min\left(\frac{x-a}{b-a}, \frac{c-x}{c-b}\right), 0\right)$,

Tập mờ hình thang: $\mu_F(x) = \max(\min(\frac{x-a}{b-a}, \frac{d-x}{d-c}, 1), 0)$,

Hàm thuộc Gauss: $\mu_F(x) = e^{-\frac{(b-x)^2}{2\sigma^2}}$,

trong đó, $a, b, c, d, \sigma, \dots$ là các tham số của hàm thuộc tương ứng.

Trong các ứng dụng, các hàm thuộc hình tam giác, hình thang và Gauss thường được sử dụng rộng rãi để biểu diễn số mờ. Tuy nhiên, trong luận án này, vai trò quan trọng nhất của lý thuyết tập mờ không nằm ở bản thân từng dạng hàm thuộc, mà ở chỗ nó mở ra khả năng chuyển thông tin ngôn ngữ mơ hồ thành đối tượng có thể tính toán được. Đồng thời, cũng chính từ đây xuất hiện giới hạn phương pháp luận đầu tiên: ngữ nghĩa của từ ngôn ngữ thường bị gán với lựa chọn hàm thuộc và các tham số thiết kế, nên khó bảo đảm tính nhất quán hình thức giữa ngôn ngữ tự nhiên và biểu diễn toán học [53, 52, 54].

1.2.1.2. Biến ngôn ngữ

Để tiến gần hơn tới việc tính toán với từ, Zadeh tiếp tục đề xuất khái niệm biến ngôn ngữ [116]. Nếu tập mờ cung cấp công cụ biểu diễn tính không chắc chắn, thì biến ngôn ngữ cho phép chính các từ hoặc phát biểu ngôn ngữ trở thành giá trị của biến, thay vì chỉ là các giá trị số. Đây là bước phát triển quan trọng, bởi trong nhiều bài toán MCGDM, chuyên gia thực sự không đánh giá bằng số, mà bằng các từ như thấp, trung bình, cao, rất cao hoặc quan trọng, rất quan trọng. Khái niệm biến ngôn ngữ được Zadeh hình thức hóa như sau:

Định nghĩa 1.2. [116] Một biến ngôn ngữ được mô tả bởi bộ 5 thành phần $(\mathbb{X}, T(\mathbb{X}), \mathbb{U}, \mathbb{R}, \mathbb{M})$, trong đó $\mathbb{X}$: tên biến, $T(\mathbb{X})$: tập các giá trị ngôn ngữ của biến $\mathbb{X}$, $\mathbb{U}$: không gian các giá trị số của biến, $\mathbb{R}$: quy tắc cú pháp sinh ra các từ ngôn ngữ $T(\mathbb{X})$, $\mathbb{M}$: tập các luật ngữ nghĩa nhằm gán ngữ nghĩa của mỗi nhãn ngôn ngữ cho một tập mờ trên nền $\mathbb{U}$.

Ví dụ 1.1. Xét biến ngôn ngữ X có tên AGE, với miền tham chiếu $U = [0, 110]$. Miền giá trị ngôn ngữ của X được biểu diễn bởi tập các từ ngôn ngữ: $T(AGE) = \{very\ old, old, little\ old, little\ young, young, very\ young, \dots\}$. Để gán ngữ nghĩa định lượng cho một từ ngôn ngữ, ta sử dụng một ánh xạ ngữ nghĩa M từ các từ ngôn ngữ sang các tập mờ trên U . Chẳng hạn, với từ ngôn ngữ “old”, ta định nghĩa:

$$M(old) = \{(u, \mu_{old}(u)) \mid u \in [0, 110]\},$$

trong đó hàm thuộc $\mu_{old}(u)$ được chọn theo dạng:

$$\mu_{old}(u) = \max\left(\min\left(1, \frac{u-50}{20}\right), 0\right).$$

Cách xây dựng này phản ánh trực giác ngữ nghĩa: mức độ *old* bắt đầu tăng rõ rệt sau ngưỡng khoảng 50 tuổi và tiến dần đến mức *extremely old* khi tuổi tăng đủ lớn.

Ngữ nghĩa của các từ ngôn ngữ khác trong $T(AGE)$ có thể tính thông qua các tập mờ của các giá trị phần tử sinh (*old, young*) bởi các phép toán tương ứng với các gia tử tác động (*very, rather, little, ...*) thường được mô hình hóa bởi các phép tăng cường (phép bình phương – CON), hoặc làm giảm (phép căn bậc hai - DIL) trên hàm thuộc [116]. Nhờ đó, ta có thể xây dựng ngữ nghĩa cho *very old, rather old, little old, ...* một cách hệ thống dựa trên $M(old)$.

Ngoài ra, các biểu thức ngôn ngữ có chứa liên từ logic như and, or, not, ... được xử lý thông qua các toán tử chuẩn của logic mờ, tương ứng với t-norm, t-conorm và negation [2]. Điều này cho phép triển khai các phép tổng hợp và suy luận trên thông tin ngôn ngữ theo cơ chế hình thức, nhưng vẫn giữ được ý nghĩa gần với suy nghĩ tự nhiên của con người. Từ các nghiên cứu về biến ngôn ngữ, các tác giả [116, 4], đã đưa ra các đặc trưng nền tảng sau:

- Tính phổ quát về cấu trúc miền giá trị: Mặc dù các biến ngôn ngữ khác nhau có thể có các giá trị phần tử sinh khác nhau (*long/short, high/low, old/young*), nhưng ý nghĩa về cấu trúc ngữ nghĩa của miền giá của chúng vẫn được giữ. Nói cách khác, giữa các miền giá trị ngôn ngữ của các biến khác nhau có thể thiết lập một quan hệ dạng *đẳng cấu*, sai khác nhau bởi giá trị phần tử sinh.

- Tính độc lập ngữ cảnh của gia tử và liên từ: Ý nghĩa ngôn ngữ của các gia tử (*very, rather, little, ...*) và các liên từ (and, or, not, ...) mang tính ổn định tương đối, ít phụ thuộc vào ngữ cảnh cụ thể. Ngược lại, các giá trị phần tử sinh (như *old, young*) lại phụ thuộc đáng kể vào ngữ cảnh của biến đang xét. Do đó, khi xây dựng mô hình hình thức cho gia tử và liên từ, ta có thể ưu tiên mô hình hóa ở cấp cấu trúc chung mà không cần phụ thuộc trực tiếp vào miền tham chiếu cụ thể của từng biến.

Hai đặc trưng trên cho phép chúng ta có thể sử dụng một tập gia tử và xây dựng một cấu trúc toán học thống nhất cho miền giá trị của nhiều biến ngôn ngữ khác nhau. Với cách tiếp cận này, biến ngôn ngữ thiết lập cầu nối giữa ngôn ngữ tự nhiên và mô hình tính toán [54]. Tuy nhiên, trong phần lớn các mô hình dựa trên biến ngôn ngữ truyền thống, ngữ nghĩa của từ vẫn được xác định chủ yếu thông qua việc lựa chọn các hàm thuộc thích hợp. Điều đó làm xuất hiện sự phụ thuộc đáng kể vào kinh nghiệm của người thiết kế, và dẫn đến khả năng cùng một tập nhãn ngôn ngữ nhưng có thể được biểu diễn theo nhiều cách khác nhau [52, 54, 3]. Chính giới hạn này tạo ra nhu cầu phải phát triển các mô hình tính toán với từ có cấu trúc biểu diễn và xử lý ngữ nghĩa chặt chẽ hơn.

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

1.2.2.1. Mô hình hóa ngôn ngữ mờ và tính toán dựa trên nguyên lý mở rộng

Trên nền tảng tập mờ và biến ngôn ngữ, một hướng phát triển tự nhiên là thực hiện tính toán trực tiếp trên các hàm thuộc của từ ngôn ngữ. Trong các mô hình này, mỗi nhãn ngôn ngữ được gắn với một tập mờ trên miền số, và các phép kết nhập được xây dựng trên cơ sở nguyên lý mở rộng của Zadeh và các phép toán mờ tương ứng [30, 116]. Cách tiếp cận này có ý nghĩa lớn ở chỗ nó cho phép thực hiện tổng hợp, so sánh và suy luận trong những tình huống mà thông tin đầu vào mang tính mơ hồ và không chắc chắn.

Một cách tổng quát, quy trình kết nhập ngôn ngữ theo nguyên lý mở rộng có thể xem như quá trình biến đổi từ tập nhãn ngôn ngữ S^n sang một tập mờ trên miền số, rồi từ đó xấp xỉ ngược về một nhãn ngôn ngữ trong S . Về mặt hình thức, điều này có thể mô tả bởi sơ đồ sau:

$$S^n \xrightarrow{\tilde{F}} F(\mathcal{R}) \xrightarrow{app_1(\cdot)} S. \quad (1.4)$$

Trong đó, S^n ký hiệu tích Đề-các bậc n của tập thuật ngữ ngôn ngữ S , $\tilde{F}$ là một toán tử kết nhập được định nghĩa trên các tập mờ dựa trên nguyên lý mở rộng, $F(\mathcal{R})$ là tập các tập mờ xác định trên tập số thực $\mathcal{R}$, $app_1(\cdot) : F(\mathcal{R}) \rightarrow S$ là một hàm xấp xỉ ngôn ngữ, thực hiện tìm kiếm một nhãn $s_i \in S$ sao cho khoảng cách ngữ nghĩa (thường là khoảng cách Euclid hoặc Hamming giữa các hàm thuộc) đến số mờ kết quả là nhỏ nhất.

Ý nghĩa của sơ đồ này là kết quả trung gian của phép toán thường là một tập mờ mới trên miền số thực, không còn trùng khớp với một từ ngôn ngữ ban đầu S . Vì vậy, mô hình buộc phải thực hiện bước xấp xỉ ngôn ngữ hoặc giải mờ để đưa kết quả quay về miền nhãn gần nhất trong S . Vấn đề phương pháp luận nổi bật của hướng tiếp cận này là sự phụ thuộc mạnh vào lựa chọn hàm thuộc và kỹ thuật xấp xỉ; hệ quả là tính nhất quán ngữ nghĩa của kết quả có thể suy giảm đáng kể khi phải thực hiện nhiều bước tính toán liên tiếp [49, 3, 48].

1.2.2.2. Mô hình tính toán ngôn ngữ dạng ký hiệu

Để giảm độ phức tạp của việc thao tác trực tiếp trên hàm thuộc, các mô hình tính toán ngôn ngữ dạng ký hiệu được đề xuất theo hướng gán cho mỗi nhãn ngôn ngữ một chỉ số số học và thực hiện phép tính trên miền chỉ số, sau đó kết quả được ánh xạ ngược về nhãn ngôn ngữ tương ứng thông qua một bước xấp xỉ hoặc làm tròn [27]. Khi đó, tập

từ ngôn ngữ $S = \{s_0, \dots, s_g\}$ được sắp thứ tự tuyến tính, sao cho $s_i < s_j$, khi và chỉ khi $i < j$; kết quả tính toán là một giá trị thực trên đoạn $[0, g]$, sau đó được ánh xạ ngược về nhãn tương ứng qua một bước làm tròn hoặc xấp xỉ:

$$S^n \xrightarrow{C} [0, g] \xrightarrow{app_2(\cdot)} \{0, 1, \dots, g\} \xrightarrow{index^{-1}} S. \quad (1.5)$$

Trong đó:

S^n : là tích Đề-các của tập từ ngôn ngữ (biểu diễn n đánh giá ngôn ngữ đầu vào);

C : là một toán tử kết nhập ngôn ngữ ký hiệu trả về một giá trị thực $\alpha \in [0, g]$ đại diện cho chỉ số tổng hợp.

$app_2(\cdot): [0, g] \rightarrow \{0, 1, \dots, g\}$ là một hàm xấp xỉ (làm tròn) để chuyển đổi chỉ số thực α về chỉ số nguyên gần nhất.

$index^{-1}$: Ánh xạ ngược từ chỉ số về nhãn ngôn ngữ tương ứng trong S .

Ưu điểm của hướng tiếp cận này là đơn giản và thuận tiện về mặt tính toán. Tuy nhiên, nó ngầm giả định rằng khoảng cách ngữ nghĩa giữa các nhãn là đồng đều, trong khi điều đó thường không phù hợp với cách con người cảm nhận ngôn ngữ tự nhiên. Hơn nữa, việc làm tròn ở các bước trung gian hoặc cuối cùng dễ gây mất mát thông tin và làm suy giảm tính nhất quán ngữ nghĩa của kết quả [49, 3, 48]. Vì vậy, dù có ý nghĩa như một bước chuyển từ tính toán mờ sang tính toán ký hiệu, mô hình này chưa phản ánh đầy đủ cấu trúc ngữ nghĩa nội tại của các từ trong bài toán MCGDM.

1.2.2.3. Mô hình 2-tuple ngôn ngữ

Để khắc phục sự mất mát thông tin do làm tròn trong mô hình ký hiệu, Herrera và cộng sự đã đề xuất mô hình 2-tuple ngôn ngữ [49]. Trong mô hình này, mỗi đánh giá được biểu diễn bởi cặp (s_i, α) , trong đó s_i là nhãn ngôn ngữ và α là tham số dịch chuyển ký hiệu. Về bản chất, α lưu giữ sai lệch giữa kết quả tính toán trên miền số với nhãn gần nhất, nhờ đó giảm đáng kể sai số xấp xỉ và cho phép kết quả đầu ra vẫn duy trì được khả năng diễn giải ngôn ngữ.

Định nghĩa 1.3. [49] Giả sử $S = \{s_0, \dots, s_g\}$ là một tập từ ngôn ngữ có thứ tự với số lượng phần tử lẻ $(g + 1)$. Một biểu diễn 2-tuple được định nghĩa là cặp (s_i, α) , trong đó:

$s_i \in S$: Là nhãn ngôn ngữ;

$\alpha \in [-0.5, 0.5)$: Là giá trị dịch chuyển ký hiệu, thể hiện khoảng cách từ giá trị thực của kết quả tính toán đến chỉ số i của nhãn s_i .

Cơ chế chuyển đổi giữa miền giá trị số thực $\beta \in [0, g]$ (kết quả của phép kết nhập) và miền 2-tuple ngôn ngữ được thực hiện thông qua cặp ánh xạ Δ và Δ^{-1} như sau:

Ánh xạ thuận (Δ): Cho một giá trị $\beta \in [0, g]$ là kết quả của một phép kết nhập ký hiệu, ánh xạ $\Delta: [0, g] \rightarrow S \times [-0.5, 0.5]$ được xác định bởi:

$$\Delta(\beta) = (s_i, \alpha) = \begin{cases} s_i, & i = \text{round}(\beta) \\ \alpha = \beta - i, & \alpha \in [-0.5, 0.5]. \end{cases} \quad (1.6)$$

Trong đó, $\text{round}(\cdot)$ là phép làm tròn đến số nguyên gần nhất.

Ánh xạ nghịch (Δ^{-1}): Cho một 2-tuple ngôn ngữ (s_i, α) , giá trị số tương đương $\beta \in [0, g]$ được khôi phục thông qua ánh xạ:

$$\Delta^{-1}(s_i, \alpha) = i + \alpha = \beta. \quad (1.7)$$

Mô hình 2-tuple được xem là một bước tiến quan trọng trong tính toán với từ, vì nó cho phép thực hiện tính toán mà không phải giải mờ theo nghĩa truyền thống. Tuy nhiên, dưới góc nhìn của luận án này, nó vẫn còn ba giới hạn nền tảng:

Thứ nhất, giả định về tính đồng đều: Mô hình 2-tuple thực hiện dựa trên giả định rằng các từ ngôn ngữ được phân bố đều và đối xứng trên miền giá trị (khoảng cách ngữ nghĩa giữa s_i và s_{i+1} là hằng số với mọi i). Điều này mâu thuẫn với nhận thức tự nhiên của con người, nơi các từ ngữ thường có độ bao phủ và khoảng cách ngữ nghĩa không đồng nhất (ví dụ: khoảng cách từ *tốt* đến *rất tốt* thường hẹp hơn từ *trung bình* đến *tốt*).

Thứ hai, thiếu cấu trúc đại số nội tại: Việc tính toán trong mô hình 2-tuple thực chất vẫn là tính toán trên giá trị số thực (thông qua chỉ số $i + \alpha$) sau đó ánh xạ ngược lại. Nó thiếu một cơ sở toán học để xử lý các quan hệ ngữ nghĩa phức tạp như quan hệ sinh, quan hệ đối xứng hay tính so sánh được giữa các từ không cùng thang đo.

Thứ ba, hạn chế trong biểu diễn độ mờ: Tham số α chỉ biểu diễn độ lệch so với tâm, chưa phản ánh được độ đo tính mờ hay độ không chắc chắn của đánh giá. Do đó, nó chưa đủ mạnh để mô hình hóa các trạng thái do dự của chuyên gia trong bài toán MCGDM nhóm.

Trong các mô hình này, hiệu quả biểu diễn phụ thuộc lớn vào độ hạt của thang đo ngôn ngữ. Thang đo càng chi tiết thì khả năng phân biệt các mức độ đánh giá càng cao, nhưng đồng thời làm tăng độ phức tạp trong mô hình hóa và tính toán. Thực tiễn nghiên cứu cho thấy, các thang đo ngôn ngữ thường được thiết kế rời rạc với số lượng từ ngôn ngữ cố định như thang Likert 5 hoặc 7 mức được sử dụng phổ biến trong khoảng 7±2 mức [83], phù hợp với nhận thức của con người [84]. Từ những phân tích trên, nhu cầu

đặt ra là cần phát triển một mô hình biểu diễn ngôn ngữ mở rộng được xây dựng trên một không gian có cấu trúc đại số tuyến tính thay vì chỉ là tập sắp thứ tự đơn thuần. Hướng tiếp cận này không chỉ cho phép định lượng ngữ nghĩa trên các thang đo không đều mà còn cung cấp các phép toán kết nhập bảo toàn được cấu trúc ngữ nghĩa, làm cơ sở vững chắc cho các bài toán ra quyết định động và suy giảm ngữ nghĩa theo thời gian.

1.2.3. Lý thuyết Đại số gia tử

Trong ngôn ngữ tự nhiên, các từ không chỉ dùng để biểu đạt thông tin mà còn hàm chứa ngữ nghĩa nội tại và thiết lập quan hệ thứ tự ngữ nghĩa với nhau. Đối với một biến ngôn ngữ X , miền từ vựng $Dom(X)$ không phải là một tập hợp rời rạc thuần túy, mà là một cấu trúc có trật tự, trong đó các từ được hình thành từ các phần tử sinh dưới tác động của các gia tử ngôn ngữ. Chẳng hạn, từ hai phần tử sinh như: *old* và *young* dưới tác động của các gia tử như *very* và *little* có thể tạo thành tập từ như *very old*, *little old*, *little young*, *very young*, v.v. và các từ này có thể được sắp xếp theo một trật tự ngữ nghĩa nhất định như: $very\ old \leq old \leq little\ old \leq little\ young \leq young \leq very\ young$ (dấu $\leq$ là ký hiệu cho quan hệ thứ tự được cảm sinh bởi ngữ nghĩa vốn có của các từ ngôn ngữ trên tập từ của cấu trúc HA). Đây chính là điểm mà các cách tiếp cận dựa trên lý thuyết tập mờ và biến ngôn ngữ còn chưa xử lý một cách tự nhiên và nhất quán khi cần tổng hợp, so sánh và tính toán trực tiếp trên các từ ngôn ngữ.

Xuất phát từ yêu cầu đó, Nguyễn Cát Hồ và cộng sự đã đề xuất lý thuyết HA như một cấu trúc toán học nhằm mô hình hóa tường minh trật tự ngữ nghĩa và cơ chế biến đổi ngữ nghĩa do các gia tử tác động lên các từ ngôn ngữ [53]. Theo hướng tiếp cận HA, miền giá trị của biến ngôn ngữ X được xem như một đại số có cấu trúc, trong đó các từ ngôn ngữ được sinh ra từ các phần tử sinh dưới tác động của các gia tử, đồng thời vẫn duy trì quan hệ thứ tự ngữ nghĩa nội tại. Đây là lý do HA được luận án lựa chọn làm nền tảng hình thức để phát triển khung phương pháp luận MCGDM.

Trên nền tảng đó, lý thuyết HA tiếp tục được phát triển qua nhiều công trình và đã hình thành các khái niệm quan trọng phục vụ trực tiếp cho tính toán với từ, gồm: cấu trúc HA tuyến tính, độ đo tính mờ, hàm định lượng ngữ nghĩa, hệ khoảng tương tự và mô hình 4-tuple ngữ nghĩa [52, 53, 54]. Các khái niệm này không chỉ có ý nghĩa lý thuyết, mà còn là cơ sở trực tiếp cho việc xây dựng thang đo ngôn ngữ, xác định khoảng cách ngữ nghĩa, kết nhập đánh giá và phát triển các mô hình ở Chương 2, Chương 3 và Chương 4.

1.2.3.1. Định nghĩa và một số tính chất của HA

Định nghĩa 1.4. [53] Một đại số gia tử $\mathcal{AX}$ 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 $0 \leq g^- \leq W \leq g^+ \leq 1$.

Mỗi phần tử $x \in X$ được gọi là một từ ngôn ngữ. Nếu X và H đều là các tập sắp thứ tự tuyến tính thì $\mathcal{AX} = (X, G, H, \leq)$ được gọi là HA tuyến tính. Hơn nữa, nếu được bổ sung thêm hai gia tử tới hạn là σ và ϕ với ngữ nghĩa là cận trên đúng và cận dưới đúng của tập $H(x)$ khi tác động lên x , thì ta được HA tuyến tính đầy đủ, ký hiệu $\mathcal{AX}^* = (X, G, H, \sigma, \phi, \leq)$. Trong luận án này, chỉ xét HA tuyến tính; vì vậy, từ đây về sau, thuật ngữ HA được hiểu là HA tuyến tính.

Tập gia tử $H = H^- \cup H^+$, trong đó H^- là tập gia tử âm có tác dụng làm giảm ngữ nghĩa của phần tử sinh, H^+ là tập các gia tử dương có tác dụng làm tăng ngữ nghĩa của phần tử sinh. Không mất tính tổng quát, có thể giả thiết

$$H^- = \{h_{-1} < h_{-2} < \dots < h_{-q}\} \text{ và } H^+ = \{h_1 < h_2 < \dots < h_p\}.$$

Khi một gia tử $h \in H$ tác động lên một giá trị ngôn ngữ $x \in X$, sẽ tạo ra một giá trị ngôn ngữ mới, ký hiệu là hx . Với mỗi $x \in X$, ký hiệu $H(x)$ là tập tất cả các giá trị ngôn ngữ được sinh từ x bởi các gia tử trong H . Mỗi phần tử $\mathcal{E} \in H(x)$ có thể được biểu diễn dưới dạng chuỗi $\mathcal{E} = h_n \dots h_1 x$, với $h_i \in H$, $n \geq 1$.

Nếu $x \in \{g^-, g^+\}$ thì biểu diễn $\mathcal{E} = h_n \dots h_1 x$ được gọi là biểu diễn chính tắc nếu $h_{j+1} \dots h_{1x} \neq h_j \dots h_{1x}$ với mọi $j = 1, \dots, n-1$. Khi đó $\mathcal{E}$ có độ dài là $n+1$, ký hiệu là $|\mathcal{E}|$ hoặc $\ell(\mathcal{E})$.

Để tiện cho việc xây dựng thang đo ngôn ngữ trong các mô hình của luận án, ký hiệu X_k là tập các từ ngôn ngữ có độ dài nhỏ hơn hoặc bằng k , còn $X_{(k)}$ là tập các từ ngôn ngữ có độ dài đúng bằng k . Chính các tập này là cơ sở để thiết lập các thang đo ngôn ngữ hữu hạn dùng trong mô hình HA-TOPSIS, mô hình 4-tuple ngữ nghĩa và mô hình MCGDM động dựa trên HA trong các Chương 2, 3 và 4.

Ví dụ 1.2. Xét biến ngôn ngữ AGE với hai phần tử sinh $\{old, young\}$, tập phần tử hằng $C = \{0, W, 1\}$, trong đó, $W = \{medium\}$, $0 = extremely\ old$, $1 = extremely\ young$, và tập gia tử $H = \{little, very\}$. Khi đó, với độ dài tối đa $k = 2$, tập X_2 gồm các từ:

$X_2 = \{extremely\ old, old, medium, young, extremely\ young\} \cup \{little\ old, very\ old, little\ young, very\ young\}$.

Ví dụ này cho thấy miền ngôn ngữ của một biến có thể được sinh có hệ thống từ các phần tử sinh và các gia tử, đồng thời bảo toàn trật tự ngữ nghĩa giữa các từ. Đây cũng là nguyên lý được sử dụng trong luận án để xây dựng các thang đo ngôn ngữ dựa trên HA cho tiêu chí, trọng số của tiêu chí và tập phương án trong các mô hình đề xuất.

Cấu trúc của $\mathcal{AX}$ được thiết lập trên cơ sở khai thác quan hệ thứ tự ngữ nghĩa vốn tồn tại tự nhiên trong miền giá trị của biến ngôn ngữ. Trên cơ sở đó, Hồ và cộng sự [52] đã chỉ ra một số tính chất cơ bản của HA tuyến tính như sau:

Định lý 1.1. [52] Cho H^- và H^+ là các tập sắp thứ tự tuyến tính của HA $\mathcal{AX} = (X, G, H, \leq)$. Khi đó:

- i) Với mỗi $u \in X$ thì $H(u)$ là một tập sắp thứ tự tuyến tính;
- ii) Nếu X được sinh từ G bởi các gia tử và G là tập sắp thứ tự tuyến tính thì X cũng là tập sắp thứ tự tuyến tính. Hơn nữa nếu $u < v$, và u, v là độc lập với nhau, tức là $u \notin H(v)$ và $v \notin H(u)$, thì $H(u) \leq H(v)$.

Tính chất này cho thấy HA không chỉ cung cấp một cơ chế sinh từ ngôn ngữ mà còn duy trì một cấu trúc thứ tự nhất quán trên toàn bộ miền ngôn ngữ. Đây là cơ sở quan trọng để các mô hình của luận án có thể thực hiện so sánh, xếp hạng và kết nhập trực tiếp trên miền từ ngôn ngữ.

Định lý 1.2. [52] Cho $x = h_n \dots h_1 u$ và $y = k_m \dots k_1 u$ là hai biểu diễn chính tắc của x và y đối với cùng phần tử gốc u . Khi đó tồn tại chỉ số $j \leq \min\{n, m\} + 1$ sao cho $h_{j'} = k_{j'}$ với mọi $j' < j$. Khi đó:

- i) $x < y$ khi và chỉ khi $h_j x_j < k_j x_j$, trong đó $x_j = h_{j-1} \dots h_1 u$;
- ii) $x = y$ khi và chỉ khi $m = n$ và $h_j x_j = k_j x_j$;
- iii) x và y là không so sánh được với nhau khi và chỉ khi $h_j x_j$ và $k_j x_j$ không so sánh được với nhau.

Định lý này cung cấp cơ sở hình thức để so sánh các từ ngôn ngữ được sinh bởi chuỗi gia tử khác nhau. Điều này đặc biệt quan trọng trong luận án, vì việc xác định trật tự, khoảng cách ngữ nghĩa và các phép kết nhập trên miền ngôn ngữ đều dựa trên khả năng so sánh nhất quán giữa các từ ngôn ngữ trong HA.

1.2.3.2. Độ đo tính mờ của giá trị ngôn ngữ

Tính mờ của một giá trị ngôn ngữ xuất phát từ thực tế rằng cùng một từ có thể được dùng để mô tả nhiều đối tượng, trạng thái hoặc hiện tượng khác nhau trong thế giới thực. Do số lượng từ ngôn ngữ là hữu hạn trong khi thế giới cần mô tả lại rất đa dạng, mỗi từ thường bao hàm nhiều mức độ và sắc thái ý nghĩa khác nhau [52].

Vì vậy, độ đo tính mờ của giá trị ngôn ngữ là một khái niệm có ý nghĩa nền tảng khi xây dựng các mô hình tính toán với từ. Khác với cách tiếp cận của lý thuyết tập mờ, nơi tính mờ thường gắn với hình dạng của hàm thuộc, trong HA tính mờ của một từ ngôn ngữ x được hiểu như là ngữ nghĩa của nó vẫn có thể được thay đổi khi tác động vào nó bằng các gia tử [52].

Theo đó, tập các từ được sinh từ một từ x bởi các gia tử, ký hiệu $H(x)$, có thể xem như biểu diễn mức độ tính mờ của x , và kích thước của tập này được lượng hóa bằng độ đo tính mờ $f\eta(x)$. Cách tiếp cận này đặc biệt quan trọng trong luận án, vì độ đo tính mờ là cơ sở để xây dựng thang đo ngôn ngữ, xác định giá trị định lượng ngữ nghĩa và phát triển mô hình 4-tuple ngữ nghĩa trong các chương sau. Ta có định nghĩa sau về độ đo tính mờ.

Định nghĩa 1.5. [52] Cho $AX^* = (X, G, H, \sigma, \phi, \preceq)$ là một HA tuyến tính đầy đủ (Com HA) của biến ngôn ngữ X . Một ánh xạ $f\eta: X \rightarrow [0,1]$ được gọi là một độ đo tính mờ của các từ ngôn ngữ trong X nếu thỏa mãn các tính chất sau:

i) $f\eta$ là một độ đo đầy đủ trên X , tức là: $f\eta(g^-) + f\eta(g^+) = 1$ và $\sum_{h \in H} f\eta(hu) = f\eta(u)$, $\forall u \in X$;

ii) $f\eta(x) = 0$, $\forall x$ thỏa mãn $H(x) = \{x\}$; đặc biệt, $f\eta(0) = f\eta(W) = f\eta(1) = 0$;

iii) $\forall x, y \in X, \forall h \in H, \mu(h) = \frac{f\eta(hx)}{f\eta(x)} = \frac{f\eta(hy)}{f\eta(y)}$, 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: $f\eta(x) = \mu(h_m) \dots \mu(h_1) f\eta(g)$, trong đó $\sum_{h \in H} \mu(h) = 1$.

Như vậy, độ đo tính mờ $f\eta(x)$ của các từ ngôn ngữ trong X được xác định hoàn toàn khi biết $f\eta(g^-)$ (hoặc $f\eta(g^+)$) và $|H| - 1$, cùng với giá trị của độ đo tính mờ của các

gia tử $\mu(h)$. Các đại lượng này chính là các tham số tính mờ độc lập của biến ngôn ngữ và giữ vai trò đầu vào ngữ nghĩa trong việc xây dựng các thang đo ngôn ngữ dựa trên HA ở các mô hình của luận án.

Từ Định nghĩa 1.5, đã chỉ ra một số tính chất quan trọng của độ đo tính mờ như sau:

Mệnh đề 1.1. [52] *Cho $f\mathfrak{m}(x)$ độ đo tính mờ của các từ ngôn ngữ và $\mu(h)$ là độ đo tính mờ của các gia tử. Khi đó:*

$$(i) f\mathfrak{m}(g^-) + f\mathfrak{m}(g^+) = 1 \text{ và } \sum_{h \in H} f\mathfrak{m}(hx) = f\mathfrak{m}(x);$$

$$(ii) \sum_{j=-q}^{-1} \mu(h_j) = \alpha, \sum_{j=1}^p \mu(h_j) = \beta, \text{ trong đó } \alpha, \beta > 0 \text{ và } \alpha + \beta = 1;$$

$$(iii) \sum_{x \in \sum_{x \in X_k}} f\mathfrak{m}(x) = 1, X_k \text{ là tập từ ngôn ngữ có độ dài nhỏ hơn hoặc bằng } k$$

$$(iv) f\mathfrak{m}(hx) = \mu(h)f\mathfrak{m}(x), \text{ và } \forall x \in X, f\mathfrak{m}(\sigma x) = f\mathfrak{m}(\phi x) = 0;$$

(v) Cho $f\mathfrak{m}(g^-)$, $f\mathfrak{m}(g^+)$ và $\mu(h)$, với $\forall h \in H$. Khi đó, với $x = h_n \dots h_1 g^\varepsilon$, $\varepsilon \in \{-, +\}$, độ đo tính mờ của x được tính theo công thức đệ quy: $f\mathfrak{m}(x) = \mu(h_n)\mu(h_{n-1}) \dots \mu(h_1) f\mathfrak{m}(g^\varepsilon)$.

Các tính chất trên cho thấy độ đo tính mờ trong HA không được gán một cách tùy ý, mà được xác định có hệ thống từ các phần tử sinh và các gia tử. Điều này giúp luận án xây dựng các thang đo ngôn ngữ có cơ sở ngữ nghĩa rõ ràng, bảo đảm tính nhất quán giữa cấu trúc ngữ nghĩa, mức độ mơ hồ và giá trị định lượng ngữ nghĩa của các từ. Trên nền tảng đó, các chương tiếp theo sử dụng độ đo tính mờ như một thành phần đầu vào để thiết lập hàm định lượng ngữ nghĩa, hệ khoảng tương tự và biểu diễn 4-tuple ngữ nghĩa phục vụ cho các mô hình MCGDM dựa trên HA.

1.2.3.3. Định lượng ngữ nghĩa của giá trị ngôn ngữ

Nếu độ đo tính mờ phản ánh mức độ biến đổi ngữ nghĩa tiềm tàng của từ, thì hàm định lượng ngữ nghĩa là cơ chế gắn mỗi từ với một giá trị số trong $[0, 1]$ theo cách vẫn bảo toàn trật tự ngữ nghĩa. Trong cách tiếp cận tập mờ, giá trị số của từ thường được xác định thông qua quá trình khử mờ của hàm thuộc [116]; cách làm đó phụ thuộc mạnh vào lựa chọn hàm thuộc và phương pháp khử mờ, nên có thể làm suy giảm tính nhất quán giữa thứ tự ngữ nghĩa và biểu diễn định lượng [48].

Theo HA, mỗi từ ngôn ngữ $x \in X$ được định lượng trực tiếp thông qua một ánh xạ $\vartheta: X \rightarrow [0, 1]$, gọi là hàm ánh xạ định lượng ngữ nghĩa (Semantic Quantifying Mapping - SQM), sao cho $\vartheta(x) \in [0, 1]$ vừa bảo toàn trật tự ngữ nghĩa, vừa phản ánh nhất quán cấu trúc đại số của miền ngôn ngữ [52]. Đây là cơ sở trực tiếp để luận án xây dựng các thang đo ngôn ngữ và thực hiện tính toán trong mô hình HA-TOPSIS, mô hình tích hợp EFAHP - 4-tuple - HA và mô hình MCGDM động dựa trên HA.

Định nghĩa 1.6. [52] Cho $AX^* = (X, G, H, \sigma, \phi, \leq)$ là một ComHA tuyến tính. Một ánh xạ $\vartheta: X \rightarrow [0, 1]$ được gọi là một SQM của 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 nghĩa: $\forall x \in X, \vartheta(\sigma x) = \infimum \vartheta(H(x))$ và $\vartheta(\phi x) = \supremum \vartheta(H(x))$.

Định nghĩa trên cho thấy hàm SQM không chỉ gán cho mỗi từ ngôn ngữ một giá trị số, mà còn duy trì được cấu trúc thứ tự và tính liên tục ngữ nghĩa của miền giá trị ngôn ngữ. Đây là yêu cầu nền tảng để các phép so sánh, đo khoảng cách và kết nhập trong các mô hình của luận án có thể được thực hiện trực tiếp trên miền ngôn ngữ.

Như vậy, tính đầy đủ của một HA có nghĩa là: với mọi tập từ $H(x)$, luôn tồn tại giới hạn cận trên và cận dưới của các tập từ $H(x)$. Tương tự, giống như dấu đại số của các số trong giải tích toán học cổ điển, các từ ngôn ngữ $x \in X$ cũng có dấu đại số được xác định như sau, gọi là hàm dấu (sign function):

Định nghĩa 1.7. [52] Hàm dấu $\zeta_g: X \rightarrow \{-1, 0, 1\}$ là một ánh xạ được định nghĩa đệ qui như sau, trong đó $h, h' \in H$ và $g \in \{g^-, g^+\}$:

- (i) $\zeta_g(g^-) = -1$, $\zeta_g(g^+) = +1$;
- (ii) $\zeta_g(hg) = -\zeta_g(g)$ nếu $hg < g$, và $\zeta_g(hg) = +\zeta_g(g)$ nếu $hg > g$;
- (iii) $\zeta_g(h'hx) = 0$ nếu $h'hx = hx$; $\zeta_g(h'hx) = -\zeta_g(hx)$, nếu $h'hx \neq hx$ và h' là âm (negative) đối với h ; ngược lại, $\zeta_g(h'hx) = \zeta_g(hx)$, nếu $h'hx \neq hx$ và h' là dương (positive) đối với h ;
- (iv) $\zeta_g(h'hx) = 0$, nếu $h'hx = hx$.

Từ hàm dấu, có thể xác định được hướng tăng hoặc giảm ngữ nghĩa của từ khi chịu tác động của gia tử

Mệnh đề 1.2. [52] Với mọi gia tử h và từ ngôn ngữ x , nếu $\zeta_g(hx) = 1$, thì $hx > x$; nếu $\zeta_g(hx) = -1$, thì $hx < x$; nếu $\zeta_g(hx) = 0$, thì $hx = x$. Các tập từ $H(x)$ có các tính chất:

$$H(x) \cap H(y) = \emptyset, \text{ với } x \notin H(x) \text{ và } y \notin H(y) \text{ và } 0 \leq H(x) \leq 1.$$

Hơn nữa:

- Nếu $\zeta_g(h_px) = +1$, thì $H(h_{-q}x) \leq H(h_{-q+1}x) \leq \dots \leq H(h_{-1}x) \leq x \leq H(h_1x) \leq H(h_2x) \leq \dots \leq H(h_px)$;
- Nếu $\zeta_g(h_px) = -1$, thì $H(h_{-q}x) \geq H(h_{-q+1}x) \geq \dots \geq H(h_{-1}x) \geq x \geq H(h_1x) \geq H(h_2x) \geq \dots \geq H(h_px)$.

Tính chất này là cơ sở để xây dựng thứ tự trên thang đo ngôn ngữ, xác định khoảng tương tự và phát triển các phép toán ngôn ngữ trong các chương sau của luận án.

Một khái niệm quan trọng trong ngôn ngữ tự nhiên là ngữ nghĩa định lượng (ngữ nghĩa số) của các từ ngôn ngữ; ta có thể hình thức hóa điều này dưới dạng một SQM được xây dựng dựa trên các tham số cho trước gồm các độ đo tính mờ của các phần tử sinh $f\mathfrak{m}(g^-)$, $f\mathfrak{m}(g^+)$ và độ đo tính mờ của các gia tử $\mu(h)$. Giá trị định lượng của giá trị ngôn ngữ x chia đoạn $f\mathfrak{m}(x)$ theo tỷ lệ $\alpha : \beta$ nếu $\zeta g(h_px) = +1$, và theo tỷ lệ $\beta : \alpha$ nếu $\zeta g(h_px) = -1$, được tính theo công thức đệ quy sau:

Định nghĩa 1.8. [52] Cho $AX^* = (X, G, H, \sigma, \phi, \leq)$ là một ComHA tuyến tính và $f\mathfrak{m}(x)$ 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) $\vartheta(W) = \theta = f\mathfrak{m}(g^-)$, $\vartheta(g^-) = \theta - \alpha f\mathfrak{m}(g^-) = \beta f\mathfrak{m}(g^-)$, $\vartheta(g^+) = \theta + \alpha f\mathfrak{m}(g^+)$;
- (ii) $\vartheta(h_jx) = \vartheta(x) + \zeta g(h_jx) \left\{ \sum_{i=\zeta g(j)}^{j-\zeta g(j)} \mu(h_i) f\mathfrak{m}(x) - \omega(h_jx) \mu(h_jx) f\mathfrak{m}(x) \right\}$ với mọi j , trong đó $-q \leq j \leq p$, $j \neq 0$, và $\omega(h_jx) = \frac{1}{2} [1 + \zeta g(h_jx) \zeta g(h_ph_jx) (\beta - \alpha)]$;
- (iii) $\vartheta(\phi g^-) = 0$, $\vartheta(\sigma g^-) = \theta = \vartheta(\phi g^+)$, $\vartheta(\sigma g^+) = 1$, và với mọi j , $-q \leq j \leq p$,
 $\vartheta(\phi h_jx) = \vartheta(x) + \zeta g(h_jx) \left\{ \sum_{i=\zeta g(j)}^{j-\zeta g(j)} \mu(h_i) f\mathfrak{m}(x) \right\} - \frac{1}{2} (1 - \zeta g(h_jx)) \mu(h_i) f\mathfrak{m}(x)$,
 $\vartheta(\sigma h_jx) = \vartheta(x) + \zeta g(h_jx) \left\{ \sum_{i=\zeta g(j)}^{j-\zeta g(j)} \mu(h_i) f\mathfrak{m}(x) \right\} + \frac{1}{2} (1 + \zeta g(h_jx)) \mu(h_i) f\mathfrak{m}(x)$.

Trong nghiên cứu của [52] đã chứng minh định nghĩa trên thỏa mãn các yêu cầu của một SQM và đảm bảo tính trừ mật của nó đối với các từ ngôn ngữ của AX^* trong đoạn $[0, 1]$. Nói cách khác, mỗi từ ngôn ngữ $x \in X$ trong HA được gắn với một giá trị định lượng ngữ nghĩa số trong đoạn $\vartheta(x) \in [0, 1]$ theo cách có cơ sở toán học, nhất quán với trật tự ngữ nghĩa và phù hợp với độ đo tính mờ của từ ngôn ngữ. Đây chính là nền tảng để luận án xây dựng các thang đo ngôn ngữ dùng cho đánh giá phương án, xác định trọng số của tiêu chí, đo khoảng cách ngữ nghĩa và biểu diễn 4-tuple ngữ nghĩa trong các mô hình đề xuất ở Chương 2, Chương 3 và Chương 4.

1.2.3.4. Hệ khoảng tương tự ngữ nghĩa

Định nghĩa 1.9 [52] Cho một độ đo tính mờ $f\mathfrak{m}$ 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 $f_{\mathfrak{M}}$. 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 .

Về bản chất, hệ khoảng tương tự ngữ nghĩa cho phép phân hoạch miền tham chiếu số $[0, 1]$ thành các khoảng ngữ nghĩa tương ứng với các từ ngôn ngữ trên thang đo. Việc xây dựng $S_k(x)$ được thực hiện trên cơ sở các khoảng tính mờ ở một bậc cao hơn $k' > k$ (thường $k' \geq k + 2$). Các khoảng này được phân hoạch thành các cụm rời nhau $C(x)$ tương ứng với từng từ $x \in X_k$; khi đó, $S_k(x)$ được xác định là hợp của các khoảng tính mờ liên tiếp tạo thành cụm bao quanh giá trị định lượng ngữ nghĩa của từ ngôn ngữ, tức $\vartheta(x)$. Do đó, hệ khoảng tương tự bậc k được xác định hoàn toàn khi chọn được cấu trúc HA và bộ tham số ngữ nghĩa như độ đo tính mờ của các phần tử sinh $f_{\mathfrak{M}}(g^-)$, gia tử $\mu(h)$ và tham số k .

Trong cách tiếp cận dựa trên HA, hệ khoảng tương tự ngữ nghĩa giữ vai trò cầu nối giữa giá trị định lượng ngữ nghĩa và miền ý nghĩa của từng từ trên thang đo ngôn ngữ. Nhờ đó, mỗi giá trị số trong $[0, 1]$ không chỉ mang ý nghĩa định lượng, mà còn gắn với khoảng lân cận ngữ nghĩa cụ thể tương ứng với một từ ngôn ngữ. Đây là cơ sở quan trọng để xây dựng biểu diễn 4-tuple ngữ nghĩa trong các mục tiếp theo, trong đó mỗi từ được đặc trưng đồng thời bởi giá trị định lượng và khoảng tương tự tương ứng.

Trong phạm vi luận án này, hệ khoảng tương tự được sử dụng trực tiếp để xây dựng các thang đo ngôn ngữ dựa trên HA cho các mô hình MCGDM. Cụ thể, ở Chương 2, nó hỗ trợ thiết lập thang đo và định vị ngữ nghĩa cho mô hình HA-TOPSIS; ở Chương 3, nó là thành phần cốt lõi của biểu diễn 4-tuple ngữ nghĩa; và ở Chương 4, nó tiếp tục làm nền tảng cho các phép kết nhập và cơ chế suy giảm ngữ nghĩa trong bài toán MCGDM động. Vì vậy, hệ khoảng tương tự ngữ nghĩa không chỉ có ý nghĩa lý thuyết, mà còn là cơ sở nền tảng trực tiếp phục vụ phát triển các mô hình của luận án.

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

Dựa trên HA, Nguyễn Cát Hồ và cộng sự đã phát triển mô hình biểu diễn 4-tuple ngữ nghĩa cho các thang đo ngôn ngữ nhằm biểu diễn đồng thời nội dung ngữ nghĩa và giá trị định lượng của các từ ngôn ngữ vốn mang tính mơ hồ, không chắc chắn [54]. Cách biểu diễn này đặc biệt hữu ích trong các bài toán MCGDM, vì cho phép thực hiện tính toán trên miền ngôn ngữ mà vẫn duy trì được khả năng diễn giải ngữ nghĩa của kết quả. Trong khuôn khổ luận án này, mô hình 4-tuple ngữ nghĩa là nền tảng trực tiếp để xây dựng thang đo ngôn ngữ cho đánh giá phương án, xác định trọng số của tiêu chí và phát triển các phép kết nhập ngôn ngữ trong các chương sau.

Định nghĩa 1.10 [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.

Khi $(r_x - \vartheta(x))$ được xem như độ lệch ngữ nghĩa giữa từ ngôn ngữ và giá trị tham chiếu r_x ; đặc biệt, nếu $r_x = \vartheta(x)$ tương ứng với biểu diễn ngữ nghĩa chuẩn của từ x . Như vậy, mô hình 4-tuple ngữ nghĩa không chỉ biểu diễn vị trí định lượng của từ trên miền $[0, 1]$, mà còn gắn từ với một khoảng tương tự ngữ nghĩa cụ thể, qua đó làm giàu khả năng biểu diễn so với cách dùng riêng giá trị số $\vartheta(x)$. Để xây dựng thang đo ngôn ngữ hữu hạn phục vụ tính toán, Nguyễn Cát Hồ và cộng sự đưa ra khái niệm thang đo ngôn ngữ đóng - trên.

Định nghĩa 1.11. [54] Cho $\mathcal{AX}^* = (X, G, \epsilon, H, \sigma, \phi, \leq)$ là một ComHA tuyến tính của biến ngôn ngữ X . Một tập từ $T = \{t_0, \dots, t_g\} \subseteq \text{Dom}(X)$ được sắp thứ tự toàn phần được gọi là thang đo ngôn ngữ đóng trên (superior-closed) ở mức độ cụ thể k nếu thỏa mãn các điều kiện:

- (i) $\forall t \in T, |t| \leq k$, và $\exists t \in T$ sao cho $|t| = k$;
- (ii) $g^-, g^+, 0, W, I \in T$;
- (iii) Nếu $y = hx \in T$ thì $x \in T$.

Do $\mathcal{AX}^*$ là một HA đầy đủ, việc $T \subseteq \text{Dom}(X)$ hàm ý rằng ngữ nghĩa định tính của các từ trong T cần được xem xét trong toàn bộ ngữ cảnh của X . Vì vậy, biểu diễn 4-tuple ngữ nghĩa phải được xác định trên toàn bộ đại số $\mathcal{AX}$.

Ba điều kiện trên lần lượt bảo đảm mức độ đặc tả của thang đo, bảo đảm sự hiện diện của các từ ngữ cốt lõi trong miền ngôn ngữ, và bảo đảm tính đóng-trên của thang đo theo cấu trúc sinh của HA. Tuy nhiên, bản thân T mới chỉ là một tập từ có trật tự, chưa đủ để hỗ trợ trực tiếp cho tính toán ngôn ngữ. Vì vậy, cần gắn cho mỗi từ trong T một biểu diễn 4-tuple ngữ nghĩa tương ứng. Trên cơ sở đó, khái niệm thang đo ngôn ngữ 4-tuple ngữ nghĩa được phát biểu như sau:

Định nghĩa 1.12. [54] Cho T là một thang đo ngôn ngữ đóng-trên ở mức đặc tả k . Khi đó, thang đo ngôn ngữ 4-tuple ngữ nghĩa kết hợp với T được xác định bởi:

$$T_{\vartheta,k} = \{(t, \vartheta(t), S_k(t), r_t) : t \in T, r_t \in S_k(t)\}.$$

Thang đo này thỏa các tính chất cơ bản sau:

(i) Với mọi $t \in T$, $(t, \vartheta(t), S_k(t), r_t)$ là một biểu diễn 4-tuple ngữ nghĩa hợp lệ, trong đó $S_k(t) \subseteq [0,1]$ biểu diễn ngữ nghĩa trong khoảng tương tự của t ; $\vartheta(t) \in [0,1]$ là giá trị định lượng ngữ nghĩa số của t trong Định nghĩa 1.11.

(ii) Các khoảng $S_k(t)$, $\forall t \in T$, tạo thành một phân hoạch của miền $[0, 1]$, trong đó mỗi khoảng $S_k(t)$ biểu diễn một khoảng tương tự ngữ nghĩa tương ứng với t . Nói cách khác, các phần tử thuộc $S_k(t)$ được xem là tương tự về mặt ngữ nghĩa với t theo một mức độ được quyết định bởi mức đặc tả k .

(iii) Với mọi $t, t' \in T$, nếu $t \leq t'$ thì $S_k(t) \leq S_k(t')$.

(iv) $\vartheta(t) \in S_k(t)$ với mọi $t \in T$.

Như vậy, thang đo ngôn ngữ 4-tuple ngữ nghĩa cho phép kết hợp đồng thời ba lớp thông tin: nhãn ngôn ngữ, giá trị định lượng ngữ nghĩa và khoảng tương tự ngữ nghĩa của từ. Đây là cấu trúc phù hợp để vừa tính toán, vừa bảo toàn khả năng diễn giải của thông tin ngôn ngữ trong các bài toán MCGDM.

Mệnh đề 1.3. [54] Cho T là một thang đo ngôn ngữ đóng-trên ở mức đặc tả k , được xây dựng dựa trên một HA cho trước $\mathcal{AX}^* = (X, G, \epsilon, H, \sigma, \phi, \preceq)$. Với các tham số độ đo tính mờ đã cho của biến ngôn ngữ X , tập:

$$T_\vartheta = \{(t, \vartheta(t), S_k(t), r_t) : t \in T, r_t \in S_k(t)\},$$

thỏa các tính chất cơ bản sau:

(i) T_ϑ cấu thành một thang đo ngôn ngữ 4-tuple ngữ nghĩa kết hợp với T ;

(ii) Mỗi khoảng $S_k(t)$ được xác định và tính toán dựa trên cấu trúc ngữ nghĩa của các từ trong $\mathcal{AX}$ như sau:

$$S_k(t) = S_{(L)}(t) \cup S_{(R)}(t) \text{ và } S_k(t) = \cup \{T(x) : x \in X_{k+2} \text{ \& } T(x) \subseteq (\vartheta(t_{L,p_L+1}), \vartheta(t_{R,p_R+1}))\}.$$

Như vậy, nếu tập mờ và biến ngôn ngữ cung cấp nền tảng ban đầu cho xử lý thông tin mơ hồ, còn các mô hình tính toán với từ như nguyên lý mở rộng, tính toán ký hiệu và 2-tuple là các bước phát triển quan trọng nhằm giảm mất mát thông tin và cải thiện khả năng thao tác trên miền ngôn ngữ, thì HA cung cấp lớp nền tảng hình thức trung tâm mà luận án cần đến. Trên cơ sở các khái niệm của HA như cấu trúc đại số ngôn ngữ, độ đo tính mờ, SQM, hệ khoảng tương tự và thang đo 4-tuple ngữ nghĩa, luận án có đủ cơ sở lý thuyết để phát triển mô hình HA-TOPSIS ở Chương 2, mô hình EFAHP - 4-tuple - HA ở Chương 3 và mô hình MCGDM động dựa trên HA ở Chương 4. Theo nghĩa đó, mục này không chỉ trình bày các khái niệm nền tảng, mà còn xác định trực tiếp cơ sở học thuật để chuyển sang phần tổng quan nghiên cứu, phân tích hạn chế và xác lập khoảng trống nghiên cứu ở các mục tiếp theo.

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

Trong bài toán MCGDM, quá trình ra quyết định thường gắn với hai nhiệm vụ trọng tâm: xác định trọng số của các tiêu chí và xếp hạng các phương án. Vì vậy, việc xác định các phương pháp nền của luận án không chỉ nhằm lựa chọn những phương pháp đã được sử dụng rộng rãi trong MCDM và MCGDM, mà còn nhằm làm rõ vai trò của chúng trong cấu trúc lời giải và khả năng phát triển trong môi trường thông tin ngôn ngữ. Trên cơ sở đó, luận án kế thừa và phát triển hai nhóm phương pháp chính: AHP/FAHP cho bài toán xác định trọng số của tiêu chí và TOPSIS/FTOPSIS cho bài toán xếp hạng phương án. Hai nhóm phương pháp này tạo nền tảng trực tiếp cho việc xây dựng các mô hình MCGDM dựa trên HA trong các chương tiếp theo.

1.3.1. Phương pháp AHP và vai trò trong xác định trọng số của tiêu chí

Phương pháp phân tích thứ bậc (AHP) do Saaty xuất là một trong những phương pháp quan trọng của MCDM và được sử dụng rộng rãi trong các mô hình MCGDM [95].

Bảng 1.1. Mô tả thang đo của phương pháp AHP [95]

Mức độ quan trọng	Định nghĩa	Giải thích
1	Tầm quan trọng ngang nhau	Hai tiêu chí đóng góp như nhau cho mục tiêu.
3	Tầm quan trọng yếu của một tiêu chí so với tiêu chí kia	Kinh nghiệm và nhận định hơi nghiêng về một tiêu chí so với tiêu chí kia.
5	Tầm quan trọng mạnh	Kinh nghiệm và nhận định nghiêng mạnh về một tiêu chí so với tiêu chí kia.
7	Tầm quan trọng mạnh hơn (rất mạnh)	Một tiêu chí được ưu tiên mạnh hơn tiêu chí kia và sự vượt trội đó đã được chứng minh trong thực tiễn.
9	Tầm quan trọng tuyệt đối	Bằng chứng ủng hộ một tiêu chí so với tiêu chí kia đạt mức cao nhất có thể khẳng định.
2, 4, 6, 8	Các giá trị trung gian giữa hai mức đánh giá liền kề	Được sử dụng khi cần thỏa hiệp.
Nghịch đảo của các giá trị khác 0 nêu trên	Nếu hoạt động i có một trong các giá trị khác 0 nêu trên khi so sánh với hoạt động j , thì j có giá trị nghịch đảo tương ứng khi so sánh với i .	
Các tỷ số	Các tỷ số phát sinh từ thang đo	Nếu cần ép buộc tính nhất quán bằng cách thu được n giá trị số để bao phủ ma trận.

Ý tưởng của phương pháp AHP là cấu trúc hóa bài toán ra quyết định thành một hệ thứ bậc gồm mục tiêu, các tiêu chí, có thể bao gồm tiêu chí con và các phương án; trên cơ sở đó, mức độ ưu tiên tương đối giữa các thành phần được xác định thông qua phán đoán so sánh cặp dựa trên thang đo gồm chín mức tương ứng với các giá trị trong Bảng 1.1. Vì vậy, AHP không chỉ hỗ trợ lựa chọn phương án, mà còn là khuôn khổ nền tảng đối với bài toán xác định trọng số của tiêu chí. Một đặc điểm có giá trị của AHP để đánh giá tính nhất quán trong các so sánh cặp số rõ của các chuyên gia. Saaty đã đề xuất công thức tính chỉ số nhất quán (Consistency Index - CI) và công thức tính tỷ lệ nhất quán (Consistency Ratio - CR) [95], như sau:

$$CI_{Saaty} = \frac{\lambda_{\max} - n}{n - 1}, \quad (1.8)$$

$$CR_{Saaty} = \frac{CI_{Saaty}}{RI(n)}, \quad (1.9)$$

trong đó, $RI(n)$ (Random Index) là chỉ số ngẫu nhiên của ma trận bậc n theo AHP [95].

Ma trận so sánh cặp số rõ của chuyên gia D^e được chấp nhận nếu thỏa điều kiện $CR_{Saaty} < 0.10$; ngược lại, chuyên gia D^e cần rà soát và điều chỉnh lại các so sánh cặp trong ma trận cho đến khi đạt yêu cầu nhất quán. Cơ chế này không loại bỏ hoàn toàn tính chủ quan, nhưng giúp phát hiện các phán đoán mâu thuẫn, từ đó yêu cầu rà soát hoặc điều chỉnh ma trận so sánh. Nhờ đó, trọng số của tiêu chí được kiểm soát ở mức nhất định về tính hợp lý nội tại. Trong bối cảnh MCGDM, cơ chế này càng có ý nghĩa vì quá trình xác định trọng số thường phải kết hợp nhiều nguồn phán đoán và nhiều quan điểm chuyên môn khác nhau.

Từ góc độ kế thừa phương pháp, AHP phù hợp với luận án ở ba khía cạnh. Thứ nhất, AHP cung cấp nguyên lý xác định trọng số của tiêu chí dựa trên so sánh cặp, là cách tiếp cận phổ biến và có khả năng diễn giải tốt trong MCDM. Thứ hai, AHP hỗ trợ biểu diễn cấu trúc phân cấp của bộ tiêu chí, phù hợp với các bài toán có nhiều tiêu chí chính và tiêu chí con. Thứ ba, AHP đặt ra yêu cầu kiểm soát tính nhất quán của phán đoán, một vấn đề quan trọng khi xây dựng mô hình xác định trọng số của tiêu chí trong môi trường có nhiều chuyên gia. Do đó, trong luận án này, AHP được xem là phương pháp nền cho giai đoạn xác định trọng số của tiêu chí, còn các mô hình mở rộng được phát triển nhằm thích ứng tốt hơn với môi trường thông tin ngôn ngữ.

Tuy nhiên, phương pháp AHP cổ điển chủ yếu được thiết lập trong điều kiện các phán đoán được biểu diễn bằng số rõ. Trong nhiều bài toán MCGDM thực tế, đặc biệt khi tiêu chí mang tính định tính, chuyên gia thường khó đưa ra một giá trị số chính xác

để biểu thị mức độ quan trọng tương đối giữa hai tiêu chí. Khi các nhận định này bị ép chuyển trực tiếp thành số rõ, một phần tính mơ hồ, tính không chắc chắn và sắc thái ngữ nghĩa của đánh giá có thể bị suy giảm. Ngoài ra, trong môi trường nhóm, việc tổng hợp các ma trận so sánh cặp và xử lý sự khác biệt giữa chuyên gia cũng đòi hỏi các cơ chế mở rộng phù hợp hơn so với AHP cổ điển. Những hạn chế trên không làm giảm giá trị nền tảng của AHP, mà chỉ ra nhu cầu mở rộng phương pháp này. FAHP được phát triển theo hướng thay thế các giá trị so sánh rõ bằng các biểu diễn mờ hoặc ngôn ngữ nhằm phản ánh tốt hơn sự không chắc chắn trong phán đoán của chuyên gia [106, 19]. Dù vậy, Wang và cộng sự [108] chỉ ra rằng phương pháp của Chang [19] có thể dẫn đến việc gán trọng số bằng 0 cho một số tiêu chí vẫn còn ý nghĩa, đồng thời bị giới hạn trong phạm vi một số dạng số mờ nhất định. Để cải thiện điểm này, Huê và cộng sự [58] đã phát triển FAHP cải tiến, trong đó sử dụng số mờ tam giác tổng quát (Generalized Triangular Fuzzy Numbers - GTFN) kết hợp với chỉ số trọng tâm (centroid index) để xác định trọng số ưu tiên của các tiêu chí một cách chính xác hơn. Mô hình FAHP cải tiến [58], gồm các bước sau:

Bước 1: Xây dựng ma trận so sánh số mờ tam giác tổng quát

$$F = (\ddot{x}_{ij})_{n \times n} = \begin{bmatrix} (1,1,1; w_{11}) & (a_{12}, b_{12}, c_{12}; w_{12}) & \dots & (a_{1n}, b_{1n}, c_{1n}; w_{1n}) \\ (a_{21}, b_{21}, c_{21}; w_{21}) & (1,1,1; w_{22}) & \dots & (a_{2n}, b_{2n}, c_{2n}; w_{2n}) \\ \dots & \dots & \dots & \dots \\ (a_{n1}, b_{n1}, c_{n1}; w_{n1}) & (a_{n2}, b_{n2}, c_{n2}; w_{n2}) & \dots & (1,1,1; w_{nm}) \end{bmatrix} \quad (1.10)$$

Trong đó: $\ddot{x}_{ij} = (a_{ij}, b_{ij}, c_{ij}; w_{ij})$, $\ddot{x}_{ij}^{-1} = (\frac{1}{c_{ij}}, \frac{1}{b_{ij}}, \frac{1}{a_{ij}}; w_{ij})$, $i, j = 1, \dots, n$; $i \neq j$.

Bước 2: Tính phạm vi tổng hợp mờ

Phạm vi tổng hợp mờ, $\mathcal{S}_i$, được tính như sau:

$$\begin{aligned} \mathcal{S}_i &= (g_i, h_i, k_i; \min(w_{ij})) \\ &= \left(\frac{\sum_{j=1}^n a_{ij}}{\sum_{j=1}^n a_{ij} + \sum_{k=1, k \neq i}^n \sum_{j=1}^n c_{kj}}, \frac{\sum_{j=1}^n b_{ij}}{\sum_{i=1}^n \sum_{j=1}^n b_{ij}}, \frac{\sum_{j=1}^n c_{ij}}{\sum_{j=1}^n c_{ij} + \sum_{k=1, k \neq i}^n \sum_{j=1}^n a_{kj}}; \min(w_{ij}) \right) \end{aligned} \quad (1.11)$$

Trong đó: $\sum_{j=1}^n M_{g_i}^j = \left(\sum_{j=1}^n a_{ij}, \sum_{j=1}^n b_{ij}, \sum_{j=1}^n c_{ij}, \min(w_{ij}) \right)$,
 $i, j = 1, 2, \dots, n$.

Tiếp theo tính chỉ số trọng tâm của phạm vi tổng hợp mờ $\mathcal{S}_i$ sử dụng phương pháp được đề xuất bởi (25). Cho $\mathcal{S}_1, \mathcal{S}_2, \dots, \mathcal{S}_n$ là các giá trị phạm vi tổng hợp mờ. Trọng tâm của mỗi số mờ $\mathcal{S}_i$, ký hiệu là $C_i = (\bar{x}_{\mathcal{S}_i}, \bar{y}_{\mathcal{S}_i})$, với $i = 1, 2, \dots, n$, được tính như sau:

$$\bar{x}_{\mathcal{S}_i} = (g_i + h_i + k_i)/3; \bar{y}_{\mathcal{S}_i} = \min(w_{ij})/3 \quad (1.12)$$

Khoảng cách giữa các trọng tâm $C_i = (\bar{x}_{\mathcal{S}_i}, \bar{y}_{\mathcal{S}_i})$, với $i = 1, 2, \dots, n$, và điểm cực tiểu $G = (x_{\min}, y_{\min})$, trong đó $x_{\min} = \min(g_i)$, và $y_{\min} = \min(w_{ij})$, $\omega = \min(w_{ij})$ được tính như sau:

$$D(\mathcal{S}_i, G) = \sqrt{(\bar{x}_{\mathcal{S}_i} - x_{\min})^2 + (\bar{y}_{\mathcal{S}_i} - \frac{\omega}{3} y_{\min})^2} \quad (1.13)$$

Bước 3: Xác định véc-tơ trọng số $W = (w_1, \dots, w_n)^T$ tương ứng với ma trận so sánh mờ

$$w_i = \frac{D(\mathcal{S}_i, C_i)}{\sum_{i=1}^n D(\mathcal{S}_i, C_i)} = \frac{\sqrt{(\bar{x}_{\mathcal{S}_i} - x_{\min})^2 + (\bar{y}_{\mathcal{S}_i} - \frac{\omega}{3} y_{\min})^2}}{\sum_{i=1}^n \sqrt{(\bar{x}_{\mathcal{S}_i} - x_{\min})^2 + (\bar{y}_{\mathcal{S}_i} - \frac{\omega}{3} y_{\min})^2}} \quad (1.14)$$

trong đó: $i = 1, \dots, n$.

Mặc dù các tiếp cận FAHP cải tiến [58], đã góp phần nâng cao độ chính xác trong xác định trọng số, Tuy nhiên, phương pháp này được phát triển dựa trên tập mờ nên cũng gặp một số hạn chế về lựa chọn hàm thuộc và tính toán mờ đã chỉ ra trong mục 1. Tính cấp thiết của đề tài. Bên cạnh đó, cơ chế kiểm soát tính nhất quán trong nhiều mô hình FAHP về cơ bản vẫn được kế thừa hoặc điều chỉnh từ AHP cổ điển, trong khi bản chất của các đánh giá ngôn ngữ đòi hỏi một khuôn khổ nhất quán phù hợp hơn với cấu trúc ngữ nghĩa của miền từ. Những giới hạn đó cho thấy AHP/FAHP cần phải tiếp tục cải tiến theo hướng trên một nền tảng ngôn ngữ có cơ sở hình thức chặt chẽ hơn.

Đối với luận án này, vai trò của AHP không nằm ở việc áp dụng nguyên dạng AHP cổ điển, mà ở việc kế thừa nguyên lý so sánh cặp, xác định véc-tơ trọng số và kiểm soát tính nhất quán để phát triển các mô hình MCGDM trong môi trường thông tin ngôn ngữ. Trên nền tảng đó, luận án tiếp cận FAHP, mô hình 4-tuple ngữ nghĩa và HA như các hướng mở rộng nhằm xử lý tốt hơn đánh giá mờ, đánh giá ngôn ngữ và yêu cầu bảo toàn cấu trúc ngữ nghĩa của thông tin đầu vào [53, 52, 54]. Vì vậy, AHP được đặt ở vị trí phương pháp nền: cung cấp quy trình xác định trọng số của tiêu chí, trong khi Chương 3 và Chương 4 sẽ phát triển cơ chế biểu diễn, kiểm soát nhất quán và tính toán phù hợp hơn với bài toán MCGDM ngôn ngữ theo hướng tiếp cận HA.

1.3.2. Phương pháp TOPSIS và vai trò trong xếp hạng phương án

TOPSIS là một trong những phương pháp MCDM được sử dụng rộng rãi để đánh giá, so sánh và xếp hạng các phương án theo nhiều tiêu chí [59]. Khác với AHP, vốn thường được dùng để xác định trọng số của các tiêu chí, TOPSIS tập trung vào giai đoạn tổng hợp đánh giá và xếp hạng phương án. Nguyên lý cơ bản của TOPSIS là một phương án được xem là có tối ưu khi đồng thời gần PIS và xa NIS [18, 110]. Nhờ vậy, TOPSIS cung cấp một cơ chế xếp hạng trực quan, dễ diễn giải và phù hợp với các bài toán có nhiều tiêu chí khác nhau về đơn vị đo, chiều hướng tối ưu và mức độ quan trọng.

Về mặt phương pháp luận, TOPSIS thiết lập một quy trình tính toán hệ thống thông qua các bước chuẩn hóa dữ liệu, tích hợp trọng số của tiêu chí và sử dụng các phép đo khoảng cách để xác định giá trị ưu tiên cho các phương án [110, 111].

Trong phương pháp TOPSIS truyền thống các giá trị được ước lượng bởi số thực và chỉ có một người ra quyết định. Đầu vào của mô hình bao gồm m phương án và n tiêu chí đánh giá. Phương pháp TOPSIS được thực hiện theo trình tự các bước dưới đây:

Bước 1. Xây dựng ma trận đánh giá các phương án

$$A = (a_{ij})_{m \times n} = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{bmatrix} \quad (1.15)$$

Trong đó: a_{ij} là mức độ đánh giá bằng số thực của chuyên gia cho phương án A_i theo tiêu chí C_j ; $i = 1, 2, \dots, m$; và $j = 1, 2, \dots, n$

Bước 2. Thiết lập ma trận quyết định chuẩn hóa dữ liệu

Ma trận $A = (a_{ij})_{m \times n}$ được chuẩn hóa thành $\tilde{A} = (\tilde{a}_{ij})_{m \times n}$, trong đó:

$$\tilde{a}_{ij} = \frac{x_{ij}}{\sqrt{\sum_{k=1}^m a_{ik}^2}}; i = 1, 2, \dots, m; j = 1, 2, \dots, n. \quad (1.16)$$

Bước 3. Xác định ma trận chuẩn hóa có trọng số

$$Y_{ij} = \tilde{a}_{ij} \times \tilde{w}_j; i = 1, 2, \dots, m; j = 1, 2, \dots, n.$$

Y_{ij} được xác định bằng cách nhân từng cột của ma trận chuẩn hóa $\tilde{a}_{ij}$ với trọng số $\tilde{w}_j$ tương ứng của tiêu chí C_j . $\tilde{w}_j$ được đo lường bởi DM bằng cách ước lượng các giá trị cho các tiêu chí sao cho $\sum_{k=1}^m \tilde{w}_j = 1$.

Bước 4. Xác định PIS (ℓ^+) và NIS (ℓ^-) được tính theo công thức sau:

$$\mathcal{B}^+ = \{\mathcal{B}_1^+, \mathcal{B}_2^+, \dots, \mathcal{B}_n^+\} \quad (1.17)$$

$$\mathcal{B}^- = \{\mathcal{B}_1^-, \mathcal{B}_2^-, \dots, \mathcal{B}_n^-\} \quad (1.18)$$

$$\begin{cases} \mathcal{B}_j^+ = \max(\mathcal{B}_{ij}); i = 1, \dots, m \text{ nếu tiêu chí } C_j \text{ là tiêu chí về lợi ích} \\ \mathcal{B}_j^- = \min(\mathcal{B}_{ij}); i = 1, \dots, m \text{ nếu tiêu chí } C_j \text{ là tiêu chí về chi phí} \end{cases}$$

Bước 5. Tính khoảng cách (d_i^+ và d_i^-) của phương án A_i đến ($\mathcal{B}^+$) và ($\mathcal{B}^-$) được tính như sau:

$$d_i^+ = \left(\sum_{j=1}^n (\mathcal{B}_{ij} - \mathcal{B}_j^+)^2 \right)^{\frac{1}{2}}; i = 1, \dots, m \quad (1.19),$$

$$d_i^- = \left(\sum_{j=1}^n (\mathcal{B}_{ij} - \mathcal{B}_j^-)^2 \right)^{\frac{1}{2}}; i = 1, \dots, m \quad (1.20)$$

Bước 6. Tính hệ số gần và xếp hạng các phương án

$$CC_i = \frac{d_i^-}{(d_i^- + d_i^+)}; 0 \leq s_i \leq 1; i = 1, \dots, m \quad (1.21)$$

Phương án có hệ số CC_i , $i = 1, \dots, m$ cao nhất là phương án gần với PIS ($\mathcal{B}^+$) được ưu tiên xếp thứ nhất. Tiếp tục xếp hạng các phương án lựa chọn theo thứ tự giảm dần hệ số CC_i .

Ưu điểm chính của TOPSIS là quy trình tính toán rõ ràng, khả năng xử lý đồng thời nhiều tiêu chí, kết quả dễ diễn giải và thuận lợi cho so sánh đối sánh giữa các phương án. Tuy nhiên, TOPSIS truyền thống chủ yếu làm việc với dữ liệu đầu vào là số rõ và trọng số của tiêu chí đã được biết trước hoặc gán trực tiếp. Trong nhiều bài toán thực tế, nhất là MCGDM trong môi trường thông tin ngôn ngữ, đánh giá của chuyên gia thường mang tính mơ hồ, không chắc chắn và được biểu đạt bằng các từ ngôn ngữ hơn là bằng số rõ. Khi đó, việc chuyển trực tiếp các đánh giá ngôn ngữ thành số thực có thể làm suy giảm một phần nội dung ngữ nghĩa của thông tin đầu vào [49, 54, 48]. Đồng thời, trong bài toán MCGDM, việc kết nhập ý kiến chuyên gia và xử lý khác biệt quan điểm cũng đòi hỏi các mở rộng phù hợp hơn so với TOPSIS truyền thống.

Những hạn chế trên là cơ sở cho sự phát triển của FTOPSIS và các mô hình xử lý thông tin ngôn ngữ [21, 114]. Các tiếp cận này mở rộng TOPSIS bằng cách biểu diễn đánh giá dưới dạng số mờ hoặc giá trị ngôn ngữ, qua đó nâng cao khả năng xử lý bất định [111]. Tuy nhiên, các mô hình FTOPSIS hiện có vẫn bộc lộ một số hạn chế như khoảng cách giữa các phương án và các giải pháp lý tưởng vẫn được xác định trên các biểu diễn số mờ, sau đó phải giải mờ hoặc xấp xỉ để có thể so sánh và xếp hạng. Quy trình này có thể tạo ra sự không trùng khớp giữa ý nghĩa ngôn ngữ ban đầu và kết quả

tính toán trung gian, nhất là khi khoảng cách số học giữa các biểu diễn mờ không phản ánh đúng khoảng cách ngữ nghĩa giữa các từ. Bên cạnh đó, việc xác định PIS/NIS trong mô hình FTOPSIS thường phụ thuộc vào không gian biểu diễn đã chọn, nên thiếu một cơ sở hình thức thống nhất để bảo đảm rằng trật tự ưu tiên cuối cùng thực sự tương thích với trật tự ngữ nghĩa của các đánh giá chuyên gia.

Do đó, với mục tiêu bảo toàn tốt hơn trật tự và nội dung ngữ nghĩa của đánh giá chuyên gia, luận án kế thừa tư tưởng xếp hạng của TOPSIS, đồng thời phát triển hướng tiếp cận dựa trên HA để xây dựng thang đo ngôn ngữ, định lượng ngữ nghĩa, xác định PIS/NIS và đo khoảng cách trên một miền ngữ nghĩa thống nhất hơn [53, 52, 54]. Đây là cơ sở phương pháp luận để Chương 2 phát triển mô hình HA-TOPSIS và tiếp tục mở rộng sang bài toán MCGDM động trong Chương 4 của luận án.

1.3.3. Mối liên hệ giữa AHP và TOPSIS trong các mô hình MCGDM

Trong các mô hình MCDM và MCGDM, xác định trọng số của tiêu chí và xếp hạng phương án là hai giai đoạn có mối quan hệ hữu cơ, hỗ trợ lẫn nhau để đạt được kết quả tối ưu [57, 117]. AHP và mở rộng mờ FAHP thể hiện ưu thế vượt trội trong việc khai thác tri thức chuyên gia thông qua các ma trận so sánh cặp, từ đó suy ra các trọng số phản ánh chính xác tầm quan trọng tương đối của các tiêu chí [57, 117, 111]. Dựa trên tập trọng số này, TOPSIS và FTOPSIS đóng vai trò là cơ chế tổng hợp, thực hiện xếp hạng dựa trên nguyên lý một phương án lý tưởng phải có khoảng cách ngắn nhất PIS và xa nhất tới NIS [21, 48]. Chính sự bổ sung chức năng này các mô hình tích hợp như AHP-TOPSIS và FAHP-FTOPSIS đã trở thành những khung phương pháp luận phổ biến nhất trong các bài toán MCGDM ngôn ngữ [110, 117]. Các mô hình lai này được ứng dụng thành công trong nhiều lĩnh vực như lựa chọn nhà cung cấp bền vững, đánh giá rủi ro và ưu tiên đầu tư [110, 48]. Sự kết hợp này không chỉ tận dụng được khả năng kiểm soát tính nhất quán trong phán đoán của AHP/FAHP, mà còn phát huy thế mạnh phân tích khoảng cách đa chiều của TOPSIS/FTOPSIS, tạo nên một quy trình MCGDM khép kín, minh bạch và có độ tin cậy khoa học cao.

Tuy nhiên, một hạn chế đáng chú ý của các mô hình tích hợp dựa trên tập mờ truyền thống là sự thiếu thống nhất về nền tảng biểu diễn thông tin xuyên suốt quy trình ra quyết định. Trong nhiều mô hình, các phán đoán so sánh cặp được biểu diễn bằng số rõ hoặc số mờ, trọng số sau đó thường phải giải mờ, trong khi đánh giá phương án và khoảng cách tới các giải pháp lý tưởng lại được xử lý trên một không gian biểu diễn

khác, dẫn đến nguy cơ sai lệch kết quả [63, 111]. Về mặt phương pháp luận, việc thiết lập sự tương thích hoàn toàn giữa các hàm thuộc số học và nội hàm của từ ngôn ngữ trong quá trình tổng hợp đánh giá vẫn là một thách thức lớn [93, 63, 11]. Theo hướng này, Shu và cộng sự nhấn mạnh nhu cầu cấp thiết về việc hạn chế mất mát thông tin trong các bước hợp nhất thông tin ngôn ngữ [99]. Đồng thời, các nghiên cứu cũng chỉ ra rằng ngay cả trong môi trường 2-tuple, quy trình đánh giá tính nhất quán và xác định trọng số của tiêu chí cần được thiết kế liên thông để vừa bảo toàn tri thức ưu tiên ban đầu của chuyên gia, vừa nâng cao độ tin cậy của thứ hạng cuối cùng [63]. Từ đó, có thể thấy rằng hạn chế lớn nhất của các mô hình tích hợp hiện nay nằm ở sự thiếu nhất quán ngữ nghĩa giữa các giai đoạn xử lý; đây chính là cơ sở khoa học để luận án lựa chọn FAHP và TOPSIS làm hai hướng tiếp cận chính, đồng thời phát triển chúng theo một khuôn khổ thống nhất hơn dựa trên lý thuyết HA.

Trong bối cảnh hiện nay, yêu cầu đặt ra đối với quy trình MCGDM không chỉ dừng lại ở việc tích hợp kỹ thuật giữa AHP/FAHP và TOPSIS/FTOPSIS, mà quan trọng hơn là phải đảm bảo sự liên thông và thống nhất ngữ nghĩa xuyên suốt từ bước xác định trọng số của tiêu chí đến bước đánh giá và xếp hạng phương án [57]. Sự đứt gãy hoặc mất mát thông tin trong quá trình chuyển đổi giữa các biểu diễn ngôn ngữ và các giá trị tính toán trung gian thường dẫn đến kết quả thiếu chính xác và khó diễn giải [63, 99; 49]. Chính vì vậy, lý thuyết HA được luận án lựa chọn làm nền tảng chủ đạo nhờ khả năng mô hình hóa trực tiếp cấu trúc thứ tự và nội hàm ngữ nghĩa của miền từ, giúp bảo toàn trật tự ngữ nghĩa tự nhiên của các phán đoán chuyên gia [53, 52, 54]. Việc kế thừa AHP/FAHP cho bài toán xác định trọng số của tiêu chí và TOPSIS/FTOPSIS cho bài toán xếp hạng là hệ quả tất yếu của cấu trúc bài toán MCGDM ngôn ngữ hiện nay, nơi tri thức chuyên gia và phương thức đo khoảng cách tới các giải pháp lý tưởng cần được kết hợp hài hòa [57, 117]; đồng thời luận án áp dụng mô hình 4-tuple ngữ nghĩa để giải quyết yêu cầu bảo toàn thông tin trong quá trình tính toán [54].

Trên cơ sở đó, luận án phát triển mô hình HA-TOPSIS cho bài toán xếp hạng phương án ở Chương 2, xây dựng mô hình tích hợp EFAHP - 4-tuple - HA ở Chương 3 và mở rộng sang L-FAHP - L-FTOPSIS động dựa trên HA ở Chương 4. Cách tiếp cận này nhằm duy trì sự thống nhất ngữ nghĩa xuyên suốt từ xác định trọng số đến xếp hạng phương án trong môi trường thông tin ngôn ngữ.

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

1.4.1. Các nghiên cứu về MCDM truyền thống và tiếp cận mở rộng dựa trên lý thuyết tập mờ

Các phương pháp MCDM truyền thống được hình thành nhằm hỗ trợ lựa chọn, đánh giá và xếp hạng phương án trong điều kiện có nhiều tiêu chí, nhưng về bản chất chủ yếu được xây dựng cho dữ liệu đầu vào là số rõ [13, 24]. Trong nhóm này, các phương pháp như AHP [95], ANP [96], TOPSIS [59], VIKOR [68], ELECTRE [94] hay PROMETHEE [109] đã cung cấp những khuôn khổ hình thức quan trọng cho các nhiệm vụ như xác định trọng số của tiêu chí, so sánh phương án, thiết lập quan hệ ưu tiên và xếp thứ hạng tổng thể [24, 92]. Ưu điểm của các phương pháp này là đặt nền tảng cho phương pháp luận của MCDM và sau đó được mở rộng sang MCGDM thông qua các cơ chế tổng hợp ý kiến của nhiều chuyên gia, tổng hợp trọng số và kết nhập ma trận đánh giá [13, 14]. Tuy nhiên, khi dữ liệu đầu vào là thông tin ngôn ngữ, các phương pháp truyền thống bộc lộ giới hạn rõ rệt bởi giả thiết một không gian đánh giá số rõ và một cấu trúc thông tin tương đối ổn định [93].

Để giải quyết các hạn chế trên, các tiếp cận dựa trên lý thuyết tập mờ được phát triển nhằm mở rộng khả năng của MCDM và MCGDM sang môi trường bất định. Trên nền tảng tập mờ và biến ngôn ngữ của Zadeh [115, 116], nhiều mô hình dựa trên số mờ và các dạng mở rộng của tập mờ đã được xây dựng để biểu diễn các đánh giá mơ hồ của chuyên gia [93]. Trong dòng nghiên cứu này, FAHP [19, 108, 25, 58] giữ vị trí nổi bật trong hướng tiếp cận xác định trọng số của tiêu chí, bởi nó kế thừa cơ chế so sánh cặp của AHP [95] nhưng thay biểu diễn số rõ bằng số mờ để phản ánh tốt hơn các phán đoán bất định [24, 92]. Tương ứng, FTOPSIS [21] giữ vị trí nổi bật trong xếp hạng phương án, bởi nó mở rộng ý tưởng giải pháp lý tưởng của TOPSIS sang không gian mờ và trở thành một công cụ được sử dụng rộng rãi trong các bài toán MCGDM [13, 92].

Cùng với đó, các hướng mở rộng dựa trên tập mờ trực cảm [45], tập mờ do dự [113, 44], tập mờ Pythagorean [8, 9, 34], tập neutrosophic [105, 69], tập mờ T-spherical [22] và một số hướng tiếp cận gần đây như tập từ ngôn ngữ xác suất (PLTS) cho phép mô tả đánh giá bằng nhiều thuật ngữ ngôn ngữ kèm xác suất hoặc mức khả năng tương ứng, phù hợp khi chuyên gia có thể lượng hóa mức độ tin cậy đối với từng nhãn ngôn ngữ [78]; tập từ ngôn ngữ mờ do dự (HFLTS) tập trung biểu diễn trạng thái do dự giữa nhiều thuật ngữ, qua đó phản ánh tốt hơn tình huống chuyên gia chưa thể lựa chọn một giá trị ngôn ngữ duy nhất [37]. Z-number ngôn ngữ mở rộng tư tưởng Z-number vào môi trường ngôn ngữ bằng cách mô hình hóa đồng thời nội dung đánh giá và độ tin cậy của

đánh giá [74] tiếp tục làm tăng đáng kể khả năng biểu diễn và xử lý bất định trong nhiều lĩnh vực khác nhau [13, 93]. Các hướng nghiên cứu này cho thấy sự phát triển mạnh của MCGDM ngôn ngữ theo chiều biểu diễn bất định. Tuy nhiên, chúng không phải là nền tảng phương pháp luận chính của luận án này, vì trọng tâm của luận án là xử lý trực tiếp thông tin ngôn ngữ trên cơ sở cấu trúc ngữ nghĩa, trật tự ngữ nghĩa, độ đo tính mờ và hàm định lượng ngữ nghĩa của HA.

Thành tựu quan trọng nhất của các nghiên cứu này là mở rộng đáng kể phạm vi áp dụng của MCDM và MCGDM từ dữ liệu số rõ sang dữ liệu mơ hồ và không chắc chắn [13]. Qua đó, các phương pháp này đã giải quyết hiệu quả nhiều bài toán thực tiễn như lựa chọn nhà cung cấp [92, 40, 70, 24], đánh giá rủi ro công nghệ [85, 35], hay đánh giá chuyển đổi số của doanh nghiệp [118, 73, 16, 75, 56, 80]. Đặc biệt, mô hình FAHP - FTOPSIS đã trở thành một quy trình tích hợp phổ biến trong MCGDM, trong đó FAHP đảm nhiệm giai đoạn lượng hóa tri thức chuyên gia để xác định trọng số của tiêu chí, còn FTOPSIS đảm nhiệm giai đoạn xếp hạng phương án [10, 64, 87, 91]. Tuy nhiên, xét từ góc độ tính toán với từ [116], phần lớn các mô hình thuộc nhóm này vẫn tiếp cận bài toán theo hướng mờ hóa thông tin ngôn ngữ. Tức là, ngữ nghĩa của các từ đánh giá chủ yếu được biểu diễn gián tiếp qua các hàm thuộc hoặc các tham số số mờ, trong khi cấu trúc ngữ nghĩa nội tại và sự sai biệt trong nhận thức cá nhân về từ ngữ chưa được mô hình hóa một cách tường minh [54, 93]. Chính giới hạn phương pháp luận này làm xuất hiện nhu cầu chuyển dịch sang các mô hình tính toán ngôn ngữ có cấu trúc biểu diễn chặt chẽ, nơi trọng tâm là bảo toàn tối đa nội dung ngữ nghĩa và tính nhất quán của thông tin trong toàn bộ quá trình tính toán [54, 74, 93].

1.4.2. Các nghiên cứu dựa trên các mô hình ngôn ngữ 2-tuple và 4-tuple

Sự phát triển của các mô hình tính toán ngôn ngữ xuất phát từ nhận thức rằng thông tin đánh giá trong MCGDM không chỉ bất định mà còn mang bản chất ngôn ngữ, do đó việc chỉ mờ hóa các từ đánh giá chưa đủ để bảo toàn đầy đủ ý nghĩa của chúng [93, 54]. Trên cơ sở các ý tưởng tính toán với từ của Zadeh [116] và các tiếp cận ngôn ngữ ký hiệu sau đó, một dòng nghiên cứu mới được hình thành nhằm thực hiện tính toán trực tiếp hơn trên miền nhân ngôn ngữ, giảm phụ thuộc vào quá trình giải mờ truyền thống và hạn chế mất mát thông tin do làm tròn hoặc xấp xỉ [27, 49].

Trong dòng phát triển này, mô hình 2-tuple của Herrera được xem là một bước tiến quan trọng [49]. Ý nghĩa của 2-tuple không chỉ ở kỹ thuật biểu diễn bằng cặp (s_i, α) gồm s_i là nhãn ngôn ngữ và α là tham số dịch chuyển, mà còn ở việc tạo ra một cơ chế tính toán cho phép duy trì thông tin ngữ nghĩa tốt hơn so với mô hình dựa trên chỉ số

ngôn ngữ đơn thuần [49, 54]. Nhờ đó, nhiều nghiên cứu đã áp dụng 2-tuple để phát triển các mô hình MCDM và MCGDM ngôn ngữ [36, 39, 63, 5, 7]; cũng như các mô hình tích hợp với AHP, TOPSIS và các phương pháp ra quyết định khác [36, 28, 5, 7, 99]. Do đó, 2-tuple là bước tiến quan trọng theo hướng tiếp cận tính toán với từ ngôn ngữ, vì nó tạo ra cầu nối giữa biểu diễn ký hiệu và yêu cầu tính toán mà vẫn giữ khả năng diễn giải bằng ngôn ngữ ở đầu ra [49, 54]. Tuy nhiên, cùng với sự phát triển của các bài toán MCGDM phức tạp hơn, các giới hạn của 2-tuple cũng dần bộc lộ. Về bản chất, nhiều mô hình 2-tuple vẫn dựa vào giả định phân bố tương đối đều của các nhãn trên một thang đo ngôn ngữ đã được xác định trước, trong khi khoảng cách ngữ nghĩa giữa các từ trong ngôn ngữ tự nhiên không phải lúc nào cũng đồng đều [52, 54]. Mặt khác, cơ chế tính toán của 2-tuple vẫn phụ thuộc vào không gian chỉ số và tham số dịch chuyển, nên chưa cung cấp đầy đủ một khuôn khổ hình thức để mô hình hóa quan hệ giữa nhãn ngôn ngữ, mức độ mờ và vị trí ngữ nghĩa của từ trong cùng một cấu trúc thống nhất.

Chính từ yêu cầu đó, các nghiên cứu tiếp tục hướng đến những mô hình biểu diễn giàu ngữ nghĩa hơn, trong đó mô hình 4-tuple ngữ nghĩa do Nguyễn Cát Hồ và cộng sự đề xuất tạo ra một cầu nối giữa biểu diễn ngữ nghĩa và khả năng tính toán, vì mỗi từ ngôn ngữ được mô tả bằng một cấu trúc bốn thành phần $(x, S_k(x), \vartheta(x), r_x)$, có ý nghĩa nổi bật [54]. So với 2-tuple, 4-tuple ngữ nghĩa không chỉ nhằm giảm mất mát thông tin trong tính toán, mà còn gắn trực tiếp nhãn ngôn ngữ với giá trị định lượng ngữ nghĩa, khoảng tương tự ngữ nghĩa và tham chiếu ngữ nghĩa trong cùng một cấu trúc biểu diễn [54]. Điều này cho phép tăng khả năng bảo toàn thông tin ngôn ngữ, đồng thời tạo nền cho việc tổng hợp, so sánh và kết nhập các đánh giá trong một không gian có quan hệ ngữ nghĩa rõ hơn. Vì vậy, nếu 2-tuple là bước chuyển quan trọng từ mờ hóa thông tin ngôn ngữ sang biểu diễn ngôn ngữ có kiểm soát, thì 4-tuple ngữ nghĩa có thể được xem là một bước tiến tiếp theo theo hướng cấu trúc hóa sâu hơn trực biểu diễn ngôn ngữ. Dẫu vậy, trong giai đoạn này, nhiều mô hình vẫn còn thiếu một nền tảng đại số đủ mạnh để thống nhất giữa sinh tập từ ngôn ngữ, trật tự ngữ nghĩa, định lượng ngữ nghĩa và tính toán với từ. Nhu cầu đó dẫn sang hướng tiếp cận dựa trên lý thuyết HA.

1.4.3. Các nghiên cứu dựa trên HA cho bài toán MCDM và MCGDM

Trong các hướng nghiên cứu MCDM và MCGDM ngôn ngữ, tiếp cận dựa trên HA có ý nghĩa quan trọng vì nó không chỉ cung cấp một công cụ biểu diễn ngôn ngữ, mà còn đặt toàn bộ miền giá trị ngôn ngữ trên một cơ sở hình thức chặt chẽ hơn [53]. So với các tiếp cận dựa trên tập mờ hoặc chỉ số ngôn ngữ, từ ngôn ngữ trong HA không chỉ được xem như một lớp nhãn, mà là một miền có cấu trúc nội tại và có thể được xử lý

trực tiếp trên chính cấu trúc đó [52, 53, 54]. Để làm rõ hơn khoảng trống nghiên cứu hiện tại, cần xem xét chi tiết cách thức HA được ứng dụng trong các nghiên cứu tiêu biểu [3, 33, 54, 89].

Thứ nhất, đối với nhóm nghiên cứu về biểu diễn ngôn ngữ và xếp hạng phương án [54, 33]: Nghiên cứu của [54] đã bước đầu sử dụng HA để xây dựng thang đo ngôn ngữ 4-tuple ngữ nghĩa, liên kết nhãn ngôn ngữ với khoảng tương tự, giá trị định lượng và tham chiếu ngữ nghĩa, nhằm phục vụ việc đánh giá và xếp hạng các phương án. Tương tự, nghiên cứu [33] đã mở rộng các toán tử kết nhập thông tin ngôn ngữ dựa trên hàm SQM. Mặc dù các nghiên cứu này đã chứng minh được khả năng của HA trong việc bảo toàn trật tự ngữ nghĩa, nhưng chúng chủ yếu giải quyết bài toán ở khâu tính toán điểm đánh giá cuối cùng của phương án. Một hạn chế lớn là các mô hình này đều thiết lập dựa trên giả định trọng số của các tiêu chí đã được cho trước hoặc được gán trực tiếp một cách chủ quan, chưa được xác định từ tri thức chuyên gia qua quy trình kiểm soát tính nhất quán.

Thứ hai, đối với nhóm nghiên cứu về các mô hình đánh giá và kết nhập tổng hợp [3, 89]: Các nghiên cứu này áp dụng HA để xây dựng các mô hình đánh giá trên miền từ, chuyển đổi trực tiếp các đánh giá ngôn ngữ thành giá trị định lượng để tổng hợp kết quả. Tuy nhiên, các nghiên cứu này lại thiếu vắng quy trình kiểm soát độ tin cậy của thông tin đầu vào. Việc thiếu một cơ chế kiểm tra tính nhất quán (Consistency Ratio) khi các chuyên gia đưa ra đánh giá so sánh cặp ngôn ngữ đã làm giảm đáng kể độ tin cậy của mô hình khi áp dụng với các bài toán có quy mô tiêu chí lớn.

Các nghiên cứu [3, 33, 54, 89] cho thấy, hướng nghiên cứu dựa trên HA tuy nhiều tiềm năng nhưng vẫn còn phát triển phân tán chưa hình thành được một khung phương pháp luận thống nhất từ biểu diễn ngôn ngữ, xác định trọng số của tiêu chí đến xếp hạng phương án. Mặc dù HA đã hỗ trợ tốt cho việc xếp hạng phương án, nhưng chưa có nghiên cứu nào phát triển mô hình TOPSIS trong môi trường ngôn ngữ dựa trên HA, trong đó xác định PIS/NIS, khoảng cách ngữ nghĩa, hệ số gần và thứ hạng được xác định trên cùng miền ngữ nghĩa. Bên cạnh đó trong các nghiên cứu trên, trọng số của tiêu chí thường được giả định là đã biết trước hoặc được gán trực tiếp từ chuyên gia, thay vì được xác định thông qua một quy trình có kiểm soát tính nhất quán trên chính không gian ngữ nghĩa dựa trên HA; và chủ yếu được phát triển trong bối cảnh tĩnh. Từ những khoảng trống đó, xu hướng tích hợp các thành phần của MCGDM và mở rộng bài toán sang bối cảnh động trở thành hướng phát triển tiếp theo của MCGDM dựa trên HA. Đây chính là tiền đề để luận án phát triển các mô hình mới trong các Chương 2, 3 và 4.

1.4.4. Các nghiên cứu về MCGDM động trong môi trường ngôn ngữ

So với các mô hình MCGDM tĩnh, các nghiên cứu về MCGDM động trong môi trường ngôn ngữ ra đời muộn hơn, nhưng có ý nghĩa thực tiễn ngày càng rõ khi nhiều bài toán ra quyết định không diễn ra tại một thời điểm đơn lẻ mà lặp lại qua nhiều chu kỳ đánh giá [17]. Trong những bối cảnh như đánh giá mức độ sẵn sàng chuyển đổi số, lựa chọn phương án đầu tư công nghệ, theo dõi hiệu quả tổ chức hoặc quản trị rủi ro, quyết định hiện tại thường chịu ảnh hưởng không chỉ từ dữ liệu tại thời điểm đang xét mà còn từ thông tin lịch sử, từ sự thay đổi của tập phương án, bộ tiêu chí và từ sự dịch chuyển trong quan điểm của nhóm chuyên gia [17, 29]. Chính vì vậy, MCGDM động được hình thành để xử lý một lớp vấn đề mới: không chỉ tổng hợp đánh giá đa tiêu chí và đa chuyên gia, mà còn phải xem xét tích hợp dữ liệu lịch sử với dữ liệu hiện tại, cũng như khả năng thích ứng với sự thay đổi của tập phương án, tập tiêu chí và hội đồng chuyên gia theo thời gian vào cấu trúc của quá trình ra quyết định [17, 29].

Phần lớn các nghiên cứu hiện có về MCGDM động vẫn phát triển dựa trên tập mờ hoặc các mở rộng của tập mờ. Nhiều mô hình tập trung vào việc cập nhật trọng số theo thời gian, tích hợp đánh giá ở nhiều thời điểm [31, 32, 91], xử lý dữ liệu khuyết thiếu hoặc xây dựng cơ chế suy giảm ảnh hưởng của dữ liệu lịch sử [29, 112, 11, 60], trong đó các cấu hình tích hợp FAHP và FTOPSISIS động là một hướng đáng chú ý [31, 32, 91]. Những kết quả này cho thấy lĩnh vực đã bắt đầu chuyển từ bài toán MCGDM ngôn ngữ tĩnh sang việc xem xét quá trình ra quyết định động theo thời gian. Về mặt ứng dụng, đây là bước phát triển cần thiết và có giá trị, bởi nó giúp các mô hình ra quyết định phản ánh tốt hơn thực tiễn vận động của dữ liệu và bối cảnh đánh giá.

Tuy vậy, xét về phương pháp luận, nhiều mô hình MCGDM động hiện nay vẫn mới chỉ mở rộng tính động trên nền các phép toán số mờ hoặc các cơ chế trung bình có trọng số, trong khi chưa thật sự hình thức hóa đầy đủ việc tích hợp ngữ nghĩa theo thời gian [31, 32, 91]. Nói cách khác, yếu tố động thường được biểu diễn ở mức điều chỉnh trọng số thời gian, cập nhật ma trận đánh giá hoặc làm suy giảm ảnh hưởng của dữ liệu cũ, nhưng chưa được gắn chặt với một không gian ngôn ngữ có cấu trúc thống nhất cho cả biểu diễn, kết nhập và diễn giải kết quả. Vì vậy, các mô hình này tuy có giá trị ứng dụng rõ rệt nhưng vẫn để lại thách thức lớn ở chỗ làm thế nào để tích hợp dữ liệu lịch sử và dữ liệu hiện tại mà vẫn bảo toàn được logic ngữ nghĩa của các đánh giá. Đây chính là điểm nổi bật giữa dòng nghiên cứu MCGDM động với hướng tiếp cận dựa trên HA, nơi có tiềm năng phát triển các phép kết nhập ngôn ngữ và các cơ chế suy giảm ngữ nghĩa trên một nền tảng hình thức rõ ràng hơn.

Từ tổng quan trên có thể thấy rằng nghiên cứu về MCDM và MCGDM ngôn ngữ đã phát triển qua nhiều hướng tiếp cận kế tiếp nhau, từ nhóm dựa trên tập mờ và số mờ, sang nhóm dựa trên các mô hình ngôn ngữ như 2-tuple và 4-tuple ngữ nghĩa, rồi đến nhóm dựa trên HA và gần đây là các mở rộng theo hướng động. Mỗi dòng nghiên cứu ra đời đều xử lý một bài toán cụ thể, đồng thời tạo ra những kết quả có giá trị ở một hoặc một số thành phần của quy trình ra quyết định. Tuy nhiên, chính sự phát triển mạnh nhưng phân tán đó cho thấy các kết quả hiện có vẫn chưa thật sự hội tụ thành một khung phương pháp luận thống nhất, trong khi cần liên kết xuyên suốt hướng biểu diễn ngôn ngữ, từ xác định trọng số của tiêu chí đến xếp hạng phương án có tích hợp dữ liệu theo thời gian. Đây là cơ sở trực tiếp để tiếp tục xác định ba khoảng trống mà luận án tập trung giải quyết trong mục 1.5 dưới đây.

1.5. Khoảng trống nghiên cứu và định hướng luận án

Trên cơ sở tổng quan các công trình liên quan ở mục 1.4, có thể nhận thấy bài toán MCDM/MCGDM trong môi trường thông tin ngôn ngữ đã được nghiên cứu theo nhiều hướng tiếp cận. Các phương pháp như AHP [95], TOPSIS [59] tạo nền tảng cho xác định trọng số của tiêu chí và xếp hạng phương án; FAHP [106, 19], FTOPSIS [21] và các mô hình tích hợp dựa trên lý thuyết tập mờ góp phần xử lý thông tin bất định; các mô hình 2-tuple [49] và các mô hình MCGDM dựa trên HA [3, 54, 33, 89] nhấn mạnh yêu cầu bảo toàn thông tin ngôn ngữ; trong khi HA cung cấp cơ sở đại số để mô hình hóa miền từ có trật tự ngữ nghĩa, độ đo tính mờ và SQM [52, 53, 54]. Tuy nhiên, các cách tiếp cận trên còn bộc lộ một số hạn chế như đã chỉ ra trong (mục 1. Tính cấp thiết của đề tài) và còn phân tán theo từng thành phần của quy trình ra quyết định. Do đó, từ những phân tích trên, ba khoảng trống khoa học chủ yếu của luận án được xác định là xếp hạng phương án, xác định trọng số của tiêu chí và mở rộng bài toán sang bối cảnh động có tích hợp dữ liệu lịch sử cần tập trung làm rõ như sau:

Khoảng trống nghiên cứu thứ nhất, liên quan đến bài toán xếp hạng phương án trong MCGDM ngôn ngữ. Các mô hình TOPSIS [59] và FTOPSIS [21] đã được sử dụng rộng rãi nhờ ý tưởng so sánh phương án với PIS/NIS. Tuy nhiên, trong nhiều mô hình hiện có, việc xác định PIS/NIS, đo khoảng cách và tính hệ số gần thường được thực hiện trên các biểu diễn số thực, số mờ hoặc thông qua các bước giải mờ, ánh xạ xấp xỉ trở lại nhãn ngôn ngữ [49]. Cách tiếp cận này thuận tiện cho tính toán, nhưng có thể làm suy giảm khả năng diễn giải ngữ nghĩa, khi khoảng cách số học không phản ánh đầy đủ cấu trúc ngữ nghĩa thứ tự vốn có của từ ngôn ngữ, ở đây thiếu một phương pháp hình thức đầy

đủ và chặt chẽ dựa trên cấu trúc đại số [48, 49, 54]. Trong khi đó, HA cung cấp cơ sở toán học dựa trên tiên đề cho việc xây dựng thang đo ngôn ngữ, xác định quan hệ thứ tự, độ đo tính mờ và SQM [52, 53]. Bên cạnh đó, các nghiên cứu MCGDM dựa trên HA hiện có mới chủ yếu sử dụng HA như một công cụ biểu diễn hoặc kết nhập, chưa khai thác đầy đủ cấu trúc trật tự ngữ nghĩa của HA để hình thành một quy trình TOPSIS ngôn ngữ cho xác định PIS/NIS, đo khoảng cách ngữ nghĩa và tính hệ số gần trong cùng một không gian ngữ nghĩa [3, 54, 33, 89]. Từ khoảng trống trên, câu hỏi nghiên cứu thứ nhất được đặt ra như sau: Làm thế nào để xây dựng mô hình TOPSIS ngôn ngữ dựa trên HA có khả năng xác định PIS/NIS, đo khoảng cách ngữ nghĩa và tính hệ số gần của các phương án trên cùng một miền ngữ nghĩa, qua đó bảo đảm tính nhất quán trong xếp hạng phương án của bài toán MCGDM ngôn ngữ? Câu hỏi này được hiện thực hóa thành đóng góp mới thứ nhất của luận án trong Chương 2.

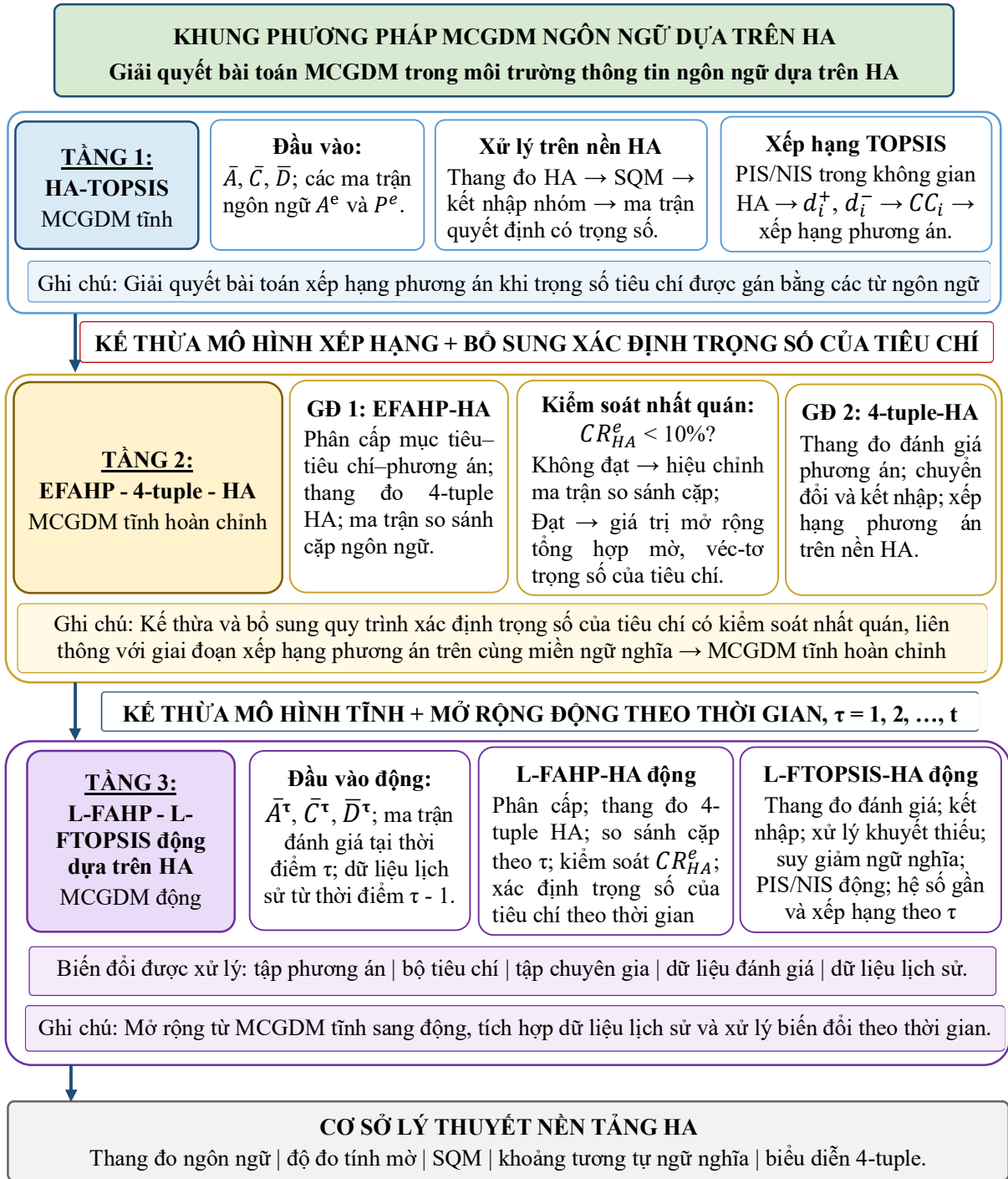
Khoảng trống nghiên cứu thứ hai, liên quan đến xác định trọng số của tiêu chí và kiểm soát tính nhất quán trong môi trường ngôn ngữ dựa trên HA. Trong MCGDM, trọng số của tiêu chí ảnh hưởng trực tiếp đến kết quả tổng hợp và xếp hạng [110, 111, 117]. Khi chuyên gia sử dụng từ ngôn ngữ để so sánh cặp tiêu chí, phán đoán có thể mang tính chủ quan, mơ hồ và không hoàn toàn nhất quán. FAHP đã được dùng rộng rãi cho so sánh cặp mờ để xác định trọng số của tiêu chí [106, 19, 58], nhưng các thao tác tổng hợp mờ và giải mờ có thể làm giảm tính liên thông ngữ nghĩa giữa giai đoạn xác định trọng số của tiêu chí và giai đoạn đánh giá xếp hạng phương án [49, 54, 57]. Trong khi đó, một số mô hình HA hiện có trọng số của tiêu chí thường giả thiết trước hoặc để chuyên gia gán trực tiếp, chưa hình thành đầy đủ một quy trình xác định trọng số của tiêu chí có kiểm soát tính nhất quán trên nền biểu diễn ngữ nghĩa dựa trên HA [3, 54, 33, 89]. Từ khoảng trống này, câu hỏi nghiên cứu thứ hai được xác định như sau: Làm thế nào để xây dựng mô hình MCGDM tích hợp FAHP cải tiến với mô hình 4-tuple ngữ nghĩa dựa trên HA nhằm xác định trọng số của tiêu chí có kiểm soát tính nhất quán, đồng thời liên thông ngữ nghĩa với giai đoạn kết nhập đánh giá và xếp hạng phương án? Câu hỏi này được cụ thể hóa thành đóng góp mới thứ hai của luận án trong Chương 3.

Khoảng trống nghiên cứu thứ ba, liên quan đến mở rộng MCGDM ngôn ngữ dựa trên HA từ bối cảnh tĩnh sang bối cảnh động. Trong thực tiễn, nhiều quá trình ra quyết định lặp lại qua nhiều thời điểm; tập phương án có thể được bổ sung hoặc loại bỏ, bộ tiêu chí có thể điều chỉnh, hội đồng chuyên gia có thể thay đổi, dữ liệu đánh giá và trọng số của tiêu chí có thể được cập nhật. Các mô hình MCGDM động dựa trên lý thuyết mờ đã bước đầu xem xét yếu tố thời gian, nhưng việc tích hợp dữ liệu lịch sử với dữ liệu

hiện tại trong một không gian ngôn ngữ có cấu trúc ngữ nghĩa thống nhất vẫn là vấn đề cần tiếp tục nghiên cứu [32, 23, 29]. Đối với hướng tiếp cận HA, phần lớn được xây dựng cho bối cảnh tĩnh, ra quyết định tại một thời điểm [3, 54, 33, 89]. Nếu chỉ xử lý dữ liệu hiện tại, mô hình có thể làm giảm khả năng kế thừa thông tin; ngược lại, nếu duy trì nguyên vẹn ảnh hưởng của dữ liệu quá khứ mà không xét đến tính lỗi thời, kết quả đánh giá có thể bị sai lệch. Hạn chế này càng trở nên rõ hơn khi các phép kết nhập ngôn ngữ trên thang đo 4-tuple ngữ nghĩa mới chỉ dừng ở các dạng trung bình số học hoặc trung bình có trọng số [3, 54], chưa khai thác đầy đủ cấu trúc 4-tuple ngữ nghĩa trong kết nhập ở nhiều thời điểm. Từ khoảng trống trên, câu hỏi nghiên cứu thứ ba được hình thành như sau: Làm thế nào để xây dựng khung MCGDM động trong môi trường thông tin ngôn ngữ dựa trên HA có khả năng xử lý sự thay đổi của phương án, tiêu chí, chuyên gia và dữ liệu đánh giá qua thời gian, đồng thời tích hợp dữ liệu lịch sử, xử lý dữ liệu khuyết thiếu và duy trì cơ chế suy giảm ngữ nghĩa nhất quán? Câu hỏi này được hiện thực hóa thành đóng góp mới thứ ba của luận án trong Chương 4.

Từ ba khoảng trống và ba câu hỏi nghiên cứu nêu trên, luận án triển khai ba định hướng nghiên cứu có quan hệ kế thừa và mở rộng thể hiện trong Hình 1.1. Thứ nhất, phát triển mô hình HA-TOPSIS nhằm xếp hạng phương án trên cơ sở thang đo ngôn ngữ, SQM, xác định PIS/NIS, đo khoảng cách ngữ nghĩa và tính hệ số gần trên cùng miền ngữ nghĩa. Thứ hai, phát triển mô hình EFAHP - 4-tuple - HA nhằm xác định trọng số của tiêu chí có kiểm soát tính nhất quán, đồng thời kết nhập đánh giá và xếp hạng phương án trên cùng khuôn khổ ngữ nghĩa. Thứ ba, phát triển khung phương pháp xây dựng mô hình L-FAHP - L-FTOPSIS động dựa trên HA nhằm xử lý thay đổi của tập phương án, bộ tiêu chí, hội đồng chuyên gia, dữ liệu đánh giá, dữ liệu khuyết thiếu và dữ liệu lịch sử theo cơ chế suy giảm ngữ nghĩa.

Như vậy, ba khoảng trống nghiên cứu, tương ứng với ba câu hỏi nghiên cứu, hiện thực ở ba đóng góp mới và thể hiện trong ba Chương phương pháp của luận án có quan hệ kế thừa và mở rộng, không chỉ về hình thức cấu trúc mà còn về logic khoa học. Chương 2 trả lời câu hỏi nghiên cứu thứ nhất bằng mô hình HA-TOPSIS; Chương 3 trả lời câu hỏi nghiên cứu thứ hai bằng mô hình EFAHP - 4-tuple - HA; Chương 4 trả lời câu hỏi nghiên cứu thứ ba bằng mô hình L-FAHP - L-FTOPSIS động dựa trên HA. Trong phạm vi luận án, các mô hình này không nhằm bao quát toàn bộ mọi vấn đề của MCGDM ngôn ngữ, mà tập trung góp phần phát triển một hướng tiếp cận có cơ sở hình thức cho biểu diễn, định lượng, kết nhập và xếp hạng phương án trong môi trường thông tin ngôn ngữ dựa trên HA.



Hình 1.1. Khung phương pháp MCGDM trong môi trường ngôn ngữ dựa trên HA.

1.6. Kết luận chương 1

Chương 1 đã hoàn thành vai trò của một chương tổng quan nghiên cứu và cơ sở lý thuyết theo đúng định hướng của luận án, nhưng giá trị của chương không dừng ở việc tập hợp và hệ thống hóa tài liệu. Quan trọng hơn, chương đã định vị rõ bài toán MCGDM trong môi trường thông tin ngôn ngữ như một đối tượng nghiên cứu có yêu cầu phương pháp luận riêng, trong đó dữ liệu đánh giá không chỉ mơ hồ và không chắc chắn, mà còn mang cấu trúc ngữ nghĩa cần được bảo toàn trong toàn bộ quá trình xác định trọng số, kết nhập thông tin và xếp hạng phương án. Từ cách đặt vấn đề đó, Chương 1 đã xác định HA không chỉ như một công cụ biểu diễn ngôn ngữ, mà như nền tảng hình thức cho việc xây dựng một khung MCGDM ngôn ngữ có tính nhất quán cao hơn về mặt ngữ nghĩa.

Về phương pháp luận, chương đã hệ thống hóa các nền tảng lý thuyết chủ yếu của lĩnh vực, từ các phương pháp MCDM/MCGDM truyền thống, các hướng tiếp cận dựa trên lý thuyết tập mờ, FAHP, FTOPSIS, các mô hình tính toán với từ, mô hình 2-tuple, mô hình 4-tuple ngữ nghĩa đến lý thuyết HA. Trên cơ sở đó, Chương 1 không chỉ chỉ ra ý nghĩa của các kết quả đã có, mà còn phân tích những giới hạn còn tồn tại khi xử lý trực tiếp thông tin ngôn ngữ, khi xác định trọng số của tiêu chí, khi bảo đảm tính liên thông ngữ nghĩa giữa các bước tính toán, và khi mở rộng bài toán sang bối cảnh động có tích hợp dữ liệu lịch sử. Như vậy, Chương 1 đã xác định ba khoảng trống nghiên cứu cốt lõi và chuyển hóa chúng thành ba hướng phát triển phương pháp luận của luận án.

Theo đó, Chương 2 phát triển mô hình HA-TOPSIS nhằm giải quyết khoảng trống về xếp hạng phương án trong MCGDM ngôn ngữ trên cơ sở định lượng ngữ nghĩa, đo khoảng cách ngữ nghĩa và xếp hạng phương án dựa trên thang đo ngôn ngữ được xây dựng theo hướng tiếp cận HA. Chương 3 phát triển mô hình EFAHP - 4-tuple - HA nhằm giải quyết khoảng trống về xác định trọng số của tiêu chí và tăng cường sự thống nhất ngữ nghĩa giữa giai đoạn xác định trọng số với giai đoạn đánh giá, xếp hạng phương án. Trên nền tảng đó, Chương 4 tiếp tục mở rộng sang mô hình L-FAHP - L-FTOPSIS động dựa trên HA thông qua các phép kết nhập ngôn ngữ và cơ chế suy giảm ngữ nghĩa để tích hợp dữ liệu lịch sử với đánh giá hiện tại.

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

Trong chương 2, luận án phát triển mô hình TOPSIS ngôn ngữ dựa trên HA (HA-TOPSIS) nhằm giải quyết bài toán MCGDM trong môi trường thông tin ngôn ngữ có tính bất định. Trên cơ sở giá trị định lượng ngữ nghĩa của các từ ngôn ngữ trong HA, mô hình đề xuất xác định PIS/NIS, tính khoảng cách ngữ nghĩa và hệ số gần để xếp hạng các phương án một cách nhất quán trên miền ngôn ngữ. Sau đó, mô hình được áp dụng cho bài toán lựa chọn nhà cung cấp nhằm kiểm chứng tính khả thi và hiệu quả của cách tiếp cận đề xuất.

Nội dung chính của chương 2 đã được công bố trong công trình [CT1].

2.1. Giới thiệu

Như đã phân tích ở Chương 1, khoảng trống nghiên cứu thứ nhất của luận án nằm ở việc còn thiếu một mô hình TOPSIS ngôn ngữ được phát triển nhất quán trên nền tảng HA để xử lý trực tiếp bài toán xếp hạng phương án trong MCGDM ngôn ngữ. Theo logic đó, Chương 2 sẽ phát triển mô hình thứ nhất của luận án, tập trung vào hướng tiếp cận xếp hạng phương án trong môi trường thông tin ngôn ngữ, qua đó cụ thể hóa hướng tiếp cận HA đối với một trong những thành phần quan trọng của bài toán MCGDM trước khi chuyển sang hướng xác định trọng số của tiêu chí ở Chương 3 và mở rộng sang bối cảnh động ở Chương 4.

Trong nhiều bài toán MCGDM, các chuyên gia thường đánh giá phương án bằng các từ ngôn ngữ thay vì bằng các giá trị số chính xác [62, 38]. Khi đó, bài toán của riêng Chương 2 được đặt ra là: với dữ liệu đầu vào gồm tập phương án, tập tiêu chí, tập chuyên gia, trọng số của các tiêu chí được gán bằng các từ ngôn ngữ, cùng ma trận đánh giá ngôn ngữ của các chuyên gia trên thang đo được xây dựng dựa trên HA, cần xây dựng một mô hình có khả năng tổng hợp và xếp hạng phương án sao cho đầu ra là thứ tự ưu tiên của các phương án vẫn phản ánh được nội dung ngữ nghĩa của thông tin đánh giá ban đầu [78, 93]. Vấn đề trọng tâm ở đây không chỉ là tính ra một hệ số gần theo quy trình TOPSIS [59], mà là phải thiết lập được một cơ chế xác định PIS/NIS, đo khoảng cách và suy ra thứ hạng trên một không gian biểu diễn ngôn ngữ có cơ sở hình thức rõ ràng [13]. TOPSIS do Hwang & Yoon đề xuất là một trong những phương pháp nền tảng của MCDM nhờ ý tưởng xếp hạng trực quan: phương án tốt là phương án gần với PIS

và xa với NIS [59, 111]. Tuy nhiên, TOPSIS truyền thống được xây dựng chủ yếu cho dữ liệu đầu vào là số rõ; vì vậy, khi áp dụng cho bài toán MCGDM trong môi trường thông tin ngôn ngữ, mô hình này bộc lộ giới hạn ngay từ lớp biểu diễn dữ liệu [21]. Các mở rộng như FTOPSIS đã góp phần xử lý thông tin bất định, bằng cách mô hình hóa các từ ngữ ngôn ngữ thông qua số mờ tam giác [21].

Tuy nhiên, các cách tiếp cận này vẫn chưa thật sự thỏa đáng đối với bài toán xử lý trực tiếp ngữ nghĩa. Dữ liệu đánh giá thường được gán cho các hàm thuộc, khiến cấu trúc ngữ nghĩa của từ bị chuyển hóa thành một biểu diễn số học phụ thuộc vào việc lựa chọn mô hình mờ và các tham số đi kèm, từ đó dẫn đến nguy cơ mất mát thông tin và giảm độ chính xác [49, 54]. Tiếp đó, việc xác định PIS/NIS trong nhiều mô hình FTOPSIS thực chất vẫn được thực hiện trên miền định lượng trung gian hơn là trên chính miền ngôn ngữ, nên chưa phản ánh đầy đủ trật tự ngữ nghĩa nội sinh của thang đo [63]. Tương tự, độ đo khoảng cách được dùng để so sánh phương án với hai giải pháp lý tưởng chủ yếu là khoảng cách giữa các biểu diễn số mờ thường không tương thích hoàn toàn với khoảng cách ngữ nghĩa giữa các từ đánh giá [63]. Hệ quả là hệ số gần và thứ hạng phương án có thể bị ảnh hưởng đáng kể bởi cơ chế biểu diễn và quá trình giải mờ thay vì xuất phát trực tiếp từ nội dung ngữ nghĩa của dữ liệu đầu vào [63, 111].

Xuất phát từ định hướng đó, Chương 2 tập trung phát triển mô hình HA-TOPSIS cho bài toán MCGDM trong môi trường thông tin ngôn ngữ, trong đó trọng tâm là xây dựng thang đo ngôn ngữ dựa trên HA, thiết lập cơ chế định lượng ngữ nghĩa cho các đánh giá ngôn ngữ, xác định PIS và NIS trên miền ngôn ngữ, đo khoảng cách ngữ nghĩa từ các phương án tới PIS/NIS, và từ đó tính hệ số gần để xếp hạng phương án. Trong toàn bộ cấu trúc luận án, mô hình này đảm nhiệm chức năng giải quyết bài toán xếp hạng phương án ở bối cảnh tĩnh, tức là một thành phần phương pháp luận cơ bản nhưng chưa bao quát bài toán xác định trọng số của tiêu chí và chưa xét đến yếu tố động theo thời gian. Chính vì vậy, kết quả của Chương 2 vừa có tính độc lập về phương pháp, vừa tạo nền trực tiếp cho Chương 3, nơi hướng tiếp cận xác định trọng số của tiêu chí được phát triển trong mô hình tích hợp EFAHP - 4-tuple - HA, và cho Chương 4, nơi mô hình MCGDM được mở rộng sang bối cảnh động dựa trên HA.

Trên cơ sở đó, các phần tiếp theo của chương sẽ lần lượt trình bày mô hình HA-TOPSIS và quy trình giải bài toán, ví dụ thực nghiệm minh họa, phân tích độ nhạy và so sánh kết quả, qua đó làm rõ cơ sở hình thức, khả năng tính toán và giá trị ứng dụng của hướng tiếp cận được đề xuất.

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

(i) Dữ liệu đầu vào:

Cho $\bar{A} = \{A_1, A_2, \dots, A_m\}$ ($m \geq 2$) là tập hữu hạn các phương án, $\bar{C} = \{C_1, C_2, \dots, C_n\}$ ($n \geq 2$) là tập hữu hạn các tiêu chí, và $\bar{D} = \{D_1, D_2, \dots, D_h\}$ ($h \geq 2$) là tập hữu hạn người ra quyết định (chuyên gia).

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. Trong đó, $L_{(j)k}(\bar{A})$ thang đo ngôn ngữ dùng để đánh giá các phương án $A_i \in \bar{A}$ theo tiêu chí $C_j \in \bar{C}$ và $L_k(\bar{C})$ thang đo ngôn ngữ để 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 $\leq k$; $i = 1, \dots, m$; $j = 1, 2, \dots, n$).

Đối với mỗi chuyên gia $D^e \in \bar{D}$ ($e = 1, \dots, h$), thực hiện:

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

$$A^e = (a_{ij}^e)_{m \times n} = \begin{pmatrix} a_{11}^e & a_{12}^e & \dots & a_{1n}^e \\ a_{21}^e & a_{22}^e & \dots & a_{2n}^e \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1}^e & a_{m2}^e & \dots & a_{mn}^e \end{pmatrix} \quad (2.1)$$

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 $i = 1, \dots, m$; $j = 1, \dots, n$; $e = 1, \dots, h$.

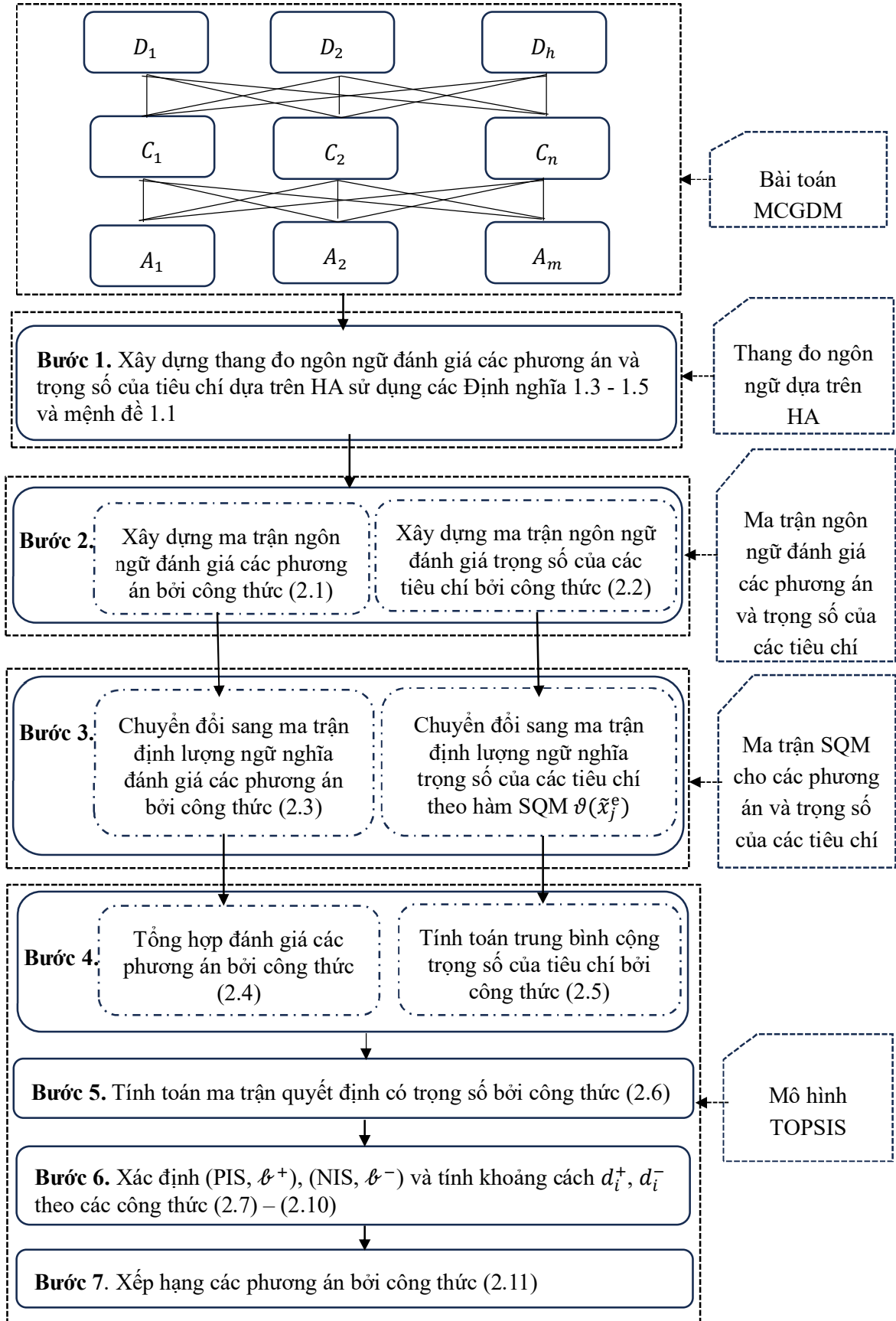
(2) Xây dựng ma trận đánh giá trọng số của các tiêu chí theo công thức (2.2):

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

Trong đó, $\tilde{x}_j^e \in L_k(\bar{C})$ ($j = 1, \dots, n$; $e = 1, \dots, h$) 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 $\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_n)^T$ của các tiêu chí, thỏa các điều kiện: $0 \leq \tilde{w}_j \leq 1$ và $\sum_{j=1}^n \tilde{w}_j = 1$.

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

Hình 2.1 trình bày quy trình ra quyết định của mô hình TOPSIS ngôn ngữ dựa trên HA (HA-TOPSIS) gồm có các bước sau:



Hình 2.1. Quy trình MCGDM của mô hình HA-TOPSIS

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.

Để đánh giá các phương án theo từng tiêu chí trong môi trường ngôn ngữ dựa trên HA, hội đồng chuyên gia thống nhất thiết lập cấu trúc HA cho từng tiêu chí $C_j \in \bar{C}$.

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 . Cụ thể, với mỗi tiêu chí $C_j \in \bar{C}$, xét một HA tuyến tính:

$$\mathcal{AX}_{(j)} = (X, G, H, \leq)$$

Trong đó: X là tập các từ ngôn ngữ; $G = \{g^-, g^+\}$ là tập các phần tử sinh, (ví dụ, $g^- = \text{bad}, g^+ = \text{good}$); $H = H^- \cup H^+$ là tập các gia tử, với $H^- = \{h_{-q}, h_{-q+1}, \dots, h_{-1}\}$ là tập gồm q gia tử âm và $H^+ = \{h_1, h_2, \dots, h_p\}$ là tập gồm p gia tử dương, (ví dụ, $H^- = \text{Little}, H^+ = \text{Very}$); $\leq$ là quan hệ thứ tự ngữ nghĩa cảm sinh trên X .

Các tham số tính mờ của HA được xác định như sau: Độ đo mờ của các phần tử sinh được ký hiệu bởi $f_m(g^-)$ và $f_m(g^+)$, với ràng buộc $f_m(g^+) = 1 - f_m(g^-)$. Độ đo tính mờ của các gia tử được ký hiệu bởi $\mu(h_{-q}), \mu(h_{-q+1}), \dots, \mu(h_{-1}), \mu(h_1), \dots, \mu(h_p)$, thỏa mãn điều kiện chuẩn hóa: $\mu(h_p) = 1 - \sum_{i=-q}^{p-1} \mu(h_i)$.

Trên cơ sở bộ tham số ngữ nghĩa của HA $(HA, f_m(g^+), \mu(h), k)$, 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, lần lượt được tạo bằng cách kết hợp các phần tử sinh với các gia tử hình thành cơ chế tự động sinh ra tập các từ ngôn ngữ bằng: $X_{(j)k} = \{x \in X_j : |x| \leq k\}$ và $X_k = \{x \in X : |x| \leq k\}$ tương ứng trên hai thang đo $L_{(j)k}(\bar{A})$ và $L_k(\bar{C})$, bảo toàn thứ tự ngữ nghĩa và phù hợp với ngữ cảnh đánh giá của các tiêu chí.

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 ($e = 1, 2, \dots, h$) sử dụng các từ ngôn ngữ $x \in X_{(j)k}$ trên thang đo ngôn ngữ $L_{(j)k}(\bar{A})$ (được tạo trong Bước 1) 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 $L_k(\bar{C})$ (được tạo trong Bước 1) 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. Cụ thể:

Với mỗi phần tử $a_{ij}^e \in X_{(j)k}$ trong ma trận A^e , giá trị định lượng ngữ nghĩa của nó được xác định bởi $\bar{a}_{ij}^e = \vartheta(a_{ij}^e)$, với $\vartheta: X_j \rightarrow [0, 1]$, trong đó $\vartheta(x)$ là giá trị định lượng ngữ nghĩa số của từ ngôn ngữ $x \in X_{(j)k}$ trên thang đo $L_{(j)k}(\bar{A})$. Từ đó ma trận định lượng ngữ nghĩa $\tilde{A}^e$ được xác định theo công thức (2.3) sau:

$$\tilde{A}^e = (\bar{a}_{ij}^e)_{m \times n} = \begin{pmatrix} \bar{a}_{11}^e & \bar{a}_{12}^e & \cdots & \bar{a}_{1n}^e \\ \bar{a}_{21}^e & \bar{a}_{22}^e & \cdots & \bar{a}_{2n}^e \\ \vdots & \vdots & \ddots & \vdots \\ \bar{a}_{m1}^e & \bar{a}_{m2}^e & \cdots & \bar{a}_{mn}^e \end{pmatrix}. \quad (2.3)$$

Trong đó, $i = 1, \dots, m; j = 1, \dots, n; e = 1, \dots, h$.

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.

Để thu được đánh giá tập thể của các phương án, các ma trận $\tilde{A}^e$ được tổng hợp theo nhóm thành một ma trận quyết định nhóm thống nhất. Với giả thiết các chuyên gia có vai trò như nhau (trọng số chuyên gia bằng nhau), ma trận đánh giá tổng hợp nhóm, ký hiệu $A\tilde{A}$ được xây dựng bằng cách tính trung bình cộng đánh giá các phương án từ tất cả các chuyên gia, 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 bằng cách tính trung bình cộng 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 $\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_n)^T$, trong đó, $\tilde{w}_j = \frac{\tilde{u}_j}{\sum_{k=1}^n \tilde{u}_k}$, $j = 1, \dots, n$, với $0 \leq \tilde{w}_j \leq 1$ 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 quyết định nhóm $A\tilde{A}$ và $\tilde{w}_j$ là véc-tơ trọng số của tiêu chí $C_j \in \tilde{C}$.

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

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

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

và

$$\mathcal{B}^- = \{\mathcal{B}_1^-, \mathcal{B}_2^-, \dots, \mathcal{B}_n^-\}. \quad (2.8)$$

+ Đối với tiêu chí lợi ích (benefit criteria), PIS là giá trị lớn nhất: $\mathcal{B}_j^+ = \sum_{i=1}^m \max \mathcal{B}_{ij}$, và NIS là giá trị nhỏ nhất: $\mathcal{B}_j^- = \sum_{i=1}^m \min \mathcal{B}_{ij}$.

+ Đối với tiêu chí chi phí (cost criteria): $\mathcal{B}_j^+ = \sum_{i=1}^m \min \mathcal{B}_{ij}$, và NIS là giá trị nhỏ nhất: $\mathcal{B}_j^- = \sum_{i=1}^m \max \mathcal{B}_{ij}$.

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

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

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

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

Hệ số gần, CC_i , ($i = 1, \dots, m$) biểu thị mức độ tối ưu của phương án A_i càng gần (PIS, $\mathcal{B}^+$) và càng xa (NIS, $\mathcal{B}^-$) thì càng tốt, được tính theo công thức (2.11):

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

Dựa trên giá trị của CC_i 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.

Giả mã của HA-TOPSIS được trình bày trong thuật toán (2.1) dưới đây:

Thuật toán 2.1. Thuật toán mô hình ra quyết định HA-TOPSIS.

Algorithm HA-TOPSIS ($\bar{A}, \bar{C}, \bar{D}, HA, fm(g^-), \mu(h), k$)

Input:

- Tập các phương án $\bar{A} = \{A_1, A_2, \dots, A_m\}$, (với $m \geq 2$);
- Tập các tiêu chí $\bar{C} = \{C_1, C_2, \dots, C_n\}$, (với $n \geq 2$);
- Tập người ra quyết định $\bar{D} = \{D_1, D_2, \dots, D_h\}$, (với $h \geq 2$);
- Bộ tham số $(HA_C, fm_C(g^-), \mu_C(h), k_C)$ // xây dựng thang đo $L_k(\bar{C})$;
- Bộ tham số $(HA_A, fm_A(g^-), \mu_A(h), k_A)$ // xây dựng thang đo $L_{(j)k}(\bar{A})$;
- Ma trận quyết định ngôn ngữ A^e đánh giá xếp hạng các phương án; Ma trận P^e đánh giá trọng số quan trọng của các tiêu chí.

Output: Thứ hạng của các phương án $A_i \in \bar{A}$.

Begin

1. for each $C_j \in \bar{C}$ do
2. *Thiết lập cấu trúc HA tuyến tính; xác định $fm(g^+)$, $\mu(h)$, chọn k ; xây dựng thang đo ngôn ngữ $L_{(j)k}(\bar{A})$, $L_k(\bar{C})$ theo Định nghĩa 1.4- 1.8 và Mệnh đề 1.1*
3. end
4. for each $e = 1, \dots, h$ do
5. *Mỗi DM xây dựng ma trận ra đánh giá ngôn ngữ $A^e = (a_{ij}^e)_{m \times n}$ theo từng tiêu chí bởi công thức (2.1);*
6. end
7. for each $e = 1, \dots, h$ do
8. *Mỗi DM xây dựng ma trận ra quyết định $P^e = (\tilde{x}_j^e)_{n \times h}$ đánh giá mức độ quan trọng tương đối của tiêu chí $C_j \in \bar{C}$ bởi công thức (2.2);*
9. end
10. for each $e = 1, \dots, h$ do
11. for each $i = 1, \dots, m$ do

```

12.      for each  $j = 1, \dots, n$  do
13.          Định lượng ngữ nghĩa bằng hàm SQM của HA theo công thức (2.3)
14.      end
15.  end
16. end
17. for each  $i = 1, \dots, m$  do
18.     for each  $j = 1, \dots, n$  do
19.         Tổng hợp ma trận đánh giá các phương án theo công thức (2.4)
20.     end
21. end
22. for each  $j = 1, \dots, n$  do
23.     Trọng số của tiêu chí  $C_j \in \bar{C}$  được tổng hợp bởi công thức (2.5)
24. end
25. for each  $i = 1, \dots, m$  do
26.     for each  $j = 1, \dots, n$  do
27.          $\gamma_{ij} = u_{ij} \times \tilde{w}_j$  // Ma trận quyết định có trọng số bởi công thức (2.6)
28.     end
29. end
30. Xác định  $(PIS, \&^+)$  và  $(NIS, \&^-)$  theo các công thức (2.7) và (2.8);
31. for each  $i = 1, \dots, m$  do
32.     Tính khoảng cách  $d_i^+, d_i^-$  từ  $A_i$  đến  $\&^+, \&^-$  theo công thức (2.9) và (2.10);
33. end
34. for each  $i = 1, \dots, m$  do
35.     Tính hệ số gần,  $CC_i$  của  $A_i$  theo công thức (2.11);
36. end
37. return Rank  $(A_i \in \bar{A})$  // thứ hạng của các phương án

```

End

Bảng 2.1. Độ phức tạp của thuật toán thực hiện các công thức trong HA-TOPSIS

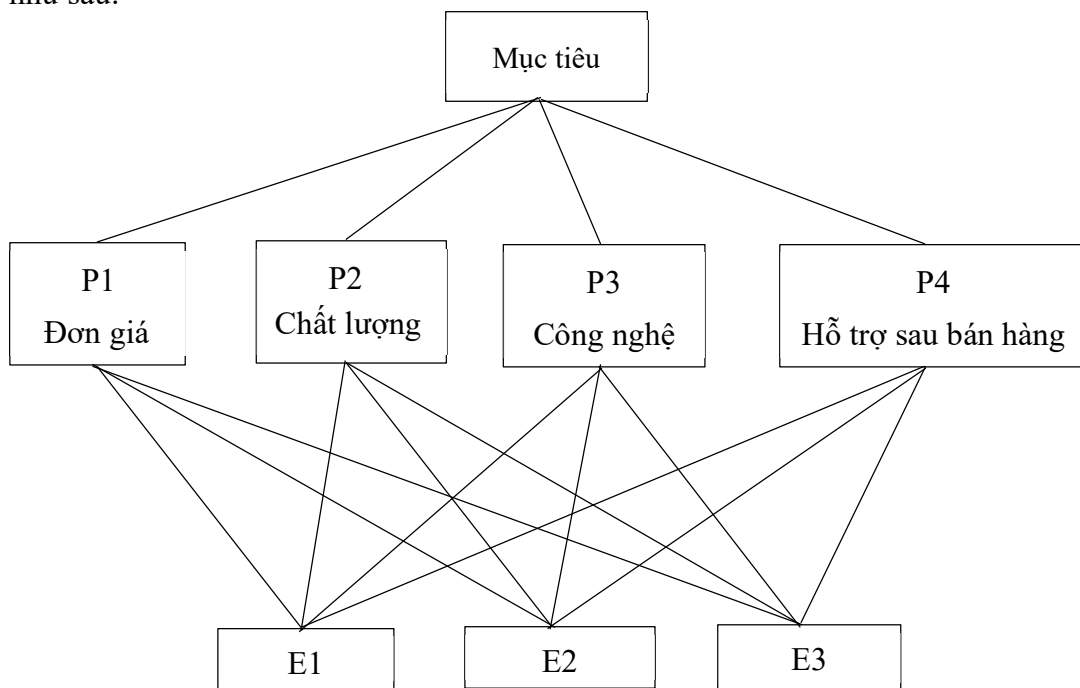
STT	Công thức	Nội dung tính toán	Độ phức tạp
1	(2.1)	Xây dựng ma trận đánh giá ngôn ngữ của một chuyên gia trên m phương án và n tiêu chí	$O(m \times n)$ cho mỗi chuyên gia; $O(h \times m \times n)$ cho toàn bộ h chuyên gia.
2	(2.2)	Xây dựng ma trận đánh giá mức độ quan trọng của n tiêu chí cho một chuyên gia	$O(n)$ cho mỗi chuyên gia; $O(h \times n)$ cho toàn bộ h chuyên gia.
3	(2.3)	Định lượng ngữ nghĩa ma trận đánh giá các phương án bằng hàm SQM	$O(m \times n)$ cho mỗi chuyên gia; $O(h \times m \times n)$ cho toàn bộ h chuyên gia.
4	(2.4)	Tổng hợp ma trận quyết định nhóm từ h ma trận đánh giá cá thể	$O(h \times m \times n)$
5	(2.5)	Tổng hợp trọng số của tiêu chí từ h chuyên gia và tính véc-tơ trọng số	$O(h \times n)$
6	(2.6)	Tính ma trận quyết định có trọng số	$O(m \times n)$
7	(2.7)	Xác định PIS trên n tiêu chí	$O(m \times n)$
8	(2.8)	Xác định NIS trên n tiêu chí	$O(m \times n)$
9	(2.9)	Tính khoảng cách từ mỗi A_i đến PIS	$O(m \times n)$
10	(2.10)	Tính khoảng cách từ mỗi A_i đến NIS	$O(m \times n)$
11	(2.11)	Tính hệ số gần cho m phương án	$O(m)$

Từ Bảng 2.1 có thể thấy các bước có chi phí lớn nhất của HA-TOPSIS là các thao tác xử lý ma trận đánh giá của h chuyên gia trên m phương án và n tiêu chí, bao gồm xây dựng ma trận đánh giá, định lượng ngữ nghĩa bằng hàm SQM và tổng hợp ma trận quyết định nhóm, với độ phức tạp $O(h \times m \times n)$. Các bước còn lại có độ phức tạp thấp hơn hoặc cùng bậc theo m và n , gồm $O(h \times n)$, $O(m \times n)$ và $O(m)$. Do đó, độ phức tạp thời gian tổng thể của quy trình HA-TOPSIS có thể viết là $O(h \times m \times n) + O(h \times n) + O(m \times n) + O(m)$, và được rút gọn thành $O(h \times m \times n)$. Như vậy, chi phí tính toán của mô hình tăng tuyến tính theo số chuyên gia, số phương án và số tiêu chí; điều này cho thấy mô hình có tính khả thi đối với các bài toán MCGDM ngôn ngữ quy mô vừa. Cần lưu ý rằng độ phức tạp trên chỉ phản ánh chi phí xử lý thuật toán sau khi dữ liệu đánh giá đã được cung cấp, không bao gồm chi phí tổ chức khảo sát và thu thập ý kiến chuyên gia.

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).

Cấu trúc phân cấp của bài toán MCGDM được thể hiện trong Hình 2.2. Sử dụng mô hình HA-TOPSIS đề xuất để giải quyết bài toán này và quy trình tính toán được tóm tắt như sau:



Hình 2.2. Cấu trúc phân cấp của bài toán lựa chọn nhà cung cấp

Bước 1: Xây dựng các thang đo ngôn ngữ cho bốn tiêu chí nêu trên và trọng số của các tiêu chí.

Hội đồng chuyên gia thống nhất xây dựng cấu trúc HA tuyến tính, các tham số mờ và độ dài tối đa của từ ngôn ngữ $k = 2$ phù hợp với ngữ cảnh của từng tiêu chí và trọng số quan trọng của các tiêu chí, như tóm tắt trong Bảng 2.2, dưới đây:

Bảng 2.2. Thiết lập cấu trúc HA và các tham số mờ xây dựng thang đo ngôn ngữ cho các tiêu chí và trọng số của tiêu chí

Tiêu chí	HA	Phần tử sinh		Các từ hằng			Gia tử		Độ đo tính mờ	
		g^-	g^+	0	W	1	h_{-1}	h_1	$fm(g^+)$	$\mu(h_1)$
P1	$\mathcal{A}X_1$	Cheap (C)	Expensive (E)	Extremely Cheap	Medium (M)	Extremely Expensive	Little (L)	Very (V)	0.5	0.5
P2	$\mathcal{A}X_2$	Bad (B)	Good (G)	Extremely Bad	Medium (M)	Extremely Good	Little (L)	Very (V)	0.525	0.548
P3	$\mathcal{A}X_3$	Outdated (O)	Advanced (A)	Extremely Outdated	Medium (M)	Extremely Advanced	Little (L)	Very (V)	0.572	0.55
P4	$\mathcal{A}X_4$	Dissatisfied (D)	Satisfied (S)	Extremely Dissatisfied	Medium (M)	Extremely Satisfied	Little (L)	Very (V)	0.374	0.473
$\tilde{W}$	$\mathcal{A}X_5$	Unimportant (U)	Important (I)	Extremely Unimportant	Medium (M)	Extremely Important	Little (L)	Very (V)	0.5	0.5

Trên cơ sở bộ tham số ngữ nghĩa của HA ($HA, fm(g^+), \mu(h_1), 2$) được chọn trong Bảng 2.2, HA sẽ tự động nội sinh tập gồm chín từ ngôn ngữ theo từng tiêu chí $C_j \in \bar{C}$ và trọng số của tiêu chí ($\tilde{W}$) tương ứng với thang đo $L_{(j)k}(\bar{A})$ và $L_k(\bar{C})$.

Tiêu chí P1, với bộ tham số ngữ nghĩa HA được chọn trên Bảng 2.2, HA sẽ tự động nội sinh ra thang đo ngôn ngữ $L_{(1)2}(\bar{A})$ gồm chín từ ngôn ngữ dưới đây:

$L_{(1)2} = \{\text{Extremely Cheap} < \text{Very Cheap} < \text{Cheap} < \text{Little Cheap} < \text{Medium} < \text{Little Expensive} < \text{Expensive} < \text{Very Expensive} < \text{Extremely Expensive}\}$

Thang đo $L_{(1)2}(\bar{A})$ được viết tắt $L_{(1)2}(\bar{A}) = \{EC, VC, C, LC, M, LE, E, VE, EE\}$

Tương tự, các thang đo ngôn ngữ còn lại của ba tiêu chí (P2, P3, P4) và trọng số của các tiêu chí ($\tilde{W}$) lần lượt được tạo và viết tắt như sau:

$$L_{(2)2}(\bar{A}) = \{EB, VB, B, LB, M, LG, G, VG, EG\},$$

$$L_{(3)2}(\bar{A}) = \{EO, VO, O, LO, M, LA, A, VA, EA\},$$

$$L_{(4)2}(\bar{A}) = \{ED, VD, D, LD, M, LS, S, VS, ES\},$$

$$L_2(\bar{C}) = \{EU, VU, U, LU, M, LI, I, VI, EI\}.$$

Bước 2. Xây dựng ma trận đánh giá các nhà cung cấp theo bốn tiêu chí và trọng số của các tiêu chí

Mỗi chuyên gia D^e sử dụng các từ ngôn ngữ trong bốn thang đo ngôn ngữ ($L_{(1)2}(\bar{A})$, $L_{(2)2}(\bar{A})$, $L_{(3)2}(\bar{A})$, $L_{(4)2}(\bar{A})$) để đánh giá ba nhà cung cấp (E1, E2, E3) theo bốn tiêu chí (P1, P2, P3, P4) bởi công thức (2.1). Kết quả thể hiện trong các Bảng (2.3 - 2.6):

Bảng 2.3. Đánh giá các nhà cung cấp theo tiêu chí P1

Nhà cung cấp	Chuyên gia đánh giá				
	D ₁	D ₂	D ₃	D ₄	D ₅
E1	LE	E	LE	E	LE
E2	M	LE	LC	M	M
E3	E	LE	E	E	LE

Bảng 2.4. Đánh giá các nhà cung cấp theo tiêu chí P2

Nhà cung cấp	Chuyên gia đánh giá				
	D ₁	D ₂	D ₃	D ₄	D ₅
E1	G	VG	EG	VG	G
E2	G	M	LG	VG	LG
E3	VG	LG	VG	EG	EG

Bảng 2.5. Đánh giá các nhà cung cấp theo tiêu chí P3

Nhà cung cấp	Chuyên gia đánh giá				
	D ₁	D ₂	D ₃	D ₄	D ₅
E1	VA	VA	EA	EA	A
E2	A	A	LA	M	LA
E3	LA	M	A	VA	VH

Bảng 2.6. Đánh giá các nhà cung cấp theo tiêu chí P4

Nhà cung cấp	Chuyên gia đánh giá				
	D ₁	D ₂	D ₃	D ₄	D ₅
E1	S	VS	S	ES	VS
E2	LD	M	LS	M	LS
E3	S	S	LS	VS	VS

Tương tự, mỗi chuyên gia D^e sử dụng các từ ngôn ngữ trên thang đo ngôn ngữ $L_2(\bar{C})$ để đánh giá trọng số của các tiêu chí bởi công thức (2.2). Kết quả thể hiện trong Bảng (2.7) dưới đây:

Bảng 2.7. Đánh giá trọng số quan trọng của các tiêu chí

Các tiêu chí	Chuyên gia đánh giá				
	D ₁	D ₂	D ₃	D ₄	D ₅
P1	VI	EI	VI	EI	VI
P2	VI	I	LI	I	VI
P3	I	VI	I	EI	VI
P4	M	LU	LI	M	I

Bước 3: Tính toán giá trị định lượng ngữ nghĩa số của từ ngôn ngữ

Sử dụng công thức (2.3) để chuyển đổi từ ngôn ngữ sang giá trị định lượng nghĩa của các từ và các công thức (2.4) - (2.5) để tổng hợp đánh giá nhà cung cấp của các chuyên gia dựa trên bốn tiêu chí, trọng số quan trọng của các tiêu chí được thể hiện trong các Bảng 2.8 và 2.9 dưới đây:

Bảng 2.8. Trung bình đánh giá của các nhà cung cấp theo các tiêu chí

Tiêu chí	Nhà cung cấp	Chuyên gia đánh giá					Trung bình
		D ₁	D ₂	D ₃	D ₄	D ₅	
P1	E1	0.625	0.750	0.625	0.750	0.625	0.675
	E2	0.500	0.625	0.375	0.500	0.500	0.500
	E3	0.750	0.625	0.750	0.750	0.625	0.700
P2	E1	0.712	0.842	1.000	0.842	0.712	0.822
	E2	0.712	0.475	0.605	0.842	0.605	0.648
	E3	0.842	0.605	0.842	1.000	1.000	0.858
P3	E1	0.827	0.827	1.000	1.000	0.685	0.868
	E2	0.685	0.685	0.570	0.428	0.570	0.588
	E3	0.570	0.428	0.685	0.827	1.000	0.702
P4	E1	0.823	0.916	0.823	1.000	0.916	0.896
	E2	0.470	0.626	0.719	0.626	0.719	0.632
	E3	0.823	0.823	0.719	0.916	0.916	0.840

Bảng 2.9. Véc-tơ trọng số của các tiêu chí

Tiêu chí	Chuyên gia đánh giá					Trọng số
	D ₁	D ₂	D ₃	D ₄	D ₅	
P1	0.875	1	0.875	1	0.875	0.298
P2	0.875	0.75	0.625	0.75	0.875	0.250
P3	0.75	0.875	0.75	1	0.875	0.274
P4	0.5	0.375	0.625	0.5	0.75	0.177

Bước 4. Tính ma trận quyết định có trọng số

Sử dụng công thức (2.6) để xây dựng ma trận quyết định định lượng ngữ nghĩa có trọng số như trình bày trong Bảng 2.10 dưới đây:

Bảng 2.10. Ma trận quyết định có trọng số

Nhà cung cấp/ Tiêu chí	P1	P2	P3	P4
E1	0.201	0.205	0.238	0.159
E2	0.149	0.162	0.161	0.112
E3	0.209	0.214	0.192	0.149

Bước 5. Tính ℓ^+ , ℓ^- , d^+ và d^-

Sử dụng công thức (2.7) và (2.8), xác định PIS, ℓ^+ và NIS, ℓ^- . Sau đó, tính khoảng cách d_i^+ và d_i^- từ nhà cung cấp thứ i đến PIS ℓ^+ và NIS ℓ^- bằng các công thức (2.9) và (2.10), như trình bày trong Bảng 2.11.

Bảng 2.11. 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

Bước 6. Tính CC_i và xếp hạng nhà cung cấp

Hệ số gần giữa các nhà cung cấp được tính theo công thức (2.11). Dựa trên giá trị CC_i , tập các nhà cung cấp được xếp hạng theo thứ tự, và nhà cung cấp E1 là phương án tối ưu nhất. Kết quả được trình bày trong Bảng 2.12, dưới đây:

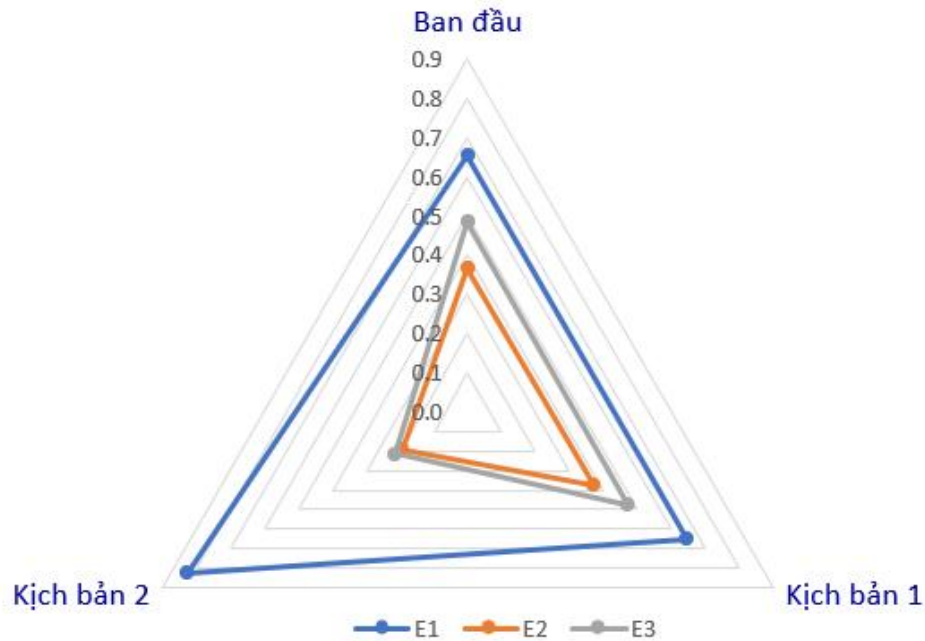
Bảng 2.12. Hệ số gần và thứ hạng nhà cung cấp

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

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 nhằm kiểm chứng tính ổn định và mức độ tin cậy của kết quả xếp hạng trong mô hình HA-TOPSIS, thông qua việc khảo sát sự biến động của kết quả đầu ra khi điều chỉnh các tham số đầu vào. Khác với cách tiếp cận thay đổi trọng số một cách cơ học, nghiên cứu này điều chỉnh trực tiếp các tham số độ đo tính mờ của các phần tử sinh $f_{\eta}(U)$, $f_{\eta}(I)$ (*Unimportant/ Important*) và các gia tử $\mu(L)$, $\mu(V)$ (*Little/ Very*) lệch hẳn về hai phía trên thang đo ngôn ngữ $L_2(\bar{C})$ dùng để xác định trọng số của tiêu chí, theo hai kịch bản đặc trưng âm - dương của miền ngữ nghĩa HA như sau:



Hình 2.3. Kết quả phân tích độ nhạy mô hình TOPSIS ngôn ngữ dựa trên HA

Trong kịch bản 1: Với $f_{\eta}(U) = 0.4$, $f_{\eta}(I) = 0.6$, $\mu(L) = 0.6$, $\mu(V) = 0.4$, ngữ nghĩa có xu hướng lệch hẳn về phía các mức quan trọng cao hơn do phần tử sinh (Important) có mức mờ lớn hơn và các gia tử đều có tác động tương đối mạnh, dẫn đến việc tăng khả năng phân giải giữa các mức trọng số.

Ngược lại, ở kịch bản 2: Với $f_{\eta}(U) = 0.6$, $f_{\eta}(I) = 0.4$, $\mu(L) = 0.4$, $\mu(V) = 0.6$, độ mờ lệch hẳn về phía (Unimportant) và tác động của gia tử (Little) giảm, làm cho thang ngôn ngữ mềm hơn ở vùng trọng số thấp - trung bình, từ đó có thể làm thu hẹp chênh lệch trọng số giữa các tiêu chí và giảm biên độ khác biệt giữa các nhà cung cấp.

Bảng 2.13. Trọng số của các tiêu chí của các kịch bản

Tiêu chí	Ban đầu	Kịch bản 1	Kịch bản 2
P1	0.298	0.308	0.288
P2	0.250	0.253	0.248
P3	0.274	0.283	0.267
P4	0.177	0.157	0.196

Kết quả ở Bảng 2.13 cho thấy việc thay đổi các tham số ngữ nghĩa chỉ làm trọng số các tiêu chí dịch chuyển trong biên độ hợp lý. Cụ thể:

Trong kịch bản 1, P1 và P3 tăng nhẹ (P1: 0.298 $\rightarrow$ 0.308; P3: 0.274 $\rightarrow$ 0.283), trong khi P4 giảm (0.177 $\rightarrow$ 0.157); ngược lại ở Kịch bản 2, P4 tăng đáng kể (0.177 $\rightarrow$ 0.196) còn P1 và P3 giảm nhẹ (P1: 0.298 $\rightarrow$ 0.288; P3: 0.274 $\rightarrow$ 0.267). Sự thay đổi này phản ánh đúng tác động của việc điều chỉnh cấu trúc ngữ nghĩa trên thang đo, nhưng không làm thay đổi thứ hạng tổng thể giữa các tiêu chí.

Bảng 2.14. 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

Tương ứng với đó, Bảng 2.14 và Hình 2.3 cho thấy hệ số gần của các nhà cung cấp có biến động về giá trị, nhưng thứ hạng cuối cùng vẫn được bảo toàn trong cả ba nhà cung cấp. Cụ thể, E1 luôn đạt giá trị cao nhất (ban đầu 0.654; kịch bản 1: 0.647; kịch bản 2: 0.825), E3 đứng thứ hai (tương ứng 0.485; 0.474; 0.217) và E2 thấp nhất (0.364; 0.372; 0.192). Như vậy, kết quả trong Bảng 2.15 cho thấy thứ hạng $E1 > E3 > E2$ được duy trì ổn định ở cả ba trường hợp ban đầu, kịch bản 1 và kịch bản 2.

Bảng 2.15. Thứ hạng của nhà cung cấp thu được trong các kịch bản

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

Kết quả này cho thấy mô hình HA-TOPSIS có độ ổn định cao trước nhiễu tham số ngữ nghĩa, trong đó các thay đổi của $fm(\cdot)$ và $\mu(\cdot)$ chủ yếu ảnh hưởng đến biên độ định lượng của trọng số và hệ số gần, nhưng không làm thay đổi thứ hạng của ba nhà cung cấp. Đây là bằng chứng quan trọng cho thấy việc xây dựng thang đo ngôn ngữ dựa trên HA và sử dụng SQM để ánh xạ mỗi từ ngôn ngữ với một giá trị trong $[0, 1]$ theo cách 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.

2.4.2. So sánh kết quả và thảo luận

Để đánh giá kết quả của mô hình HA-TOPSIS [CT1] trong phạm vi bài toán lựa chọn nhà cung cấp, mục này đối sánh với FTOPSIS [21] và mô hình 4-tuple ngữ nghĩa [54] trên các phương diện: thứ hạng phương án, cơ chế hình thành kết quả, vai trò của trọng số của tiêu chí và độ ổn định của thứ hạng trước biến động tham số ngữ nghĩa. Do ba mô hình không sử dụng cùng một đại lượng đầu ra, HA-TOPSIS và FTOPSIS sử dụng hệ số gần, còn mô hình 4-tuple ngữ nghĩa sử dụng tổng điểm, các giá trị số trong Bảng 2.16 không được xem là cùng một thang đo để so sánh trực tiếp về độ lớn. Vì vậy, phân tích tập trung vào thứ hạng, ý nghĩa nội tại của kết quả trong từng mô hình và mức độ phù hợp của cơ chế tính toán với dữ liệu ngôn ngữ.

Bảng 2.16. So sánh kết quả xếp hạng giữa HA-TOPSIS, FTOPSIS và 4-tuple ngữ nghĩa

Phương án	HA-TOPSIS [CT1]		FTOPSIS [21]		4-tuple ngữ nghĩa [54]	
	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 tại Bảng 2.16 cho thấy cả ba mô hình đều xếp E1 ở vị trí thứ nhất. HA-TOPSIS và mô hình 4-tuple ngữ nghĩa cùng cho thứ hạng $E1 > E3 > E2$, trong khi FTOPSIS cho $E1 > E2 > E3$. Như vậy, kết quả lựa chọn E1 là nhất quán trên bộ dữ liệu đang xét, còn khác biệt chỉ xuất hiện ở thứ hạng của E2 và E3. Sự trùng khớp giữa HA-TOPSIS và 4-tuple ngữ nghĩa phản ánh sự tương đồng về kết quả xếp hạng trong ví dụ này, nhưng không đồng nghĩa với hai mô hình tương đương về cơ chế tính toán hoặc giá trị đầu ra. Tương tự, sự khác biệt với FTOPSIS cho thấy các cơ chế biểu diễn và lượng hóa khác nhau có thể dẫn đến khác biệt cục bộ đối với các phương án tầm trung nơi có kết quả tiệm cận nhau.

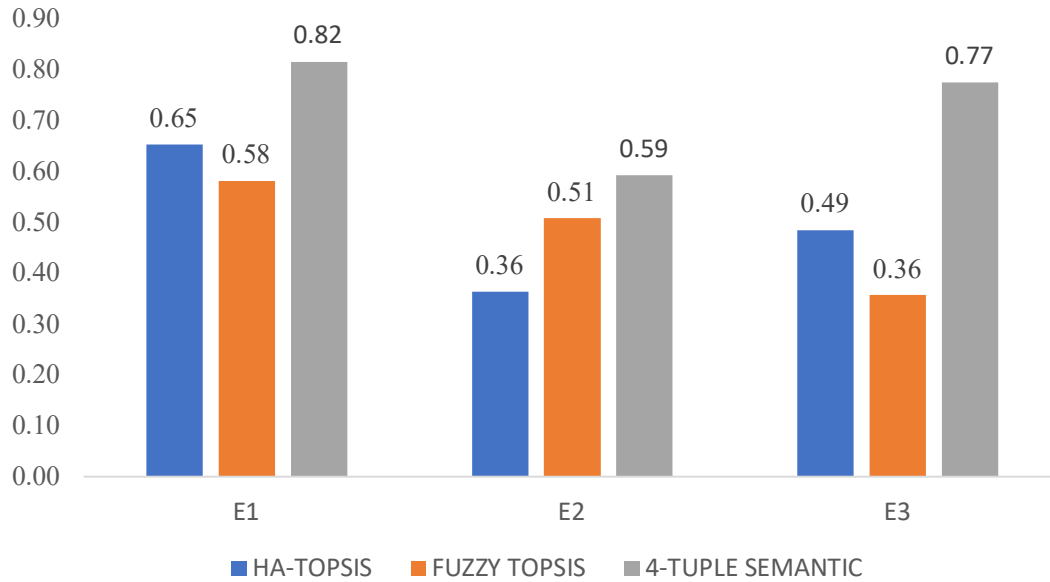
Mức độ tương đồng thứ hạng giữ các mô hình đối sánh trong Bảng 2.16 được mô tả bằng hệ số tương quan thứ hạng Spearman [100] và mức độ đồng thuận Kendall [67] khi không có đồng hạng. Dựa trên Bảng 2.16 các hệ số tương đồng thứ hạng của ba phương án được tính và tổng hợp trong Bảng 2.17 dưới đây:

Bảng 2.17. Hệ số tương quan thứ hạng và mức độ đồng thuận giữa các mô hình đối sánh

Cặp mô hình đối sánh	Hệ số Spearman (ρ)	Hệ số Kendall (κ)
HA-TOPSIS [CT1] và FTOPSIS [21]	0.50	0.33
HA-TOPSIS [CT1] và 4-tuple ngữ nghĩa [54]	1.00	1.00
FTOPSIS [21] và 4-tuple ngữ nghĩa [54]	0.50	0.33

Kết quả các mô hình HA-TOPSIS và 4-tuple ngữ nghĩa đạt hệ số tương quan thứ hạng và mức độ đồng thuận tuyệt đối thể hiện qua các giá trị $\rho = 1.00$, $\kappa = 1.00$, phản ánh sự trùng khớp hoàn toàn về thứ hạng trong bộ dữ liệu là minh chứng thực nghiệm quan trọng cho thấy nền tảng HA cung cấp một khung tham chiếu ngữ nghĩa thống nhất, giúp kết quả ra quyết định đạt độ ổn định cao, ít phụ thuộc vào biến động kỹ thuật của thuật toán xếp hạng cụ thể. Ngược lại đạt mức độ tương quan thấp hơn với hai cặp đối sánh còn lại với FTOPSIS đạt $\rho = 0.50$ và $\kappa = 0.33$, cho thấy mô hình đề xuất đã khắc phục hiệu quả các nhiễu loạn ngữ nghĩa mà các phương pháp mờ truyền thống gặp phải.

Đối với HA-TOPSIS, các Bảng 2.11 và 2.12 cho thấy E1 có khoảng cách nhỏ nhất tới PIS (0.053) và khoảng cách lớn nhất tới NIS (0.100), nên đạt hệ số gần 0.654. E3 có các khoảng cách tương ứng 0.076 và 0.071, đạt hệ số gần 0.485; E2 có khoảng cách tới PIS lớn nhất (0.104) và khoảng cách tới NIS nhỏ nhất (0.060), nên hệ số gần chỉ đạt 0.364. Cấu trúc khoảng cách này dẫn tới thứ hạng $E1 \succ E3 \succ E2$ được trình bày trong Bảng 2.16. Kết quả trên có thể được giải thích thêm từ véc-tơ trọng số của tiêu chí và ma trận quyết định có trọng số trong mục 2.3. Trọng số của bốn tiêu chí lần lượt là $P1 = 0.298$, $P3 = 0.274$, $P2 = 0.250$ và $P4 = 0.177$; trong đó P1 và P3 có ảnh hưởng lớn nhất. E1 đạt mức đánh giá có trọng số cao ở $P3 = 0.238$, $P2 = 0.205$ và $P4 = 0.159$. E3 có giá trị cao ở $P2 = 0.214$, nhưng thấp hơn E1 tại $P3 = 0.192$ và $P4 = 0.149$. E2 có lợi thế về P1 = 0.149 khi P1 được xử lý theo hướng tiêu chí chi phí, song các giá trị tại $P2 = 0.162$, $P3 = 0.161$ và $P4 = 0.112$ đều thấp hơn hai phương án còn lại. Vì vậy, thứ hạng của HA-TOPSIS phản ánh đồng thời hướng ưu tiên của từng tiêu chí, trọng số của tiêu chí và vị trí của phương án so với PIS/NIS, thay vì chỉ dựa vào tổng điểm thành phần.



Hình 2.4. Kết quả so sánh giữa HA-TOPSIS với FTOPSIS và 4-tuple ngữ nghĩa

Đối với FTOPSIS [21], cho các hệ số gần 0.58, 0.51 và 0.36, tương ứng với thứ hạng là $E1 > E2 > E3$. Kết quả này vẫn xác định E1 là phương án tốt nhất, nhưng hoán đổi vị trí của E2 và E3 so với HA-TOPSIS. Về mặt ý nghĩa, điều này cho thấy hai mô hình cùng bảo toàn được thứ hạng phương án nổi trội của dữ liệu, tức E1 là nhà cung cấp tốt nhất, song tồn tại sự khác biệt trong cơ chế lượng hóa các phương án tầm trung có kết quả gần nhau. Sự sai khác này phản ánh hạn chế cố hữu của phương pháp mờ truyền thống: việc ánh xạ từ nhãn ngôn ngữ sang số mờ tam giác và thực hiện các bước giải mờ trung gian thường gây ra hiện tượng mất mát thông tin ngữ nghĩa và làm suy giảm độ phân biệt giữa các phương án có mức đánh giá gần nhau. Trong khi HA-TOPSIS tính toán trực tiếp trên miền định lượng ngữ nghĩa của từ ngôn ngữ, FTOPSIS phụ thuộc vào việc lựa chọn hàm thuộc và phương pháp khử mờ, dẫn đến sai số tích lũy có thể làm thay đổi thứ hạng của phương án như trường hợp giữa E2 và E3 trong bài toán này.

Mô hình 4-tuple ngữ nghĩa [54], kết quả trong Bảng 2.16 cho tổng điểm của ba phương án lần lượt là $E1 = 0.82$, $E3 = 0.77$ và $E2 = 0.59$, tương ứng với thứ hạng $E1 > E3 > E2$. Thứ hạng này trùng khớp với HA-TOPSIS cho thấy hai cách tiếp cận dựa trên HA duy trì cùng một thứ hạng trong ví dụ này. Tuy nhiên, hai kết quả không được hình thành theo cơ chế tương đương: mô hình 4-tuple ngữ nghĩa sử dụng tổng điểm và giả định các tiêu chí có trọng số bằng nhau, trong khi HA-TOPSIS sử dụng véc-tơ trọng số không đồng đều, xác định PIS/NIS và tính hệ số gần. Vì vậy, sự tương đồng về thứ hạng chỉ phản ánh tính nhất quán của kết luận trên bộ dữ liệu đang xét, không đồng nghĩa với sự tương đương về mô hình hoặc giá trị đầu ra.

Kết quả phân tích độ nhạy tại mục 2.4.1 bổ sung bằng chứng về độ ổn định và độ tin cậy của mô hình HA-TOPSIS. Khi các tham số độ đo tính mờ của phần tử sinh và gia tử trên thang đo trọng số được điều chỉnh theo hai kịch bản, trọng số của tiêu chí và hệ số gần đều thay đổi. Chẳng hạn, ở kịch bản 1, trọng số P1 và P3 lần lượt là 0.308 và 0.283, còn ở kịch bản 2 giảm xuống 0.288 và 0.267; hệ số gần của E1, E2 và E3 cũng biến động giữa các kịch bản. Dù vậy, thứ hạng thứ hạng $E1 > E3 > E2$ vẫn được duy trì. Đây là bằng chứng thực nghiệm quan trọng cho thấy kết quả của HA-TOPSIS không chỉ hợp lý tại một cấu hình tham số duy nhất mà còn có độ ổn định đáng kể trước nhiễu của các tham số ngữ nghĩa. Ngược lại, sự thay đổi thứ hạng giữa E2 và E3 trong phép đối sánh với FTOPSIS cho thấy hướng tiếp cận mờ truyền thống có thể nhạy hơn với cách gán hàm thuộc và phép đo khoảng cách.

Tóm lại, thông qua các kết quả trên, mô hình HA-TOPSIS đã kế thừa logic của TOPSIS [59] nhưng triển khai các bước định lượng ngữ nghĩa, xác định PIS/NIS, đo khoảng cách và tính hệ số gần trên miền ngữ nghĩa dựa trên HA. Trong phạm vi bài toán lựa chọn nhà cung cấp, mô hình cho cùng phương án xếp thứ nhất với hai phương pháp đối sánh, trùng hoàn toàn thứ hạng với mô hình 4-tuple ngữ nghĩa và duy trì thứ hạng khi các tham số ngữ nghĩa được thay đổi theo hai kịch bản của Mục 2.4.1. Các kết quả này hỗ trợ tính hợp lý và độ ổn định thứ hạng của HA-TOPSIS trong phạm vi thực nghiệm và tạo tiền đề để luận án tiếp tục phát triển cơ chế xác định trọng số tiêu chí có kiểm soát nhất quán tại Chương 3 và mở rộng sang MCGDM động tại Chương 4.

2.5. Kết luận chương 2

Chương 2 đã phát triển mô hình TOPSIS ngôn ngữ dựa trên HA (HA-TOPSIS) nhằm giải quyết bài toán MCGDM trong môi trường thông tin ngôn ngữ. Điểm cốt lõi của mô hình là xây dựng thang đo ngôn ngữ nội sinh dựa trên HA, thông qua bộ tham số ngữ nghĩa gồm $(HA, f\eta(g^-), \mu(h), k)$. Từ các tham số này, áp dụng hàm SQM để ánh xạ giá trị định lượng ngữ nghĩa số $\vartheta(x) \in [0, 1]$ cho mỗi từ ngôn ngữ, đồng thời bảo toàn thứ tự ngữ nghĩa và nhất quán với ngữ nghĩa định tính. Nhờ đó, toàn bộ thang đo ngôn ngữ được sinh đồng bộ từ cùng một bộ tham số, giảm phụ thuộc vào thiết kế hàm thuộc và khử mờ như trong tiếp cận tập mờ truyền thống. Trên cơ sở đó, chương 2 đã thiết lập quy trình ra quyết định HA-TOPSIS gồm xây dựng ma trận quyết định ngôn ngữ, chuyển đổi sang ma trận định lượng ngữ nghĩa, xác định PIS/NIS, tính khoảng cách và hệ số gần để xếp hạng phương án. Kết quả thực nghiệm trên bài toán lựa chọn nhà cung cấp cho thấy mô hình đề xuất xác định được phương án tối ưu một cách rõ ràng. Phân tích độ nhạy cho thấy thứ hạng các nhà cung cấp được duy trì ổn định khi thay đổi các tham số ngữ nghĩa của thang đo, qua đó khẳng định độ tin cậy kết quả của mô hình. So sánh với các mô hình 4-tuple ngữ nghĩa và FTOPSIS cho thấy HA-TOPSIS có ưu thế hơn về tính nhất quán ngữ nghĩa, độ ổn định và khả năng diễn giải kết quả.

Trên nền tảng đó, Chương 3 sẽ tiếp tục nghiên cứu bài toán xác định trọng số của tiêu chí trong môi trường HA, thông qua mô hình MCGDM tích hợp FAHP cải tiến với 4-tuple ngữ nghĩa dựa trên HA để nâng cao độ tin cậy và khả năng diễn giải của quá trình xếp hạng các phương án.

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

Chương 3 phát triển mô hình EFAHP - 4-tuple - HA nhằm giải quyết đồng thời hai vấn đề cốt lõi của bài toán MCGDM ngôn ngữ: xác định trọng số của tiêu chí và xếp hạng phương án. Trọng tâm của chương là xây dựng cơ chế kiểm soát tính nhất quán của ma trận so sánh cặp ngôn ngữ để xác định véc-tơ trọng số của tiêu chí có thể được kiểm soát tốt hơn về tính nhất quán, đồng thời áp dụng mô hình 4-tuple ngữ nghĩa để biểu diễn, kết nhập và xử lý đánh giá ngôn ngữ một cách nhất quán hơn về mặt ngữ nghĩa. Mô hình đề xuất sau đó được áp dụng cho bài toán đánh giá mức độ sẵn sàng chuyển đổi số của 16 SMEs bán lẻ tại Việt Nam nhằm kiểm chứng tính khả thi và giá trị ứng dụng của hướng tiếp cận.

Nội dung chính của chương 3 đã được công bố trong công trình [CT2].

3.1. Giới thiệu

Chương 3 tập trung giải quyết khoảng trống nghiên cứu thứ hai đã được xác định tại Chương 1, liên quan đến xác định trọng số của tiêu chí có kiểm soát tính nhất quán trong MCGDM ngôn ngữ và duy trì sự liên thông ngữ nghĩa giữa giai đoạn xác định trọng số của tiêu chí với giai đoạn kết nhập, đánh giá và xếp hạng phương án. Nếu Chương 2 phát triển HA-TOPSIS để xử lý bài toán xếp hạng phương án trên miền ngữ nghĩa dựa trên HA, thì Chương 3 bổ sung thành phần xác định trọng số của tiêu chí từ phán đoán chuyên gia và liên kết kết quả này với quá trình đánh giá phương án trong cùng một khuôn khổ 4-tuple ngữ nghĩa.

Trong MCGDM, trọng số của tiêu chí phản ánh mức độ quan trọng tương đối của các thuộc tính quyết định và tác động trực tiếp đến kết quả tổng hợp, lựa chọn hoặc xếp hạng phương án. Khi trọng số được xác định từ so sánh cặp của chuyên gia, các phán đoán có thể mang tính chủ quan, mơ hồ và không hoàn toàn nhất quán, nhất là khi được biểu diễn bằng từ ngôn ngữ [110, 111, 117]. Vì vậy, xác định trọng số của tiêu chí cần được tổ chức thành một quy trình có cơ sở phương pháp luận, vừa khai thác tri thức chuyên gia, vừa kiểm soát tính nhất quán trước khi sử dụng các trọng số ở các bước tiếp theo. AHP cung cấp cơ chế xác định trọng số dựa trên so sánh cặp [95], còn FAHP mở rộng cách tiếp cận này để xử lý các phán đoán bất định hoặc khó biểu diễn bằng số rõ [106, 19, 58]. Tuy nhiên, trong MCGDM ngôn ngữ, các đánh giá thường phải được chuyển sang số mờ hoặc hàm thuộc trước khi tính toán, làm cho kết quả phụ thuộc vào cách thiết kế thang đo và tham số biểu diễn [63, 99]. Các bước tổng hợp mờ, so sánh

mức độ mở rộng, tính trọng tâm hoặc giải mờ cũng có thể làm suy giảm mối liên hệ trực tiếp với ngữ nghĩa ban đầu của phán đoán [49, 54]. Mặt khác, cơ chế kiểm soát tính nhất quán trong nhiều biến thể FAHP vẫn chủ yếu được kế thừa hoặc điều chỉnh từ AHP trên miền số, trong khi dữ liệu đầu vào của bài toán đang xét là các phán đoán ngôn ngữ [111, 117]. Điều này đặt ra yêu cầu xác định trọng số có kiểm soát tính nhất quán và đồng thời duy trì liên thông ngữ nghĩa với giai đoạn xếp hạng

Trong bối cảnh đó, hướng tiếp cận dựa trên HA đã cung cấp cơ sở hình thức cho xây dựng thang đo, xác định trật tự ngữ nghĩa, định lượng và kết nhập đánh giá [53, 52, 54]. Tuy nhiên, như đã phân tích tại Mục 1.4.3 và Mục 1.5, trong một số mô hình MCGDM dựa trên HA, trọng số của tiêu chí vẫn được giả định biết trước, gán trực tiếp bởi chuyên gia hoặc sử dụng dưới dạng trọng số bằng nhau [54, 3, 33, 89]. Các cách tiếp cận này chưa hình thành đầy đủ một quy trình suy ra trọng số từ so sánh cặp ngôn ngữ có cơ chế kiểm soát tính nhất quán, đồng thời liên kết kết quả trọng số với quá trình kết nhập và xếp hạng trên cùng nền tảng ngữ nghĩa. Đây là vấn đề phương pháp luận trực tiếp mà Chương 3 tiếp tục xử lý

Lý thuyết HA mô hình hóa miền giá trị ngôn ngữ như một cấu trúc có trật tự, trong đó các phần tử sinh kết hợp với các gia tử tạo ra các từ ngôn ngữ mới [53, 52]. Độ đo tính mờ và hàm SQM của HA cho phép ánh xạ giá trị ngôn ngữ sang miền số nhưng vẫn bảo toàn quan hệ thứ tự và nội dung ngữ nghĩa cơ bản của từ [53]. Trên nền tảng đó, mô hình 4-tuple ngữ nghĩa với cấu trúc biểu diễn mỗi từ ngôn ngữ được mô tả bằng một cấu trúc bốn thành phần $(x, S_k(x), \vartheta(x), r_x)$, liên kết nhãn ngôn ngữ, giá trị định lượng ngữ nghĩa, khoảng tương tự ngữ nghĩa và giá trị tham chiếu trong một cấu trúc thống nhất [54]. Cách biểu diễn này phù hợp với yêu cầu liên kết kết quả xác định trọng số của tiêu chí với đánh giá phương án mà không phải chuyển đổi lặp lại giữa nhiều miền biểu diễn.

Từ các cơ sở trên, Chương 3 phát triển mô hình MCGDM tích hợp FAHP cải tiến với mô hình 4-tuple ngữ nghĩa theo hướng tiếp cận HA, ký hiệu EFAHP - 4-tuple - HA. Mô hình được tổ chức theo hai giai đoạn. Giai đoạn thứ nhất, EFAHP-HA sử dụng ma trận so sánh cặp ngôn ngữ để xác định trọng số của tiêu chí và xây dựng cơ chế kiểm soát tính nhất quán của phán đoán trong môi trường HA. Giai đoạn thứ hai, mô hình 4-tuple ngữ nghĩa dựa trên HA được sử dụng để biểu diễn, kết nhập đánh giá và xếp hạng phương án bằng các trọng số đã xác định. Cấu trúc này được thiết kế nhằm tạo sự liên thông giữa tri thức chuyên gia ở bước xác định trọng số của tiêu chí và thông tin đánh giá ở bước xếp hạng. Các công thức, chỉ số và tính chất của cơ chế kiểm soát tính nhất quán được trình bày tại Mục 3.2, còn quy trình tính toán đầy đủ được phát triển trong các mục tiếp theo.

Để kiểm chứng mô hình, Chương 3 áp dụng mô hình EFAHP - 4-tuple - HA cho bài toán đánh giá mức độ sẵn sàng chuyển đổi số của 16 SMEs trong lĩnh vực bán lẻ tại Việt Nam, với 7 tiêu chí chính và 34 tiêu chí con. Phần thực nghiệm tại Mục 3.4 được sử dụng để kiểm tra khả năng triển khai của quy trình hai giai đoạn; Mục 3.5 tiếp tục phân tích độ nhạy và đối sánh kết quả với HA-TOPSIS [CT1] và FAHP - FTOPSIS [58, 21]. Qua đó, Chương 3 xem xét tính ổn định và khả năng diễn giải của mô hình đề xuất trong phạm vi bài toán thực nghiệm, làm cơ sở đánh giá mức độ đáp ứng của hướng tiếp cận đối với khoảng trống nghiên cứu đã đặt ra.

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

Trong mô hình EFAHP - 4-tuple - HA, ma trận so sánh cặp không chỉ là công cụ tổng hợp ý kiến chuyên gia mà còn là dữ liệu đầu vào để suy ra véc-tơ trọng số của tiêu chí. Vì vậy, nếu các phán đoán so sánh cặp không được kiểm soát về tính nhất quán thì sai lệch không chỉ dừng ở bước xác định trọng số, mà còn lan truyền sang giai đoạn kết nhập, đánh giá và xếp hạng phương án [110, 117]. Theo nghĩa đó, kiểm soát tính nhất quán trong Chương 3 không phải là một bước phụ, mà là khâu then chốt nối giữa tri thức chuyên gia, cơ chế xác định trọng số của tiêu chí và quy trình tích hợp của toàn mô hình.

Khác với Chương 2, nơi trọng tâm là xếp hạng phương án trên miền ngôn ngữ, mục này tập trung riêng vào độ tin cậy của thông tin đầu vào dùng để xác định trọng số của tiêu chí. Sự tách bạch đó là cần thiết về phương pháp luận, bởi trong bài toán MCGDM, trọng số của tiêu chí giữ vai trò chi phối mức đóng góp của từng tiêu chí vào kết quả cuối cùng; do đó, một véc-tơ trọng số chỉ có ý nghĩa khi nó được suy ra từ một quy trình ra quyết định đủ nhất quán.

Trong AHP cổ điển [95], tính nhất quán được hiểu trên ma trận số rõ dương đối nghịch đảo; ma trận đạt nhất quán hoàn hảo khi tồn tại véc-tơ trọng số $\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_n)^T$ sao cho $u_{ij}^e = \frac{\tilde{w}_i}{\tilde{w}_j}$ với mọi i, j , khi đó $\lambda_{\max} = n$ và chỉ số nhất quán được xác định bởi công thức (1.8). Cơ chế này chỉ phù hợp khi đối tượng so sánh là các số rõ. Tuy nhiên, khi chuyển sang FAHP và rộng hơn là các mô hình ngôn ngữ, bản chất của ma trận so sánh cặp đã thay đổi: phán đoán không còn là một giá trị số duy nhất, mà chứa mức độ mơ hồ, độ do dự và quan hệ ngữ nghĩa giữa các nhãn đánh giá. Bởi vậy, việc sao chép trực tiếp tỷ lệ nhất quán (CR) được xác định bởi công thức (1.9) của AHP cổ điển vào môi trường ngôn ngữ sẽ làm suy giảm ý nghĩa của phép kiểm tra, vì mốc $\lambda_{\max} = n$ không phản ánh dung sai ngữ nghĩa vốn có của các phán đoán [95].

Xét theo đối tượng biểu diễn, có thể phân biệt ba lớp nhất quán. Thứ nhất, nhất quán số rõ trong AHP cổ điển là nhất quán đại số trên ma trận số rõ. Thứ hai, nhất quán trong FAHP là nhất quán trên các đại lượng mờ, thường vẫn phải quy chiếu về một lớp biểu diễn số để kiểm tra. Thứ ba, trong môi trường HA và 4-tuple ngữ nghĩa, nhất quán còn phải bảo toàn trật tự ngữ nghĩa và cấu trúc đại số của miền giá trị, nghĩa là không chỉ kiểm tra độ lệch số học mà còn phải xem xét việc tỷ số ưu tiên suy diễn được có còn nằm trong khoảng tương tự ngữ nghĩa chấp nhận được hay không.

Do đó, vấn đề đặt ra ở đây không chỉ là kiểm tra nhất quán theo nghĩa của AHP cổ điển, mà là kiểm tra nhất quán trong một không gian biểu diễn ngôn ngữ dựa trên HA. Cơ chế được đề xuất trong luận án kế thừa tư tưởng của AHP ở chỗ vẫn đo độ lệch của ma trận thông qua giá trị riêng lớn nhất, nhưng điều chỉnh mốc chuẩn theo dung sai ngữ nghĩa của từ trung hòa trên thang đo ngôn ngữ. Nhờ vậy, phép kiểm tra vừa giữ được liên hệ với tư tưởng của AHP và FAHP, vừa thích ứng với bản chất ngôn ngữ của bài toán MCGDM.

Định nghĩa 3.1. Xét chuyên gia $D^e (e = 1, \dots, h)$, với ma trận so sánh cặp ngôn ngữ

$$P^e = (x_{ij}^e)_{n \times n},$$

trong đó, $x_{ij}^e \in X_k$ biểu thị mức độ quan trọng tương đối của tiêu chí i so với tiêu chí j , $x_{ii}^e = w$ và w là từ ngôn ngữ trung hòa trên thang đo ngôn ngữ L_k . Gọi $\tilde{w}^e = (\tilde{w}_1^e, \tilde{w}_2^e, \dots, \tilde{w}_n^e)^T$, $\tilde{w}_i^e > 0$, là véc-tơ trọng số suy ra từ ma trận P^e ; gọi $\phi: [0, 1] \rightarrow \mathbb{R}_+$ là ánh xạ từ miền chuẩn hóa ngữ nghĩa sang thang tỷ số của AHP. Khi đó, ma trận P^e được gọi là HA-nhất quán nếu với mọi i, j ta có,

$$\phi^{-1}\left(\frac{\tilde{w}_i^e}{\tilde{w}_j^e}\right) \in S_k(x_{ij}^e), \forall i, j. \quad (3.1)$$

Điều kiện (3.1) cho thấy trong môi trường HA, tính nhất quán không còn được hiểu như sự trùng khít một giá trị số duy nhất, mà là sự phù hợp của tỷ số ưu tiên suy diễn được với khoảng tương tự ngữ nghĩa của phán đoán ban đầu. Đây là điểm phân biệt cốt lõi giữa ma trận HA-ngôn ngữ và ma trận số rõ của AHP cổ điển.

Để lượng hóa mức độ nhất quán, mỗi phần tử $x_{ij}^e \in P^e$ được ánh xạ sang giá trị định lượng ngữ nghĩa $\vartheta(x_{ij}^e)$ và chuyển sang biểu diễn số dương đối nghịch đảo, thu được ma trận:

$$\mathcal{R}^e = (u_{ij}^e)_{n \times n}, \text{ trong đó } u_{ij}^e > 0, u_{ii}^e = \vartheta(M), u_{ji}^e = 1/u_{ij}^e. \quad (3.2)$$

Ký hiệu $\lambda_{\max}^e = \lambda_{\max}(\mathcal{R}^e)$ là giá trị riêng lớn nhất của ma trận này.

Gọi $\vartheta(w) \in (0, 1)$ là giá trị định lượng ngữ nghĩa của từ trung hòa w theo hàm SQM của HA, và đặt:

$$\delta = 1 - \vartheta(w) \in (0, 1).$$

Trong đó, đại lượng δ được xem như dung sai ngữ nghĩa toàn cục của thang đo, phản ánh thực tế rằng quanh từ trung hòa luôn tồn tại một vùng lân cận ngữ nghĩa thay vì một điểm đúng tuyệt đối. Khi đó, mốc HA-nhất quán được xác định bởi:

$$\lambda_0^{HA} = n + \delta = n + (1 - \vartheta(w)). \quad (3.3)$$

Trên cơ sở đó, luận án đề xuất chỉ số nhất quán trong môi trường HA cho ma trận P^e của chuyên gia D^e như sau:

$$CI_{HA}^e = \frac{\lambda_{\max}^e - n - (1 - \vartheta(w))}{n - (1 - \vartheta(w))} = \frac{\lambda_{\max}^e - (n + \delta)}{n - \delta}. \quad (3.4)$$

So với CI_{Saaty} của [95], công thức (3.4) không thay đổi nguyên lý đo độ lệch bằng $\lambda_{\max}^e$, nhưng thay mốc nhất quán n bằng mốc mềm $n + \delta$. Luận án đã phát triển CI_{HA}^e trong môi trường HA thông qua việc nhúng trực tiếp dung sai ngữ nghĩa của giá trị định lượng ngữ nghĩa phần tử trung hòa vào đo độ lệch nhất quán của ma trận so sánh cặp ngôn ngữ. Nói cách khác, chỉ số CI_{HA}^e đề xuất kế thừa tư tưởng của AHP và FAHP, đồng thời thích nghi với đặc trưng của miền ngữ nghĩa dựa trên HA.

Vì CI_{HA}^e vẫn chịu ảnh hưởng bởi kích cỡ ma trận, tỷ lệ nhất quán được chuẩn hóa như sau:

$$CR_{HA}^e = \frac{CI_{HA}^e}{RI(n)}, \text{ với } e = 1, \dots, h; n \geq 3 \quad (3.5)$$

trong đó $RI(n)$ (Random Index) là chỉ số ngẫu nhiên của ma trận bậc n theo AHP cổ điển [95], được sử dụng ở đây 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

Trong luận án này, ma trận so sánh cặp của chuyên gia D^e được chấp nhận nếu $CR_{HA}^e < 0.10$. Nếu điều kiện này không đạt, chuyên gia D^e phải rà soát và điều chỉnh lại các so sánh cặp trong ma trận P^e cho đến khi đạt yêu cầu nhất quán.

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 (3.4) 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 $\lambda_{\max}^e$;
- (iv) Công thức (3.4) 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 $\lambda_{\max}^e$, $\vartheta(w)$ 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).

Chứng minh:

- (i) Chứng minh tính xác định tốt của CI_{HA}^e :

Ta cần chứng minh mẫu số của công thức (3.4) luôn dương

$$\text{Vì } \delta = 1 - \vartheta(w) \in (0, 1) \text{ và } n \geq 2 \text{ nên } n - \delta > 0.$$

Do đó, mẫu số của công thức (3.4) luôn dương, nên CI_{HA}^e được xác định tốt trên miền áp dụng.

- (ii) Chứng minh chuẩn hóa theo mốc HA-nhất quán

Ta cần chứng minh rằng, nếu $\lambda_{\max}^e = n + \delta = n + (1 - \vartheta(w))$ thì $CI_{HA}^e = 0$.

Thật vậy, thay $\lambda_{\max}^e = n + \delta$ vào (3.4), ta có

$$CI_{HA}^e = \frac{(n + \delta) - (n + \delta)}{n - \delta} = \frac{0}{n - \delta} = 0.$$

Suy ra, công thức (3.4) được chuẩn tại mốc HA-nhất quán

- (iii) Chứng minh tính đơn điệu theo mức độ không nhất quán: $CI_{HA}^e = 0$ tại mốc $\lambda_0^{HA} = n + (1 - \vartheta(w))$.

Ta cần chứng minh rằng với n và w cố định, CI_{HA}^e là hàm tăng theo $\lambda_{\max}^e$.

Từ (3.4), lấy đạo hàm theo $\lambda_{\max}^e$, ta được:

$$\frac{\partial CI_{HA}^e}{\partial \lambda_{\max}^e} = \frac{1}{n - \delta}$$

Theo (i), $n - \delta > 0$, do đó

$$\frac{\partial CI_{HA}^e}{\partial \lambda_{\max}^e} > 0.$$

Vì vậy, CI_{HA}^e là hàm tăng nghiêm ngặt theo $\lambda_{\max}^e$. Tính chất này bảo toàn ý tưởng của AHP cổ điển, ma trận càng lệch khỏi mốc nhất quán thì chỉ số nhất quán càng lớn.

(iv) *Chứng minh nhúng trực tiếp $\vartheta(w)$ trên thang đo ngôn ngữ dựa trên HA vào công thức (3.4).*

Ta cần chứng minh rằng công thức (3.4) phụ thuộc trực tiếp vào $\vartheta(w)$, hay tương đương là $\delta = 1 - \vartheta(w)$. Từ (3.4),

$$CI_{HA}^e = \frac{\lambda_{\max}^e - n - \delta}{n - \delta}$$

Ta thấy δ xuất hiện ở cả tử số: $\lambda_{\max}^e - n - \delta$, và mẫu số: $n - \delta$.

Điều đó cho thấy, giá trị định lượng ngữ nghĩa của từ trung hòa $\vartheta(w)$ không chỉ mang tính mô tả mà được tích hợp trực tiếp vào cơ chế đo mức độ nhất quán. Xét độ nhạy theo $\vartheta(w)$, ta có:

Viết CI_{HA}^e theo $\vartheta(w)$:

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

Lấy đạo hàm theo $\vartheta(w)$:

$$\frac{\partial CI_{HA}^e}{\partial \vartheta(w)} = \frac{2n - \lambda_{\max}^e}{(n - 1 + \vartheta(w))^2} \quad (3.6)$$

Vì $(n - 1 + \vartheta(w))^2 > 0$, dấu của đạo hàm phụ thuộc vào $2n - \lambda_{\max}^e$.

Xét ba trường hợp:

Trường hợp 1: $\lambda_{\max}^e < 2n$, khi đó,

$$\frac{\partial CI_{HA}^e}{\partial \vartheta(w)} > 0$$

Nghĩa là, $\vartheta(w)$ tăng thì CI_{HA}^e tăng; tương đương, δ giảm thì phép kiểm tra chặt hơn.

Trường hợp 2: $\lambda_{\max}^e = 2n$, khi đó,

$$\frac{\partial CI_{HA}^e}{\partial \vartheta(w)} = 0$$

Chỉ số không nhạy với $\vartheta(w)$ tại đúng điểm biên này.

Trường hợp 3: $\lambda_{\max}^e > 2n$, khi đó,

$$\frac{\partial CI_{HA}^e}{\partial \vartheta(w)} < 0$$

Đây là miền ma trận rất kém nhất quán, thường nằm ngoài miền giá trị của phương pháp AHP và FAHP.

(v) *Chứng minh tính liên tục và ổn định*

Bước 1: Chứng minh tính liên tục

Từ (3.4), CI_{HA}^e là hàm hữu tỉ theo $\lambda_{\max}^e$ và $\vartheta(w)$, với mẫu số $n - \delta = n - 1 + \vartheta(w)$ luôn dương theo tính chất (i). Vì vậy, CI_{HA}^e liên tục trên miền

$$\Omega = \{(\lambda_{\max}^e, \vartheta(w)): \lambda_{\max}^e > 0, \vartheta(w) \in (0,1)\}.$$

Bước 2: Chứng minh ổn định cục bộ theo nghĩa Lipschitz

Xét một miền compact: $\mathcal{K} = [n, \Lambda] \times [m_0, m_1]$, $0 < m_0 < m_1 < 1$.

Trên $\mathcal{K}$, ta có

$$\left| \frac{\partial CI_{HA}^e}{\partial \lambda_{\max}^e} \right| = \frac{1}{n - 1 + \vartheta(w)} \leq \frac{1}{n - 1 + m_0},$$

và từ công thức (3.6),

$$\left| \frac{\partial CI_{HA}^e}{\partial \vartheta(w)} \right| = \frac{|2n - \lambda_{\max}^e|}{(n - 1 + \vartheta(w))^2} \leq \frac{\max_{\lambda \in [n, \Lambda]} |2n - \lambda|}{(n - 1 + m_0)^2}.$$

Do đó, hai đạo hàm riêng đều bị chặn trên $\mathcal{K}$. Theo định lý giá trị trung bình, tồn tại hằng số $L > 0$ sao cho với mọi $(\lambda_1, \vartheta(w_1)), (\lambda_2, \vartheta(w_2)) \in \mathcal{K}$,

$$|CI_{HA}^e(\lambda_1, \vartheta(w_1)) - CI_{HA}^e(\lambda_2, \vartheta(w_2))| \leq L(|\lambda_1 - \lambda_2| + |\vartheta(w_1) - \vartheta(w_2)|)$$

Suy ra CI_{HA}^e có tính ổn định cục bộ theo nghĩa Lipschitz.

Vì vậy, thay đổi nhỏ của $\lambda_{\max}^e$ không làm CI_{HA}^e biến động bất thường; sai lệch nhỏ trong lượng hoá ngữ nghĩa của $\vartheta(w)$ chỉ gây thay đổi nhỏ cho kết quả; điều này rất quan trọng trong bài toán MCGDM xếp hạng các phương án và phân tích độ nhạy.

(vi) *Chứng minh khả năng diễn giải vùng dung sai ngữ nghĩa*

Ta xét dấu của CI_{HA}^e trong (3.4).

Vì mẫu số $n - \delta > 0$, nên dấu của CI_{HA}^e phụ thuộc hoàn toàn vào tử số

$$\lambda_{\max}^e - (n + \delta)$$

Xét ba trường hợp:

Trường hợp 1: $\lambda_{\max}^e < n + \delta$.

Khi đó, $\lambda_{\max}^e - (n + \delta) < 0$ nên $CI_{HA}^e < 0$.

Điều này được diễn giải là sai lệch của ma trận nằm trong vùng dung sai ngữ nghĩa.

Trường hợp 2: $\lambda_{\max}^e = n + \delta$

Khi đó, $CI_{HA}^e = 0$. Ma trận nằm đúng tại mốc HA-nhất quán.

Trường hợp 3: $\lambda_{\max}^e > n + \delta$

Khi đó, $CI_{HA}^e > 0$. Ma trận đã vượt quá dung sai ngữ nghĩa, và mức vượt này được lượng hoá bằng giá trị dương của CI_{HA}^e .

Vì vậy,

$$CI_{HA}^e \begin{cases} < 0, & \lambda_{\max} < n + \delta, \\ = 0, & \lambda_{\max} = n + \delta, \\ > 0, & \lambda_{\max} > n + \delta. \end{cases}$$

Điều này cho phép diễn giải trực tiếp ba trạng thái của ma trận so sánh cặp: nằm trong vùng dung sai ngữ nghĩa, đúng tại mốc HA-nhất quán, hoặc vượt quá dung sai ngữ nghĩa. Vì vậy, CI_{HA}^e không chỉ là một đại lượng đại số mà còn là một chỉ báo có ý nghĩa ngữ nghĩa rõ ràng về mức độ nhất quán của ma trận so sánh cặp trong môi trường HA.

Tóm lại, Có thể thấy điểm mới của cơ chế đề xuất nằm ở chỗ phép kiểm tra nhất quán không còn dựa trên một mốc đại số cố định như trong AHP cổ điển, mà được điều chỉnh bởi đặc trưng ngữ nghĩa của thang đo HA thông qua $\vartheta(w)$. Phần kế thừa của phương pháp là cách đo độ lệch bằng $\lambda_{\max}$ và chuẩn hóa theo $RI(n)$; phần cải tiến là đưa dung sai ngữ nghĩa của từ trung hòa vào mốc chuẩn; còn phần mới do môi trường HA và 4-tuple ngữ nghĩa đặt ra là yêu cầu diễn giải nhất quán theo khoảng tương tự ngữ nghĩa thay vì theo một điểm số duy nhất. Nhờ đó, cơ chế kiểm soát tính nhất quán được đề xuất vừa bảo đảm tính chặt chẽ về toán học, vừa tăng khả năng diễn giải và độ phù hợp khi xử lý các ma trận so sánh cặp của bài toán MCGDM trong môi trường HA.

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

Dựa trên cơ chế định lượng ngữ nghĩa của các từ ngôn ngữ trong HA, mô hình 4-tuple ngữ nghĩa và cơ chế kiểm soát tính nhất quán của ma trận so sánh cặp ngôn ngữ, luận án phát triển mô hình tích hợp EFAHP - 4-tuple - HA nhằm giải quyết bài toán MCGDM trong môi trường thông tin ngôn ngữ. Mô hình này cho phép xác định trọng số của tiêu chí một cách nhất quán trên miền ngôn ngữ, đồng thời tổng hợp và xếp hạng các phương án trên thang đo 4-tuple ngữ nghĩa theo hướng bảo toàn cấu trúc ngữ nghĩa của dữ liệu đánh giá.

(i) Dữ liệu đầu vào:

Cho $\bar{A} = \{A_1, A_2, \dots, A_m\}$ ($m \geq 2$) là tập hữu hạn các phương án khả thi, $\bar{C} = \{C_1, C_2, \dots, C_n\}$ ($n \geq 2$) là tập hữu hạn các tiêu chí, và $\bar{D} = \{D_1, D_2, \dots, D_h\}$ ($h \geq 2$) là tập hữu hạn người ra quyết định (chuyên gia).

Hội đồng chuyên gia $\bar{D}$ thống nhất xây dựng hai thang đo 4-tuple ngữ nghĩa dựa trên HA, ký hiệu $L_k(\bar{C})$ và $L_k(\bar{A})$. Trong đó, thang đo $L_k(\bar{C})$ để đánh giá tầm quan trọng tương đối giữa các tiêu chí, còn $L_k(\bar{A})$ dùng để đánh giá xếp hạng phương án $A_i \in \bar{A}$ theo tiêu chí $C_j \in \bar{C}$. Đối với mỗi chuyên gia $D^e \in \bar{D}$ (với $e = 1, 2, \dots, h$), thực hiện:

(1) *Xây dựng ma trận so sánh cặp ngôn ngữ để xác định trọng số của các tiêu chí, theo công thức (3.7):*

$$P^e = (x_{ij}^e)_{n \times n} = \begin{pmatrix} x_{11}^e & x_{12}^e & \dots & x_{1n}^e \\ & x_{22}^e & \dots & x_{2n}^e \\ & & \ddots & \vdots \\ & & & x_{nn}^e \end{pmatrix} \quad (3.7)$$

Trong đó, $x_{ij}^e \in L_k(\bar{C})$ biểu diễn mức độ quan trọng tương đối của tiêu chí C_i so với tiêu chí C_j theo đánh giá của chuyên gia D^e . Các phần tử trên đường chéo chính x_{ii}^e là phần tử trung hòa của thang đo ngôn ngữ $L_k(\bar{C})$, 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 P^e chỉ được chấp nhận nếu đạt tỷ lệ nhất quán $CR_{HA}^e < 0.10$; ngược lại, chuyên gia D^e phải rà soát và điều chỉnh lại các so sánh cặp ngôn ngữ trong ma trận P^e cho đến khi đạt yêu cầu nhất quán.

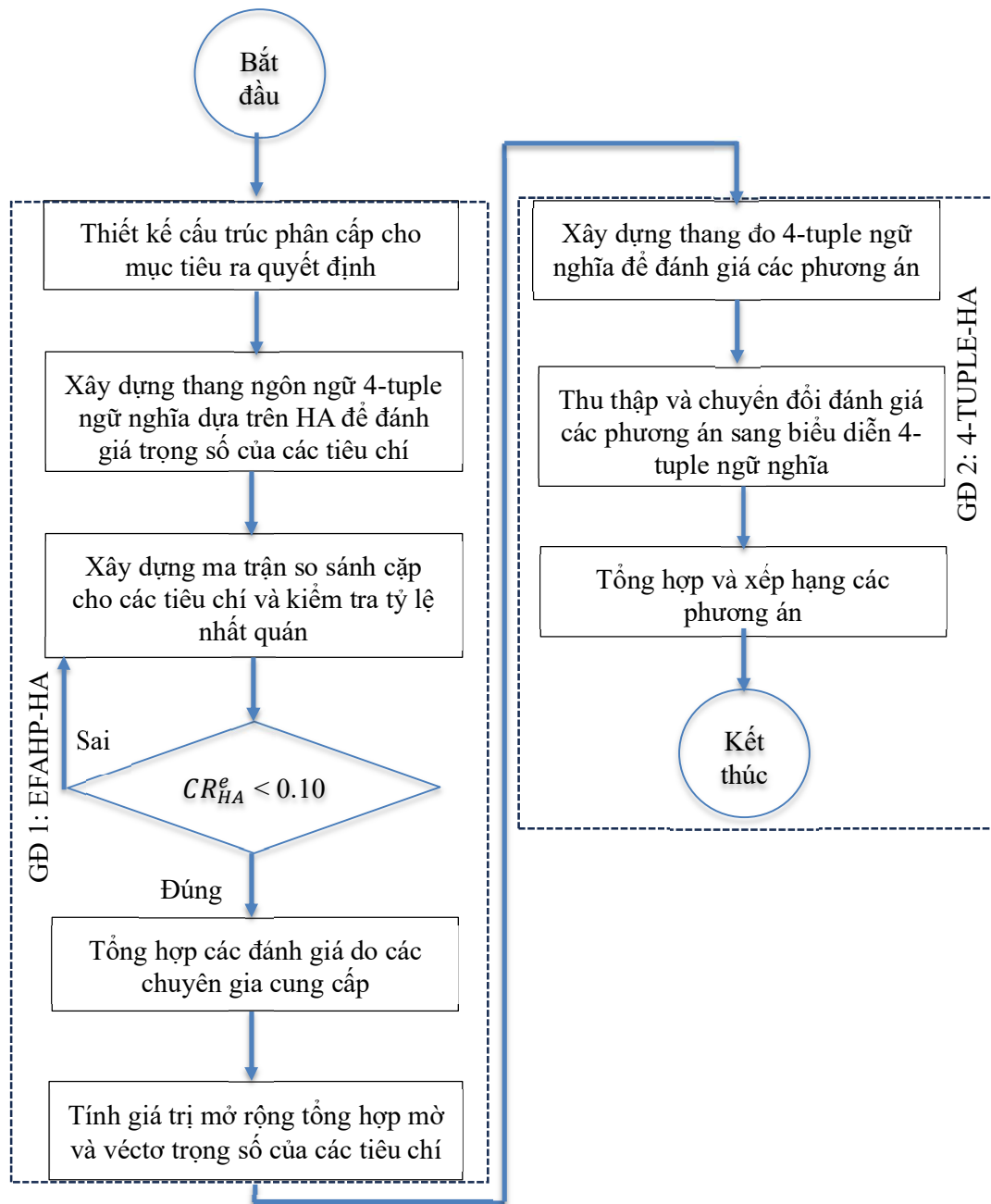
(2) *Xây dựng ma trận đánh giá các phương án theo công thức (3.8)*

$$A^e = (a_{ij}^e)_{m \times n} = \begin{pmatrix} a_{11}^e & a_{12}^e & \dots & a_{1n}^e \\ a_{21}^e & a_{22}^e & \dots & a_{2n}^e \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1}^e & a_{m2}^e & \dots & a_{mn}^e \end{pmatrix} \quad (3.8)$$

Trong đó, $a_{ij}^e \in L_k(\bar{A})$ là đánh giá ngôn ngữ của chuyên gia $D^e \in \bar{D}$ đối với phương án thứ $A_i \in \bar{A}$ theo tiêu chí $C_j \in \bar{C}$, với $e = 1, \dots, h; i = 1, \dots, m; j = 1, \dots, n$. Ma trận này phản ánh mức độ đáp ứng của từng phương án đối với từng tiêu chí trong bối cảnh đánh giá.

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

Mô hình đề xuất trong Hình 3.1 thực hiện theo hai giai đoạn. Giai đoạn 1, phát triển EFAHP-HA dựa trên phương pháp FAHP [58], để xác định trọng số của tiêu chí thông qua so sánh cặp ngôn ngữ và kiểm soát tính nhất quán trong phán đoán của các chuyên gia 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 [54]. Quy trình MCGDM với mô hình EFAHP - 4-tuple - HA gồm các bước sau:



Hình 3.1. Quy trình MCGDM của mô hình EFAHP - 4-tuple - HA

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

Trước hết, bài toán MCGDM được mô hình hóa dưới dạng cấu trúc phân cấp 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. Cấu trúc này được xác lập trên cơ sở kết hợp giữa cơ sở lý thuyết và tri thức chuyên gia, nhằm bảo đảm tính logic giữa các mức phân cấp và sự phù hợp với ngữ cảnh ứng dụng.

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

Để đánh giá tầm quan trọng tương đối giữa các tiêu chí, hội đồng chuyên gia thống nhất xây dựng một ComHA tuyến tính và lựa chọn bộ tham số ngữ nghĩa $(HA, f\eta(g^-))$ hoặc $f\eta(g^+)$, $\mu(h)$, k , từ đó thang đo ngôn ngữ 4-tuple ngữ nghĩa $L_k(\bar{C})$ được xây dựng như sau:

Cho một ComHA tuyến tính, $\mathcal{A}X^* = (X, G, E, H, \sigma, \phi, \leq)$, trong đó, $G = \{g^-, g^+\}$ là tập các từ của phần tử sinh, $E = \{0, W, 1\}$ là tập các từ hằng, và $H = H^- \cup H^+$, trong đó $H^- = \{h_{-q}, h_{-q+1}, \dots, h_{-1}\}$ là tập hợp q các gia tử âm và $H^+ = \{h_1, h_2, \dots, h_p\}$ là tập hợp p các gia tử dương.

Độ đo tính mờ của các từ phần tử sinh được ký hiệu $f\eta(g^-)$ và $f\eta(g^+)$, trong đó $f\eta(g^+) = 1 - f\eta(g^-)$. Độ đo tính mờ của các gia tử được ký hiệu $\mu(h_{-q})$, $\mu(h_{-q+1})$, ..., $\mu(h_{-1})$; $\mu(h_1)$, ..., $\mu(h_p)$, và thỏa mãn điều kiện sau: $\mu(h_p) = 1 - \sum_{i=-q}^{p-1} \mu(h_i)$.

Trên cơ sở bộ tham số ngữ nghĩa $(HA, f\eta(g^-), \mu(h), k)$ của HA, tập các từ ngôn ngữ $X_k \in L_k(\bar{C})$ được tạo bằng cách kết hợp các phần tử sinh với các gia tử hình thành cơ chế tự động sinh ra tập các từ ngôn ngữ sao cho $|x| \leq k$. Với mọi $x \in X_k$ và $h \in H$, độ đo tính mờ của từ ngôn ngữ sinh bởi gia tử được xác định đệ quy theo công thức (3.9):

$$f\eta(hx) = \mu(h) f\eta(x) \quad (3.9)$$

Ví dụ 3.1. Hãy xem xét một ComHA tuyến tính, $\mathcal{A}X^* = (X, G, E, H, \sigma, \phi, \leq)$ của biến ngôn ngữ Criteria, trong đó, $G = (Unimportant, Important)$, $E = (Extremely Unimportant, Medium, Extremely Important)$, $H = H^- \cup H^+ = (Little, Very)$, (với $q = p = 1$), theo đó các từ ngôn ngữ bao gồm: Unimportant, Important, Extremely Unimportant, Medium, and Extremely Important, kết hợp với các gia tử: Little và Very. Các từ này lần lượt được viết tắt là: U, I, EU, M, EI, L, và V. Theo Mệnh đề 1.1, độ đo tính mờ của các từ phần tử sinh và gia tử được xác định như sau:

Cấu trúc HA và tham số ngữ nghĩa của HA, được xác định dựa trên tri thức chuyên gia hoặc dữ liệu thực nghiệm phù hợp với ngữ cảnh ứng dụng cụ thể. Chẳng hạn, trong bối cảnh đánh giá mức độ chuyển đổi số của SMEs tại Việt Nam, hội đồng chuyên gia thống nhất chọn cấu trúc HA và tham số ngữ nghĩa như sau:

$$f\eta(U) = 0.489, f\eta(I) = 1 - f\eta(U); \mu(V) = 0.472, \mu(L) = 1 - \mu(V), k = 2.$$

Khi đó, tập từ ngôn ngữ $X_2 \in L_2(\bar{C})$ được nội sinh dựa trên HA như sau:

$$L_2(\bar{C}) = \{EU, VU, U, LU, M, LI, I, VI, EI\}.$$

Sử dụng Định nghĩa 1.5 và Mệnh đề 1.1 tính độ đo tính mờ cho các từ ngôn ngữ $x \in X_2$ trên thang đo 4-tuple ngữ nghĩa $L_2(\bar{C})$ được tổng hợp trong Bảng 3.2 dưới đây:

Bảng 3.2. Độ đo tính mờ của các từ ngôn ngữ $x \in X_2$ trong thang đo $L_2(\bar{C})$

Thang đo ngôn ngữ ($L_2(\bar{C})$)	EU	VU	U	LU	M	LI	I	VI	EI
$f\eta(x)$	0	0.231	0.489	0.258	0.000	0.270	0.511	0.241	0

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 $\vartheta(x) \in [0, 1]$. Để thuận tiện cho việc tính toán nghịch đảo trong ma trận so sánh cặp (tránh gặp phải chia cho 0) ở bước tiếp theo, các giá trị này được chuẩn hóa từ đoạn $[0, 1]$ sang đoạn $[0.1, 1]$.

Ví dụ 3.2. Xét lại ComHA tương ứng với biến ngôn ngữ Criteria trong Ví dụ 3.1. Độ đo tính mờ của các từ phân tử sinh và gia tử được xác định là: $f\eta(U) = 0.489, f\eta(I) = 1 - f\eta(U)$ và $\mu(V) = 0.472, \mu(L) = 1 - \mu(V)$, với $q = p = 1$.

Áp dụng các Định nghĩa 1.6 đến 1.8 và Mệnh đề 1.2, mỗi từ ngôn ngữ $x \in X_2$ trong thang đo $L_2(\bar{C})$ được ánh xạ một giá trị định lượng ngữ nghĩa $\vartheta(x) \in [0, 1]$ và được chuẩn hóa về đoạn $[0.1, 1]$, như được trình bày trong Bảng 3.3.

Bảng 3.3. Giá trị định lượng ngữ nghĩa của từ ngôn ngữ $x \in X_2$ trong thang đo $L_2(\bar{C})$

Thang đo ngôn ngữ ($L_2(\bar{C})$)	EU	VU	U	LU	M	LI	I	VI	EI
$\vartheta(x)$	0.1	0.198	0.308	0.430	0.540	0.654	0.783	0.897	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 $L_k(\bar{C})$.

Ví dụ 3.3. Xét lại ComHA liên quan đến biến ngôn ngữ “Criteria” trong Ví dụ 3.1 và Ví dụ 3.2, trong đó: $f_{\mathfrak{M}}(U) = 0.489, f_{\mathfrak{M}}(I) = 1 - f_{\mathfrak{M}}(U); \mu(V) = 0.472, \mu(L) = 1 - \mu(V)$, và $q = p = 1, k = 2$. Các giá trị định lượng ngữ nghĩa số của chín từ ngôn ngữ được cho trong Bảng 3.2 và Bảng 3.3. Áp dụng các Định nghĩa từ 1.9 đến 1.12 và Mệnh đề 1.3, để tính khoảng tương tự ngữ nghĩa $S_2(x)$ trong đoạn $[0.1, 1]$ cho mỗi từ $x \in X_2$ trên thang đo $L_2(\bar{C})$ như sau:

$$S_{(2)}(EU) = [0.1, 0.146), S_{(2)}(VU) = [0.146, 0.256), S_{(2)}(U) = [0.256, 0.366), S_{(2)}(LU) = [0.366, 0.488), S_{(2)}(M) = [0.488, 0.594), S_{(2)}(LI) = [0.594, 0.722), S_{(2)}(I) = [0.722, 0.837), S_{(2)}(VI) = [0.837, 0.952), S_{(2)}(EI) = [0.952, 1].$$

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:

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

Ví dụ 3.4. Xét lại ComHA liên quan đến biến ngôn ngữ Criteria trong Ví dụ 3.1 đến Ví dụ 3.3, với các tham số: $f_{\mathfrak{M}}(U) = 0.489, f_{\mathfrak{M}}(I) = 1 - f_{\mathfrak{M}}(U); \mu(V) = 0.472, \mu(L) = 1 - \mu(V), k = 2$ và $q = p = 1$. Các giá trị định lượng ngữ nghĩa và các khoảng tương tự của chín từ ngôn ngữ được trình bày trong các Bảng 3.2 – 3.3. Áp dụng các Định nghĩa các Định nghĩa từ 1.9 đến 1.12 và Mệnh đề 1.3 ta thu được một thang đo $L_2(\bar{C})$, trong đó các từ $x \in X_2$ được biểu diễn dưới dạng 4-tuple ngữ nghĩa như sau:

$$L_2(\bar{C}) = \{(EU, 0.1, [0.1, 0.146), 0.1), (VU, 0.198, [0.146, 0.256), 0.198), (U, 0.308, [0.256, 0.366), 0.308), (LU, 0.430, [0.366, 0.488), 0.430), (M, 0.540, [0.488, 0.594), 0.540), (LI, 0.655, [0.594, 0.722), 0.655), (I, 0.783, [0.722, 0.837), 0.783), (VI, 0.898, [0.837, 0.952), 0.898), (EI, 1, [0.952, 1], 1)\}.$$

Như vậy, mỗi từ ngôn ngữ trong thang đo $L_k(\bar{C})$ không chỉ được xác định bởi nhãn ngôn ngữ x , mà còn được đặc trưng bởi giá trị định lượng ngữ nghĩa $\vartheta(x)$, khoảng tương tự ngữ nghĩa $S_k(x)$ và giá trị tham chiếu r_x . Cách biểu diễn này cho phép bảo toàn đồng thời trật tự ngữ nghĩa, mức độ mơ hồ và khả năng tính toán trên miền ngôn ngữ, qua đó tạo cơ sở trực tiếp cho bước xây dựng ma trận so sánh cặp và xác định trọng số của tiêu chí trong môi trường HA.

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

Giả sử tại một mức của cấu trúc có n tiêu chí. Mỗi chuyên gia $D^e (e = 1, \dots, h)$, sử dụng các từ ngôn ngữ $x \in X_k$ trên thang đo ngôn ngữ $L_k(\bar{C})$ đã được tạo trong Bước 2 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 (3.7).

Trong ma trận P^e , 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:

$$x_{ii}^e = w, (i = 1, \dots, n).$$

Cách biểu diễn này khác với AHP cổ điển [95] và các phiên bản FAHP dựa trên lý thuyết mờ [19, 58] thường sử dụng các giá trị số cố định như số 1 hay số mờ tam giác $(1, 1, 1)$ hoặc GTFN $(1, 1, 1; 1)$ để biểu diễn mức độ quan trọng bằng nhau. Trong phần ứng dụng và thực nghiệm các mô hình đề xuất trong Chương 3 và Chương 4 sử dụng trực tiếp từ trung hòa M trên hai thang đo $L_2(\bar{C})$, $L_2(\bar{A})$ và $L_k(\bar{C}^T)$, $L_k(\bar{A}^T)$ để diễn đạt mức độ quan trọng như: *bằng nhau (equal)*, *xấp xỉ bằng nhau (approximately equal)* hoặc *gần như tương đương (almost the same)* qua đó phản ánh phù hợp hơn bản chất ngôn ngữ và bất định phán đoán của chuyên gia.

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 GTFN $u_{ij}^e = (L_{ij}^e, C_{ij}^e, R_{ij}^e; \omega_{ij}^e)$, Từ đó, ma trận so sánh cặp ngôn ngữ được chuyển thành ma trận số mờ tam giác tổng quát bởi công thức (3.10) dưới đây:

$$\begin{aligned} \mathcal{R}^e &= (u_{ij}^e)_{n \times n} \\ &= \begin{bmatrix} (L_{11}^e, C_{11}^e, R_{11}^e; \omega_{11}^e) & (L_{12}^e, C_{12}^e, R_{12}^e; \omega_{12}^e) & \dots & (L_{1n}^e, C_{1n}^e, R_{1n}^e; \omega_{1n}^e) \\ (\frac{1}{R_{21}^e}, \frac{1}{C_{21}^e}, \frac{1}{L_{21}^e}; \omega_{21}^e) & (L_{22}^e, C_{22}^e, R_{22}^e; \omega_{22}^e) & \dots & (L_{2n}^e, C_{2n}^e, R_{2n}^e; \omega_{2n}^e) \\ \dots & \dots & \dots & \dots \\ (\frac{1}{R_{n1}^e}, \frac{1}{C_{n1}^e}, \frac{1}{L_{n1}^e}; \omega_{n1}^e) & (\frac{1}{R_{n2}^e}, \frac{1}{C_{n2}^e}, \frac{1}{L_{n2}^e}; \omega_{n2}^e) & \dots & (L_{nn}^e, C_{nn}^e, R_{nn}^e; \omega_{nn}^e) \end{bmatrix} \quad (3.10) \end{aligned}$$

Trong đó, $u_{ij}^e = (L_{ij}^e, C_{ij}^e, R_{ij}^e; \omega_{ij}^e)$, $(u_{ij}^e)^{-1} = (\frac{1}{R_{ij}^e}, \frac{1}{C_{ij}^e}, \frac{1}{L_{ij}^e}; \omega_{ij}^e)$, $(e = 1, \dots, h; i, j = 1, \dots, n \text{ và } i \neq j)$; 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.

Như vậy, R^e là ma trận dương đối nghịch đảo, làm cơ sở cho bước kiểm tra nhất quán và xác định trọng số của tiêu chí. Bảng 3.4 trình bày tập các giá trị GTFN tương ứng với các từ ngôn ngữ $x \in X_2$ trên thang đo $L_2(\bar{C})$ dùng để đánh giá mức độ quan trọng giữa các tiêu chí, dưới đây:

Bảng 3.4. Bảng giá trị GTFN của các từ ngôn ngữ $x \in X_2$ trên thang đo $L_2(\bar{C})$

TT	Từ ngôn ngữ	GTFNs dựa trên HA	GTFNs nghịch đảo dựa trên HA
1	EU	(0.1, 0.1, 0.198; 0.7)	(1/0.198, 1/0.1, 1/0.1; 0.7)
2	VU	(0.1, 0.198, 0.308; 0.7)	(1/0.308, 1/0.198, 1/0.1; 0.7)
3	U	(0.198, 0.308, 0.430; 0.8)	(1/0.430, 1/0.308, 1/0.198; 0.8)
4	LU	(0.308, 0.430, 0.540; 0.8)	(1/0.540, 1/0.430, 1/0.308; 0.8)
5	M	(0.430, 0.540, 0.654; 1.0)	(1/0.654, 1/0.540, 1/0.430; 1.0)
6	LI	(0.540, 0.654, 0.783; 0.9)	(1/0.783, 1/0.654, 1/0.540; 0.9)
7	I	(0.654, 0.783, 0.897; 0.9)	(1/0.897, 1/0.783, 1/0.654; 0.9)
8	VI	(0.783, 0.897, 1; 1.0)	(1, 1/0.897, 1/0.783; 1.0)
9	EI	(0.897, 1, 1; 1.0)	(1, 1, 1/0.897; 1.0)

Sử dụng công thức (3.5) 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 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 so sánh cặp ngôn ngữ 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 (3.11):

$$AP = (\tilde{u}_{ij})_{n \times n} = \begin{pmatrix} \tilde{u}_{11} & \tilde{u}_{12} & \cdots & \tilde{u}_{1n} \\ \tilde{u}_{21} & \tilde{u}_{22} & \cdots & \tilde{u}_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{u}_{n1} & \tilde{u}_{n2} & \cdots & \tilde{u}_{nn} \end{pmatrix}, \quad (3.11)$$

trong đó: $\tilde{u}_{ij} = \left(\frac{\sum_{e=1}^h L_{ij}^e}{h}, \frac{\sum_{e=1}^h C_{ij}^e}{h}, \frac{\sum_{e=1}^h R_{ij}^e}{h}; \frac{\sum_{e=1}^h \omega_{ij}^e}{h} \right)$, với $e = 1, \dots, h$; $i, j = 1, \dots, n$ 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 (3.12) như sau:

$$\mathbb{P}_i = (\mathbb{D}_i, \mathbb{K}_i, \mathbb{T}_i; \min(\omega_{ij}))$$

$$= \left(\frac{\sum_{j=1}^n L_{ij}}{\sum_{i=1}^n L_{ij} + \sum_{k=1, k \neq i}^n \sum_{j=1}^n R_{kj}}, \frac{\sum_{j=1}^n C_{ij}}{\sum_{i=1}^n \sum_{j=1}^n C_{ij}}, \frac{\sum_{j=1}^n R_{ij}}{\sum_{j=1}^n R_{ij} + \sum_{k=1, k \neq i}^n \sum_{j=1}^n L_{kj}}; \min(\omega_{ij}) \right) \quad (3.12)$$

với $i, j = 1, \dots, n$ và $i \neq j$, 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 $\mathbf{m} = (\bar{\mathbb{A}}_{\mathbb{P}_i}, \bar{\mathcal{L}}_{\mathbb{P}_i})$, với $i = 1, 2, \dots, n$, và được tính theo công thức (3.13):

$$\bar{\mathbb{A}}_{\mathbb{P}_i} = \frac{(\mathbb{D}_i + \mathbb{K}_i + \mathbb{T}_i)}{3}, \quad \bar{\mathcal{L}}_{\mathbb{P}_i} = \frac{\min(\omega_{ij})}{3}. \quad (3.13)$$

Khoảng cách từ trọng tâm $\mathbf{m} = (\bar{\mathbb{A}}_{\mathbb{P}_i}, \bar{\mathcal{L}}_{\mathbb{P}_i})$ đến điểm tham chiếu $\mathbf{c} = (\mathbb{A}_{min}, \mathcal{L}_{min})$ được tính theo công thức (3.14):

$$D(\mathbb{P}_i, \mathbf{c}) = \sqrt{(\bar{\mathbb{A}}_{\mathbb{P}_i} - \mathbb{A}_{min})^2 + \left(\bar{\mathcal{L}}_{\mathbb{P}_i} - \frac{\psi}{3} \mathcal{L}_{min}\right)^2} \quad (3.14)$$

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

Cuối cùng, véc-tơ trọng số $\tilde{\mathbf{w}} = (\tilde{w}_1, \dots, \tilde{w}_n)^T$ tương ứng với ma trận so sánh mờ được tính như công thức (3.15) sau:

$$\tilde{w}_i = \frac{D(\mathbb{P}_i, \mathbf{c}_i)}{\sum_{i=1}^n D(\mathbb{P}_i, \mathbf{c}_i)} = \frac{\sqrt{(\bar{\mathbb{A}}_{\mathbb{P}_i} - \mathbb{A}_{min})^2 + \left(\bar{\mathcal{L}}_{\mathbb{P}_i} - \frac{\psi}{3} \mathcal{L}_{min}\right)^2}}{\sum_{i=1}^n \sqrt{(\bar{\mathbb{A}}_{\mathbb{P}_i} - \mathbb{A}_{min})^2 + \left(\bar{\mathcal{L}}_{\mathbb{P}_i} - \frac{\psi}{3} \mathcal{L}_{min}\right)^2}} \quad (3.15)$$

với $i = 1, \dots, n$.

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 xây dựng một ComHA tuyến tính và lựa chọn bộ tham số ngữ nghĩa phù hợp ($f\eta(g^-)$ hoặc $f\eta(g^+)$, $\mu(h)$, k) để tạo một thang đo ngôn ngữ 4-tuple dựa trên HA dành cho tập phương án, ký hiệu là $L_k(\bar{A})$. 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 $L_k(\bar{A})$ được tạo trong Bước 6 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 (3.8).

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 ngôn ngữ 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

$$g(\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_p) = \Delta \left(\frac{1}{p} \sum_{j=1}^p r_j \right), \quad (3.16)$$

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:

$$\Upsilon(\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_p) = \Delta\left(\sum_{j=1}^p \tilde{w}_j r_j\right) \quad (3.17)$$

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

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

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.

Thuật toán 3.1. Thuật toán mô hình tích hợp EFAHP - 4-tuple - HA

Algorithm EFAHP - 4-tuple - HA ($\bar{A}, \bar{C}, \bar{D}, HA_C, f\mathfrak{m}_C(g^-), \mu_C(h), k_C, HA_A, f\mathfrak{m}_A(g^-), \mu_A(h), k_A$).

Input:

- Tập các phương án $\bar{A} = \{A_1, A_2, \dots, A_m\}$, (với $m \geq 2$);
- Tập các tiêu chí $\bar{C} = \{C_1, C_2, \dots, C_n\}$, (với $n \geq 2$);
- Tập người ra quyết định $\bar{D} = \{D_1, D_2, \dots, D_h\}$, (với $h \geq 2$);
- Bộ tham số ($HA_C, f\mathfrak{m}_C(g^-), \mu_C(h), k_C$)// xây dựng thang đo $L_k(\bar{C})$;
- Bộ tham số ($HA_A, f\mathfrak{m}_A(g^-), \mu_A(h), k_A$)// xây dựng thang đo $L_k(\bar{A})$;
- Ma trận quyết định ngôn ngữ A^e đánh giá xếp hạng các phương án; Ma trận quyết định ngôn ngữ P^e đánh giá trọng số quan trọng của các tiêu chí.

Output:

Thứ hạng của các phương án $A_i \in \bar{A}$.

Begin

// Giai đoạn 1: EFAHP-HA xác định trọng số quan trọng của các tiêu chí

1. Thiết kế cấu trúc phân cấp MCDM (mục tiêu – tiêu chí – phương án).

2. Xây dựng thang đo ngôn ngữ 4-tuple $L_k(\bar{C})$ dựa trên HA bằng cách lựa chọn bộ tham số $ComHA$, $f\eta_C(g^-)$, $\mu_C(h)$, (k_C) theo các Định nghĩa từ 1.4 đến 1.12 cùng các Mệnh đề 1.1 đến 1.3.
3. for each $D_e \in \bar{D}$ do
 - Xây dựng ma trận so sánh cặp ngôn ngữ P^e bởi công thức (3.7)
 - Chuyển P^e sang ma trận GTFN bởi công thức (3.10)
 - Tính chỉ số nhất quán CI_{HA}^e theo công thức (3.4)
 - Tính tỷ lệ nhất quán CR_{HA}^e theo công thức (3.5)
 - while $CR_{HA}^e \geq 0.10$ do
 - Yêu cầu chuyên gia D^e hiệu chỉnh lại ma trận P^e
 - Cập nhật lại P^e , CI_{HA}^e , CR_{HA}^e theo các công thức (3.7), (3.10), (3.4)
- end while
- end for
4. Tổng hợp các ma trận so sánh cặp của tất cả các chuyên gia thành ma trận tổng hợp nhóm AP theo công thức (3.11).
5. for each $C_j \in \bar{C}$ do
 - Tính giá trị mở rộng tổng hợp mờ $\mathbb{P}_i$ theo công thức (3.12).
 - Tính trọng tâm m_i theo công thức (3.13).
 - Tính khoảng cách $D(\mathbb{P}_i, C)$ theo công thức (3.14).
- end for
6. Chuẩn hóa để thu được vector trọng số $\tilde{w} = (\tilde{w}_1, \dots, \tilde{w}_n)^T$ theo (3.15).
- // Giai đoạn 2: 4-tuple - HA đánh giá và xếp hạng phương án bằng 4-tuple - HA
7. Xây dựng thang đo 4-tuple $L_k(\bar{A})$ dựa trên HA bằng cách lựa chọn bộ tham số $ComHA$, $f\eta_A(g^-)$, $\mu_A(h)$, (k_A) bởi các Định nghĩa từ (1.4) - (1.12)
8. for each $D_e \in \bar{D}$ do
 - Thu thập ma trận $A^e = (u_{ij}^e)_{m \times n}$ theo công thức (3.8);
 - Chuyển $u_{ij}^e \rightarrow 4\text{-tuple } (x, S_k(x), \vartheta(x), r_x)$ bởi Định nghĩa 1.10.
- end for
9. for each $A_i \in \bar{A}$ do
 - for each $C_j \in \bar{C}$ do
 - Tổng hợp đánh giá của h chuyên gia bằng phép trung bình cộng g để xây dựng ma trận tổng hợp $AA = (\bar{a}_{ij})_{m \times n}$ theo công thức (3.16).

end for

Tính điểm tổng hợp có trọng số cho phương án A_i trên toàn bộ tập tiêu chí $\bar{C}$ bằng phép kết nhập trung bình có trọng số và véc-tơ trọng số theo công thức (3.17)

end for

10. *Xếp hạng các phương án theo $r(A_i)$ giảm dần (điểm lớn nhất xếp hạng 1).*

11. return Rank (A_i), với $A_i \in \bar{A}$

End

Bảng 3.5. Độ phức tạp của thuật toán thực hiện các công thức trong mô hình EFAHP - 4-tuple - HA

STT	Công thức	Nội dung tính toán	Độ phức tạp
1	(3.7)	Xây dựng các ma trận so sánh cặp ngôn ngữ cho h chuyên gia, mỗi ma trận có kích thước $n \times n$.	$O(h \times n^2)$
2	(3.8)	Xây dựng các ma trận đánh giá phương án cho h chuyên gia, mỗi ma trận có kích thước $m \times n$.	$O(h \times m \times n)$
3	(3.9)	Tính độ đo tính mờ của một từ ngôn ngữ sinh bởi gia tử; nếu áp dụng cho toàn bộ tập từ X_k thì có thể xem là $O(X_k)$.	$O(1)$
4	(3.10)	Chuyển toàn bộ các ma trận so sánh cặp ngôn ngữ sang ma trận GTFN; mỗi phần tử được xử lý trong thời gian hằng số.	$O(h \times n^2)$
5	(3.11)	Tổng hợp các ma trận so sánh cặp của h chuyên gia thành một ma trận nhóm kích thước $n \times n$.	$O(h \times n^2)$
6	(3.12)	Tính các giá trị mở rộng tổng hợp mờ cho toàn bộ n tiêu chí; mỗi tiêu chí cần các phép cộng và chuẩn hóa trên các hàng và phần bù của ma trận nhóm.	$O(n^3)$
7	(3.13)	Tính chỉ số trọng tâm cho n số mờ tổng hợp, mỗi trọng tâm được xác định bằng số lượng hữu hạn phép toán số học.	$O(n)$
8	(3.14)	Tính khoảng cách từ n trọng tâm đến điểm tham chiếu; mỗi khoảng cách có chi phí hằng số.	$O(n)$

STT	Công thức	Nội dung tính toán	Độ phức tạp
9	(3.15)	Chuẩn hóa n khoảng cách để thu được véc-tơ trọng số của tiêu chí.	$O(n)$
10	(3.16)	Tổng hợp đánh giá của h chuyên gia trên toàn bộ ma trận quyết định nhóm kích thước $m \times n$.	$O(h \times m \times n)$
11	(3.17)	Tính điểm tổng hợp có trọng số cho m phương án trên n tiêu chí.	$O(m \times n)$

Ký hiệu: m là số phương án; n là số tiêu chí; h là số chuyên gia. Độ phức tạp dưới đây được xác định theo chi phí thời gian để tính toàn bộ đại lượng tương ứng với từng công thức trong một lần thực hiện mô hình.

Nhận xét: Độ phức tạp thời gian chi phối của mô hình EFAHP - 4-tuple - HA là $O(\max(h \times n^2, h \times m \times n, n^3, m \times n))$. Trong trường hợp thông thường khi số tiêu chí nhỏ hơn đáng kể so với số phương án, chi phí tổng hợp đánh giá phương án $O(h \times m \times n)$ thường là thành phần chi phối; ngược lại, nếu số tiêu chí lớn thì bước tính giá trị mở rộng tổng hợp mở theo công thức (3.12) có thể trở thành nút thắt với độ phức tạp $O(n^3)$.

3.4. Ứng dụng mô hình đề xuất vào đánh giá mức độ sẵn sàng chuyển đổi số của SMEs trong lĩnh vực bán lẻ tại Việt Nam

Để kiểm chứng tính khả thi và giá trị ứng dụng của mô hình đề xuất, luận án áp dụng mô hình EFAHP - 4-tuple - HA vào bài toán đánh giá mức độ sẵn sàng chuyển đổi số của SMEs trong lĩnh vực bán lẻ tại Việt Nam. Lĩnh vực bán lẻ được lựa chọn vì đây là một trong những ngành chịu tác động mạnh của chuyển đổi số, gắn với sự thay đổi nhanh của hành vi tiêu dùng, phương thức bán hàng và mô hình kinh doanh.

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, bốn 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 và được trình bày trong Bảng A1.1 ở phần Phụ lục A [1]; 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 $L_k(\bar{C})$ trong Bảng 3.7 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 $L_k(\bar{A})$ trong Bảng 3.8.

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. Nội dung thực hiện cụ thể của từng giai đoạn được trình bày như sau:

Bước 1. Xây dựng thang đo ngôn ngữ 4-tuple ngữ nghĩa dựa trên HA để đánh giá mức độ quan trọng của các tiêu chí chính và tiêu chí con

Trên cơ sở bộ tiêu chí gồm 7 tiêu chí chính và 34 tiêu chí con, hội đồng chuyên gia thống nhất lựa chọn cấu trúc ComHA tuyến tính và bộ tham số ngữ nghĩa trình bày trong Bảng 3.6 để xây dựng thang đo ngôn ngữ 4-tuple dựa trên HA dùng cho đánh giá mức độ quan trọng của các tiêu chí.

Bảng 3.6. Thiết lập cấu trúc ComHA và bộ tham số ngữ nghĩa

ComHA	Phân tử sinh		Các từ hằng			Gia tử		Độ đo tính mờ		Độ dài tối đa của từ ngôn ngữ
	g^-	g^+	0	W	1	h_{-1}	h_1	$f\eta(g^-)$	$\mu(h_{-1})$	k
$\hat{A}_X$	U	I	EU	M	EI	L	V	0.489	0.528	2

Từ các Định nghĩa 1.4 - 1.12 cùng các Mệnh đề 1.1 - 1.3, thang đo ngôn ngữ $L_2(\bar{C})$ được thiết lập dưới dạng 4-tuple ngữ nghĩa, trong đó mỗi từ ngôn ngữ được biểu diễn bởi một bộ 4-tuple $(x, \vartheta(x), S_k(x), r_x)$, cho phép bảo toàn đồng thời nội dung ngữ nghĩa và mức độ bất định của đánh giá. Kết quả xây dựng thang đo ngôn ngữ được trình bày trong Bảng 3.7 dưới đây:

Bảng 3.7. Thang đo ngôn ngữ 4-tuple xác định trọng số của tiêu chí ($L_2(\bar{C})$)

Từ ngôn ngữ	4-tuple ngữ nghĩa
EU	(EU, 0.1, [0.1, 0.146], 0.1)
VU	(VU, 0.198, [0.146, 0.256], 0.198)
U	(U, 0.308, [0.256, 0.366], 0.308)
LU	(LU, 0.430, [0.366, 0.488], 0.430)
M	(M, 0.540, [0.488, 0.594], 0.540)
LI	(LI, 0.655, [0.594, 0.722], 0.655)
I	(I, 0.783, [0.722, 0.837], 0.783)
VI	(VI, 0.898, [0.837, 0.952], 0.898)
EI	(EI, 1, [0.952, 1], 1))

Bước 2. Xây dựng, kiểm tra tính nhất quán và tổng hợp ma trận so sánh cặp

Dựa trên bộ tiêu chí trong Bảng A1.1 ở phần Phụ lục A và thang đo $L_2(\bar{C})$, mỗi chuyên gia xây dựng các ma trận so sánh cặp ngôn ngữ để đánh giá mức độ quan trọng tương đối giữa các tiêu chí chính và các tiêu chí con theo công thức (3.7).

Các ma trận này sau đó được chuyển sang các ma trận GTFN theo công thức (3.10), làm cơ sở để tính chỉ số nhất quán CI_{HA}^e và tỷ lệ nhất quán CR_{HA}^e theo các công thức (3.4) và (3.5). Kết quả cho thấy tất cả các ma trận đều thỏa mãn điều kiện $CR_{HA}^e < 0.10$, chứng tỏ các phán đoán của chuyên gia đạt yêu cầu về tính nhất quán.

Trên cơ sở đó, các ma trận cá nhân được tổng hợp bằng phép trung bình cộng theo công thức (3.11) để thu được ma trận so sánh cặp mờ nhóm AP . Kết quả tổng hợp ma trận so sánh mờ trung bình của bảy tiêu chí chính thể hiện trong các Bảng A1.2 ở phần Phụ lục A đã phản ánh tương đối nhất quán đánh giá tổng hợp của hội đồng chuyên gia về tầm quan trọng tương đối giữa các tiêu chí. Tương tự với tất cả các ma trận so sánh mờ trung bình ở 34 tiêu chí con thuộc bảy tiêu chí chính cũng đạt yêu cầu nhất quán.

Bước 3. Xác định các giá trị mở rộng tổng hợp mờ và véc-tơ trọng số của các Tiêu chí chính, tiêu chí con

Từ các ma trận AP , các giá trị mở rộng tổng hợp mờ của 7 tiêu chí chính và 34 tiêu chí được tính bởi các công thức (3.12) - (3.14), sau đó chuẩn hóa theo công thức (3.15) để thu được các véc-tơ trọng số tương ứng.

Kết quả tính toán cho thấy toàn bộ bộ tiêu chí được lượng hóa nhất quán trong cùng một khuôn khổ ngữ nghĩa dựa trên HA, qua đó phản ánh rõ mức độ quan trọng tương đối giữa các Tiêu chí chính cũng như giữa các tiêu chí con trong từng tiêu chí chính.

Các giá trị mở rộng tổng hợp mờ của 7 tiêu chí chính và 34 tiêu chí con được trình bày trong Bảng A1.3, còn các véc-tơ trọng số tương ứng được trình bày trong Bảng A1.4

Kết quả ở Bảng A1.4 cho thấy tiêu chí chính Năng lực con người và tổ chức (P7) có trọng số lớn nhất, tiếp theo là Quản lý rủi ro và an ninh mạng (P6), trong khi các tiêu chí chính còn lại có mức độ quan trọng phân bố tương đối đồng đều.

Kết quả này phù hợp với đặc thù chuyển đổi số của SMEs bán lẻ tại Việt Nam, nơi năng lực con người, khả năng thích ứng của tổ chức và năng lực bảo đảm an toàn công nghệ giữ vai trò quan trọng. Ở cấp độ tiêu chí con, các trọng số thu được là cơ sở định lượng cho giai đoạn tổng hợp đánh giá và xếp hạng doanh nghiệp ở bước tiếp theo.

Bước 4. Xây dựng thang đo ngôn ngữ 4-tuple để đánh giá mức độ sẵn sàng chuyển đổi số của SMEs bán lẻ tại Việt Nam

Căn cứ Quyết định số 2158/QĐ-BTTTT ngày 07/11/2023 của Bộ Thông tin và Truyền thông về khung đánh giá mức độ sẵn sàng chuyển đổi số của SMEs theo 5 cấp độ, luận án sử dụng thang đo năm mức để thu thập đánh giá của 16 SMEs bán lẻ tại Việt Nam đối với bộ tiêu chí đã thiết lập trong Bảng A1.1 ở phần Phụ lục A [1].

Để bảo đảm tính nhất quán ngữ nghĩa trong toàn bộ quy trình tính toán với từ ngôn ngữ, hội đồng chuyên gia thống nhất xây dựng một ComHA tuyến tính $AX^* = (X, G, \epsilon, H, \sigma, \phi, \preceq)$ với tập phần tử sinh $G = \{\text{Low}, \text{High}\}$, tập từ hằng $\epsilon = \{\text{Extremely Low}, \text{Medium}, \text{Extremely High}\}$, cùng bộ tham số ngữ nghĩa $f_m(\text{Low}) = 0.518$, $\mu(\text{Little}) = 0.495$ và $k = 1$.

Trên cơ sở các Định nghĩa 1.4 - 1.12, cùng các Mệnh đề 1.1 - 1.3, thang đo ngôn ngữ 4-tuple $L_2(\bar{A})$ được xây dựng dựa trên HA để đánh giá mức độ sẵn sàng chuyển đổi số của SMEs.

Kết quả thu được là thang đo gồm năm mức ngôn ngữ được biểu diễn 4-tuple ngữ nghĩa tương ứng trong Bảng 3.8 dưới đây:

Bảng 3.8. Thang đo ngôn ngữ 4-tuple để đánh giá các phương án $L_2(\bar{A})$

TT	Từ ngôn ngữ	4-tuple ngữ nghĩa
1	Extremely Low (EL)	(EL, 0.1, [0.1, 0.219], 0.1)
2	Low (L)	(L, 0.336, [0.219, 0.450], 0.336)
3	Medium (M)	(M, 0.566, [0.450, 0.675], 0.566)
4	High (H)	(H, 0.781, [0.675, 0.889], 0.781)
5	Extremely High (EH)	(EH, 1, [0.889, 1], 1)

Trong thang đo mỗi mức đánh giá không chỉ được biểu diễn bằng một nhãn ngôn ngữ, mà còn gắn với giá trị định lượng ngữ nghĩa, khoảng tương tự ngữ nghĩa và giá trị tham chiếu trong miền chuẩn hóa $[0.1, 1]$. Nhờ đó, dữ liệu khảo sát của doanh nghiệp vừa bảo toàn được khả năng diễn giải, vừa có thể được xử lý định lượng trong các phép kết nhập và tổng hợp có trọng số ở các bước tiếp theo. Cách tiếp cận này bảo đảm sự tương thích giữa giai đoạn xác định trọng số bằng EFAHP-HA và giai đoạn xếp hạng doanh nghiệp bằng mô hình 4-tuple ngữ nghĩa dựa trên HA, qua đó duy trì tính nhất quán xuyên suốt của toàn bộ mô hình đề xuất.

Bước 5. Thu thập và chuyển đổi các đánh giá mức độ chuyển đổi số của 16 SMEs sang biểu diễn dạng 4-tuple ngữ nghĩa

Dựa trên thang đo ngôn ngữ 4-tuple đã xây dựng ở Bước 4, dữ liệu khảo sát được thu thập từ 16 SMEs bán lẻ tại Việt Nam cho toàn bộ 34 tiêu chí thuộc 7 tiêu chí chính của mô hình. Mỗi doanh nghiệp tự đánh giá mức độ sẵn sàng chuyển đổi số theo từng tiêu chí bằng một trong năm mức ngôn ngữ trong Bảng 3.8, từ đó hình thành ma trận đánh giá ngôn ngữ A^e theo công thức (3.8).

Tiếp theo, mỗi phần tử trong ma trận A^e được chuyển sang biểu diễn 4-tuple ngữ nghĩa theo các Định nghĩa 1.10 -1.12, tức dưới dạng $(x, \vartheta(x), S_{(k)}(x), r_x)$. Việc chuyển đổi này được thực hiện thống nhất cho toàn bộ 16 SMEs và 34 tiêu chí, bảo đảm mọi đánh giá đầu vào đều được biểu diễn trong cùng một không gian ngữ nghĩa chuẩn hóa dựa trên HA.

Bước 6. Tổng hợp kết quả và xếp hạng mức độ sẵn sàng chuyển đổi số của 16 SMEs bán lẻ tại Việt Nam

Sau khi dữ liệu khảo sát được chuyển sang biểu diễn 4-tuple ngữ nghĩa, các đánh giá của từng doanh nghiệp theo từng tiêu chí được tổng hợp bằng phép kết nhập trung bình cộng theo công thức (3.16).

Tiếp đó, điểm tổng hợp cuối cùng của mỗi doanh nghiệp được xác định bằng phép kết nhập trung bình có trọng số theo công thức (3.17), trong đó véc-tơ trọng số của các tiêu chí chính và tiêu chí con được lấy từ kết quả của Bước 3. Theo cách tiếp cận này, mức độ sẵn sàng chuyển đổi số tổng thể của mỗi doanh nghiệp phản ánh đồng thời mức độ đáp ứng thực tế trên từng tiêu chí và tầm quan trọng tương đối của tiêu chí đó trong cấu trúc phân cấp của bài toán. Kết quả đánh giá có trọng số theo từng tiêu chí của 16 SMEs được trình bày trong Bảng A1.5, còn tổng điểm và thứ hạng cuối cùng được trình bày trong Bảng 3.9 dưới đây:

Bảng 3.9. Tổng hợp điểm số và xếp hạng mức độ chuyển đổi số cho 16 SMEs trong lĩnh vực bán lẻ tại Việt Nam

Doanh nghiệp	Tổng điểm	Mô hình 4-tuple ngữ nghĩa				Thứ hạng
		x	$S_k(x)$	$\mathcal{G}(x)$	r_x	
E1	0.3522	L	[0.219, 0.450)	0.3360	0.3522	2
E2	0.3930	L	[0.219, 0.450)	0.3360	0.3930	1
E3	0.2318	L	[0.219, 0.450)	0.3360	0.2318	16
E4	0.3380	L	[0.219, 0.450)	0.3360	0.3380	5
E5	0.3446	L	[0.219, 0.450)	0.3360	0.3446	3
E6	0.3381	L	[0.219, 0.450)	0.3360	0.3381	4
E7	0.2876	L	[0.219, 0.450)	0.3360	0.2876	10
E8	0.3250	L	[0.219, 0.450)	0.3360	0.3250	6
E9	0.2956	L	[0.219, 0.450)	0.3360	0.2956	8
E10	0.3010	L	[0.219, 0.450)	0.3360	0.3010	7
E11	0.2767	L	[0.219, 0.450)	0.3360	0.2767	13
E12	0.2826	L	[0.219, 0.450)	0.3360	0.2826	11
E13	0.2770	L	[0.219, 0.450)	0.3360	0.2770	12
E14	0.2950	L	[0.219, 0.450)	0.3360	0.2950	9
E15	0.2765	L	[0.219, 0.450)	0.3360	0.2765	14
E16	0.2761	L	[0.219, 0.450)	0.3360	0.2761	15

Kết quả trong Bảng 3.9 cho thấy mô hình EFAHP - 4-tuple - HA có khả năng tổng hợp và xếp hạng cho 16 SMEs một cách nhất quán cả về mặt định lượng và ngữ nghĩa. Mặc dù các doanh nghiệp đều được ánh xạ vào cùng một nhãn ngôn ngữ L (Low), mô hình vẫn phân biệt được mức độ khác nhau giữa các doanh nghiệp thông qua giá trị tham chiếu r_x . Nhờ đó, kết quả không chỉ dừng ở mức phân lớp ngôn ngữ tổng quát mà còn cho phép xếp hạng chi tiết ngay trong cùng một lớp ngôn ngữ. Trong số 16 SMEs khảo sát, E2 đạt điểm cao nhất (0.3930), tiếp theo là E1 (0.3522), E5 (0.3446), E6 (0.3381) và E4 (0.3380), trong khi E3 có điểm thấp nhất (0.2318). Điều này cho thấy mức độ sẵn sàng chuyển đổi số của tập SMEs bán lẻ tham gia khảo sát nhìn chung còn ở mức thấp, nhưng ngay trong nội bộ tập này vẫn tồn tại sự phân hóa tương đối rõ giữa nhóm dẫn đầu, nhóm tầm trung và nhóm xếp cuối.

Về phương pháp luận, kết quả thực nghiệm khẳng định ưu điểm của mô hình đề xuất trong môi trường HA: vừa bảo toàn được khả năng diễn giải tự nhiên của thông tin ngôn ngữ, vừa duy trì được độ nhạy cần thiết để nhận diện các khác biệt định lượng nhỏ

giữa các đối tượng được đánh giá. Về mặt thực tiễn, kết quả trong Bảng 3.9 cho thấy SMEs khảo sát tuy chưa vượt qua ngưỡng “Low” về mức độ sẵn sàng chuyển đổi số, nhưng trong nội bộ nhóm này đã xuất hiện sự khác biệt tương đối rõ về mức độ chuẩn bị và triển khai chuyển đổi số giữa nhóm dẫn đầu, nhóm tầm trung và nhóm cuối. Nhóm doanh nghiệp có điểm cao hơn, tiêu biểu là E2, E1 và E5, cho thấy khả năng đã thiết lập được nền tảng tổ chức, quy trình và năng lực triển khai số tốt hơn so với phần còn lại. Ngược lại, các doanh nghiệp có điểm thấp như E3, E16, E15 và E11. Điều đó cho thấy nhu cầu tiếp tục đầu tư và cải thiện có hệ thống đối với các năng lực số cốt lõi của SMEs bán lẻ tại Việt Nam trong giai đoạn tới.

3.5. Phân tích độ nhạy và so sánh kết quả

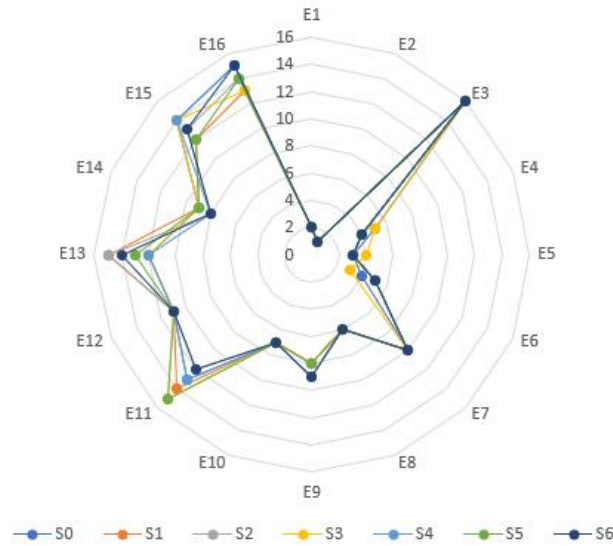
3.5.1. Phân tích độ nhạy

Để đánh giá độ ổn định và độ tin cậy kết quả của mô hình đề xuất, luận án thực hiện phân tích độ nhạy theo hai kịch bản nhiễu trên ma trận so sánh cặp giữa bảy tiêu chí chính. Trong đó:

Ở kịch bản 1, từ kịch bản ban đầu S_0 , các kịch bản $S_1 - S_6$ được xây dựng bằng cách lần lượt tăng một mức ưu tiên của từng tiêu chí chính $P_1, \dots, P_7$ so với sáu Tiêu chí chính còn lại trên thang đo ngôn ngữ $L_2(\bar{C})$, nhằm xem xét riêng rẽ mức độ ảnh hưởng của từng Tiêu chí chính đến kết quả xếp hạng.

Ở kịch bản 2, toàn bộ các giá trị so sánh cặp của bảy tiêu chí chính được gây nhiễu đồng đều bằng cách tăng lên một cấp độ trên cùng thang đo $L_2(\bar{C})$, qua đó đánh giá mức độ ổn định của mô hình trước biến động hệ thống của bộ trọng số.

Thiết kế này cho phép kiểm tra đồng thời độ nhạy cục bộ và độ nhạy toàn cục của mô hình, từ đó làm rõ khả năng duy trì tính nhất quán và độ tin cậy của kết quả xếp hạng trong bài toán đánh giá mức độ sẵn sàng chuyển đổi số của 16 SMEs bán lẻ tại Việt Nam. Kết quả phân tích được tổng hợp trong Bảng A1.6 và minh họa trực quan trong Hình 3.2 và Hình 3.3 dưới đây:



Hình 3.2. Phân tích độ nhạy theo kịch bản nhiều đơn tiêu chí

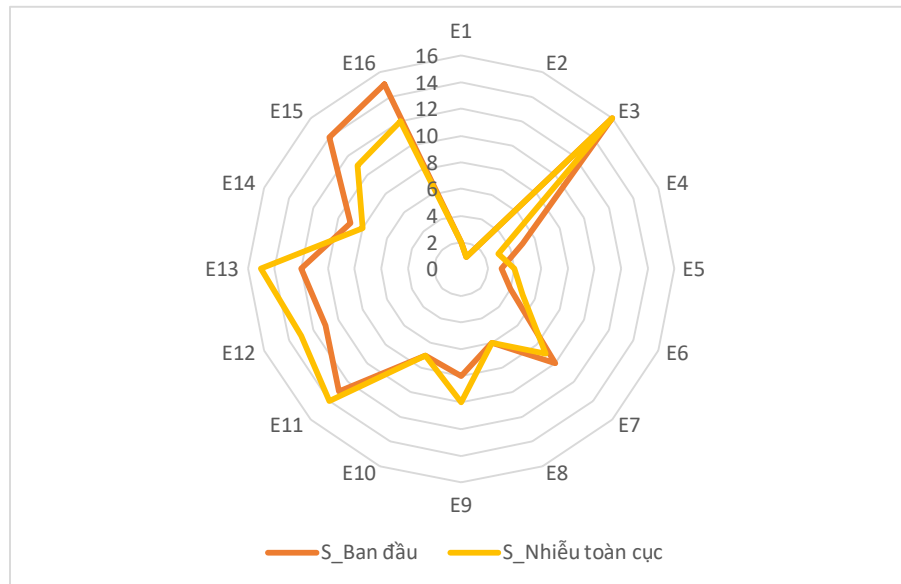
Kết quả của kịch bản 1: thể hiện trong Bảng A1.6 và Hình 3.2 cho thấy mô hình có độ ổn định cao đối với nhóm doanh nghiệp đứng đầu và nhóm doanh nghiệp xếp cuối. Cụ thể, doanh nghiệp E2 luôn duy trì vị trí hạng nhất trong tất cả các kịch bản, với tổng điểm chỉ dao động nhẹ trong khoảng 0.389 - 0.400; doanh nghiệp E1 cũng giữ ổn định ở hạng 2 trong toàn bộ các kịch bản; tương tự, E3 luôn xếp cuối ở hạng 16. Điều này cho thấy kết quả xếp hạng của mô hình không bị đảo lộn trước các điều chỉnh nhỏ trong mức độ ưu tiên của các Tiêu chí chính, qua đó phản ánh độ ổn định và độ ổn định cao của mô hình đối với các phương án có ưu thế hoặc bất lợi rõ rệt.

Biến động xếp hạng chủ yếu xuất hiện ở nhóm doanh nghiệp tầm trung, nơi các tổng điểm tương đối gần nhau. Cụ thể, E4 dao động giữa hạng 4 và 5; E5 và E6 hoán đổi giữa hạng 3, 4 và 5 ở một số kịch bản; E9 và E14 thay đổi qua lại giữa hạng 8 và 9. Mức dịch chuyển rõ hơn xuất hiện ở E11, E13, E15 và E16; trong đó E13 có thể giảm từ hạng 12 xuống hạng 15, còn E16 có thể tăng từ hạng 15 lên hạng 13. Tuy nhiên, các thay đổi này vẫn nằm trong phạm vi hẹp và không làm thay đổi cấu trúc xếp hạng tổng thể. Kết quả này cho thấy mô hình vừa đủ nhạy để nhận diện các khác biệt nhỏ giữa SMEs có năng lực tương đương, vừa đủ ổn định để tránh tạo ra các biến động cực đoan khi đầu vào chỉ thay đổi nhẹ.

Về phương pháp luận, kết quả phân tích theo kịch bản nhiều đơn tiêu chí cho thấy mô hình đề xuất đạt được sự cân bằng hợp lý giữa độ ổn định và độ nhạy. Độ ổn định thể hiện ở việc các doanh nghiệp đứng đầu và xếp cuối được nhận diện nhất quán qua tất cả các kịch bản; trong khi đó, độ nhạy thể hiện ở khả năng phân biệt tinh trong nhóm doanh nghiệp xếp hạng tầm trung, nơi khoảng cách điểm số nhỏ và sự thay đổi trọng số

của một tiêu chí chính có thể làm thay đổi quan hệ thứ hạng cục bộ. Đây là đặc điểm quan trọng đối với bài toán MCGDM sử dụng thông tin ngôn ngữ, vì trong thực tế, các đánh giá chuyên gia luôn chứa đựng tính mơ hồ và dao động nhất định trong nhận thức về tầm quan trọng của tiêu chí. Nhờ dựa trên cấu trúc ngữ nghĩa nội sinh của HA, mô hình không chỉ bảo toàn được trật tự ngữ nghĩa của các đánh giá mà còn hạn chế sự méo lệch kết quả khi bộ trọng số bị nhiễu.

Xét theo kịch bản 2: nhiễu toàn cục, kết quả minh họa trong Hình 3.3 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. Một số dịch chuyển như E16 tăng hạng hoặc E13 giảm hạng cho thấy kết quả xếp hạng có phản ứng nhất định trước biến động hệ thống của trọng số, nhưng mức độ thay đổi vẫn không đủ lớn để phá vỡ cấu trúc xếp hạng chung. Điều này cung cấp thêm bằng chứng rằng mô hình đề xuất có tính ổn định cao không chỉ trước các nhiễu cục bộ theo từng tiêu chí chính mà còn trước các biến dạng có tính hệ thống trong toàn bộ ma trận so sánh cặp.



Hình 3.3. Phân tích độ nhạy theo kịch bản nhiễu toàn cục

Tóm lại, kết quả phân tích độ nhạy khẳng định rằng 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. Đây là một minh chứng quan trọng cho thấy mô hình phù hợp với các bài toán MCGDM trong môi trường ngôn ngữ mơ hồ và không chắc chắn, nơi vừa cần bảo đảm khả năng diễn giải ngữ nghĩa, vừa cần duy trì tính ổn định của kết quả ra quyết định.

3.5.2. So sánh kết quả và thảo luận

Nhằm đánh giá mức độ tương đồng và làm rõ các khác biệt trong kết quả xếp hạng, mục này đối sánh ba hướng tiếp cận trên cùng bộ dữ liệu 16 SMEs bán lẻ tại Việt Nam: EFAHP - 4-tuple - HA [CT2], FAHP-FTOPSIS [58, 21] và HA-TOPSIS [CT1]. Việc đối sánh tập trung vào thứ hạng cuối cùng, mức độ phân hóa giữa các phương án, cơ chế xác định trọng số của tiêu chí, cách biểu diễn và xử lý thông tin ngôn ngữ, đồng thời liên hệ với kết quả phân tích độ nhạy của mô hình đề xuất tại Mục 3.5.1. Do FAHP-FTOPSIS và HA-TOPSIS sử dụng hệ số gần, còn EFAHP - 4-tuple - HA sử dụng tổng điểm kết nhập, các giá trị số trong Bảng 3.10 không nằm trên cùng một thang đo và không được so sánh trực tiếp về độ lớn. Vì vậy, nội dung thảo luận chủ yếu dựa trên thứ hạng và cơ chế hình thành kết quả của từng mô hình.

Bảng 3.10. So sánh kết quả giữa mô hình EFAHP - 4-tuple - HA với FAHP - FTOPSIS và HA-TOPSIS

Doanh nghiệp	FAHP-FTOPSIS [58, 21]		HA-TOPSIS [CT1]		EFAHP - 4-tuple - HA [CT2]	
	Hệ số gần	Thứ hạng	Hệ số gần	Thứ hạng	Tổng điểm	Thứ hạng
E1	0.5813	4	0.0139	2	0.3522	2
E2	0.7010	1	0.0212	1	0.3930	1
E3	0.5206	11	0.0078	16	0.2318	16
E4	0.5183	12	0.0123	4	0.3380	5
E5	0.4764	13	0.0133	3	0.3446	3
E6	0.6081	2	0.0111	6	0.3381	4
E7	0.3388	16	0.0110	8	0.2876	10
E8	0.6051	3	0.0121	5	0.3250	6
E9	0.5476	8	0.0102	12	0.2956	8
E10	0.4724	14	0.0104	11	0.3010	7
E11	0.5646	7	0.0108	9	0.2767	13
E12	0.5650	6	0.0107	10	0.2826	11
E13	0.4392	15	0.0100	13	0.2770	12
E14	0.5725	5	0.0110	7	0.2950	9
E15	0.5256	10	0.0100	14	0.2765	14
E16	0.5406	9	0.0094	15	0.2761	15

Bảng 3.10 cho thấy cả ba mô hình đều xác định E2 ở vị trí thứ nhất. Tuy nhiên, sự phân hóa bắt đầu xuất hiện rõ rệt ở nhóm SMEs kế cận và nhóm tầm trung. EFAHP - 4-tuple - HA xếp 5 SMEs dẫn đầu theo thứ tự $E2 > E1 > E5 > E6 > E4$; HA-TOPSIS

cho $E2 > E1 > E5 > E4 > E8$; còn FAHP-FTOPSIS cho $E2 > E6 > E8 > E1 > E14$. EFAHP - 4-tuple - HA và HA-TOPSIS cùng xếp E3 ở vị trí cuối, trong khi FAHP-FTOPSIS xếp E7 ở vị trí 16 và E3 ở vị trí 11. Hai mô hình dựa trên HA trùng khớp vị trí xếp hạng tại 6 SMEs: E1, E2, E3, E5, E15 và E16, đồng thời có cùng 3 vị trí dẫn đầu là E2, E1 và E5. Kết quả này cho thấy mức độ tương đồng thứ hạng giữa EFAHP - 4-tuple - HA và HA-TOPSIS cao hơn so với sự tương đồng của từng mô hình này với FAHP-FTOPSIS trên bộ dữ liệu đang xét.

Để lượng hóa mức độ tương đồng thứ hạng giữa các mô hình đối sánh, luận án sử dụng hệ số tương quan thứ hạng Spearman (ρ) [100] và Kendall (κ) [67], khi không có hạng trùng trong từng mô hình. Kết quả được tổng hợp tại Bảng 3.11 dưới đây:

Bảng 3.11. Hệ số tương quan thứ hạng và mức độ đồng thuận giữa các mô hình đối sánh.

Cặp mô hình đối sánh	Hệ số Spearman (ρ)	Hệ số Kendall (κ)
EFAHP - 4-tuple - HA [CT2] và HA-TOPSIS [CT1]	0.9059	0.8000
EFAHP - 4-tuple - HA [CT2] và FAHP-FTOPSIS [58, 21]	0.3912	0.3167
HA-TOPSIS [CT1] và FAHP-FTOPSIS [58, 21]	0.4382	0.3500

Cặp mô hình EFAHP - 4-tuple - HA và HA-TOPSIS đạt chỉ số tương quan thứ hạng và mức độ đồng thuận ở mức rất cao với $\rho = 0.9059$ và $\kappa = 0.8000$ là minh chứng thực nghiệm khẳng định rằng: khi cùng dựa trên nền tảng HA, các mô hình ra quyết định đạt được tính nhất quán rất cao bất kể cơ chế xếp hạng là kết nhập trực tiếp hay đo khoảng cách. Ngược lại cặp EFAHP - 4-tuple - HA và FAHP-FTOPSIS thấp hơn rõ rệt, với $\rho = 0.3912$ và $\kappa = 0.3167$; và cặp HA-TOPSIS và FAHP-FTOPSIS có $\rho = 0.4382$ và $\kappa = 0.3500$ chỉ ra rằng các phương pháp dựa trên tập mờ truyền thống, qua bước giải mờ hay chuyển đổi số mờ sang số thực để tính khoảng cách thường gây ra hiện tượng suy giảm hoặc mất mát thông tin ngôn ngữ ban đầu.

Đối với mô hình EFAHP - 4-tuple - HA [CT2]: Cấu trúc xếp hạng được hình thành qua hai giai đoạn liên thông. Ở giai đoạn 1, trọng số của 7 tiêu chí chính và 34 tiêu chí con được suy ra từ các ma trận so sánh cặp ngôn ngữ bằng EFAHP-HA; các phán đoán được kiểm soát tính nhất quán theo cơ chế phát triển tại Mục 3.2 để tính CI_{HA}^e và CR_{HA}^e thỏa mãn $CR_{HA}^e < 0.10$ trước khi sử dụng để xác định véc-tơ trọng số của tiêu chí. Ở giai đoạn thứ hai, đánh giá của SMEs được biểu diễn bằng 4-tuple ngữ nghĩa, kết nhập theo công thức (3.16), sau đó tổng hợp có trọng số để tính điểm theo công thức (3.17) và xếp

hạng. Theo Bảng 3.9, 16 SMEs trong mẫu khảo sát đều được ánh xạ vào nhãn L (Low), nhưng vẫn có các giá trị định lượng và giá trị tham chiếu khác nhau. Nhờ đó, mô hình duy trì được khả năng phân biệt thứ hạng giữa các doanh nghiệp ngay cả khi chúng cùng thuộc một lớp ngôn ngữ. Chẳng hạn, E2 đạt tổng điểm 0.3930, tiếp theo là E1 với 0.3522 và E5 với 0.3446; trong khi E6 và E4 có điểm rất gần nhau, tương ứng 0.3381 và 0.3380. Những chênh lệch nhỏ này cho thấy giá trị định lượng trong biểu diễn 4-tuple được sử dụng để sắp thứ tự chi tiết trong cùng lớp ngôn ngữ, thay vì đồng nhất các phương án chỉ vì chúng có cùng nhãn đánh giá, phản ánh chính xác sự khác biệt tinh tế trong năng lực chuyển đổi số giữa SMEs.

Đối với FAHP-FTOPSIS [58, 21]: E2 cũng được xác định ở vị trí thứ nhất với hệ số gần 0.7010, nhưng cấu trúc thứ hạng còn lại khác đáng kể so với EFAHP - 4-tuple - HA. E6 và E8 lần lượt đứng thứ hai và thứ ba, trong khi E5 từ hạng 3 ở mô hình đề xuất xuống hạng 13, E4 từ hạng 5 xuống hạng 12 và E8 từ hạng 6 lên hạng 3. FAHP-FTOPSIS sử dụng FAHP [58] để xác định trọng số và FTOPSIS [21] để xếp hạng trên không gian biểu diễn mờ, trong đó các đánh giá được xử lý thông qua số mờ, giải pháp lý tưởng, khoảng cách và hệ số gần. Những sai khác này cho thấy phương pháp dựa trên tập mờ truyền thống vẫn nhận diện đúng SME đứng đầu, nhưng thiếu ổn định hơn ở nhóm SMEs tầm trung và cận trên. Nguyên nhân chủ yếu của sự khác biệt nêu trên xuất phát từ nền tảng biểu diễn, tính toán mờ và khử mờ; có thể làm suy giảm một phần cấu trúc ngữ nghĩa ban đầu, đặc biệt khi các phương án có mức đánh giá tương đối gần nhau. Vì vậy, kết quả của nó có thể nhạy hơn đối với cách lựa chọn hàm thuộc, dạng số mờ và quy tắc giải mờ, dẫn tới sự biến động thứ hạng ở những phương án không cách biệt lớn về chất lượng tổng thể.

Thứ ba, đối với HA-TOPSIS [CT1]: thứ hạng $E2 > E1 > E5$ ở ba vị trí đầu trùng với EFAHP - 4-tuple - HA. HA-TOPSIS sử dụng thang đo ngôn ngữ và hàm SQM của HA để định lượng các đánh giá, xác định PIS/NIS, đo khoảng cách và tính hệ số gần. Về trọng số của tiêu chí, các chuyên gia đánh giá trực tiếp mức độ quan trọng của từng tiêu chí bằng các từ ngôn ngữ; các đánh giá này được định lượng và kết nhập để hình thành véc-tơ trọng số. Cách xác định này không sử dụng ma trận so sánh cặp và cơ chế kiểm soát tính nhất quán như EFAHP-HA. Do đó, khác biệt giữa hai mô hình không chỉ nằm ở quy tắc xếp hạng - khoảng cách tới PIS/NIS đối với HA-TOPSIS và kết nhập 4-tuple có trọng số đối với EFAHP - 4-tuple - HA - mà còn ở cách hình thành trọng số của tiêu chí. Do đó, các sai khác thứ hạng của SMEs: E4, E6, E9, E10, E11 và một số phương án khác, được xem là khoảng trống phương pháp luận của HA-TOPSIS: mặc dù đã xử lý tốt bài toán xếp hạng trên miền ngôn ngữ, mô hình này chưa tích hợp một cơ chế xác

định trọng số của tiêu chí nội sinh có kiểm soát tính nhất quán như EFAHP-HA. Vì vậy, mô hình EFAHP - 4-tuple - HA có thể được xem là phù hợp với logic phát triển của luận án từ giải quyết bài toán xếp hạng sang bổ sung cơ chế xác định trọng số có kiểm soát tính nhất quán

Xét tổng thể Bảng 3.10, sự tương đồng giữa EFAHP - 4-tuple – HA [CT2] và HA-TOPSIS [CT1] tập trung ở nhóm dẫn đầu và nhóm xếp cuối, trong khi các sai khác chủ yếu xuất hiện ở nhóm tầm trung. Các SMEs: E4, E6, E7, E9, E10, E11, E12, E13 và E14 có những thay đổi thứ hạng với mức độ khác nhau giữa hai mô hình. Điều này phù hợp với đặc điểm của dữ liệu thực nghiệm, khi nhiều SMEs có tổng điểm tương đối gần nhau và cùng thuộc nhãn L (Low). Trong trường hợp đó, thay đổi về trọng số của tiêu chí hoặc quy tắc tổng hợp có thể dẫn đến hoán đổi cục bộ mà không làm thay đổi các phương án có ưu thế hoặc bất lợi rõ rệt. So với hai mô hình dựa trên HA, mô hình FAHP-FTOPSIS dựa trên lý thuyết tập mờ tạo ra nhiều thay đổi thứ hạng hơn trên cùng bộ dữ liệu. Kết quả phân tích độ nhạy tại Mục 3.5.1 bổ sung bằng chứng về độ ổn định của chính mô hình EFAHP - 4-tuple - HA trong phạm vi các kịch bản đã khảo sát. Khi gây nhiễu đơn tiêu chí và nhiễu toàn cục trên ma trận so sánh cặp của bảy tiêu chí chính, các SMEs dẫn đầu E2, E1 và SME xếp cuối E3 duy trì vị trí tương đối ổn định; biến động chủ yếu tập trung ở nhóm SMEs có điểm gần nhau. Kết quả này cho thấy mô hình có phản ứng trước thay đổi của trọng số nhưng không làm đảo lộn cấu trúc xếp hạng tổng thể trong các kịch bản được thiết kế; củng cố luận điểm rằng lợi thế của HA không chỉ nằm ở biểu diễn ngôn ngữ, mà còn ở khả năng duy trì cấu trúc ngữ nghĩa xuyên suốt từ giai đoạn xác định trọng số đến giai đoạn tổng hợp và xếp hạng.

Tóm lại, trong phạm vi thực nghiệm Chương 3, EFAHP - 4-tuple - HA cho kết quả thống nhất với hai mô hình đối sánh ở phương án đứng đầu và có mức tương đồng thứ hạng đáng kể với HA-TOPSIS, đặc biệt ở ba vị trí dẫn đầu và các phương án cuối bảng. Điểm khác biệt phương pháp luận của mô hình đề xuất nằm ở việc bổ sung cơ chế xác định trọng số của tiêu chí có kiểm soát tính nhất quán với biểu diễn, kết nhập và xếp hạng 4-tuple trên nền HA. Kết quả phân tích độ nhạy cho thấy cấu trúc xếp hạng của mô hình duy trì tương đối ổn định trong các kịch bản nhiễu đã khảo sát. Các kết quả này hỗ trợ tính khả thi và khả năng diễn giải của mô hình trong bài toán đánh giá mức độ sẵn sàng chuyển đổi số của 16 SMEs bán lẻ tại Việt Nam. Do đó, mô hình EFAHP - 4-tuple - HA có thể được xem là một mở rộng hợp lý, nhất quán và có giá trị phương pháp luận cho bài toán MCGDM ngôn ngữ mờ hồ và không chắc chắn.

3.6. Kết luận chương 3

Chương 3 đã đề xuất mô hình MCGDM tích hợp FAHP cải tiến với 4-tuple ngữ nghĩa dựa trên HA nhằm đánh giá mức độ sẵn sàng chuyển đổi số của SMEs bán lẻ tại Việt Nam. Mô hình kết hợp cơ chế xác định trọng số của tiêu chí bằng phương pháp EFAHP kèm cơ chế kiểm soát tính nhất quán trong môi trường HA, đồng thời sử dụng biểu diễn 4-tuple ngữ nghĩa để kết nhập và xếp hạng phương án trên miền ngôn ngữ. Cách tiếp cận này góp phần bảo toàn cấu trúc ngữ nghĩa của thông tin đầu vào, giảm mất mát thông tin trong tính toán, đồng thời nâng cao tính nhất quán, khả năng diễn giải và độ ổn định của kết quả.

Kết quả thực nghiệm trên 16 SMEs cho thấy mô hình đề xuất tích hợp hiệu quả các đánh giá ngôn ngữ của chuyên gia và doanh nghiệp, phản ánh hợp lý sự khác biệt giữa các phương án ngay cả khi chúng cùng thuộc một lớp ngôn ngữ đánh giá. Phân tích độ nhạy và đối sánh với mô hình HA-TOPSIS và FAHP - FTOPSIS tiếp tục cho thấy tính khả thi, độ ổn định và khả năng diễn giải của mô hình theo hướng tiếp cận HA trong bài toán MCGDM sử dụng thông tin ngôn ngữ.

Trên cơ sở đó, Chương 4 sẽ mở rộng từ mô hình MCGDM tĩnh sang MCGDM động dựa trên HA, nhằm mô hình hóa sự biến thiên theo thời gian của tập phương án, bộ tiêu chí và hội đồng chuyên gia; đồng thời xây dựng các phép kết nhập ngôn ngữ mới trong quá trình tổng hợp nhóm và kết nhập giữa thông tin hiện tại và lịch sử mà vẫn bảo đảm trật tự ngữ nghĩa và tính nhất quán định lượng trên thang đo ngôn ngữ HA, đồng thời đảm bảo tính ổn định, khả năng diễn giải và độ tin cậy của kết quả trong điều kiện bất định và biến động.

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

Trong chương này, luận án đề xuất một số lý thuyết mở rộng của HA như phép kết nhập ngôn ngữ nhị phân và phép kết nhập với dữ liệu khuyết thiếu có cơ chế suy giảm ngữ nghĩa phù hợp trên thang đo ngôn ngữ 4-tuple ngữ nghĩa nhằm giải quyết một số vấn đề về bộ tiêu chí, tập phương án, tập chuyên gia có thể thay đổi theo thời gian và dữ liệu lịch sử của bài toán MCGDM động trong môi trường ngôn ngữ dựa trên HA.

Nội dung chính của chương 4 đã được công bố trong công trình [CT3].

4.1. Giới thiệu

Chương 4 là bước phát triển mô hình tiếp theo của luận án nhằm mở rộng khung MCGDM dựa trên HA từ bối cảnh tĩnh sang bối cảnh động. Nếu Chương 2 tập trung vào việc phát triển mô hình HA-TOPSIS để xếp hạng phương án trên miền ngôn ngữ, và Chương 3 tập trung vào cơ chế xác định trọng số của tiêu chí có kiểm soát tính nhất quán trong mô hình EFAHP - 4-tuple - HA, thì Chương 4 hướng tới một yêu cầu phương pháp luận rộng hơn: xử lý quá trình ra quyết định khi dữ liệu đánh giá, phương án, tiêu chí hoặc hội đồng chuyên gia có thể thay đổi qua nhiều thời điểm. Theo nghĩa đó, Chương 4 không lặp lại các mô hình tĩnh đã xây dựng, mà kế thừa chúng để phát triển một khung phương pháp xây dựng mô hình MCGDM động có khả năng tích hợp dữ liệu lịch sử với đánh giá hiện tại trong môi trường thông tin ngôn ngữ.

Sự cần thiết của Chương 4 xuất phát trực tiếp từ khoảng trống nghiên cứu thứ ba đã được xác định ở Chương 1. Trong nhiều bài toán thực tiễn, trong đó các bài toán đánh giá theo chu kỳ như đánh giá mức độ sẵn sàng chuyển đổi số [73, 16, 74], lựa chọn giải pháp công nghệ [26, 6], hoặc xếp hạng phương án đầu tư [10, 90], quyết định cuối cùng không chỉ phụ thuộc vào dữ liệu tại một thời điểm đơn lẻ. Tập phương án có thể được bổ sung hoặc loại bỏ; bộ tiêu chí có thể được điều chỉnh theo mục tiêu đánh giá [23]; hội đồng chuyên gia cũng có thể thay đổi do điều kiện tổ chức, chuyên môn hoặc phạm vi tham gia thực tế [32]; trọng số của tiêu chí có thể thay đổi khi ưu tiên của tổ chức thay đổi và ma trận đánh giá có thể được cập nhật khi xuất hiện dữ liệu mới. Vì vậy, MCGDM động cần được hiểu không phải là việc lặp lại độc lập nhiều bài toán MCGDM tĩnh, mà là một lớp bài toán trong đó thông tin ra quyết định thay đổi theo thời gian và kết quả tại thời điểm hiện tại được xem xét trong quan hệ với dữ liệu lịch sử [17, 29].

Trong môi trường thông tin ngôn ngữ, đặc trưng động của bài toán làm phát sinh thêm những khó khăn về phương pháp luận. Các đánh giá của chuyên gia thường được biểu đạt bằng các giá trị ngôn ngữ, khi các đánh giá này thay đổi theo thời gian, việc tích hợp dữ liệu lịch sử với dữ liệu hiện tại không thể chỉ dựa trên phép trung bình số học hoặc gán trọng số thời gian cho các đại lượng số trung gian. Cách tiếp cận như vậy có thể làm suy giảm mối liên hệ giữa biểu diễn định lượng và nội dung ngữ nghĩa ban đầu, trong khi kết quả tổng hợp nằm gần ranh giới giữa các khoảng tương tự ngữ nghĩa hoặc khi dữ liệu tại một số thời điểm bị khuyết thiếu [23, 32]. Do đó, bài toán MCGDM động trong môi trường ngôn ngữ đòi hỏi một phép kết nhập có khả năng bảo toàn trật tự ngữ nghĩa, xử lý dữ liệu khuyết thiếu và phản ánh hợp lý mức độ suy giảm ảnh hưởng của dữ liệu lịch sử theo thời gian [91, 29].

Các nghiên cứu hiện có về MCGDM động trong môi trường bất định đã bước đầu giải quyết một số khía cạnh của bài toán này, chủ yếu dựa trên lý thuyết tập mờ, số mờ và các toán tử kết nhập theo thời gian [23, 29, 31, 32, 91]. Tuy nhiên, khi đặt trong yêu cầu xử lý trực tiếp thông tin ngôn ngữ, các mô hình này vẫn còn một số hạn chế. Thứ nhất, kết quả tính toán thường phụ thuộc vào lựa chọn hàm thuộc, tham số mờ hóa, phương pháp giải mờ hoặc xấp xỉ, làm suy giảm tính toàn vẹn ngữ nghĩa của dữ liệu đầu vào [32, 23]. Thứ hai, cơ chế tích hợp dữ liệu lịch sử thường được xây dựng bằng trọng số thời gian ngoại sinh, nhưng chưa làm rõ đầy đủ sự suy giảm về mặt ngữ nghĩa của dữ liệu lịch sử [29, 32]. Thứ ba, trong nhiều trường hợp, mô hình chưa xử lý một cách nhất quán các tình huống dữ liệu khuyết thiếu, phương án mới xuất hiện, tiêu chí thay đổi hoặc hội đồng chuyên gia không ổn định qua các thời điểm đánh giá [23, 29]. Những hạn chế này cho thấy nhu cầu phát triển một mô hình MCGDM động có cơ sở ngữ nghĩa chặt chẽ hơn, do đó luận án đã lựa chọn HA làm nền tảng biểu diễn và định lượng thông tin ngôn ngữ.

Trong bối cảnh đó, HA cung cấp cơ sở toán học phù hợp để tiếp cận yêu cầu trên [52, 53]. Trong HA, miền giá trị của biến ngôn ngữ được tổ chức như một cấu trúc có trật tự ngữ nghĩa; các phần tử sinh và gia tử cho phép hình thành các từ ngôn ngữ mới; độ đo tính mờ và SQM cho phép ánh xạ từ ngôn ngữ sang miền số mà vẫn bảo toàn quan hệ thứ tự và nội dung ngữ nghĩa cơ bản của từ [52, 53, 54]. Trên nền tảng đó, mô hình 4-tuple ngữ nghĩa cho phép biểu diễn đồng thời nhãn ngôn ngữ, khoảng tương tự ngữ nghĩa, giá trị định lượng ngữ nghĩa và giá trị tham chiếu của từ [54]. Đây là cơ sở quan trọng để xây dựng các phép kết nhập ngôn ngữ, xử lý sự dịch chuyển ngữ nghĩa của đánh giá và thiết lập cơ chế suy giảm ngữ nghĩa cho dữ liệu lịch sử trong bài toán MCGDM động.

Mặc dù HA đã được ứng dụng trong một số mô hình MCDM và MCGDM ngôn ngữ, nhưng phần lớn các mô hình hiện có chủ yếu được phát triển trong bối cảnh tĩnh [3, 33, 54, 89]. Trong các mô hình đó, dữ liệu đánh giá thường được giả định là ổn định tại một thời điểm; trọng số của tiêu chí, tập phương án và tập chuyên gia thường được xác định trước; còn các phép kết nhập trên thang đo 4-tuple ngữ nghĩa chủ yếu dựa trên trung bình số học hoặc trung bình có trọng số. Cách tiếp cận này có ý nghĩa trong bài toán tĩnh, nhưng chưa đủ để mô hình hóa đầy đủ sự biến đổi của thông tin theo thời gian, chưa khai thác đầy đủ cấu trúc ngữ nghĩa – đại số của HA trong kết nhập nhiều giai đoạn, và chưa hình thành cơ chế tích hợp dữ liệu lịch sử theo hướng suy giảm ngữ nghĩa. Do đó, việc mở rộng MCGDM dựa trên HA sang bối cảnh động là bước phát triển cần thiết để hoàn thiện khung phương pháp luận của luận án.

Xuất phát từ các phân tích trên, Chương 4 phát triển khung phương pháp xây dựng mô hình MCGDM động trong môi trường ngôn ngữ dựa trên HA với ba trọng tâm chính. Thứ nhất, phát triển phép kết nhập ngôn ngữ nhị phân mới trên thang đo 4-tuple ngữ nghĩa, nhằm kết nhập các đánh giá ngôn ngữ trong điều kiện cần bảo toàn quan hệ ngữ nghĩa và một số tính chất đại số cơ bản. Thứ hai, xây dựng phép kết nhập với dữ liệu khuyết thiếu và cơ chế suy giảm ngữ nghĩa để tích hợp dữ liệu lịch sử với đánh giá hiện tại, qua đó phản ánh hợp lý sự giảm dần ảnh hưởng của thông tin quá khứ mà không làm đứt gãy cấu trúc ngữ nghĩa của đánh giá. Thứ ba, đề xuất mô hình hai giai đoạn L-FAHP - L-FTOPSIS động dựa trên HA, trong đó L-FAHP - HA được sử dụng để xác định trọng số của tiêu chí trong môi trường ngôn ngữ, còn L-FTOPSIS - HA được sử dụng để kết nhập, đánh giá và xếp hạng phương án qua nhiều thời điểm. Mô hình đề xuất được kiểm chứng, phân tích độ nhạy, đối sánh với mô hình liên quan thông qua bài toán đánh giá mức độ chuyển đổi số của 10 SMEs tại Việt Nam theo 4 tiêu chí chính và 14 tiêu chí con qua 3 thời điểm, nhằm xem xét khả năng bảo toàn ngữ nghĩa, độ ổn định của thứ hạng và mức độ thích ứng của mô hình trước sự biến đổi của dữ liệu theo thời gian.

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

Một trong những cách diễn giải ngữ nghĩa quan trọng của từ ngôn ngữ là biểu diễn chúng dưới dạng tập mờ. Trên cơ cấu trúc HA, Nguyễn Cát Hồ và cộng sự đã đề xuất khái niệm khung nhận thức ngôn ngữ (Linguistic Frame of Cognition – LFoC), về bản chất có thể xem như một thang đo ngôn ngữ [55]. Từ đó, ngữ nghĩa của các từ ngôn ngữ trong LFoC được biểu diễn thông qua ngữ nghĩa đơn thể hạt dưới dạng tập mờ tam giác.

Trong chương này, để xây dựng ngữ nghĩa của các từ ngôn ngữ trên thang đo ngôn ngữ L_k theo cách tiếp cận phân hoạch bằng các tập mờ tam giác, luận án sử dụng một quy trình sinh ngữ nghĩa dựa trên các tính chất đại số của HA. Quy trình này bảo đảm đồng thời ba yêu cầu cơ bản: bảo toàn thứ tự ngữ nghĩa của các từ ngôn ngữ, phản ánh đúng sự biến đổi ngữ nghĩa do tác động của gia tử, và duy trì tính nhất quán định lượng thông qua hàm SQM của HA. Quy trình này được gọi là phân hoạch ngữ nghĩa tập mờ tam giác, được mô tả như sau:

Xét mỗi từ ngôn ngữ $x_i \in X_k$, với $X_k = \{x_1, x_2, \dots, x_n\}$. Ngữ nghĩa tập mờ của x_i được biểu diễn bởi một số mờ tam giác (l, m, u) , trong đó các tham số l, m, u được xác định như sau:

- (i) Nếu $i = 1$ thì $l = 0, m = 0, u = \vartheta(x_2)$;
- (ii) Nếu $i = n$ thì $l = \vartheta(x_{n-1}), m = 1, u = 1$;
- (iii) Nếu $1 < i < n$ thì $l = \vartheta(x_{i-1}), m = \vartheta(x_i), u = \vartheta(x_{i+1})$.

Trong đó, $\vartheta(x_i)$ là giá trị định lượng ngữ nghĩa của từ ngôn ngữ $x_i \in X_k$

Như vậy, tập các tập mờ tương ứng với các từ ngôn ngữ $x_i \in X_k$ trên thang đo L_k tạo thành một phân hoạch mờ của miền ngữ nghĩa được sinh bởi cấu trúc HA. Gọi F_{L_k} là tập các tập mờ gắn với các từ trong thang đo L_k .

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, $\ddot{x} = (x, \vartheta(x), S_k(x), r_x) \in L_k$, ánh xạ $f_z(\ddot{x})$ được xác định bởi,

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

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

Ví dụ 4.1. Xét một ComHA, $AX^* = (X, G, \epsilon, H, \sigma, \phi, \leq)$, trong đó $G = \{\text{Unimportant (U), Important (I)}\}$, $\epsilon = \{0, w, 1\} = \{\text{Extremely Unimportant (EU), Medium (M), Extremely Important (EI)}\}$, $H = \{\text{Little (L), Very (V)}\}$, với $\mu(L) = 0.528, f\eta(U) = 0.489, \mu(V) = 1 - \mu(L) = 1 - 0.528 = 0.472$.

Với $k = 2$, ta có $X_2 = \{EU < VU < U < LU < M < LI < I < VI < EI\}$.

(i) Áp dụng Mệnh đề 1.1, tính độ đo tính mờ của các từ ngôn ngữ $x_i \in X_2$ được xác định như sau:

$$f_{\mathcal{M}}(EU) = 0;$$

$$f_{\mathcal{M}}(VU) = \mu(V) \times f_{\mathcal{M}}(U) = 0.472 \times 0.489 = 0.231;$$

$$f_{\mathcal{M}}(U) = 0.489;$$

$$f_{\mathcal{M}}(LU) = \mu(L) \times f_{\mathcal{M}}(U) = 0.528 \times 0.489 = 0.258;$$

$$f_{\mathcal{M}}(M) = 0;$$

$$f_{\mathcal{M}}(LI) = \mu(L) \times f_{\mathcal{M}}(I) = 0.528 \times 0.511 = 0.270;$$

$$f_{\mathcal{M}}(I) = 1 - f_{\mathcal{M}}(U) = 1 - 0.489 = 0.511;$$

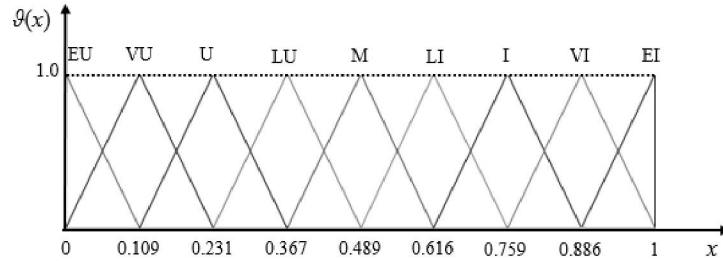
$$f_{\mathcal{M}}(VI) = \mu(V) \times f_{\mathcal{M}}(I) = 0.472 \times 0.511 = 0.241;$$

$$f_{\mathcal{M}}(EI) = 0.$$

(ii) Tiếp theo áp dụng các Định nghĩa 1.5 đến 1.8 cùng các Mệnh đề 1.1 và 1.2, giá trị định lượng ngữ nghĩa của các từ ngôn ngữ trong X_2 được xác định là:

$$\mathcal{Q}(EU) = 0, \mathcal{Q}(VU) = 0.109, \mathcal{Q}(U) = 0.231, \mathcal{Q}(LU) = 0.367, \mathcal{Q}(M) = 0.489, \mathcal{Q}(LI) = 0.616, \mathcal{Q}(I) = 0.759, \mathcal{Q}(VI) = 0.886, \mathcal{Q}(EI) = 1.$$

(iii) Từ các giá trị định lượng ngữ nghĩa này, TFSs tương ứng với các từ ngôn ngữ $x_i \in X_2$ được xác định theo công thức (4.1), qua đó hình thành tập F_{L_2} như minh họa trong Hình 4.1:



Hình 4.1. TFSs của các từ ngôn ngữ $x \in X_2$ của F_{L_2}

Kết quả cho thấy việc biểu diễn các từ ngôn ngữ $x \in X_k$ bằng TFSs dựa trên HA đã tạo cơ sở hình thức cho việc xử lý và tính toán với thông tin ngôn ngữ. Trên cơ sở đó, vấn đề tiếp theo là xây dựng các phép kết nhập phù hợp để tổng hợp nhiều đánh giá ngôn ngữ mà vẫn bảo toàn được trật tự ngữ nghĩa, tính nhất quán định lượng và khả năng diễn giải của kết quả. Vì vậy, phần tiếp theo của luận án sẽ trình bày các phép kết nhập ngôn ngữ trên thang đo 4-tuple ngữ nghĩa, kể cả trong trường hợp dữ liệu khuyết thiếu, làm nền tảng cho quá trình tổng hợp thông tin trong mô hình MCGDM động theo hướng tiếp cận HA.

4.3. Các phép kết nhập ngôn ngữ

Từ các ma trận so sánh cặp và ma trận đánh giá các phương án do các chuyên gia cung cấp, việc xác định trọng số của tiêu chí và xếp hạng các phương án đều đòi hỏi một bước quan trọng là kết nhập các đánh giá ngôn ngữ. Xuất phát từ tính chất của phép cộng trên tập số thực, luận án đề xuất một phép kết nhập ngôn ngữ nhị phân trên thang đo 4-tuple ngữ nghĩa L_k , nhằm bảo đảm kết quả tổng hợp vừa phù hợp với trực giác ngữ nghĩa, vừa duy trì được các tính chất đại số cần thiết cho bài toán MCGDM động.

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

Ta nhận thấy rằng, với hai số thực bất kỳ a và b , tổng $a + b$ thỏa mãn ba trường hợp cơ bản sau:

- (i) Nếu $a, b \geq 0$ thì $a + b \geq a$ và $a + b \geq b$;
- (ii) Nếu $a, b \leq 0$ thì $a + b \leq a$ và $a + b \leq b$;
- (iii) Nếu $a \leq 0$ và $b \geq 0$ thì $a \leq a + b \leq b$ hoặc nếu $b \leq 0$ và $a \geq 0$ thì $b \leq a + b \leq a$.

Trên cơ sở đó, phép kết nhập ngôn ngữ nhị phân được xây dựng bằng cách kết hợp đặc trưng âm - dương của miền ngữ nghĩa HA với tư tưởng của toán tử uninorm trong logic mờ, ký hiệu là: $\oplus$.

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.

Tương đương, bằng cách đưa vào hàm dấu $\zeta g\eta(\cdot)$, được định nghĩa trên các lỗi của ngữ nghĩa định lượng là:

$$\bar{r}_x = \vartheta(x) - r_w, \bar{r}_y = \vartheta(y) - r_w, \bar{r}_z \in [-1, 0, +1],$$

trong đó $r_w = \vartheta(w)$, thì ba trường hợp trên có thể diễn giải theo dấu của $\bar{r}_x$ và $\bar{r}_y$. Cụ thể, trường hợp 1 tương ứng với $\zeta g\eta(\bar{r}_x) = \zeta g\eta(\bar{r}_y) = -1$; Trường hợp 2 tương ứng với $\zeta g\eta(\bar{r}_x) = \zeta g\eta(\bar{r}_y) \in \{0, +1\}$; và Trường hợp 3 bao gồm tất cả các giá trị còn lại, ít nhất một trong hai giá trị $\zeta g\eta(\bar{r}_x), \zeta g\eta(\bar{r}_y)$ là (-1) và giá trị còn lại thuộc $\{0, +1\}$. Do đó, ba nhánh của công thức (4.2) tạo thành một phân hoạch đầy đủ và loại trừ lẫn nhau của $L_k \times L_k$, bảo đảm toán tử $\oplus$ được xác định tốt cho mọi cặp từ ngôn ngữ.

Kết quả của phép kết nhập $\bar{x} \oplus \bar{y}$ là một từ ngôn ngữ $z \in X_k$, sao cho $r_z \in S_k(z) \subseteq L_k$, còn $\vartheta(z)$ biểu diễn giá trị định lượng ngữ nghĩa tương ứng của từ z . Nhờ vậy, phép kết nhập được xây dựng không chỉ bảo toàn tính toán trên miền định lượng, mà còn duy trì được cấu trúc ngữ nghĩa của thang đo ngôn ngữ trong môi trường HA.

Ví dụ 4.2. Xét lại ComHA trong Ví dụ 4.1, sử dụng các Định nghĩa 1.9 - 1.12 và Mệnh đề 1.3 để tính các khoảng tương tự ngữ nghĩa $S_2(x)$ của mỗi từ $x \in X_2$ trên thang đo ngôn ngữ $L_2(\bar{C})$ như sau:

$$\begin{aligned} S_{(2)}(EU) &= [0.1, 0.146], S_{(2)}(VU) = [0.146, 0.256], S_{(2)}(U) = [0.256, 0.366], \\ S_{(2)}(LU) &= [0.366, 0.488], S_{(2)}(M) = [0.488, 0.594], S_{(2)}(LI) = [0.594, 0.722], \\ S_{(2)}(I) &= [0.722, 0.837], S_{(2)}(VI) = [0.837, 0.952], S_{(2)}(EI) = [0.952, 1]. \end{aligned}$$

Áp dụng các Định nghĩa 1.9 đến 1.12 và Mệnh đề 1.3 thu được một thang đo 4-tuple ngữ nghĩa $L_2(\bar{C})$ như trong Bảng 4.1 dưới đây:

Bảng 4.1. Thang đo ngôn ngữ 4-tuple ngữ nghĩa $L_2(\bar{C})$

Từ ngôn ngữ	4-tuple ngữ nghĩa
EU	(EU, 0, [0, 0.051), 0)
VU	(VU, 0.109, [0.051, 0.173), 0.109)
U	(U, 0.231, [0.173, 0.295), 0.231)
LU	(LU, 0.367, [0.295, 0.431), 0.367)
M	(M, 0.489, [0.431, 0.549), 0.489)
LI	(LI, 0.616, [0.549, 0.692), 0.616)
I	(I, 0.759, [0.692, 0.819), 0.759)
VI	(VI, 0.886, [0.819, 0.946), 0.886)
EI	(EI, 1, [0.946, 1], 1)

Trên cơ sở đó, xét một số trường hợp minh họa cho phép kết nhập ngôn ngữ nhị phân theo Định nghĩa 4.2.

(i) Với $\check{x} = (U, 0.231, [0.173, 0.295), 0.231)$ và $\check{y} = (LU, 0.367, [0.295, 0.431), 0.367)$, ta có $x < w, y < w$. Do đó, áp dụng biểu thức nhánh thứ 1 của công thức (4.2):

$$r_z = fm(g^-) \times T\left(\frac{r_x}{fm(g^-)}, \frac{r_y}{fm(g^-)}\right) = 0.489 \times \frac{0.231}{0.489} \times \frac{0.367}{0.489}$$

$$= 0.173 \in S_{(2)}(U) = [0.173, 0.295), \text{ nên kết quả thuộc về từ ngôn ngữ } U.$$

Suy ra, $\check{z} = \check{x} \oplus \check{y} = (U, 0.231, [0.173, 0.295), 0.173)$.

(ii) Với $\check{x} = (I, 0.759, [0.692, 0.819), 0.759)$, và $\check{y} = (VI, 0.886, [0.819, 0.946), 0.886)$, ta có $x > w, y > w$. Do đó, áp dụng biểu thức ở nhánh thứ 2 của công thức (4.2):

$$r_z = 0.489 + 0.511 \times S\left(\frac{0.27}{0.511}, \frac{0.397}{0.511}\right) =$$

$$= 0.489 + 0.511 \times \frac{0.27}{0.511} + 0.511 \times \frac{0.397}{0.511} - 0.511 \times \frac{0.27}{0.511} \times \frac{0.397}{0.511}$$

$$= 0.946 \in S_{(2)}(EI) = [0.946, 1]$$

Suy ra, $\check{z} = \check{x} \oplus \check{y} = (EI, 1, [0.946, 1], 0.946)$.

(iii) Với $\check{x} = (VU, 0.109, [0.051, 0.173), 0.109)$ và $\check{y} = (LI, 0.616, [0.549, 0.692), 0.616)$, ta có $x < w < y$. Khi đó, áp dụng biểu thức ở nhánh thứ 3 của công thức (4.2):

$$r_z = \frac{fm(LI)}{fm(VU) + fm(LI)} \times \vartheta(VU) + \frac{fm(VU)}{fm(VU) + fm(LI)} \times \vartheta(LI)$$

$$= \frac{0.231}{0.231 + 0.270} \times 0.109 + \frac{0.270}{0.231 + 0.270} \times 0.616$$

$= 0.342 \in S_{(2)}(LU) = [0.295, 0.431]$, nên kết quả thuộc từ ngôn ngữ LU .

Suy ra, $\tilde{z} = \tilde{x} \oplus \tilde{y} = (LU, 0.367, [0.295, 0.431], 0.342)$.

Các kết quả trên cho thấy phép kết nhập đề xuất luôn đưa kết quả trở lại miền biểu diễn 4-tuple ngữ nghĩa của thang đo $L_2(\tilde{C})$, đồng thời bảo toàn được sự nhất quán giữa giá trị tham chiếu r_z , khoảng tương tự ngữ nghĩa $S_2(z)$ và từ ngôn ngữ kết quả. Nhờ đó, phép toán không chỉ có ý nghĩa tính toán trên miền định lượng mà còn duy trì được khả năng diễn giải ngôn ngữ trong môi trường HA.

Mệnh đề 4.1. Với mọi $\tilde{x}, \tilde{y} \in L_k$, phép kết nhập $\oplus$ được định nghĩa trong Định nghĩa 4.2 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 . Cụ thể, nếu $x, y, z \geq w$ hoặc $x, y, z \leq w$ thì $(\tilde{x} \oplus \tilde{y}) \oplus \tilde{z} = \tilde{x} \oplus (\tilde{y} \oplus \tilde{z})$.

Chứng minh:

Theo các Định nghĩa 1.3 - 1.7, ta có $w = fm(g^-) = \theta$ và $fm(g^-) = 1 - \theta$. Gọi $r_x = \vartheta(x)$, $r_y = \vartheta(y)$ và $r_w = \vartheta(w)$. Do biểu diễn 4-tuple ngữ nghĩa và ánh xạ định lượng ngữ nghĩa đều bảo toàn thứ tự ngữ nghĩa, việc chứng minh các tính chất của phép kết nhập $\oplus$ trên thang đo ngôn ngữ L_k được quy về chứng minh các tính chất tương ứng của thành phần số thực $r_z = F(r_x, r_y)$ được xác định bởi công thức (4.2), như sau:

(i) Tính đóng trên thang đo ngôn ngữ L_k

Để chứng minh phép kết nhập có tính đóng trên thang đo L_k , ta chỉ cần chỉ ra rằng $r_z \in [0, 1]$, trong đó r_z được tính theo công thức (4.2), có ba trường hợp

Trường hợp 1: Với $x, y < w$, khi đó r_z được tính theo nhánh thứ 1 trong công thức (4.2). r_z được xác định thông qua hàm $T: [0, 1]^2 \rightarrow [0, 1]$. Vì T ánh xạ mọi cặp giá trị trong $[0, 1]^2$ vào $[0, 1]$.

Suy ra, $r_z \in [0, 1]$. Điều này chứng minh rằng r_z nằm trong khoảng $[0, 1]$.

Trường hợp 2: Với $r_x, r_y \geq w$, khi đó r_z được tính theo nhánh 2 trong công thức (4.2),

$$r_z = \theta + (1 - \theta) S\left(\frac{r_x - \theta}{1 - \theta}, \frac{r_y - \theta}{1 - \theta}\right),$$

trong đó, $S: [0,1] \times [0,1] \rightarrow [0,1]$ là một toán tử kết nhập dạng uninorm, ánh xạ mọi cặp giá trị trong $[0, 1]^2$ về một giá trị trong $[0, 1]$. Vì $0 \leq r_z \leq 1$, ta có $0 \leq \frac{r_x - \theta}{1 - \theta} \leq 1$. Tương tự, $0 \leq \frac{r_y - \theta}{1 - \theta} \leq 1$. Vì miền xác định của S đều nằm trong khoảng $[0, 1]$, suy ra: $0 \leq S\left(\frac{r_x - \theta}{1 - \theta}, \frac{r_y - \theta}{1 - \theta}\right) \leq 1$. Nhân với số không âm $(1 - \theta)$ và cộng thêm θ (với $0 \leq \theta \leq 1$),

suy ra, $0 \leq \theta + (1 - \theta) \times 0 \leq r_z \leq \theta + (1 - \theta) \times 1 \leq 1$.

Do đó, $r_z \in [0, 1]$, điều cần chứng minh.

Trường hợp 3: Với $x > w > y$ hoặc $x < w < y$ (trường hợp hỗn hợp), khi đó r_z được tính theo nhánh thứ 3 trong công thức (4.2).

Ta có $fm(x) > 0$ và $fm(y) > 0$, $\Rightarrow fm(x) + fm(y) > 0$, đồng thời:

$$\alpha = \frac{fm(x)}{fm(x) + fm(y)}, \beta = \frac{fm(y)}{fm(x) + fm(y)}.$$

Khi đó $\alpha + \beta = 1$. Công thức có thể viết lại thành: $r_z = \alpha r_x + \beta r_y$, nghĩa là r_z là một tổ hợp lồi của r_x và r_y . Vì $r_x, r_y \in [0, 1]$, và $\alpha, \beta > 0$ với $\alpha + \beta = 1$, các tính chất cơ bản của tổ hợp lồi suy ra: $\min(r_x, r_y) \leq r_z \leq \max(r_x, r_y)$. Vì cả hai đầu mút đều thuộc $[0,1]$ nên $r_z \in [0, 1]$.

Suy ra, $r_z \in [0, 1]$ điều cần chứng minh.

(ii) *Chứng minh tính giao hoán của phép kết nhập*

Xét công thức (4.2), giá trị tham chiếu $r_x \oplus_y$ được xác định theo ba trường hợp tùy thuộc vào vị trí của x và y so với phần tử trung hòa w .

Trường hợp 1: $x, y < w$. Khi đó, được tính theo nhánh thứ 1 của công thức (4.2),

ta có

$$r_x \oplus_y = r_w T\left(\frac{r_x}{r_w}, \frac{r_y}{r_w}\right),$$

trong đó, $T(a, b) = a \cdot b$ là t-norm dạng tích. Do T có tính giao hoán, nên

$$T\left(\frac{r_x}{r_w}, \frac{r_y}{r_w}\right) = T\left(\frac{r_y}{r_w}, \frac{r_x}{r_w}\right),$$

suy ra,

$$r_x \oplus y = r_y \oplus x.$$

Vì vậy, phép kết nhập có tính giao hoán trong trường hợp này.

Trường hợp 2: $x, y \geq w$. Đặt

$$u = \frac{r_x - r_w}{1 - r_w}, v = \frac{r_y - r_w}{1 - r_w}.$$

Khi đó, được tính theo nhánh thứ hai của công thức (4.2),

$$r_x \oplus y = r_w + (1 - r_w) S(u, v),$$

trong đó $S(a, b) = a + b - a \times b$ là t-conorm dạng tổng xác suất. Vì S có tính giao hoán, nên $S(u, v) = S(v, u)$, do đó việc hoán đổi vị trí của r_x và r_y không làm thay đổi giá trị $r_x \oplus y$. Suy ra

$$r_x \oplus y = r_y \oplus x.$$

Do đó, phép kết nhập là giao hoán trong trường hợp này.

Trường hợp 3: x và y nằm ở hai phía khác nhau của w . Không mất tính tổng quát, giả sử $x < w < y$. Khi đó, được tính theo nhánh thứ ba của công thức (4.2),

$$r_x \oplus y = \frac{fm(y)}{fm(x) + fm(y)} r_x + \frac{fm(x)}{fm(x) + fm(y)} r_y$$

Nếu hoán đổi x và y , ta được:

$$r_y \oplus x = \frac{fm(x)}{fm(x) + fm(y)} r_y + \frac{fm(y)}{fm(x) + fm(y)} r_x$$

Hai biểu thức trên là tương đương về mặt đại số, do cùng có mẫu số $fm(x) + fm(y)$ và chỉ khác nhau ở thứ tự các hạng tử. Vì vậy,

$$r_x \oplus y = r_y \oplus x.$$

Suy ra phép kết nhập cũng có tính giao hoán trong trường hợp hỗn hợp.

Tổng hợp cả ba trường hợp, phép kết nhập $\oplus$ có tính giao hoán trên L_k .

(iii) Chứng minh tính đơn điệu của phép kết nhập

Cần chứng minh rằng nếu $\check{x}_1 \leq \check{x}_2$ thì $\check{x}_1 \oplus \check{y} \leq \check{x}_2 \oplus \check{y}$.

Do ánh xạ định lượng ngữ nghĩa ϑ bảo toàn thứ tự, chỉ cần chứng minh rằng nếu $r_{x_1} \leq r_{x_2}$ thì $r_{x_1} \oplus y \leq r_{x_2} \oplus y$. Xét ba trường hợp của công thức (4.2):

Trường hợp 1: $x, y < w$ (cả hai nằm phía âm của w). Khi đó, nhánh thứ 1 của công thức (4.2) được áp dụng:

$$r_{x \oplus y} = r_w T\left(\frac{r_x}{r_w}, \frac{r_y}{r_w}\right)$$

trong đó, $T(a, b) = a \cdot b$ là t -norm dạng tích (product t -norm). Với y cố định, hàm $r_x \mapsto T\left(\frac{r_x}{r_w}, \frac{r_y}{r_w}\right)$ là hàm tăng trên $[0, r_w]$.

Do đó, nếu $r_{x_1} \leq r_{x_2}$ thì $r_{x_1} \oplus y \leq r_{x_2} \oplus y$. Lập luận tương tự cũng áp dụng cho đối số thứ hai nhờ tính đối xứng của T .

Trường hợp 2: $x, y \geq w$ (cả hai nằm phía không âm của w): Khi đó, nhánh thứ 2 của công thức (4.2) được sử dụng. Xác định các giá trị đã chuẩn hóa:

$$u_x = \frac{r_x - r_w}{1 - r_w}, u_y = \frac{r_y - r_w}{1 - r_w},$$

Suy ra: $r_{x \oplus y} = r_w + (1 - r_w) S(u_x, u_y)$, trong đó, $S(a, b) = a + b - ab$ là tổng xác suất (t -conorm). Vì S là hàm tăng theo từng đối số trên khoảng $[0, 1]$, nên nếu $r_{x_1} \leq r_{x_2}$ thì $r_{x_1} \oplus y \leq r_{x_2} \oplus y$.

Do đó, tính đơn điệu được đảm bảo trong trường hợp này.

Trường hợp 3: x và y nằm về hai phía khác nhau so với w (hoặc một trong hai phần tử bằng w) (trường hợp hỗn hợp): Khi đó nhánh thứ ba của công thức (4.2) được áp dụng và có dạng như sau:

$$r_{x \oplus y} = \frac{fm(x) r_x + fm(y) r_y}{fm(x) + fm(y)},$$

trong đó $fm(\cdot)$ là độ đo tính mờ đã được định nghĩa trước. Với y cố định, biểu thức trên là một hàm tuyến tính (lồi) theo r_x với hệ số

$$\frac{fm(x)}{fm(x) + fm(y)} > 0.$$

Vì vậy, $r_{x \oplus y}$ tăng theo r_x . Tương tự $r_{x \oplus y}$ cũng tăng theo r_y do tính đối xứng

Tóm lại, Tổng hợp cả ba trường hợp và sử dụng tính chất bảo toàn thứ tự của ngữ nghĩa, suy ra phép kết nhập là đơn điệu theo từng đối số trên miền L_k .

(iv) *Tồn tại phần tử trung hòa* $w \in L_k$

Gọi w là từ ngôn ngữ trung hoà trên thang đo L_k , với $r_w = \vartheta(w)$ và $fm(w) = 0$.

Trường hợp 1: $x \geq w$. Khi đó, cặp (x, w) được tính theo nhánh thứ 2 của công thức (4.2). Vì $u_w = \frac{r_w - r_w}{1 - r_w} = 0$, nên:

$$r_x \oplus w = r_w + (1 - r_w) S\left(\frac{r_x - r_w}{1 - r_w}, 0\right)$$

Do $S(a, b) = a$, suy ra

$$r_x \oplus w = r_w + (1 - r_w) \frac{r_x - r_w}{1 - r_w} = r_x$$

Trường hợp 2: $x \leq w$. Khi đó cặp (x, w) được tính theo nhánh hợp hỗn hợp của công thức (4.2). Vì $fm(w) = 0$, ta có:

$$r_x \oplus w = \frac{fm(x) r_x + fm(w) r_w}{fm(x) + fm(w)} = \frac{fm(x) r_x}{fm(x)} = r_x.$$

Như vậy, trong mọi trường hợp đều có: $r_x \oplus w = r_x$.

Kết hợp với tính giao hoán của phép kết nhập, suy ra

$$w \oplus x = x \oplus w = x.$$

Do đó, w là phần tử trung hòa (hay phần tử đơn vị) của phép kết nhập.

(v) *Phép kết nhập có tính kết hợp bộ phận khi các toán hạng cùng nằm một phía so với phần tử trung hòa w .*

Tính chất này chỉ được thỏa mãn trong trường hợp tất cả các toán hạng đều nằm cùng một phía của phần tử trung hòa w , tức là khi,

$$x, y, z \leq w \text{ hoặc } x, y, z \geq w,$$

Khi đó, cần chứng minh rằng:

$$(\ddot{x} \oplus \ddot{y}) \oplus \dot{z} = \ddot{x} \oplus (\ddot{y} \oplus \dot{z}).$$

Thật vậy, nếu cả ba từ ngôn ngữ đều nằm về phía âm của w , tức $x, y, z < w$, thì công thức (4.2) luôn vận hành trong nhánh thứ nhất. Khi đó, phép kết nhập được quy về phép hợp thành của t-norm dạng tích T trên các giá trị định lượng ngữ nghĩa đã chuẩn hóa. Vì T là phép toán kết hợp trên $[0, 1]$, nên ta có,

$$r_{(x \oplus y) \oplus z} = r_{x \oplus (y \oplus z)}.$$

Tương tự, nếu cả ba từ ngôn ngữ đều nằm về phía không âm của w , tức $x, y, z \geq w$, thì công thức (4.2) luôn thực hiện trong nhánh thứ hai. Khi đó, phép kết nhập được quy về phép hợp thành của t-conorm S trên các giá trị đã chuẩn hóa. Do S cũng là phép toán kết hợp trên $[0, 1]$,

suy ra: $r_{(x \oplus y) \oplus z} = r_{x \oplus (y \oplus z)}$.

Vì biểu diễn 4-tuple ngữ nghĩa bảo toàn trật tự và ánh xạ ngược từ giá trị tham chiếu r_x về miền ngôn ngữ là xác định, nên từ đẳng thức trên miền số thực suy ra

$$(\ddot{x} \oplus \ddot{y}) \oplus \ddot{z} = \ddot{x} \oplus (\ddot{y} \oplus \ddot{z}).$$

Như vậy, phép kết nhập $\oplus$ có tính kết hợp bộ phận, nghĩa là chỉ kết hợp khi toàn bộ các toán hạng cùng nằm trong một nhánh của công thức (4.2), tức cùng ở phía âm hoặc cùng ở phía không âm của phần tử trung hòa w .

Cần lưu ý rằng tính kết hợp nói trên không còn được bảo đảm trên toàn thang đo L_k . Khi các toán hạng nằm ở hai phía khác nhau của w , công thức (4.2) phải chuyển sang nhánh hỗn hợp, khi đó phép kết nhập không còn quy về một phép toán kết hợp duy nhất như T hoặc S . Điều này phản ánh đúng đặc trưng ngữ nghĩa của thang đo 4-tuple trong môi trường HA, nơi phần tử trung hòa w đóng vai trò ranh giới giữa hai miền ngữ nghĩa âm và dương. Vì vậy, trong các trường hợp hỗn hợp, để bảo đảm tính xác định và khả năng tái lập của kết quả, các toán hạng được sắp xếp theo thứ tự tăng dần trước khi thực hiện phép kết nhập lặp. Quy ước này không làm thay đổi khả năng áp dụng của phép kết nhập trong thực tiễn, nhưng giúp duy trì sự nhất quán của quá trình tổng hợp.

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

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

Nói cách khác, phép kết nhập của n từ ngôn ngữ được xây dựng bằng cách áp dụng lặp lại phép kết nhập ngôn ngữ nhị phân $\oplus$ trong Định nghĩa 4.2 theo thứ tự đã sắp xếp. Sau mỗi bước kết nhập, kết quả trung gian vẫn là một 4-tuple ngữ nghĩa thuộc L_k , nhờ đó bảo đảm tính đóng và tính nhất quán ngữ nghĩa trong toàn bộ quá trình tổng hợp.

Định nghĩa 4.3 cho phép mở rộng tự nhiên phép kết nhập nhị phân trên thang đo L_k sang trường hợp tổng hợp nhiều đánh giá ngôn ngữ trong bài toán hay MCDM hay MCGDM. Đồng thời, là cơ sở để xây dựng các phép kết nhập trong bối cảnh dữ liệu động, khi tại một số thời điểm có thể xuất hiện các đánh giá bị thiếu hoặc không quan sát được. Trên cơ sở đó, phần tiếp theo sẽ trình bày phép kết nhập ngôn ngữ trong trường hợp dữ liệu khuyết thiếu, nhằm bảo đảm tính liên tục của quá trình tổng hợp thông tin trong môi trường MCGDM động.

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 hòa.

Sau khi xác định $r_{x_{agg}^{(\tau)}}$, các thành phần còn lại của 4-tuple ngữ nghĩa $\check{x}_{agg}^{(\tau)}$ 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.

Chứng minh:

(i). Chứng minh, tính bảo toàn miền giá trị: 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 $\vartheta(w)$, tức là:

$$\min\{s_{\tau-1}, \vartheta(w)\} \leq s_{\tau} \leq \max\{s_{\tau-1}, \vartheta(w)\}.$$

Từ công thức (4.6), suy ra:

$$s_{\tau} - \vartheta(w) = \lambda(s_{\tau-1} - \vartheta(w)).$$

Vì $\lambda \in [0, 1]$, ta xét hai trường hợp:

- Trường hợp 1: Nếu $s_{\tau-1} \geq \vartheta(w)$ (đánh giá lịch sử tốt hơn hoặc bằng mức trung bình).

Khi đó, $s_{\tau-1} - \vartheta(w) \geq 0$. Do, $0 \leq \lambda \leq 1$, suy ra $0 \leq \lambda(s_{\tau-1} - \vartheta(w)) \leq s_{\tau-1} - \vartheta(w)$. Cộng $\vartheta(w)$ vào cả ba vế, ta có: $\vartheta(w) \leq \vartheta(w) + \lambda(s_{\tau-1} - \vartheta(w)) \leq s_{\tau-1}$.

Suy ra, $\vartheta(w) \leq s_{\tau} \leq s_{\tau-1}$.

- *Trường hợp 2*: Nếu $s_{\tau-1} < \vartheta(w)$ (đánh giá lịch sử kém hơn mức trung bình).

Khi đó, $s_{\tau-1} - \vartheta(w) < 0$. Vì, $0 \leq \lambda \leq 1$, nên $(s_{\tau-1} - \vartheta(w)) \leq \lambda(s_{\tau-1} - \vartheta(w)) \leq 0$.

Cộng $\vartheta(w)$ vào các vế: $s_{\tau-1} \leq \vartheta(w) + (s_{\tau-1} - \vartheta(w)) \leq \vartheta(w)$.

Suy ra, $s_{\tau-1} \leq s_{\tau} \leq \vartheta(w)$.

Từ hai trường hợp trên, ta luôn có s_t nằm giữa s_{t-1} và $\vartheta(w)$.

Suy ra, $\min\{s_{\tau-1}, \vartheta(w)\} \leq s_{\tau} \leq \max\{s_{\tau-1}, \vartheta(w)\}$, điều phải chứng minh.

(ii). *Chứng minh, tính hội tụ về phân tử trung hòa*: Nếu dữ liệu bị khuyết thiếu liên tiếp trong n thời điểm và $(0 \leq \lambda < 1)$ thì $s_{\tau+n}$ sẽ hội tụ về $\vartheta(w)$ khi $n \rightarrow \infty$.

Giả sử tại thời điểm bắt đầu quá trình suy giảm t_0 , giá trị lịch sử là $s_0 = s_{t-1}$. Do không có dữ liệu mới trong các thời điểm tiếp theo, công thức suy giảm ngẫu nhiên được áp dụng lặp lại:

Tại thời điểm τ_1 (sau một bước suy giảm): $s_1 - \vartheta(w) = \lambda(s_0 - \vartheta(w))$

Tại thời điểm τ_2 (sau hai bước suy giảm): $s_2 - \vartheta(w) = \lambda(s_1 - \vartheta(w)) = \lambda[\lambda(s_0 - \vartheta(w))] = \lambda^2(s_0 - \vartheta(w))$. Bằng phương pháp quy nạp, tại thời điểm s_n , ta có:

$$s_n - \vartheta(w) = \lambda^n(s_0 - \vartheta(w))$$

Hay tương đương: $s_{\tau+n} = \vartheta(w) + \lambda^n(s_{\tau-1} - \vartheta(w))$

Lấy giới hạn khi $n \rightarrow \infty$: Do $\lambda \in [0, 1]$. Nếu $\lambda < 1$, nên $\lim_{n \rightarrow \infty} \lambda^n = 0$

Vì vậy, $\lim_{n \rightarrow \infty} s_{\tau+n} = \vartheta(w) + 0 \cdot (s_{\tau-1} - \vartheta(w)) = \vartheta(w)$, điều phải chứng minh.

Trường hợp $\lambda = 1$ cho $s_{\tau+n} = s_{\tau-1}$, không suy giảm ngẫu nhiên.

Trường hợp $\lambda = 0$ cho $s_{\tau+1} = \vartheta(w)$, quên hoàn toàn sau một bước.

Tính chất này đảm bảo hệ thống tự động quên hay suy giảm ngẫu nhiên theo thời gian và đưa đánh giá về trạng thái trung bình, tránh việc lưu giữ thiên kiến vĩnh viễn.

(iii). *Chứng minh, tính đơn điệu theo hệ số suy giảm*: với mọi $n \geq 1$, nếu $0 \leq \lambda_1 \leq \lambda_2 \leq 1$ thì $|s_{\tau+n}(\lambda_1) - \vartheta(w)| \leq |s_{\tau+n}(\lambda_2) - \vartheta(w)|$, tức là, λ càng nhỏ thì $s_{\tau+n}$ tiến về $\vartheta(w)$ càng nhanh.

Từ công thức suy giảm ngẫu nhiên (4.6) sau n bước,

$$s_{\tau+n} = \vartheta(w) + \lambda^n(s_{\tau-1} - \vartheta(w)),$$

suy ra, $s_{\tau+n}(\lambda) - \vartheta(w) = \lambda^n(s_{\tau-1} - \vartheta(w))$.

Lấy trị tuyệt đối hai vế, ta được: $|s_{\tau+n}(\lambda) - \vartheta(w)| = \lambda^n |s_{\tau-1} - \vartheta(w)|$.

Vì $0 \leq \lambda_1 \leq \lambda_2 \leq 1$ và $n \geq 1$, nên $\lambda_1^n \leq \lambda_2^n$.

Mặt khác, $|s_{\tau-1} - \vartheta(w)| \geq 0$.

Do đó, nhân cả hai vế với $|s_{\tau-1} - \vartheta(w)|$, suy ra

$$\lambda_1^n |s_{\tau-1} - \vartheta(w)| \leq \lambda_2^n |s_{\tau-1} - \vartheta(w)|.$$

Hay tương đương, $|s_{\tau+n}(\lambda_1) - \vartheta(w)| \leq |s_{\tau+n}(\lambda_2) - \vartheta(w)|$.

Vậy tính đơn điệu theo hệ số suy giảm được chứng minh. Kết quả này cho thấy λ càng nhỏ thì tốc độ suy giảm ngữ nghĩa của dữ liệu lịch sử càng lớn, và ngược lại, λ càng gần 1 thì ảnh hưởng của thông tin lịch sử được duy trì lâu hơn.

(iv). *Chứng minh tính ổn định tại điểm trung hòa (w):*

Cần chứng minh rằng nếu $s_{\tau-1} = \vartheta(w)$, thì với mọi $\lambda \in [0, 1]$, ta luôn có $s_\tau = \vartheta(w)$. Thật vậy, thay $s_{\tau-1} = \vartheta(w)$ vào công thức (4.6), ta được

$$s_\tau = \vartheta(w) + \lambda(\vartheta(w) - \vartheta(w)) = \vartheta(w) + \lambda \times 0 = \vartheta(w).$$

Vậy, $s_{\tau-1} = \vartheta(w)$ với mọi λ .

Suy ra, $s_{\tau-1} = \vartheta(w)$, điều phải chứng minh.

Ví dụ 4.3. Giả sử Doanh nghiệp A, tại thời điểm τ_2 được đánh giá mức độ sẵn sàng chuyển đổi số ở mức cao (High) với giá trị tham chiếu ngữ nghĩa $s_{\tau_2} = 0.781$. Chọn $\vartheta(w) = 0.566$, $\lambda = 0.7$. Nếu tại các thời điểm τ_3, τ_4, τ_5 , doanh nghiệp không tham gia đánh giá hoặc dữ liệu bị khuyết thiếu thì giá trị suy giảm ngữ nghĩa được tính như sau:

$$\text{Tại thời điểm } \tau_3: s_{\tau_3} = 0.566 + 0.7 \times (0.781 - 0.566) = 0.717;$$

$$\text{Tại thời điểm } \tau_4: s_{\tau_4} = 0.566 + 0.7 \times (0.7165 - 0.566) = 0.672;$$

$$\text{Tại thời điểm } \tau_5: s_{\tau_5} = 0.566 + 0.7 \times (0.672 - 0.566) = 0.640;$$

Như vậy, điểm số của doanh nghiệp giảm dần theo chuỗi:

$$s_{\tau_2} = 0.781 \rightarrow s_{\tau_3} = 0.717 \rightarrow s_{\tau_4} = 0.672 \rightarrow s_{\tau_5} = 0.640.$$

Kết quả này cho thấy khi dữ liệu bị khuyết thiếu liên tiếp trong nhiều thời điểm, ảnh hưởng của đánh giá lịch sử không bị loại bỏ ngay lập tức mà suy giảm dần một cách có kiểm soát theo thời gian, đồng thời tiệm cận về mốc trung hòa $\vartheta(w)$. Cơ chế này phù hợp với thực tế trong bài toán đánh giá chuyển đổi số, nơi năng lực và mức độ sẵn sàng của doanh nghiệp không bất biến mà có thể thay đổi theo thời gian dưới tác động của công nghệ, thị trường và môi trường quản trị.

Tóm lại, phép kết nhập với dữ liệu khuyết thiếu cung cấp một cơ chế suy giảm ngữ nghĩa nội tại, trong đó w đóng vai trò là mốc trung hòa và λ là tham số điều khiển tốc độ quên của dữ liệu lịch sử do chuyên gia xác định. Nhờ đó, mô hình vừa duy trì được dấu vết của dữ liệu quá khứ, vừa tránh việc bảo toàn tuyệt đối các đánh giá cũ khi thiếu vắng dữ liệu mới, qua đó làm cho quá trình xếp hạng trong bài toán MCGDM động trở nên hợp lý, ổn định và phản ánh đúng hơn sự biến thiên theo thời gian của đối tượng được đánh giá.

Trong thực tiễn, nhiều bài toán ra quyết định không diễn ra tại một thời điểm đơn lẻ mà cần được đánh giá định kỳ qua nhiều thời điểm kế tiếp, khi thông tin, đối tượng và bối cảnh ra quyết định có thể thay đổi theo thời gian [29, 23]. Dạng bài toán này được mô hình hóa dưới tên gọi ra MCGDM động [17].

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

Dựa trên các phép kết nhập ngôn ngữ nhị phân, phép kết nhập với dữ liệu khuyết thiếu và cơ chế suy giảm ngữ nghĩa để tích hợp dữ liệu lịch sử với dữ liệu hiện tại trên thang đo biểu diễn 4-tuple ngữ nghĩa; luận án đề xuất mô hình ra quyết định L-FAHP - L-FTOPSIS động dựa trên HA nhằm giải quyết bài toán MCGDM động trong môi trường ngôn ngữ, trong đó tập phương án, bộ tiêu chí và hội đồng chuyên gia có thể thay đổi theo thời gian và quan tâm đến dữ liệu lịch sử.

(i) Dữ liệu đầu vào:

Xét bài toán MCGDM động trong môi trường ngôn ngữ dựa trên HA. Giả sử quá trình ra quyết định được thực hiện tại T thời điểm rời rạc, ký hiệu $T = \{1, 2, \dots, t\}$, với $t \geq 1$.

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

(1) Xét ba tập hữu hạn sau: $\bar{A}^\tau = \{A_1^\tau, A_2^\tau, \dots, A_{m_\tau}^\tau\}$, ($m_\tau \geq 2$) là tập hữu hạn các phương án khả thi; $\bar{C}^\tau = \{C_1^\tau, C_2^\tau, \dots, C_{n_\tau}^\tau\}$, ($n_\tau \geq 2$) là tập hữu hạn các tiêu chí; và $\bar{D}^\tau = \{D_1^\tau, D_2^\tau, \dots, D_{h_\tau}^\tau\}$ ($h_\tau \geq 2$) là tập hữu hạn người ra quyết định (chuyên gia).

(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 $L_k(\bar{A}^\tau)$. Trong đó, thang đo $L_k(\bar{C}^\tau)$ được sử dụng để xác định trọng số của tiêu chí $C_j^\tau \in \bar{C}^\tau$; còn thang đo $L_k(\bar{A}^\tau)$ để đánh giá mức độ đáp ứng của các phương án $A_i^\tau \in \bar{A}^\tau$ theo từng tiêu chí $C_j^\tau \in \bar{C}^\tau$.

(3) 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 các tiêu chí, theo công thức (4.7):*

$$P_{D_e^\tau}^\tau = (x_{ij}^e(\tau))_{n_\tau \times n_\tau} = \begin{pmatrix} x_{11}^e(\tau) & x_{12}^e(\tau) & \cdots & x_{1n_\tau}^e(\tau) \\ & x_{22}^e(\tau) & \cdots & x_{2n_\tau}^e(\tau) \\ & & \ddots & \vdots \\ & & & x_{nn}^e(\tau) \end{pmatrix} \quad (4.7)$$

Trong đó, mỗi phần tử $x_{ij}^e(\tau) \in L_k(\bar{C}^\tau)$ biểu diễn mức độ quan trọng tương đối của tiêu chí C_i^τ so với tiêu chí C_j^τ theo đánh giá của chuyên gia D_e^τ . 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 $e = 1, 2, \dots, h_\tau$; $\tau = 1, 2, \dots, t$; $i, j = 1, 2, \dots, n_\tau$; $i \neq j$.

Ma trận $P_{D_e^\tau}^\tau$ được chấp nhận nếu đạt tỷ lệ nhất quán $CR_{HA}^e(\tau) < 0.10$; ngược lại, chuyên gia D_e^τ cần rà soát và điều chỉnh lại các so sánh cặp ngôn ngữ trong ma trận $P_{D_e^\tau}^\tau$ cho đến khi đạt yêu cầu nhất quán.

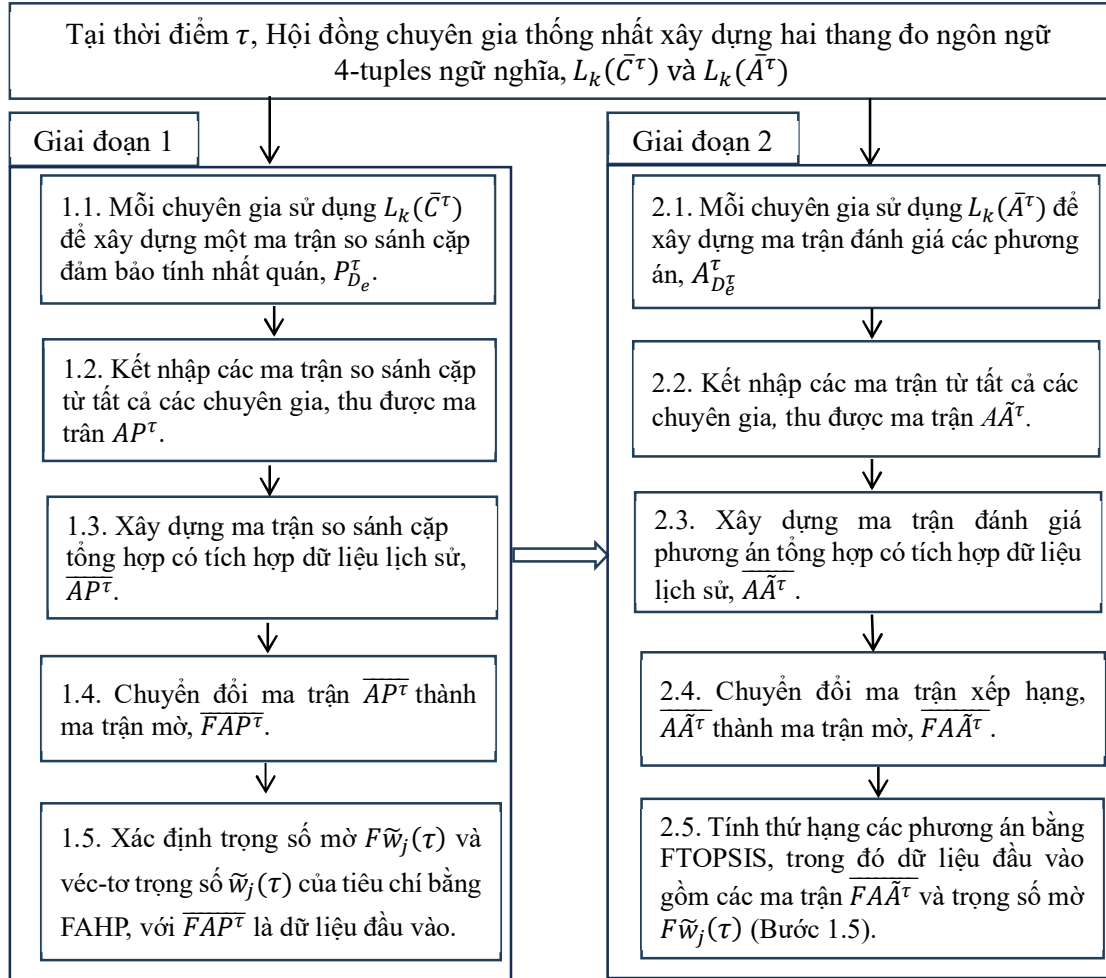
- *Xây dựng ma trận đánh giá các phương án theo công thức (4.8)*

$$A_{D_e^\tau}^\tau = (a_{ij}^e(\tau))_{m_\tau \times n_\tau} = \begin{pmatrix} a_{11}^e(\tau) & a_{12}^e(\tau) & \cdots & a_{1n_\tau}^e(\tau) \\ a_{21}^e(\tau) & a_{22}^e(\tau) & \cdots & a_{2n_\tau}^e(\tau) \\ \vdots & \vdots & \ddots & \vdots \\ a_{m_\tau 1}^e(\tau) & a_{m_\tau 2}^e(\tau) & \cdots & a_{m_\tau n_\tau}^e(\tau) \end{pmatrix} \quad (4.8)$$

Trong đó, $a_{ij}^e(\tau) \in L_k(\bar{A}^\tau)$ là đánh giá ngôn ngữ của chuyên gia $D_e^\tau \in \bar{D}^\tau$ đối với phương án $A_i^\tau \in \bar{A}^\tau$ theo tiêu chí $C_j^\tau \in \bar{C}^\tau$ tại thời điểm τ , với $e = 1, \dots, h_\tau$; $i = 1, \dots, m_\tau$; $j = 1, \dots, n_\tau$; $\tau = 1, 2, \dots, t$. Ma trận này phản ánh mức độ đáp ứng của từng phương án đối với từng tiêu chí trong bối cảnh đánh giá tại thời điểm τ .

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).



Hình 4.2. Quy trình MCGDM động trong môi trường ngôn ngữ dựa trên HA

Hình 4.2 trình bày khung tổng quát của mô hình L-FAHP - L-FTOPSIS động dựa trên HA. Quy trình ra quyết định này gồm hai giai đoạn: giai đoạn 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 phương pháp FAHP [19, 15], còn giai đoạn 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 phương pháp 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. Nhờ vậy, mô hình bảo đảm tính nhất quán ngữ nghĩa, nâng cao độ tin cậy và khả năng diễn giải của kết quả ra quyết định trong bối cảnh dữ liệu thay đổi theo thời gian.

Để thuận tiện cho việc trình bày các bước tính toán trong các mục tiếp theo, ký hiệu $P_{D_e}^\tau$ và $A_{D_e}^\tau$ lần lượt biểu thị ma trận so sánh cặp tiêu chí và ma trận đánh giá phương án do chuyên gia D_e cung cấp tại thời điểm τ . Các ma trận AP^τ và $A\tilde{A}^\tau$ lần lượt là kết quả kết nhập thông tin từ toàn bộ hội đồng chuyên gia tại thời điểm τ cho cả giai đoạn xác định trọng số của tiêu chí và đánh giá phương án. Tương ứng, $\overline{AP}^\tau$ và $\overline{A\tilde{A}}^\tau$ là các ma trận ngôn ngữ sau khi tích hợp thêm dữ liệu lịch sử từ thời điểm $\tau - 1$. Sau bước chuyển đổi ngữ nghĩa mờ thu được ma trận $\overline{FAP}^\tau$ và $\overline{FA\tilde{A}}^\tau$ là đầu vào trực tiếp trong các bước tính toán của mô hình L-FAHP-HA và L-FTOPSIS-HA. Cách ký hiệu này cho phép phân biệt rõ ba lớp xử lý thông tin trong mô hình, gồm: thông tin cá nhân chuyên gia, thông tin nhóm tại thời điểm hiện tại và thông tin nhóm sau khi tích hợp dữ liệu lịch sử.

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 sử dụng các từ ngôn ngữ $x \in X_k$ trên thang đo ngôn ngữ $L_k(\bar{C}^\tau)$ được trình bày trong Bảng 4.1 để xây dựng ma trận so sánh cặp ngôn ngữ giữa các tiêu chí, ký hiệu $P_{D_e}^\tau$, như trong công thức (4.7). Ma trận này phản ánh phán đoán của chuyên gia về mức độ quan trọng giữa từng cặp tiêu chí trong tập $\bar{C}^\tau$.

Để bảo đảm độ tin cậy của các phán đoán, ma trận $P_{D_e}^\tau$ 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 của [95] được tra trong Bảng 4.2.

Bảng 4.2. Giá trị chỉ số ngẫu nhiên ($RI(n_\tau)$)

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

Trong mục 3.2 của Chương 3, luận án đã đề xuất công thức tính chỉ số nhất quán (CI_{HA}^e) cho ma trận so sánh cặp ngôn ngữ $P_{D_e}^\tau$ do chuyên gia D_e cung cấp trong môi trường HA như công thức (4.10) dưới đây:

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

Trong đó, $\lambda_{\max}^e(\tau)$ là giá trị riêng lớn nhất của ma trận $P_{D_e}^\tau$ do chuyên gia D_e cung cấp tại thời điểm τ , n_τ là số lượng tiêu chí trong tập $\bar{C}^\tau$ tại thời điểm τ , còn $\vartheta(w)^\tau$ là giá trị định lượng ngữ nghĩa của từ trung hoà w trên thang đo ngôn ngữ $L_k(\bar{C}^\tau)$ tại thời điểm τ . Trong ví dụ thực nghiệm của Chương 4 sử dụng từ ngôn ngữ M là phần tử trung hòa của thang đo $L_2(\bar{C}^\tau)$ và $L_2(\bar{A}^\tau)$.

Để tính $CR_{HA}^e(\tau)$, trước hết ma trận ngôn ngữ $P_{D_e}^\tau$ được chuyển sang ma trận mờ ngữ nghĩa $FP_{D_e}^\tau$ 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 $P_{D_e}^\tau$ 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 $P_{D_e}^\tau$, ($e = 1, \dots, h_\tau$) 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 . Việc thực hiện kết nhập trực tiếp trên miền ngôn ngữ 4-tuple giúp bảo toàn cấu trúc thứ tự và ngữ nghĩa nội tại của các từ ngôn ngữ. Đây là cơ sở để mô hình tiếp tục tích hợp dữ liệu lịch sử trong bước tiếp theo mà vẫn duy trì được tính nhất quán ngữ nghĩa của thông tin đầu vào.

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: $\bar{C}_u^\tau = \bar{C}^{\tau-1} \cup \bar{C}^\tau$, 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, $\widehat{\oplus}$: 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 $\overline{AP^T}$ 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 $\overline{FAP^T}$ và được xác định theo công thức (4.14) dưới đây:

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

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 $i, j = 1, \dots, n_t^u, i \neq j$; hàm $fz(\cdot)$ là phép ánh xạ của tập mờ tam giác của một từ ngôn ngữ, như minh họa trong Hình 4.1.

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 tập mờ tam giác (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. Bảng 4.3 trình bày các TFSs của các từ ngôn ngữ dùng để so sánh mức độ quan trọng tương đối giữa hai tiêu chí trên thang đo $L_2(\bar{C}^T)$ trong môi trường HA. Các giá trị này được xây dựng nhất quán từ cấu trúc ngữ nghĩa của thang đo 4-tuple và được sử dụng thống nhất trong toàn bộ giai đoạn xác định trọng số.

Bảng 4.3. Các TFSs của các từ ngôn ngữ dùng để so sánh cặp giữa hai tiêu chí trên thang đo $L_2(\bar{C}^\tau)$ dựa trên HA

TT	Từ ngôn ngữ	TFS của các từ ngôn ngữ dựa trên HA	TFS nghịch đảo của từ ngôn ngữ dựa trên HA
1	EU	(0.1, 0.1, 0.198)	(1/0.198, 1/0.1, 1/0.1)
2	VU	(0.1, 0.198, 0.308)	(1/0.308, 1/0.198, 1/0.1)
3	U	(0.198, 0.308, 0.430)	(1/0.430, 1/0.308, 1/0.198)
4	LU	(0.308, 0.430, 0.540)	(1/0.540, 1/0.430, 1/0.308)
5	M	(0.430, 0.540, 0.654)	(1/0.654, 1/0.540, 1/0.430)
6	LI	(0.540, 0.654, 0.783)	(1/0.783, 1/0.654, 1/0.540)
7	I	(0.654, 0.783, 0.897)	(1/0.897, 1/0.783, 1/0.654)
8	VI	(0.783, 0.897, 1)	(1, 1/0.897, 1/0.783)
9	EI	(0.897, 1, 1)	(1, 1, 1/0.897)

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 [15].

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 mờ $\overline{FAP}^\tau$.

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, $\overline{\oplus}$: 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:

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

Giá trị $M_j(\tau)$ là đại diện số rõ của trọng số mờ $F\tilde{w}_j(\tau)$. Cuối cùng, 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 chuẩn hóa này là kết quả đầu ra của giai đoạn 1 và được sử dụng làm đầu vào cho giai đoạn 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 $L_k(\bar{A}^\tau)$ được trình bày trong Bảng 4.4, để xây dựng ma trận đánh giá ngôn ngữ cho các phương án, ký hiệu $\tilde{A}_{D_e^\tau}^\tau$, 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 . Việc sử dụng thang đo 4-tuple ngữ nghĩa dựa trên HA cho phép các chuyên gia biểu đạt đánh giá bằng ngôn ngữ tự nhiên, đồng thời vẫn bảo đảm khả năng định lượng ngữ nghĩa và tính toán ở các bước tiếp theo.

Bảng 4.4. Thang đo ngôn ngữ 4-tuple và TFS của các từ ngôn ngữ dựa trên HA dùng để xếp hạng các phương án $L_2(\bar{A}^\tau)$

TT	Từ ngôn ngữ	4-tuple ngữ nghĩa	TFS của các từ ngôn ngữ dựa trên HA
1	Extremely Low (EL)	(EL, 0.1, [0.1, 0.219], 0.1)	(0.1, 0.1, 0.336)
2	Low (L)	(L, 0.336, [0.219, 0.450], 0.336)	(0.1, 0.336, 0.566)
3	Medium (M)	(M, 0.566, [0.450, 0.675], 0.566)	(0.336, 0.566, 0.781)
4	High (H)	(H, 0.781, [0.675, 0.889], 0.781)	(0.566, 0.781, 1)
5	Extremely High (EH)	(EH, 1, [0.889, 1], 1)	(0.781, 1.000, 1)

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:

$$\tilde{a}_{ij}(\tau) = a_{ij}^1(\tau) \oplus a_{ij}^2(\tau) \oplus \dots \oplus a_{ij}^{h_\tau}(\tau) \quad (4.19)$$

trong đó, $i = 1, \dots, m_\tau; j = 1, \dots, n_\tau$.

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 . Việc thực hiện kết nhập trực tiếp trên miền ngôn ngữ 4-tuple giúp bảo toàn cấu trúc thứ tự và ngữ nghĩa nội tại của các từ ngôn ngữ, thay vì chuyển đổi sang miền số.

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_τ^u, n_τ^u . 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 $\overline{A\tilde{A}^\tau} = (\tilde{a}_{ij}^*(\tau))_{m_\tau^u \times n_\tau^u}$, được xác định theo công thức (4.20) sau:

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

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

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: $\emptyset \widehat{\oplus} \emptyset := \tilde{w} = (w, S_k(w), \mathfrak{H}(w), r_w)$. 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 τ , các tiêu chí không thuộc $\bar{C}^\tau$ không tham gia ma trận quyết định hiện tại, nên giá trị $\tilde{w}$ chỉ có ý nghĩa lưu vết cấu trúc và không ảnh hưởng đến kết quả xếp hạng tại thời điểm τ .

Công thức (4.20) cho phép mô hình thích ứng với bối cảnh MCGDM động, trong đó tập phương án $\bar{A}^\tau$ và tập tiêu chí $\bar{C}^\tau$ có thể thay đổi theo thời gian. Nhờ cơ chế này, dữ liệu lịch sử được kết nhập với dữ liệu hiện tại bị khuyết thiếu theo nguyên tắc kế thừa có kiểm soát, tránh việc loại bỏ hoàn toàn thông tin cũ, đồng thời không làm cho dữ liệu lịch sử chi phối quá mức kết quả đánh giá hiện tại.

Bước 2.4. Chuyển đổi ma trận đánh giá ngôn ngữ đã 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 phân 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 $\overline{FA\tilde{A}^\tau}$ và được xác định theo công thức (4.21) sau:

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

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 $\overline{FA\tilde{A}^\tau}$ 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 ($\mathcal{B}^+$) và NIS ($\mathcal{B}^-$) tại thời điểm τ được tính lần lượt theo các công thức (4.23) và (4.24):

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

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

trong đó, $\mathcal{B}^+ = (1, 1, 1)$, $\mathcal{B}^- = (0, 0, 0)$, 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). Trên cơ sở đó, 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.

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

Thuật toán 4.1. Thuật toán mô hình L-FAHP - L-FTOPSIS động dựa trên HA.

Algorithm L-FAHP-L-FTOPSIS-HA ($\bar{A}$, $\bar{C}$, $\bar{D}$, HA_C , $f\mathfrak{m}_C(g^-)$, $\mu_C(h)$, k_C , HA_A , $f\mathfrak{m}_A(g^-)$, $\mu_A(h)$, k_A).

Input:

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

- Ở mỗi thời điểm $\tau \in T$ ($\tau = 1, 2, \dots, t$):

- Tập các phương án $\bar{A}^\tau = \{A_1^\tau, A_2^\tau, \dots, A_{m_\tau}^\tau\}$, (với $m_\tau \geq 2$);
- Tập các tiêu chí $\bar{C}^\tau = \{C_1^\tau, C_2^\tau, \dots, C_{n_\tau}^\tau\}$, (với $n_\tau \geq 2$);
- Tập người ra quyết định $\bar{D}^\tau = \{D_1^\tau, D_2^\tau, \dots, D_{h_\tau}^\tau\}$, (với $h_\tau \geq 2$);

- Bộ tham số $(HA_C, f\mathfrak{m}_C(g^-), \mu_C(h), k_C)$ // xây dựng thang đo $L_k(\bar{C}^\tau)$;
- Bộ tham số $(HA_A, f\mathfrak{m}_A(g^-), \mu_A(h), k_A)$ // xây dựng thang đo $L_k(\bar{A}^\tau)$;

- Từ thời điểm τ - I : History_data (dữ liệu lịch sử).

Output: Thứ hạng các phương án $A_i^\tau \in \bar{A}^\tau$ tại mỗi thời điểm τ .

Begin

1. for $\tau \in T$ do

// **Giai đoạn 1: Xác định trọng số của các tiêu chí bằng L-FAHP - HA**

2. Thiết kế cấu trúc phân cấp MCDM (mục tiêu – tiêu chí – phương án).

3. Xây dựng thang đo ngôn ngữ 4-tuple $L_k(\bar{C})$ dựa trên HA bằng cách lựa chọn cấu trúc ComHA, và các tham số $f\mathfrak{m}_C(g^-)$, $\mu_C(h)$, (k_C) theo các Định nghĩa từ 1.4 đến 1.12 cùng các Mệnh đề 1.1 đến 1.3.

4. for $D_e^\tau \in \bar{D}^\tau$ do

5. Xây dựng ma trận so sánh cặp ngôn ngữ $P_{D_e}^\tau$ bởi công thức (4.7);

6. Kiểm tra tỷ lệ nhất quán $CR_{HA}^e(\tau)$ bởi công thức (4.9) và (4.10);

7. while $CR_{HA}^e(\tau) \geq 0.10$ do

8. Yêu cầu D_e^τ hiệu chỉnh lại ma trận $P_{D_e}^\tau$; quay lại bước 3-4;

9. end while

10. end for

11. Kết nhập các ma trận so sánh cặp tại thời điểm τ , thu được ma trận nhóm, AP^τ theo công thức (4.11) và (4.12);

12. if $\tau > 1$ then

13. Tính ma trận so sánh cặp tổng hợp có tích hợp dữ liệu lịch sử $\overline{AP}^\tau$ bởi công thức (4.13);

14. else

15. $\overline{AP}^\tau = AP^\tau$;

16. end if

17. Áp dụng hàm ánh xạ $fz(.)$ để chuyển đổi $\overline{AP}^\tau$ sang $\overline{FAP}^\tau$ bởi công thức (4.14);

18. for $C_j^\tau \in \bar{C}^\tau$ do

19. Tính trung bình nhân mờ RAP_j^τ theo công thức (4.15);

20. end for

21. Tính trọng số mờ của tiêu chí $F\tilde{w}_j(\tau)$ bởi công thức (4.16);

22. *Khử mờ theo COA để thu giá trị số rõ bởi công thức (4.17);*
23. *Chuẩn hóa $\{M_i(\tau)\}$ để thu véc-tơ trọng số $\tilde{w}_j(\tau)$*
- // Giai đoạn 2: Đánh giá xếp hạng phương án tại bằng mô hình L-FTOPSIS-HA.**
24. *Xây dựng thang đo ngôn ngữ 4-tuple $L_k(\bar{A})$ dựa trên HA bằng cách lựa chọn cấu trúc ComHA, và các tham số $f_{m_A}(g^-)$, $\mu_A(h)$, (k_A) theo các Định nghĩa từ 1.4 đến 1.12 cùng các Mệnh đề 1.1 đến 1.3.*
25. *for $D_e^\tau \in \bar{D}^\tau$ do*
26. *Xây dựng ma trận đánh giá ngôn ngữ cho các phương án, $\tilde{A}_{D_e^\tau}^\tau$, bởi công thức (4.8);*
27. *end for*
28. *Kết nhập các ma trận đánh giá ngôn ngữ thành ma trận nhóm $A\tilde{A}^\tau$ tại thời điểm τ , bởi công thức (4.18), (4.19);*
29. *if $\tau > 1$ then*
30. *Tính ma trận đánh giá phương án tổng hợp có tích hợp dữ liệu lịch sử $\overline{A\tilde{A}^\tau}$ bởi công thức (4.20);*
31. *else*
32. *$\overline{A\tilde{A}^\tau} = A\tilde{A}^\tau$;*
33. *end if*
34. *Chuyển đổi từ ma trận $\overline{A\tilde{A}^\tau}$ sang $\overline{FA\tilde{A}^\tau}$ theo công thức (4.21)*
35. *Tính ma trận quyết định mờ chuẩn hóa có trọng số, Υ_i^τ bởi công thức (4.22)*
36. *for $A_i^\tau \in \bar{A}^\tau$ do*
37. *Tính khoảng cách từ A_i^τ đến PIS và NIS bởi công thức (4.23) và (4.24);*
38. *end for*
39. *for $A_i^\tau \in \bar{A}^\tau$ do*
40. *Tính hệ số gần CC_i^τ bởi công thức (4.25)*
41. *end for*
42. *Xếp hạng phương án A_i^τ theo CC_i^τ giảm dần*
43. *end for*
44. *Return $\{Rank^\tau(A_i^\tau \in \bar{A}^\tau), i = 1, \dots, m_\tau; \tau = 1, 2, \dots, t\}$:*

End

Bảng 4.5. Độ phức tạp của thuật toán thực hiện các công thức trong mô hình L-FAHP - L-FTOPSIS động dựa trên HA

STT	Công thức	Nội dung tính toán	Độ phức tạp
1	(4.7)	Xây dựng các ma trận so sánh cặp ngôn ngữ cho h_t chuyên gia; mỗi ma trận có kích thước $n_t \times n_t$.	$O(h_t \times n_t^2)$
2	(4.8)	Xây dựng các ma trận đánh giá phương án cho h_t chuyên gia; mỗi ma trận có kích thước $m_t \times n_t$.	$O(h_t \times m_t \times n_t)$
3	(4.9)	Tính tỷ lệ nhất quán $CR_{HA}^e(\tau)$ cho toàn bộ h_t chuyên gia khi chỉ số $CI_{HA}^e(\tau)$ và $RI(n_t)$ đã có.	$O(h_t)$
4	(4.10)	Tính chỉ số nhất quán dựa trên giá trị riêng lớn nhất của ma trận so sánh cặp; chỉ phí chi phối là phép tính $\lambda_{\max}^e(\tau)$ cho h_t ma trận bậc n_t .	$O(h_t \times n_t^3)$
5	(4.11) - (4.12)	Kết nhập các ma trận so sánh cặp của các chuyên gia tại thời điểm τ để thu được ma trận nhóm AP^τ .	$O(h_t \times n_t^2)$
6	(4.13)	Tích hợp dữ liệu lịch sử vào ma trận so sánh cặp trên không gian tiêu chí hợp nhất; mỗi phần tử được xử lý trong thời gian hằng số.	$O((n_t^u)^2)$
7	(4.14)	Ánh xạ toàn bộ các phần tử ngôn ngữ của ma trận so sánh cặp tổng hợp sang ma trận mờ tam giác.	$O((n_t^u)^2)$
8	(4.15)	Tính trung bình hình học mờ cho toàn bộ n_t^u tiêu chí; mỗi tiêu chí cần kết hợp n_t^u phần tử.	$O((n_t^u)^2)$
9	(4.16)	Tính và chuẩn hóa các giá trị trung bình hình học mờ để thu được trọng số mờ của các tiêu chí.	$O(n_t^u)$
10	(4.17)	Khử mờ theo phương pháp COA và chuẩn hóa để thu được véc-tơ trọng số rõ cho n_t^u tiêu chí.	$O(n_t^u)$
11	(4.18) - (4.19)	Kết nhập các ma trận đánh giá phương án của h_t chuyên gia để thu được ma trận đánh giá nhóm tại thời điểm τ .	$O(h_t \times m_t \times n_t)$
12	(4.20)	Tích hợp dữ liệu lịch sử vào ma trận đánh giá phương án trên không gian phương án - tiêu chí hợp nhất.	$O(m_t^u \times n_t^u)$
13	(4.21)	Chuyển đổi ma trận đánh giá ngôn ngữ đã tích hợp lịch sử sang ma trận đánh giá mờ.	$O(m_t^u \times n_t^u)$

STT	Công thức	Nội dung tính toán	Độ phức tạp
14	(4.22)	Tính giá trị đánh giá tổng hợp có trọng số cho toàn bộ m_τ^u phương án trên n_τ^u tiêu chí.	$O(m_\tau^u \times n_\tau^u)$
15	(4.23)	Tính khoảng cách từ các phương án đến PIS trên toàn bộ $m_\tau^u \times n_\tau^u$ tiêu chí.	$O(m_\tau^u \times n_\tau^u)$
16	(4.24)	Tính khoảng cách từ các phương án đến NIS trên toàn bộ $m_\tau^u \times n_\tau^u$ tiêu chí.	$O(m_\tau^u \times n_\tau^u)$
17	(4.25)	Tính hệ số gần cho toàn bộ phương án tại thời điểm τ .	$O(m_\tau^u)$

Ký hiệu: t là số thời điểm đánh giá; tại mỗi thời điểm τ , m_τ là số phương án, n_τ là số tiêu chí, h_τ là số chuyên gia; m_τ^u và n_τ^u lần lượt là số phần tử của tập phương án và tập tiêu chí hợp nhất khi tích hợp dữ liệu lịch sử. Vì vậy, phân tích độ phức tạp cần được tách thành hai mức: chi phí tại một thời điểm đánh giá là τ và chi phí tổng thể trên toàn bộ chuỗi thời điểm đánh giá là t .

Tại một thời điểm τ , độ phức tạp thời gian chi phối của mô hình L-FAHP - L-FTOPSIS động dựa trên HA bị chi phối chủ yếu bởi hai thao tác: quá trình kiểm tra tính nhất quán của ma trận so sánh cặp bậc n_τ với độ phức tạp $O(h_\tau \times n_\tau^3)$ và quá trình kết nhập ma trận đánh giá phương án của h_τ chuyên gia với độ phức tạp $O(h_\tau \times m_\tau \times n_\tau)$, còn $\rho_\tau(m_\tau^u, n_\tau^u)$ biểu diễn nhóm chi phí gần với tích hợp dữ liệu lịch sử, xử lý không gian phương án - tiêu chí hợp nhất và các bước tính toán tiếp theo theo các bậc đã nêu trong Bảng 4.5. Do mô hình L-FAHP - L-FTOPSIS dựa trên HA được xây dựng cho bài toán MCGDM động qua nhiều thời điểm, độ phức tạp tổng thể cần được lượng hóa bằng tổng chi phí tính toán qua toàn bộ chuỗi thời điểm. Vì vậy, độ phức tạp tổng thể của mô hình trên t thời điểm đánh giá được viết dưới dạng:

$$O\left(\sum_{\tau=1}^t [h_\tau \times n_\tau^3 + h_\tau \times m_\tau \times n_\tau + \rho_\tau(m_\tau^u, n_\tau^u)]\right)$$

Trong trường hợp bỏ qua các hạng bậc thấp và giả định chi phí tra cứu, chuyển đổi trên thang đo HA hữu hạn được xem là hằng số, thành phần chính có thể diễn giải là: $O(t \times h \times \max\{n^3, m \times n\})$. Trong thực tiễn của các bài toán MCGDM, số lượng tiêu chí n_t và số lượng chuyên gia h_t thường là các hằng số nhỏ, kích thước bài toán chủ yếu tăng theo số lượng phương án m_t . Đồng thời, độ phức tạp có tham số t chỉ phản ánh chi phí xử lý thuật toán sau khi dữ liệu đánh giá đã được cung cấp; chi phí tổ chức khảo sát, lựa chọn chuyên gia và thu thập đánh giá ngôn ngữ qua nhiều thời điểm thuộc về chi phí triển khai thực nghiệm, không phải độ phức tạp thuật toán theo nghĩa tính toán. Vì vậy, mô hình L-FAHP - L-FTOPSIS động dựa trên HA có độ phức tạp tuyến tính theo m_t ,

cho thấy mô hình có độ phức tạp phù hợp với phạm vi thực nghiệm; khả năng mở rộng cần tiếp tục kiểm chứng và xử lý dữ liệu thời gian thực trong các hệ thống hỗ trợ ra quyết định động.

4.5. Ví dụ thực nghiệm

Trong mục này, mô hình ra quyết định L-FAHP - L-TOPSIS dựa trên HA được áp dụng để giải quyết bài toán đánh giá mức độ sẵn sàng chuyển đổi số của SMEs trong lĩnh vực bán lẻ tại Việt Nam, qua đó minh chứng khả năng ứng dụng và hiệu quả của mô hình đề xuất trong môi trường MCGDM động. Bài toán được lựa chọn vì đây là một bài toán MCGDM điển hình, trong đó thông tin đánh giá mang tính định tính, mơ hồ, khó lượng hóa chính xác [31, 91] và có xu hướng biến thiên theo thời gian dưới tác động của sự thay đổi công nghệ, thị trường, hành vi tiêu dùng và năng lực thích ứng của doanh nghiệp [29, 32]. Trên cơ sở tổng hợp các nghiên cứu liên quan, mức độ chuyển đổi số của SMEs bán lẻ có thể được xem xét một cách tương đối toàn diện thông qua bốn tiêu chí chính dưới đây:

Về công nghệ và năng lực số: Các doanh nghiệp cần chú trọng khả năng tích hợp và ứng dụng các nền tảng số như thương mại điện tử, Big Data, IoT và blockchain [101], cũng như năng lực triển khai trí tuệ nhân tạo (AI), điện toán đám mây và đưa ra quyết định dựa trên phân tích dữ liệu lớn [42].

Về chiến lược và tổ chức: Việc xây dựng một chiến lược chuyển đổi số rõ ràng cùng các mục tiêu và chỉ số đo lường hiệu quả (KPIs), đi kèm với cam kết và sự ủng hộ mạnh mẽ từ ban lãnh đạo, khả năng linh hoạt thích ứng và phân bổ nguồn lực hợp lý, là những yếu tố quan trọng quyết định sự thành công của chuyển đổi số [101, 107, 65].

Về trải nghiệm khách hàng: SMEs cần tăng cường năng lực bán hàng đa kênh [65], cá nhân hóa trải nghiệm khách hàng thông qua việc ứng dụng công nghệ [42], đồng thời nâng cao sự tiện lợi và tốc độ giao dịch trong các dịch vụ số để đáp ứng nhu cầu ngày càng khắt khe từ phía khách hàng [107].

Về yếu tố nhân lực và văn hóa: Việc đầu tư vào đào tạo và phát triển kỹ năng số cho nhân viên, xây dựng văn hóa đổi mới, khuyến khích chấp nhận rủi ro và điều chỉnh linh hoạt cơ cấu tổ chức phù hợp với chuyển đổi số là những điều kiện tiên quyết nhằm đảm bảo sự bền vững và hiệu quả của quá trình này [77, 107].

Nhìn chung, bộ tiêu chí này không chỉ phản ánh tình trạng và năng lực chuyển đổi số của SMEs trong lĩnh vực bán lẻ mà còn là cơ sở quan trọng để xây dựng các chiến lược phù hợp nhằm tối ưu hóa vận hành và nâng cao năng lực cạnh tranh trong môi trường kinh doanh số hóa hiện nay.

Trong ví dụ thực nghiệm này, bộ tiêu chí đánh giá gồm 4 tiêu chí chính và 14 tiêu chí con, được sử dụng để mô hình hóa bài toán đánh giá mức độ sẵn sàng chuyển đổi số của SMEs bán lẻ tại Việt Nam được trình bày trong Bảng A2.1 ở phần Phụ lục A.

Giả sử một tỉnh A tại Việt Nam có nhu cầu đánh giá và xếp hạng mức độ sẵn sàng chuyển đổi số của SMEs trong lĩnh vực bán lẻ. Sau quá trình sàng lọc, 10 doanh nghiệp (E1, E2, E3, E4, E5, E6, E7, E8, E9, E10), được chọn làm tập phương án khảo sát. Quá trình này được thực hiện tại 3 thời điểm τ_1 , τ_2 và τ_3 trên cơ sở bộ tiêu chí gồm 4 tiêu chí chính và 14 tiêu chí con được trình bày trong Bảng A2.1 ở phần Phụ lục A. Trong bối cảnh ra quyết định động, tập phương án, bộ tiêu chí và thành phần hội đồng chuyên gia có thể thay đổi theo từng thời điểm nhằm phản ánh sự biến động của môi trường ra quyết định, chẳng hạn như sự thay đổi của công nghệ, thị trường, hành vi tiêu dùng, chính sách quản lý hoặc năng lực thích ứng của doanh nghiệp. Để thực hiện quá trình đánh giá, một hội đồng chuyên gia được thành lập gồm 3 chuyên gia, ký hiệu D_1 , D_2 , D_3 . Đây là những người có kiến thức chuyên sâu và kinh nghiệm thực tiễn trong lĩnh vực chuyển đổi số SMEs tại Việt Nam.

Hội đồng chuyên gia thống nhất xây dựng hai thang đo ngôn ngữ 4-tuple ngữ nghĩa dựa trên HA, ký hiệu lần lượt là $L_2(\bar{C}^t)$ và $L_2(\bar{A}^t)$. Trong đó, thang đo $L_2(\bar{C}^t)$ được sử dụng trong giai đoạn L-FAHP-HA với bộ tham số ngữ nghĩa của HA được thiết lập như trong Ví dụ 4.1 (AX^* , $\text{fn}(U)$, $\mu(L)$, $k = 2$) để biểu diễn các đánh giá so sánh cặp ngôn ngữ về mức độ quan trọng giữa các tiêu chí trình bày trong Bảng 4.3; còn thang đo $L_2(\bar{A}^t)$ được sử dụng trong giai đoạn L-FTOPSIS-HA với tập phân tử sinh $G = \{\text{Low}, \text{High}\}$, tập gia tử $H = \{\text{Little}, \text{Very}\}$, cùng độ đo tính mờ của $\text{fn}(\text{Low}) = 0.518$, $\mu(\text{Little}) = 0.495$ và độ dài tối đa của từ ngôn ngữ $k = 1$ để biểu diễn mức độ đáp ứng của từng doanh nghiệp theo từng tiêu chí tại Bảng 4.4; đồng thời thống nhất thiết lập hệ số suy giảm ngữ nghĩa $\lambda = 0.7$.

Dựa trên hai thang đo ngôn ngữ trên, tiến hành thu thập dữ liệu thông qua phương pháp phỏng vấn bán cấu trúc đối với các nhà quản lý cấp cao và trưởng bộ phận của các doanh nghiệp được lựa chọn khảo sát theo bộ tiêu chí nêu trên. Ba chuyên gia sẽ cung cấp ma trận đánh giá so sánh cặp ngôn ngữ để xác định trọng số của tiêu chí, đồng thời cung cấp ma trận đánh giá mức độ đáp ứng của các doanh nghiệp theo từng tiêu chí tại từng thời điểm quan sát.

Quy trình tính toán thực nghiệm của mô hình ra quyết định đề xuất qua ba thời điểm (τ_1 , τ_2 và τ_3) được thực hiện như sau:

(1) Thời điểm τ_1 :

Các chuyên gia đã thực hiện đánh giá 5 doanh nghiệp (E1, E2, E3, E4, E5). Dựa trên phân tích bối cảnh hiện tại, mười một tiêu chí con được nhóm vào ba tiêu chí chính, bao gồm: (i) chiến lược và tổ chức, (ii) công nghệ và năng lực số, và (iii) trải nghiệm khách hàng, đã được xác định để phục vụ việc đánh giá.

Các Bảng (từ A2.2 đến A2.7) trong Phụ lục A trình bày ma trận so sánh mờ cho các tiêu chí và tiêu chí con, cùng với các đánh giá kết nhập của doanh nghiệp và trọng số quan trọng của các tiêu chí tại thời điểm τ_1 . Bảng 4.6 và Bảng 4.7 là kết quả tính toán của mô hình ra quyết định tại thời điểm τ_1 , bao gồm: đánh giá trung bình có trọng số của các doanh nghiệp, tính khoảng cách từ các doanh nghiệp đến PIS và NIS, cũng như hệ số gần của các doanh nghiệp.

Bảng 4.6. Đánh giá trung bình có trọng số của các doanh nghiệp tại thời điểm τ_1

Doanh nghiệp	Đánh giá trung bình có trọng số
E1	(0.047, 0.111, 0.177)
E2	(0.061, 0.133, 0.185)
E3	(0.029, 0.080, 0.149)
E4	(0.025, 0.072, 0.135)
E5	(0.028, 0.076, 0.141)

Bảng 4.7. Khoảng cách từ các doanh nghiệp tới ℓ^+ và ℓ^- tại thời điểm τ_1

Doanh nghiệp	d^+	d^-
E1	0.890	0.124
E2	0.875	0.136
E3	0.915	0.099
E4	0.924	0.090
E5	0.919	0.094

Bảng 4.8 cho thấy kết quả xếp hạng các SMEs, trong đó, doanh nghiệp E2 đạt hệ số gần lớn nhất 0.135, xếp hạng thứ nhất, tiếp theo là E1 (0.122) và E3 (0.098); trong khi E4 (0.088) là phương án có mức độ ưu tiên thấp nhất. Thứ hạng này cho thấy mô hình có khả năng nhận diện được doanh nghiệp nổi trội ngay từ giai đoạn đầu, đồng thời bảo đảm sự phân hóa hợp lý giữa các phương án có mức đánh giá tương đối gần nhau trong cùng môi trường ngôn ngữ dựa trên HA.

Bảng 4.8. Hệ số gần và thứ hạng của các doanh nghiệp tại thời điểm τ_1

Doanh nghiệp	Hệ số gần	Thứ hạng
E1	0.122	2
E2	0.135	1
E3	0.098	3
E4	0.088	5
E5	0.093	4

(2) Thời điểm τ_2 :

Tại thời điểm này, ba doanh nghiệp mới (E6, E7, E8) được bổ sung vào tập phương án, đồng thời ba tiêu chí con thuộc tiêu chí chính “*Nguồn nhân lực và văn hóa*” cũng được đưa vào khảo sát. Các bước tính toán, kết nhập đánh giá của các doanh nghiệp và tính trọng số quan trọng cho bốn tiêu chí và mười bốn tiêu chí con được trình bày tại các Bảng A2.8 và A2.9 trong Phụ lục A. Các Bảng 4.9 và 4.10 là kết quả tính toán, bao gồm: đánh giá trung bình có trọng số của các doanh nghiệp, khoảng cách đến PIS và NIS, hệ số gần, và thứ hạng của các doanh nghiệp tại thời điểm τ_2 .

Bảng 4.9. Đánh giá trung bình có trọng số của các doanh nghiệp tại thời điểm τ_2

Doanh nghiệp	Đánh giá trung bình có trọng số
E1	(0.023, 0.054, 0.111)
E2	(0.025, 0.058, 0.114)
E3	(0.019, 0.047, 0.103)
E4	(0.014, 0.037, 0.090)
E5	(0.015, 0.040, 0.092)
E6	(0.030, 0.066, 0.112)
E7	(0.022, 0.052, 0.112)
E8	(0.016, 0.043, 0.095)

Bảng 4.10. Khoảng cách từ các doanh nghiệp tới $\mathcal{B}^+$ và $\mathcal{B}^-$ tại thời điểm τ_2

Doanh nghiệp	d^+	d^-
E1	0.938	0.072
E2	0.935	0.075
E3	0.944	0.067
E4	0.954	0.057
E5	0.952	0.058
E6	0.931	0.077
E7	0.939	0.072
E8	0.949	0.061

Kết quả xếp hạng cuối cùng của các doanh nghiệp được trình bày trong Bảng 4.11, trong đó E6 vươn lên xếp hạng thứ nhất với hệ số gần 0.077, tiếp theo là E2 (0.075) và E1 (0.072); các doanh nghiệp còn lại có hệ số gần thấp hơn và phân bố thứ hạng tương đối ổn định. Kết quả này cho thấy mô hình không chỉ thích ứng tốt với sự thay đổi của tập phương án và bộ tiêu chí theo thời gian, mà còn duy trì được khả năng so sánh, xếp hạng nhất quán khi tích hợp thông tin hiện tại với dữ liệu lịch sử trong mô hình L-FAHP - L-FTOPSIS động dựa trên HA.

Bảng 4.11. Hệ số gần và thứ hạng của các doanh nghiệp tại thời điểm τ_2

Doanh nghiệp	Hệ số gần	Thứ hạng
E1	0.072	3
E2	0.075	2
E3	0.066	5
E4	0.056	8
E5	0.058	7
E6	0.077	1
E7	0.072	4
E8	0.060	6

(3) Thời điểm τ_3 :

Vì một số lý do thực tế, hai doanh nghiệp E1 và E2 được đưa ra khỏi tập phương án; đồng thời, hai doanh nghiệp mới là E9 và E10 được đưa vào mô hình. Bộ tiêu chí được giữ nguyên như ở thời điểm τ_2 , gồm 14 tiêu chí con trong 4 tiêu chí chính chính. Các Bảng A2.10 và A2.11 trong Phụ lục A lần lượt trình bày các mức đánh giá tổng hợp của doanh nghiệp và trọng số quan trọng của các tiêu chí. Tiếp theo, Bảng 4.12 và 4.13 trình bày các kết quả tính toán, bao gồm: đánh giá trung bình có trọng số, khoảng cách từ mỗi doanh nghiệp đến PIS và NIS, hệ số gần, và thứ hạng cuối của các doanh nghiệp.

Bảng 4.12. Đánh giá trung bình có trọng số của các doanh nghiệp tại thời điểm τ_3

Doanh nghiệp	Đánh giá trung bình có trọng số
E1	(0.018, 0.046, 0.094)
E2	(0.034, 0.070, 0.115)
E3	(0.011, 0.032, 0.075)
E4	(0.015, 0.038, 0.086)
E5	(0.020, 0.046, 0.095)
E6	(0.030, 0.062, 0.105)
E7	(0.029, 0.060, 0.105)
E8	(0.020, 0.046, 0.095)
E9	(0.031, 0.063, 0.105)
E10	(0.028, 0.059, 0.103)

Bảng 4.13. Khoảng cách từ các doanh nghiệp tới b^+ và b^- tại thời điểm τ_3

Doanh nghiệp	d^+	d^-
E1	0.948	0.061
E2	0.928	0.080
E3	0.961	0.047
E4	0.954	0.055
E5	0.947	0.062
E6	0.935	0.073
E7	0.936	0.072
E8	0.947	0.062
E9	0.934	0.073
E10	0.937	0.071

Tại thời điểm τ_3 , Bảng 4.14 là kết quả xếp thứ hạng tổng hợp cả quá trình đánh giá ở 3 thời điểm của 10 doanh nghiệp, trong đó kết quả xếp hạng cho thấy E2 đứng đầu với hệ số gần 0.079, tiếp theo là E9 (0.073), E6 (0.072) và E7 (0.071), trong khi E3 có hệ số gần nhỏ nhất 0.047. Điểm đáng chú ý ở thời điểm này là mặc dù một số doanh nghiệp không còn được đánh giá mới, mô hình vẫn bảo lưu và khai thác thông tin lịch sử theo chính sách lưu giữ kết hợp cơ chế suy giảm ngữ nghĩa, chẳng hạn như E2 bị đưa ra khỏi tập phương án (không được đánh giá ở thời điểm hiện tại, τ_3), nhưng hệ thống vẫn lưu trữ kết quả lịch sử xuất sắc của nó từ các thời điểm τ_1 (xếp thứ nhất) và τ_2 (xếp thứ hai) nhờ đó E2 vẫn được xác định là phương án tối ưu do có thành tích tích lũy nổi trội ở các thời điểm trước.

Bảng 4.14. Hệ số gần và thứ hạng của các doanh nghiệp tại thời điểm τ_3

Doanh nghiệp	Hệ số gần	Thứ hạng
E1	0.061	8
E2	0.079	1
E3	0.047	10
E4	0.054	9
E5	0.061	7
E6	0.072	3
E7	0.071	4
E8	0.062	6
E9	0.073	2
E10	0.070	5

Kết quả này làm rõ ưu điểm quan trọng của mô hình đề xuất: không chỉ đánh giá theo lát cắt tại một thời điểm, mà còn hỗ trợ xếp hạng tổng thể trong bối cảnh động, qua đó nâng cao tính kế thừa, độ ổn định và khả năng diễn giải của kết quả ra quyết định.

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 trên hai kịch bản nhằm kiểm tra độ ổn định và độ nhạy của mô hình ra quyết định L-FAHP - L-FTOPSIS động dựa trên HA trước các biến động có hệ thống của dữ liệu đầu vào và trước sự thay đổi đồng thời của cấu trúc ngữ nghĩa. Cụ thể, ở *kịch bản 1*, toàn bộ các phán đoán so sánh cặp ngôn ngữ được dịch tăng một mức trên thang đo $L_2(\bar{C}^r)$, qua đó mô phỏng kịch bản nhiễu loạn toàn cục nhưng vẫn bảo toàn quan hệ thứ tự tương đối giữa các phán đoán. còn *kịch bản 2*, ngoài nhiễu loạn toàn cục nói trên, mô hình còn thay đổi độ đo tính mờ của HA, nhờ đó kiểm tra khả năng ổn định của kết quả không chỉ theo dữ liệu đánh giá mà còn theo chính cơ chế biểu diễn ngữ nghĩa của mô hình.

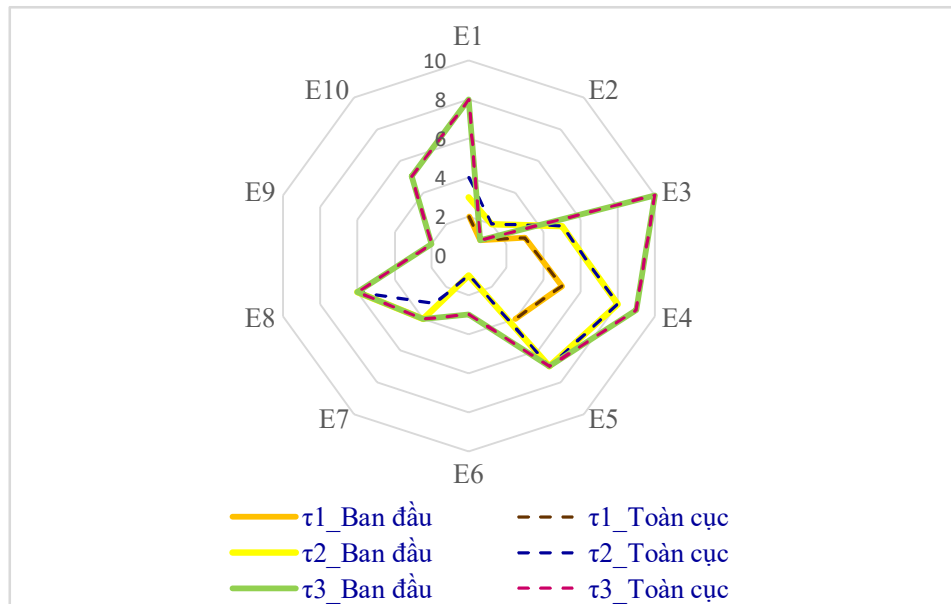
Kết quả tổng quát theo hai kịch bản được tóm tắt trong Bảng 4.15, Hình 4.3 và Hình 4.4 cho thấy thứ hạng SMEs biến động rất nhỏ, trong đó nhóm dẫn đầu và nhóm xếp cuối hầu như được bảo toàn qua cả ba thời điểm. Điều đó cung cấp bằng chứng thực nghiệm rõ ràng về độ ổn định và độ tin cậy của mô hình trong bối cảnh L-FAHP - L-FTOPSIS động dựa trên HA.

Bảng 4.15. So sánh kết quả phân tích độ nhạy trong các kịch bản tại thời điểm τ, τ_2 và τ_3

Thời điểm	Ban đầu	Kết quả phân tích độ nhạy trên 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	$E6 > E2 > \mathbf{E1} > E7 > E3 > E8 > E5 > E4$	$E6 > E2 > \mathbf{E7} > E1 > E3 > E8 > E5 > E4$	$E6 > E2 > \mathbf{E7} > E1 > E3 > E8 > E5 > E4$
τ_3	$E2 > E9 > E6 > E7 > E10 > E8 > E5 > E1 > E4 > E3$	$E2 > E9 > E6 > E7 > E10 > E8 > E5 > E1 > E4 > E3$	$E2 > E9 > E6 > E7 > E10 > \mathbf{E1} > \mathbf{E8} > E5 > E4 > E3$

(i) **Kịch bản 1:** Cho thấy mô hình có mức chống chịu rất tốt trước nhiễu loạn toàn cục của các phán đoán so sánh cặp. Theo Bảng 4.15, ở thời điểm τ_1 , kết quả ban đầu là $E2 > E1 > E3 > E5 > E4$, và thứ hạng này được giữ nguyên ở cả hai kịch bản. Điều này cho thấy cấu trúc ưu tiên giữa các phương án tại thời điểm đầu là khá bền vững trước cả nhiễu dữ liệu lẫn biến động tham số ngữ nghĩa. Về mặt định lượng, E2 là doanh nghiệp có ưu thế rõ rệt khi đạt hệ số gần 0.135, cao hơn E1 với 0.122 và E3 với 0.098, trong khi E4 chỉ đạt 0.088. Khoảng cách đủ lớn giữa phương án đứng đầu, nhóm tầm trung và xếp hạng cuối làm cho các nhiễu loạn áp dụng lên ma trận so sánh cặp không đủ mạnh để gây ra đảo hạng.

Kết quả này cho thấy tại thời điểm τ_1 , ưu thế của E2 không mang tính ngẫu nhiên mà phản ánh năng lực chuyển đổi số tốt hơn một cách tương đối ổn định ngay trong cấu hình bài toán ban đầu. Tại thời điểm τ_2 , khi không gian ra quyết định mở rộng từ 5 lên 8 SMEs, từ 3 tiêu chí chính lên 4 tiêu chí và từ 11 lên 14 tiêu chí con, thứ hạng chỉ thay đổi cục bộ: từ kết quả ban đầu $E6 > E2 > E1 > E7 > E3 > E8 > E5 > E4$ chuyển thành $E6 > E2 > E7 > E1 > E3 > E8 > E5 > E4$; tức chỉ có sự hoán đổi vị trí giữa E1 và E7, còn phương án tốt nhất E6, nhóm tầm trung và phương án cuối E4 đều được giữ nguyên. Tại τ_3 , thứ hạng ban đầu $E2 > E9 > E6 > E7 > E10 > E8 > E5 > E1 > E4 > E3$ được bảo toàn hoàn toàn trong Kịch bản 1. Diễn biến này phù hợp với Hình 4.3, nơi các đường biểu diễn giữa trường hợp ban đầu và trường hợp nhiễu loạn gần như chồng khít ở hầu hết các tiêu chí và thời điểm, cho thấy hệ số gần chỉ dao động biên nhỏ và không đủ để phá vỡ cấu trúc thứ hạng tổng thể.



Hình 4.3. Kịch bản phân tích độ nhạy với nhiễu loạn toàn cục

Ý nghĩa của kết quả trên là rất rõ. Thứ nhất, khi một sự dịch chuyển hệ thống được áp dụng đồng đều lên toàn bộ ma trận so sánh cặp, mô hình vẫn duy trì được trật tự ưu tiên của các doanh nghiệp chủ chốt. Điều đó cho thấy véc-tơ trọng số sinh ra bởi L-FAHP dựa trên HA không nhạy quá mức với nhiễu loạn toàn cục, vì cấu trúc thứ tự ngữ nghĩa của các đánh giá vẫn được bảo toàn sau khi dịch mức. Thứ hai, SMEs đứng đầu tại từng thời điểm đều là những phương án có ưu thế tương đối rõ rệt trong kết quả gốc: tại τ_1 , E2 có hệ số gần 0.135, cao hơn E1 (0.122) và E3 (0.098); tại τ_2 , E6 đạt 0.077,

cao hơn E2 (0.075) và E1 (0.072); tại τ_3 , E2 tiếp tục đứng đầu với 0.079, tiếp theo là E9 (0.073), E6 (0.072) và E7 (0.071). Điều đó cho thấy các cực trị trong cấu trúc đánh giá được bảo toàn sau gây nhiễu, và mô hình ổn định không chỉ ở mức thứ hạng mà còn ở mức khoảng cách tới giải pháp lý tưởng.

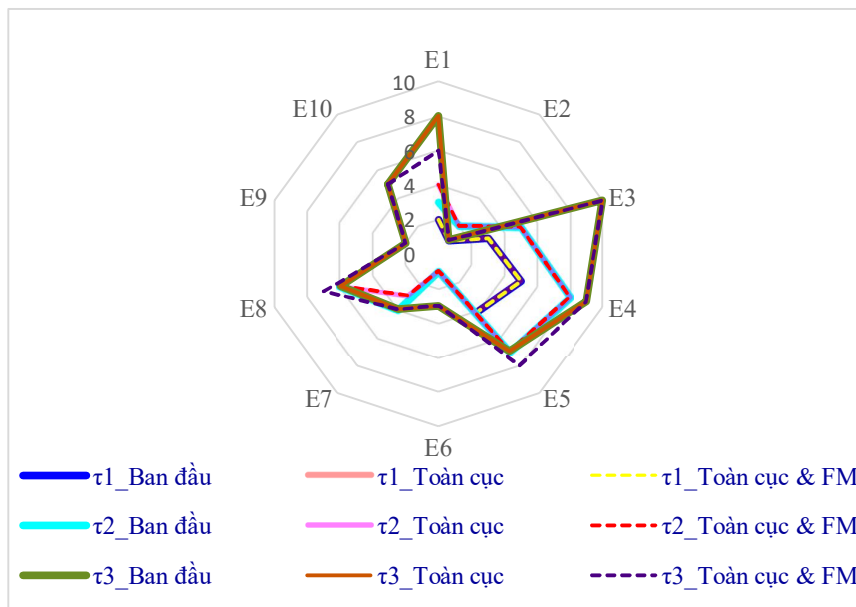
Kết quả trong kịch bản 1 phản ánh tác động đồng thời của hai thành phần trong mô hình: trọng số của tiêu chí được xác định bằng L-FAHP - HA và điểm đánh giá phương án được tổng hợp bằng L-FTOPSIS - HA đồng. Việc dịch tăng đồng đều toàn bộ các phán đoán chỉ làm thay đổi mức độ ưu tiên mà không làm thay đổi quan hệ thứ tự tương đối giữa các tiêu chí. Do đó, véc-tơ trọng số do L-FAHP dựa trên HA sinh ra vẫn duy trì được tính ổn định, và các phương án có lợi thế tích lũy trên nhiều tiêu chí tiếp tục giữ ưu thế. Chẳng hạn, tại τ_3 , E2 có $d^+ = 0.928$ nhỏ hơn và $d^- = 0.080$ lớn hơn so với phần lớn các phương án khác, từ đó đạt hệ số gần lớn nhất 0.079; ngược lại, E3 có $d^+ = 0.961$ lớn nhất và $d^- = 0.047$ nhỏ nhất, nên luôn ở cuối bảng xếp hạng. Diễn biến này phù hợp với Hình 4.3, khi các đường biểu diễn giữa trường hợp ban đầu và trường hợp nhiễu loạn gần như chồng khít ở hầu hết các tiêu chí và thời điểm.

Về phương pháp luận, kết quả phân tích theo kịch bản nhiễu loạn toàn cục cho thấy mô hình đề xuất đạt được sự cân bằng hợp lý giữa độ ổn định và độ nhạy. Độ ổn định thể hiện ở việc các doanh nghiệp đứng đầu và đứng cuối được nhận diện nhất quán qua tất cả các kịch bản; trong khi đó, độ nhạy thể hiện ở khả năng phân biệt tinh trong nhóm xếp hạng tầm trung, nơi khoảng cách điểm số nhỏ và sự thay đổi trọng số của một Tiêu chí chính có thể làm thay đổi quan hệ thứ hạng cục bộ. Đây là đặc trưng có giá trị thực tiễn, vì trong các bài toán đánh giá mức độ sẵn sàng chuyển đổi số, điều quan trọng thường là nhận diện đúng nhóm doanh nghiệp nổi trội, nhóm xếp cuối và xu hướng dịch chuyển của nhóm tầm trung.

(ii) Kịch bản 2: Là phép thử khắt khe hơn, vì ngoài việc dịch tăng toàn bộ các phán đoán so sánh cặp ngôn ngữ lên một mức trên thang $L_2(\bar{C}^r)$, mô hình còn đồng thời thay đổi độ đo tính mờ của phần tử sinh và gia tử trong cấu trúc HA về mức $(fm(g^-) = 0,5$ và $\mu(h^-) = 0,5)$. Tuy nhiên, kết quả vẫn cho thấy mô hình có độ ổn định rất cao. Theo Bảng 4.15, tại τ_1 , thứ hạng tiếp tục được giữ nguyên hoàn toàn. Tại τ_2 , kết quả Kịch bản 2 là $E6 \succ E2 \succ E7 \succ E1 \succ E3 \succ E8 \succ E5 \succ E4$, tức giống với Kịch bản 1 và chỉ khác kết quả ban đầu ở sự hoán đổi nhẹ giữa E1 và E7. Tại τ_3 , kịch bản 2 cho kết quả $E2 \succ E9 \succ E6 \succ E7 \succ E10 \succ E1 \succ E8 \succ E5 \succ E4 \succ E3$, nghĩa là chỉ có một thay

đổi cục bộ giữa E1 và E8, còn doanh nghiệp đứng đầu E2, nhóm ưu tiên cao gồm E9, E6, E7 và E10; doanh nghiệp cuối bảng là E3 vẫn không thay đổi.

Hình 4.4 cũng cho thấy các đường radar của ba trường hợp: ban đầu, nhiều toàn cục và nhiều toàn cục kết hợp thay đổi độ đo tính mờ, nhìn chung vẫn bám khá sát nhau, chỉ tách nhẹ ở một số trục. Điều đó chứng tỏ các thay đổi hợp lý của độ đo tính mờ chủ yếu làm thay đổi biên độ định lượng, nhưng không làm mất trật tự ngữ nghĩa giữa các từ ngôn ngữ và cũng không dẫn tới đảo hạng quy mô lớn. Đây là ưu thế nổi bật của mô hình dựa trên HA so với nhiều mô hình dựa trên tập mờ truyền thống, nơi thay đổi hàm thuộc hoặc tham số mờ thường có thể làm biến động đáng kể kết quả xếp hạng.



Hình 4.4. Độ nhảy của mô hình trước nhiễu loạn toàn cục và thay đổi độ đo tính mờ của HA.

Nhìn từ bối cảnh động của bài toán, ý nghĩa của thứ hạng phương án cũng khá nhất quán. Ở thời điểm τ_1 , E2 luôn đứng đầu cho thấy doanh nghiệp này đã có mức độ sẵn sàng chuyển đổi số nổi trội ngay từ cấu hình ban đầu. Ở τ_2 , sự vươn lên của E6 lên vị trí số một khi bài toán được mở rộng lên 8 SMEs, 4 tiêu chí chính và 14 tiêu chí con cho thấy SMEs này có khả năng thích ứng tốt hơn khi chiều đánh giá trở nên đầy đủ hơn, đặc biệt trong tiêu chí chính mới liên quan đến “Nguồn nhân lực và văn hóa”. Ở τ_3 , mặc dù tập phương án thay đổi khi đưa thêm E9, E10 và loại E1, E2 khỏi tập phương án hiện hành, kết quả xếp hạng tổng hợp vẫn xác định E2 là tốt nhất, tiếp theo là E9, E6, E7. Điều đó cho thấy cơ chế tích hợp dữ liệu lịch sử có suy giảm ngữ nghĩa đã phát huy tác dụng, mô hình không đánh giá SMEs chỉ theo lát cắt hiện thời, mà phản ánh được cả quỹ đạo phát triển theo thời gian. Chính vì vậy, thứ hạng đầu ra mang ý nghĩa cả quá trình đánh giá chứ không đơn thuần là thời điểm.

Tổng thể, các kết quả phân tích độ nhạy phù hợp với kỳ vọng lý thuyết và với các kết quả thực nghiệm trước đó trong luận án. Về nguyên tắc, một mô hình MCGDM ngôn ngữ có nền tảng đại số tốt cần đồng thời bảo toàn trật tự ngữ nghĩa và không nhảy quá mức trước các nhiễu loạn hợp lý. Kết quả cho thấy mô hình L-FAHP - L-FTOPSIS động dựa trên HA đáp ứng tốt cả hai yêu cầu này. Đồng thời, phát hiện đó cũng nhất quán với các kết quả ở Chương 2 và Chương 3, nơi các mô hình dựa trên HA đều cho thấy độ ổn định thứ hạng cao và khả năng diễn giải tốt hơn so với các tiếp cận tập mờ truyền thống. Vì vậy, Chương 4 có thể được xem là sự mở rộng hợp logic từ bối cảnh MCGDM tĩnh sang MCGDM động, đồng thời củng cố thêm luận điểm rằng HA là nền tảng phù hợp cho các bài toán MCGDM ngôn ngữ có dữ liệu lịch sử và thay đổi theo thời gian.

Tóm lại, phân tích độ nhạy trong hai kịch bản cho thấy 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. Nhiễu loạn toàn cục trên ma trận so sánh cặp chỉ gây ra các dịch chuyển cục bộ rất nhỏ; ngay cả khi đồng thời thay đổi độ đo tính mờ của HA, cấu trúc xếp hạng tổng thể vẫn được bảo toàn. Điều này cho thấy mô hình không chỉ ổn định trước sai lệch đánh giá của chuyên gia mà còn ổn định trước những thay đổi hợp lý của cấu trúc ngữ nghĩa, qua đó khẳng định tính khả thi và giá trị ứng dụng của phương pháp đề xuất trong các bài toán MCGDM động thực tế.

4.6.2. So sánh kết quả và thảo luận

Phân tích so sánh nhằm đối sánh mô hình L-FAHP - L-FTOPSIS động dựa trên HA với phương pháp FAHP - FTOPSIS động dựa trên lý thuyết tập mờ của [91] tại ba thời điểm τ_1 , τ_2 và τ_3 , qua đó xem xét mức độ tương đồng của kết quả xếp hạng và khác biệt trong cơ chế xử lý thông tin. Các Bảng 4.16 - 4.18 cho thấy hai mô hình nhìn chung thống nhất trong nhận diện một số SMEs nổi trội, trong khi khác biệt chủ yếu xuất hiện ở các phương án có mức đánh giá gần nhau và rõ hơn khi bài toán mở rộng theo thời gian. Mô hình [91] biểu diễn đánh giá ngôn ngữ bằng số mờ và xử lý trên miền mờ, còn mô hình đề xuất [CT3] sử dụng biểu diễn 4-tuple ngữ nghĩa dựa trên HA và kết nhập trên miền ngôn ngữ; đối với dữ liệu lịch sử, mô hình áp dụng cơ chế suy giảm ngữ nghĩa hướng giá trị lịch sử dần về phần tử trung hòa. Vì vậy, sai khác thứ hạng phản ánh sự khác nhau về cơ chế biểu diễn và tổng hợp thông tin; trong phạm vi thực nghiệm của Chương 4, kết quả thu được phù hợp với mục tiêu duy trì trật tự ngữ nghĩa và kiểm soát ảnh hưởng của dữ liệu lịch sử trong MCGDM động.

Tại thời điểm τ_1 : Kết quả tóm tắt trong Bảng 4.16 cho thấy hai mô hình thống nhất hoàn toàn ở hai vị trí dẫn đầu: FAHP - FTOPSIS động [91] xếp E2 hạng 1 với hệ số gần 0.290 và E1 hạng 2 với 0.231; mô hình đề xuất [CT3] cũng xếp E2 hạng 1 với 0.135 và E1 hạng 2 với 0.122.

Bảng 4.16. Bảng so sánh kết quả giữa hai mô hình tại thời điểm τ_1

Doanh nghiệp	FAHP–FTOPSIS động của [91]				Mô hình đề xuất [CT3]			
	d^+	d^-	Hệ số gần	Thứ hạng	d^+	d^-	Hệ số gần	Thứ hạng
E1	0.932	0.281	0.231	2	0.890	0.124	0.122	2
E2	0.862	0.351	0.290	1	0.875	0.136	0.135	1
E3	0.972	0.241	0.199	5	0.915	0.099	0.098	3
E4	0.961	0.253	0.209	4	0.924	0.090	0.088	5
E5	0.960	0.256	0.210	3	0.919	0.094	0.093	4

Tuy nhiên, sự khác biệt tập trung ở E3, E4 và E5: mô hình [91] cho thứ tự $E5 > E4 > E3$, trong khi mô hình đề xuất cho thứ tự $E3 > E5 > E4$. Thứ hạng $E2 > E1 > E3 > E5 > E4$ của mô hình đề xuất thống nhất với Mục 4.5 và được bảo toàn trong cả hai kịch bản phân tích độ nhạy ở Mục 4.6.1, cho thấy cấu trúc xếp hạng tại thời điểm đầu ổn định trước các nhiễu loạn được khảo sát. Kết quả này cho thấy mô hình dựa trên HA không làm thay đổi nhận diện phương án tối ưu, nhưng có khả năng phân biệt tốt hơn giữa các doanh nghiệp tầm trung khi hệ số gần chỉ chênh lệch nhỏ. Nguyên nhân chủ yếu là ở mô hình đối chứng, quá trình ánh xạ và giải mờ có thể làm mất khoảng cách ngữ nghĩa giữa các mức đánh giá lân cận, trong khi mô hình đề xuất duy trì tốt hơn quan hệ ngữ nghĩa trong quá trình kết nhập.

Trên cơ sở Bảng 4.16, luận án tính hệ số tương quan thứ hạng Spearman [100] đạt $\rho = 0.7000$ và mức độ đồng thuận Kendall [67] với $\kappa = 0.6000$, phản ánh mức độ tương đồng đáng kể giữa hai kết quả, mặc dù tồn tại hai cặp nghịch trong nhóm E3 - E5. Do τ_1 là thời điểm đầu, khác biệt này chủ yếu liên quan đến cách biểu diễn và kết nhập thông tin ngôn ngữ, chưa phản ánh tác động tích lũy của cơ chế suy giảm dữ liệu lịch sử. Theo đó, mô hình đề xuất với phép kết nhập ngôn ngữ nhị phân duy trì tốt hơn quan hệ ngữ nghĩa khi tổng hợp, nên E3 được đưa lên hạng 3 thay vì hạng 5. Đây là một phát hiện có giá trị học thuật, vì nó cho thấy ưu thế của tiếp cận HA không nhất thiết nằm ở việc thay đổi phương án tối ưu, mà ở khả năng tái cấu trúc hợp lý nhóm phương án cận biên, nơi sai lệch ngữ nghĩa thường ảnh hưởng mạnh nhất đến kết quả xếp hạng.

Tại thời điểm τ_2 : Khi không gian phương án được mở rộng lên 8 SMEs và bộ tiêu chí được hoàn chỉnh thành 4 tiêu chí chính với 14 tiêu chí con, kết quả trong Bảng 4.17 cho thấy FAHP - FTOPSIS động [91] và mô hình L-FAHP - L-FTOPSIS động dựa trên HA [CT3] vẫn thống nhất ở các vị trí chủ chốt, gồm E6 hạng 1, E2 hạng 2 và E4 hạng 8; sự trùng khớp này chỉ xét về thứ hạng, trong khi hệ số gần của các phương án giữa hai mô hình có giá trị khác nhau. Cụ thể, thứ hạng của FAHP - FTOPSIS động [91] là $E6 > E2 > E8 > E7 > E1 > E5 > E3 > E4$, trong khi mô hình đề xuất [CT3] cho thứ tự $E6 > E2 > E1 > E7 > E3 > E8 > E5 > E4$; theo đó, E1 tăng từ hạng 5 lên hạng 3, E3 từ hạng 7 lên hạng 5, E5 giảm từ hạng 6 xuống hạng 7 và E8 giảm từ hạng 3 xuống hạng 6.

Bảng 4.17. Bảng so sánh kết quả giữa hai mô hình tại thời điểm τ_2

Doanh nghiệp	FAHP–FTOPSIS động của [91] (trung bình cộng của τ_1 và τ_2)				Mô hình đề xuất [CT3]			
	d^+	d^-	Hệ số gần	Thứ hạng	d^+	d^-	Hệ số gần	Thứ hạng
E1	0.937	0.278	0.229	5	0.938	0.072	0.072	3
E2	0.896	0.319	0.262	2	0.935	0.075	0.075	2
E3	0.958	0.257	0.2114	7	0.944	0.067	0.066	5
E4	0.959	0.256	0.211	8	0.954	0.057	0.056	8
E5	0.949	0.268	0.220	6	0.952	0.058	0.058	7
E6	0.852	0.366	0.300	1	0.931	0.077	0.077	1
E7	0.925	0.290	0.239	4	0.939	0.072	0.072	4
E8	0.916	0.304	0.249	3	0.949	0.061	0.060	6

So sánh định lượng hai véc-tơ thứ hạng cho hệ số tương quan thứ hạng Spearman [100] với $\rho = 0.7857$ và mức độ đồng thuận Kendall [67] đạt $\kappa = 0.6429$, tương ứng 23 cặp thuận và 5 cặp nghịch trong tổng số 28 cặp; năm cặp nghịch gồm (E1, E7), (E1, E8), (E3, E5), (E3, E8) và (E7, E8), cho thấy khác biệt tập trung chủ yếu ở quan hệ thứ tự giữa các SMEs thuộc nhóm trung gian, dù E7 vẫn giữ nguyên hạng 4.

Về mặt ý nghĩa, kết quả này cho thấy khi tích hợp dữ liệu lịch sử từ τ_1 được tích hợp với đánh giá hiện tại tại τ_2 , mô hình đề xuất không chỉ phản ánh trạng thái hiện thời mà còn xem xét hợp lý hơn xu hướng phát triển tích lũy của doanh nghiệp theo thời gian. Đồng thời, việc nhóm dẫn đầu gồm E6, E2, E1, E7 có hệ số gần dao động trong khoảng hẹp 0.072 - 0.077, nhưng vẫn được phân giải thứ hạng ổn định cho thấy khả năng xử lý tốt hơn tính động của bài toán. Tính ổn định của cấu trúc xếp hạng còn được củng cố qua kết quả phân tích độ nhạy ở Mục 4.6.1, khi các kịch bản gây nhiễu cho thấy E6 và E2 tiếp tục giữ vững hai vị trí đầu, E4 tiếp tục đứng cuối, và chỉ xảy ra sự hoán đổi mang tính cục bộ giữa E1 và E7. Điều này cho thấy phương pháp AHP-FTOPSIS [91] sử dụng kỹ thuật lấy trung bình cộng đơn thuần giữa hai thời điểm τ_1 và τ_2 , đồng thời trải qua bước giải mờ làm gia tăng phép tính, dễ mất mát thông tin ngữ nghĩa tích lũy và chưa phản ánh đầy đủ quy luật kế thừa và suy giảm của dữ liệu lịch sử, trong khi mô hình đề xuất [CT3] xử lý trực tiếp trên miền ngữ nghĩa với cơ chế suy giảm ngữ nghĩa hướng về phần tử trung hòa của HA phù hợp hơn với bản chất của MCGDM động.

Tại thời điểm τ_3 : Khi bài toán cập nhật không gian đánh giá với sự xuất hiện của hai SMEs mới (E9, E10) và ngừng đánh giá mới đối với E1, E2, nhưng thông tin lịch sử của các phương án trước đó vẫn được lưu giữ và tham gia vào kết quả xếp hạng tổng hợp của mô hình đề xuất. Kết quả trong Bảng 4.18 cho thấy sự khác biệt về thứ hạng giữa hai mô hình biểu hiện rõ hơn ở nhóm dẫn đầu khi bối cảnh ra quyết định tiếp tục thay đổi.

Bảng 4.18. Bảng so sánh kết quả giữa hai mô hình tại thời điểm τ_3

Doanh nghiệp	FAHP–FTOPSIS động của [91] (trung bình của τ_1 , τ_2 , và τ_3)				Mô hình đề xuất [CT3]			
	d^+	d^-	Hệ số gần	Thứ hạng	d^+	d^-	Hệ số gần	Thứ hạng
E1	0.937	0.278	0.229	7	0.948	0.061	0.061	8
E2	0.875	0.342	0.281	4	0.928	0.080	0.079	1
E3	0.958	0.257	0.212	10	0.961	0.047	0.047	10
E4	0.943	0.273	0.225	9	0.954	0.055	0.054	9
E5	0.940	0.278	0.228	8	0.947	0.062	0.061	7
E6	0.875	0.344	0.282	3	0.935	0.073	0.072	3
E7	0.925	0.292	0.240	6	0.936	0.072	0.071	4
E8	0.907	0.313	0.256	5	0.947	0.062	0.062	6
E9	0.825	0.390	0.321	1	0.934	0.073	0.073	2
E10	0.865	0.355	0.291	2	0.937	0.071	0.070	5

Mặc dù E1 và E2 không còn dữ liệu đánh giá hiện tại (thời điểm τ_3), mô hình đề xuất L-FAHP - L-FTOPSIS động dựa trên HA [CT3] vẫn bảo lưu và khai thác thông tin lịch sử của các phương án này ở các thời điểm trước đó (τ_1 và τ_2) duy trì tập phương án gồm 10 SMEs. Trong khi mô hình đối sánh FAHP - FTOPSIS động [91] xếp các phương án mới E9 ở hạng 1 và E10 ở hạng 2, đồng thời đẩy E2 xuống hạng 4; còn mô hình đề xuất [CT3] lại ghi nhận hiệu suất tích lũy vượt trội của E2 khi đưa phương án này lên vị trí dẫn đầu với hệ số gần 0.079, tiếp theo là E9 hạng 2 với 0.073, E6 hạng 3 với 0.072, E7 hạng 4 với 0.071 và E10 hạng 5 với 0.070; ở nhóm cuối, E4 và E3 cùng giữ tương ứng hạng 9 và hạng 10 ở cả hai mô hình. Kết quả tính toán tại thời điểm τ_3 cho thấy mức độ tương đồng thứ hạng tổng thể vẫn duy trì ở mức cao, với Spearman [100] đạt $\rho = 0.8424$ và Kendall [67] $\kappa = 0.6889$. Trong tổng số 45 cặp so sánh, kết quả ghi nhận 38 cặp đồng thuận và 7 cặp nghịch thuận. Sự xuất hiện của các cặp nghịch này chủ yếu xuất phát từ khác biệt trong cơ chế xử lý dữ liệu động. Thay vì lấy trung bình cộng kết quả của ba thời điểm τ_1 , τ_2 và τ_3 như mô hình đối sánh [91], mô hình đề xuất [CT3] kết nhập dữ liệu lịch sử thông qua cơ chế suy giảm ngưỡng nghĩa hướng về phần tử trung hòa. Kết quả thực nghiệm này phản ánh rõ vai trò của cơ chế lưu giữ thông tin, cho phép quá trình tổng hợp đánh giá một cách hợp lý mức độ đóng góp của dữ liệu quá khứ mà không làm thiên lệch đánh giá hiện tại của các phương án.

Nhìn chung, kết quả đối sánh tại τ_1 , τ_2 và τ_3 cho thấy hai mô hình vừa duy trì những điểm tương đồng ở các phương án chủ chốt, vừa xuất hiện các khác biệt thứ hạng phù hợp với sự thay đổi của không gian ra quyết định và cơ chế tích hợp dữ liệu. Tại τ_1 , hai mô hình cùng xác định E2 xếp hạng 1 và E1 hạng 2; tại τ_2 , cả hai tiếp tục thống nhất ở E6 hạng 1 và E2 hạng 2, trong khi khác biệt chủ yếu xuất hiện ở nhóm SMEs tầm trung; đến τ_3 , sự khác biệt rõ hơn ở nhóm ưu tiên cao khi E9, E10 được bổ sung vào tập phương án hiện tại còn dữ liệu lịch sử của E1, E2 tiếp tục được xử lý trong kết quả tổng hợp.

Mức độ tương đồng thứ hạng qua ba thời điểm có xu hướng tăng dần, với Spearman ρ lần lượt là 0.7000, 0.7857 và 0.8424, còn Kendall κ tương ứng là 0.6000, 0.6429 và 0.6889. Mức độ tương đồng thứ hạng tăng dần trong phạm vi ví dụ thực nghiệm cung cấp thêm bằng chứng về tính ổn định của cấu trúc thuật toán. Kết quả này đồng thời phù hợp với phân tích độ nhạy tại Mục 4.6.1: thứ hạng tại τ_1 được bảo toàn trong cả hai kịch bản; tại τ_2 chỉ xuất hiện hoán đổi cục bộ giữa E1 và E7; tại τ_3 , Kịch bản 1 giữ nguyên cấu trúc thứ hạng ban đầu, còn Kịch bản 2 chỉ tạo ra thay đổi cục bộ giữa E1 và E8. Đáng chú ý, các phương án thuộc nhóm dẫn đầu và nhóm cuối nhìn chung vẫn được duy trì dưới các nhiễu loạn được khảo sát; E4 liên tục thuộc nhóm thứ hạng thấp, trong khi E3 đứng cuối tại τ_3 ở cả hai mô hình. Do đó, kết quả so sánh giữa các mô hình và kết quả phân tích độ nhạy bổ trợ cho nhau trong việc cung cấp bằng chứng thực nghiệm về độ ổn định của cấu trúc xếp hạng của [CT3] trong các kịch bản được xem xét.

Tóm lại, phân tích đối sánh tại ba thời điểm thực nghiệm cho thấy mô hình đề xuất L-FAHP - L-FTOPSIS động dựa trên HA [CT3] duy trì được khả năng nhận diện các phương án chủ chốt đồng thời xử lý được sự thay đổi của tập phương án và sự tham gia của dữ liệu lịch sử trong môi trường MCGDM động. Khác với mô hình FAHP - FTOPSIS động [91] vốn ánh xạ thông tin ngôn ngữ sang tập mờ, giải mờ thành số thực và dùng trung bình cộng, mô hình đề xuất của [CT3] thực hiện biểu diễn và kết nhập thông tin trên miền 4-tuple ngữ nghĩa dựa trên HA, kết hợp cơ chế lưu giữ và suy giảm ngữ nghĩa đối với dữ liệu lịch sử. Cách tiếp cận này hướng tới duy trì trực tiếp hơn quan hệ thứ tự ngữ nghĩa và tính kế thừa của thông tin khi môi trường ra quyết định thay đổi; đồng thời, kết quả phân tích độ nhạy cho thấy các biến động thứ hạng chủ yếu mang tính cục bộ trong phạm vi các kịch bản khảo sát. Vì vậy, các kết quả tại Mục 4.6.2 cung cấp thêm cơ sở thực nghiệm về tính khả thi, khả năng xử lý dữ liệu lịch sử và độ ổn định của mô hình đề xuất trong bài toán đánh giá mức độ sẵn sàng chuyển đổi số của SMEs bán lẻ tại Việt Nam.

4.7. Kết luận chương 4

Trong chương này, luận án đã trình bày một số lý thuyết HA 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 “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” được thực hiện với mục tiêu phát triển các mô hình MCGDM có khả năng xử lý đánh giá ngôn ngữ trên nền tảng HA. Luận án tập trung vào ba thành phần cốt lõi của quy trình MCGDM: xếp hạng phương án, xác định trọng số của tiêu chí có kiểm soát tính nhất quán và mở rộng quá trình ra quyết định sang bối cảnh động. Các phương pháp nền như TOPSIS [59], FAHP [58, 15, 19], FTOPSIS [21] và mô hình 4-tuple ngữ nghĩa [54] được kế thừa có chọn lọc và phát triển theo hướng tiếp cận HA, phù hợp với phạm vi, khoảng trống khoa học và định hướng nghiên cứu đã xác định.

Trên cơ sở tổng quan và phân tích các nghiên cứu liên quan, Chương 1 xác định ba khoảng trống khoa học. Khoảng trống thứ nhất liên quan đến xếp hạng phương án MCGDM ngôn ngữ trên một miền ngữ nghĩa có cơ sở hình thức; khoảng trống thứ hai liên quan đến xác định trọng số từ phán đoán ngôn ngữ có kiểm soát tính nhất quán và duy trì sự liên thông với giai đoạn xếp hạng; khoảng trống thứ ba liên quan đến mở rộng MCGDM dựa trên HA từ bối cảnh tĩnh sang động, có xét sự thay đổi của cấu trúc bài toán, dữ liệu lịch sử và dữ liệu khuyết thiếu. Các Chương 2, 3 và 4 lần lượt phát triển ba nhóm kết quả tương ứng theo logic kế thừa, bổ sung và mở rộng.

Thứ nhất, Chương 2 phát triển mô hình TOPSIS ngôn ngữ dựa trên HA, ký hiệu HA-TOPSIS [CT1], cho bài toán xếp hạng phương án trong MCGDM ngôn ngữ. Mô hình xây dựng thang đo ngôn ngữ dựa trên HA và sử dụng hàm SQM để định lượng các giá trị ngôn ngữ sao cho $\vartheta(x) \in [0, 1]$ theo cách bảo toàn trật tự ngữ nghĩa; từ đó xác định PIS/NIS, đo khoảng cách ngữ nghĩa, tính hệ số gần và thiết lập thứ hạng trên cùng miền định lượng ngữ nghĩa. Cách tiếp cận này kế thừa nguyên lý của TOPSIS [59] nhưng tổ chức các bước tính toán theo cấu trúc ngữ nghĩa của HA. Mô hình được kiểm chứng trên bài toán lựa chọn nhà cung cấp phần mềm, kết hợp phân tích độ nhạy và đối sánh với FTOPSIS [21] và mô hình 4-tuple ngữ nghĩa [54] 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.

Thứ hai, Chương 3 phát triển mô hình MCGDM tích hợp FAHP cải tiến với 4-tuple ngữ nghĩa theo hướng tiếp cận HA, ký hiệu EFAHP - 4-tuple - HA [CT2]. Mô hình bổ sung quy trình xác định trọng số của tiêu chí từ ma trận so sánh cặp ngôn ngữ và kiểm soát tính nhất quán của phán đoán trong môi trường HA (Chương 3, mục 3.2). Ở giai đoạn thứ nhất, EFAHP-HA được sử dụng để xác định véc-tơ trọng số sau khi kiểm tra

tính nhất quán; ở giai đoạn thứ hai, biểu diễn 4-tuple ngữ nghĩa dựa trên HA được sử dụng để biểu diễn, kết nhập đánh giá và xếp hạng phương án. Cấu trúc này tạo sự liên thông giữa xác định trọng số của tiêu chí và tổng hợp đánh giá trên cùng nền tảng ngữ nghĩa (Chương 3, mục 3.3). Mô hình được ứng dụng vào đánh giá mức độ sẵn sàng chuyển đổi số của 16 SMEs bán lẻ tại Việt Nam (Chương 3, mục 3.4). Phân tích độ nhạy và đối sánh với HA-TOPSIS [CT1] và FAHP-FTOPSIS [58, 21] cho thấy cấu trúc thứ hạng duy trì tương đối ổn định trong các kịch bản nhiều đã khảo sát và mô hình vẫn phân biệt được SMEs khi nhiều phương án cùng thuộc một lớp ngôn ngữ (Chương 3, mục 3.5). Kết quả cho thấy tính khả thi, độ tin cậy và khả năng diễn giải của mô hình trong phạm vi thực nghiệm.

Thứ ba, Chương 4 mở rộng hướng tiếp cận sang MCGDM động và phát triển khung phương pháp dựa trên HA cho quá trình ra quyết định qua nhiều thời điểm [CT3]. Trên thang đo 4-tuple ngữ nghĩa, luận án xây dựng phép kết nhập ngôn ngữ nhị phân, phép kết nhập đối với dữ liệu khuyết thiếu và cơ chế suy giảm ngữ nghĩa để tích hợp dữ liệu lịch sử với đánh giá hiện tại (Chương 4, mục 4.3). Trên nền tảng đó, mô hình L-FAHP - L-FTOPSIS động dựa trên HA được xây dựng để xác định trọng số của tiêu chí và xếp hạng phương án khi tập phương án, bộ tiêu chí, hội đồng chuyên gia hoặc dữ liệu đánh giá thay đổi theo thời gian (Chương 4, mục 4.4). Mô hình được kiểm chứng trên bài toán đánh giá mức độ sẵn sàng chuyển đổi số của 10 SMEs tại Việt Nam với 4 tiêu chí chính, 14 tiêu chí con và 3 thời điểm đánh giá (Chương 4, mục 4.5). Kết quả thực nghiệm, phân tích độ nhạy và đối sánh với phương pháp liên quan cho thấy mô hình có khả năng khai thác thông tin lịch sử, xử lý các thay đổi của cấu trúc bài toán và duy trì khả năng diễn giải ngữ nghĩa trong các kịch bản đã xây dựng (Chương 4, mục 4.6). Qua đó, Chương 4 hoàn thiện bước mở rộng từ MCGDM tĩnh sang MCGDM động trong phạm vi luận án.

Từ các kết quả trên, có thể khẳng định rằng luận án đã tạo ra một số đóng góp mới có giá trị về cả lý thuyết, phương pháp luận và ứng dụng như sau:

- Về lý thuyết, luận án đã góp phần mở rộng lý thuyết HA thông qua việc đề xuất công thức tính chỉ số nhất quán trong môi trường ngôn ngữ, đồng thời xây dựng các phép kết nhập ngôn ngữ nhị phân, phép kết nhập với dữ liệu khuyết thiếu và mô hình hóa sự suy giảm ảnh hưởng của dữ liệu lịch sử trên thang đo 4-tuple ngữ nghĩa.

- Về phương pháp luận, luận án đã hình thành một khung phương pháp luận thống nhất cho MCGDM ngôn ngữ dựa trên HA có tính kế thừa, bổ sung và mở rộng: HA-TOPSIS giải quyết thành phần xếp hạng phương án trên miền ngữ nghĩa; EFAHP - 4-tuple - HA bổ sung quy trình xác định trọng số của tiêu chí có kiểm soát tính nhất quán và liên kết với kết nhập, xếp hạng; L-FAHP - L-FTOPSIS động tiếp tục mở rộng quy trình sang nhiều thời điểm, có tích hợp dữ liệu lịch sử và xử lý dữ liệu khuyết thiếu.

- Về ứng dụng, các mô hình đề xuất được kiểm chứng trên các bài toán có ý nghĩa thực tiễn cao, cho thấy hướng tiếp cận dựa trên HA có tiềm năng ứng dụng trong lựa chọn giải pháp công nghệ, đánh giá năng lực chuyển đổi số, quản trị hệ thống thông tin và các bài toán đánh giá định kỳ trong SMEs. Kết quả nghiên cứu tạo tiền đề cho việc phát triển các công cụ hỗ trợ ra quyết định có khả năng xử lý đánh giá ngôn ngữ, tổng hợp ý kiến nhóm chuyên gia và cung cấp kết quả có khả năng diễn giải.

Bên cạnh các kết quả đạt được, luận án còn một số hạn chế cần tiếp tục được nghiên cứu. Thứ nhất, các mô hình chủ yếu được xây dựng trong phạm vi HA tuyến tính và sử dụng các bộ tham số ngữ nghĩa theo giả thiết hoặc cấu hình thực nghiệm; ảnh hưởng của các cấu trúc HA phức tạp hơn và cơ chế lựa chọn tham số tối ưu chưa được khảo sát đầy đủ. Thứ hai, phạm vi kiểm chứng thực nghiệm tuy gắn với các bài toán có ý nghĩa thực tiễn nhưng vẫn còn giới hạn về số lượng miền ứng dụng, số lượng bộ dữ liệu và quy mô bài toán. Thứ ba, khả năng mở rộng của các thuật toán đối với bài toán MCGDM quy mô lớn, dữ liệu không đồng nhất, dữ liệu cập nhật liên tục hoặc mức độ tương tác cao giữa các tiêu chí cần được đánh giá sâu hơn. Thứ tư, luận án tập trung vào mô hình, phép toán và quy trình tính toán cốt lõi, chưa phát triển thành một hệ hỗ trợ ra quyết định hoàn chỉnh để kiểm chứng trong môi trường vận hành thực tế.

Từ những hạn chế trên, nghiên cứu tiếp theo có thể tập trung vào ba hướng. Trước hết, cần nghiên cứu cơ chế hiệu chỉnh hoặc tối ưu các tham số ngữ nghĩa của HA dựa trên dữ liệu thực nghiệm và phản hồi chuyên gia, đồng thời xem xét các cấu trúc biểu diễn phong phú hơn. Tiếp theo, cần kiểm chứng các mô hình trên nhiều bộ dữ liệu và miền ứng dụng, nghiên cứu khả năng mở rộng đối với MCGDM quy mô lớn, dữ liệu không đồng nhất, khuyết thiếu hoặc cập nhật theo thời gian thực; đồng thời có thể đối sánh sâu hơn với các mô hình dựa trên tập mờ Pythagorean, tập neutrosophic hoặc mô hình ngôn ngữ đa hạt để xác định rõ điều kiện áp dụng của hướng tiếp cận HA. Cuối cùng, cần phát triển hệ hỗ trợ ra quyết định tích hợp các mô hình đã đề xuất nhằm kiểm chứng khả năng vận hành và tạo cơ sở mở rộng sang các bài toán thực tế khác như đánh giá rủi ro, quản lý dự án, y tế, giáo dục, tài chính hoặc môi trường.

Tóm lại, luận án đã đạt được mục tiêu nghiên cứu đề ra là phát triển các mô hình MCGDM trong môi trường thông tin ngôn ngữ dựa trên HA và hình thành một khung phương pháp từ xếp hạng phương án, xác định trọng số của tiêu chí có kiểm soát tính nhất quán đến xử lý quá trình ra quyết định động. Các kết quả đạt được bổ sung cơ sở phương pháp luận cho xử lý thông tin ngôn ngữ trong MCGDM, đồng thời tạo nền tảng cho việc tiếp tục hoàn thiện và triển khai các mô hình dựa trên HA trong những bối cảnh ra quyết định phức tạp hơn.

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 và 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*).

TÀI LIỆU THAM KHẢO

Tiếng Việt

- [1] Bộ Thông tin và Truyền thông (2023): Quyết định số 2158/QĐ-BTTTT ngày 07/11/2023 quy định về việc phê duyệt đề án xác định chỉ số đánh giá mức độ chuyển đổi số doanh nghiệp, hỗ trợ thúc đẩy doanh nghiệp chuyển đổi số.
- [2] Nguyễn Cát Hồ (2006), Lý thuyết tập mờ và công nghệ tính toán mềm, *Tuyển tập các bài giảng về trường thu hệ mờ và ứng dụng*, lần in thứ 2, tr. 51–92.
- [3] Nguyễn Văn Long, Hoàng Văn Thông (2011), Vấn đề kết nhập thông tin biểu diễn bằng bộ 4 với ngữ nghĩa dựa trên đại số gia tử. *Tin học và Điều khiển học*, 27(3), 241–253. <https://doi.org/10.15625/1813-9663/27/3/488>.
- [4] Lê Xuân Vinh (2006), Về một cơ sở đại số và logic cho lập luận xấp xỉ và ứng dụng, *Luận án Tiến sĩ Toán học*, Viện Công nghệ Thông tin - Viện Khoa học và Công nghệ Việt Nam.

Tiếng Anh

- [5] Abbas, F., Ali, J., Mashwani, W. K., & Syam, M. I. (2024), “An integrated group decision-making method under q-rung orthopair fuzzy 2-tuple linguistic context with partial weight information”, *PLoS ONE*, 19, e0297462. <https://doi.org/10.1371/journal.pone.0297462>
- [6] Aggarwal, M., Krishankumar, R., Ravichandran, K. S., & Hanmandlu, M. (2024), “Cloud vendor selection using choice models based on interactive criteria and varying attitudes of experts”, *Expert Systems with Applications*, 239, 122021. <https://doi.org/10.1016/j.eswa.2023.122021>
- [7] Akram, M., Noreen, U., & Deveci, M. (2023), “Enhanced ELECTRE II method with 2-tuple linguistic m-polar fuzzy sets for multi-criteria group decision making”, *Expert Systems with Applications*, 213, (Part C), 119237. <https://doi.org/10.1016/j.eswa.2022.119237>
- [8] Akram, M., Ramzan, N., & Deveci, M. (2023). Linguistic Pythagorean fuzzy CRITIC-EDAS method for multiple-attribute group decision analysis. *Engineering Applications of Artificial Intelligence*, 119, 105777. <https://doi.org/10.1016/j.engappai.2022.105777>
- [9] Alkan, N., & Kahraman, C. (2023). CODAS extension using novel decomposed Pythagorean fuzzy sets: Strategy selection for IoT-based sustainable supply chain systems. *Expert Systems with Applications*, 237, Part C, 121534. <https://doi.org/10.1016/j.eswa.2023.121534>

- [10] Almjally, A., Khan, M., & Khan, A. (2025), “Broadband technology selection based on three-way decision-making model under the tripolar fuzzy information”, *International Journal of Fuzzy Systems*, Advance online Asat, S., Bahram, S. B., & Mazyar, Z. (2025), “Supplier selection utilizing AHP and TOPSIS in a fuzzy environment based on KPIs”. *Contemporary Mathematics*, 6, 574–593. <https://doi.org/10.37256/cm.6120256152>
- [11] Bastanifar, I., Khan, K. H., Azmi, S. N., & Opera, E. (2025), “Applying dynamic TOPSIS: A multi-criteria decision-making approach to economic corridors under uncertainty—the case of IMEEC”, *Cogent Economics & Finance*, 13, 2558028. <https://doi.org/10.1080/23322039.2025.2558028>
- [12] Bjarnason, E., Åberg, P., & Ali, N. B. (2023), “Software selection in large-scale software engineering: A model and criteria based on interactive rapid reviews”, *Empirical Software Engineering*, 28, Article 51. <https://doi.org/10.1007/s10664-023-10288-w>
- [13] Boix-Cots, D., Pardo-Bosch, F., & Pujadas, P. (2023), A systematic review on multi-criteria group decision-making methods based on weights: Analysis and classification scheme. *Information Fusion*, 96, 16–36. <https://doi.org/10.1016/j.inffus.2023.03.004>
- [14] Bui, T., & Jarke, M. (1984). A DSS for cooperative multiple criteria group decision making. In *Proceedings of the 5th International Conference on Information Systems (ICIS 1984)*, Tucson, Arizona
- [15] Buckley, (1985), “J.J. Fuzzy hierarchical analysis”, *Fuzzy Sets Syst*, 17, 233–247. [https://doi.org/10.1016/0165-0114\(85\)90090-9](https://doi.org/10.1016/0165-0114(85)90090-9)
- [16] Cai, C., Tu, Y., & Li, Z. (2023), “Enterprise digital transformation and ESG performance”, *Finance Research Letters*, 58(Part D), 104692. <https://doi.org/10.1016/j.frl.2023.104692>
- [17] Campanella, G., & Ribeiro, R. A. (2011), “A framework for dynamic multiple-criteria decision making”, *Decision Support Systems*, 52(1), 52–60. <https://doi.org/10.1016/j.dss.2011.05.003>
- [18] Chakraborty, S. (2022). TOPSIS and Modified TOPSIS: A comparative analysis. *Decision Analytics Journal*, 2, 100021. <https://doi.org/10.1016/j.dajour.2021.100021>
- [19] Chang, D.-Y. (1996), “Applications of the extent analysis method on fuzzy AHP”, *European Journal of Operational Research*, 95(3), 649–655. [https://doi.org/10.1016/0377-2217\(95\)00300-2](https://doi.org/10.1016/0377-2217(95)00300-2)

- [20] Chaurasiya, R., & Jain, D. (2023), “A new algorithm on Pythagorean fuzzy-based multi-criteria decision-making and its application”, *Iranian Journal of Science and Technology, Transactions of Electrical Engineering*, 47, 871–886. <https://doi.org/10.1007/s40998-023-00600-1>
- [21] Chen, T.-C. (2000), “Extensions of the TOPSIS for group decision-making under fuzzy environment”, *Fuzzy Sets and Systems*, 114(1), 1–9. [https://doi.org/10.1016/S0165-0114\(97\)00377-1](https://doi.org/10.1016/S0165-0114(97)00377-1)
- [22] Chen, T.-Y. (2022), “A point operator-driven approach to decision-analytic modeling for multiple criteria evaluation problems involving uncertain information based on T-spherical fuzzy sets”, *Expert Systems with Applications*, 203, 117559. <https://doi.org/10.1016/j.eswa.2022.117559>
- [23] Cheng, R., Fan, J., & Wu, M. (2023), “A dynamic multi-attribute group decision-making method with R-numbers based on MEREC and CoCoSo method”, *Complex & Intelligent Systems*, 9, 6393–6426. <https://doi.org/10.1007/s40747-023-01032-4>
- [24] Dağci Yüksel, B., & Ersöz, F. (2024), “Evaluation of ERP software selection criteria with fuzzy AHP approach: An application in the metal production enterprises in the aviation industry”, *International Journal of System Assurance Engineering and Management*, 15(6), 2656–2667. <https://doi.org/10.1007/s13198-024-02287-x>
- [25] Dat, L. Q., Yu, V. F., & Chou, S.-Y. (2012), “An improved ranking method for fuzzy numbers based on the centroid-index”, *International Journal of Fuzzy Systems*, 14, 3, 413–441.
- [26] Deason, J. P., Adams, S. J., Rahman, A., Lovo, S., & Mendez, I. (2025), “A technology selection tool applying multiple criteria decision analysis for virtual care implementation”. *Mayo Clinic Proceedings: Digital Health*, 3(4), 100260. <https://doi.org/10.1016/j.mcpgdig.2025.100260>
- [27] Delgado, M., Verdegay, J. L., & Vila, M. A. (1993), “On aggregation operations of linguistic labels”, *International Journal of Intelligent Systems*, 8(3), 351–370. <https://doi.org/10.1002/int.4550080303>
- [28] Díaz, G. M., & Salvador, J. L. G. (2023), “Group decision-making model based on 2-tuple fuzzy linguistic model and AHP applied to measuring digital maturity level of organizations”, *Systems*, 11, 341. <https://doi.org/10.3390/systems11070341>

- [29] Dong, Y., Cheng, X., Xu, Z., & Ma, T. (2024), “Multi-criteria group decision-making methods with dynamic probabilistic linguistic information characterized by multiple consecutive time points”, *International Journal of Machine Learning and Cybernetics*, 15, 1277–1293. <https://doi.org/10.1007/s13042-023-01967-7>
- [30] Dubois, D., & Prade, H. (1980), “Fuzzy sets and systems: Theory and applications”, *Academic Press*. <https://doi.org/10.1016/C2013-0-07153-8>
- [31] Duc, D. A., Phuong, B. H., Van, L. H., & Diep, P. T. H. (2019). A dynamic fuzzy multiple criteria decision-making approach for lecturer performance evaluation. *Journal of Management Information and Decision Sciences*, 22(3), 250–261.
- [32] Duc, D. A., Van, L. H., Yu, V. F., Chou, S.-Y., Hien, N. V., Chi, N. T., Toan, D. V., & Dat, L. Q. (2021). A dynamic generalized fuzzy multi-criteria group decision making approach for green supplier segmentation. *PLoS ONE*, 16(1), e0245187. <https://doi.org/10.1371/journal.pone.0245187>
- [33] Duy, N. T., & Son, T. T. (2019), “A group decision-making model with comparative linguistic expression base hedge algebra”, *Advances in Engineering Research and Application (ICERA 2018)*, Conference paper. *Springer*, 31–45. https://doi.org/10.1007/978-3-030-04792-4_7
- [34] Erdebilli, B., Yilmaz, I., Aksoy, T., Hacıoglu, U., Yüksel, S., & Dinçer, H. (2023), “An interval-valued Pythagorean fuzzy AHP and COPRAS hybrid methods for the supplier selection problem”, *International Journal of Computational Intelligence Systems*, 16, 124. <https://doi.org/10.1007/s44196-023-00297-4>
- [35] Erdem, M., & Özdemir, A. (2024), “Sustainability and risk assessment of data center locations under a fuzzy environment”, *Journal of Cleaner Production*, 450, 141982. <https://doi.org/10.1016/j.jclepro.2024.141982>
- [36] Faizi, S., Sałabun, W., Shaheen, N., Rehman, A. U., & Wątróbski, J. (2021), “A novel multi-criteria group decision-making approach based on Bonferroni and Heronian mean operators under hesitant 2-tuple linguistic environment”, *Mathematics*, 9(13), 1489. <https://doi.org/10.3390/math9131489>
- [37] Ferrarini, F., Muzzioli, S., & De Baets, B. (2024). A TOPSIS analysis of regional competitiveness at European level. *Competitiveness Review*, 34(7), 52–72. <https://doi.org/10.1108/CR-01-2024-0005>

- [38] Fronte, P., Agell, N., Mesa, D., & Torrens, M. (2025), “Criteria definition for digital requirements using hesitant fuzzy linguistic terms sets: An application to the automotive industry”, *Annals of Operations Research*, 353, 147–169. <https://doi.org/10.1007/s10479-024-06449-9>
- [39] Gabriel Marín Díaz (2025). A Fuzzy-XAI Framework for Customer Segmentation and Risk Detection: Integrating RFM, 2-Tuple Modeling, and Strategic Scoring. *Mathematics*, 113, 2141. <https://doi.org/10.3390/math13132141>.
- [40] Ghorui, N., Mondal, S. P., Chatterjee, B., Ghosh, A., Pal, A., De, D., & Giri, B. C. (2023), “Selection of cloud service providers using MCDM methodology under intuitionistic fuzzy uncertainty” *Soft Computing*, 27, 2403–2423. <https://doi.org/10.1007/s00500-022-07772-8>
- [41] Gongora-Salazar, P., Rocks, S., Fahr, P., Rivero-Arias, O., & Tsiachristas, A. (2023), “The use of multicriteria decision analysis to support decision making in healthcare: An updated systematic literature review”, *Value in Health*, 26(5), 780–790. <https://doi.org/10.1016/j.jval.2022.11.007>
- [42] Gouveia, F.D.; Mamede, H.S. (2022), “Digital Transformation for SMEs in the Retail Industry”, *Procedia Comput. Sci*, 204, 671–681. <https://doi.org/10.1016/j.procs.2022.08.081>.
- [43] Greco, S., Figueira, J., & Ehrgott, M. (Eds.). (2016), *Multiple criteria decision analysis*. Springer.
- [44] Guangfen, Y., Min, R., Xiumei, H. (2023), “Multi-criteria decision-making problem based on the novel probabilistic hesitant fuzzy entropy and TODIM method”, *Alexandria Engineering Journal*. 68, 437–451. <https://doi.org/10.1016/j.aej.2023.01.014>
- [45] Gülçin, B.; Sezin, G. (2016), “Multi Criteria Group Decision Making Approach for Smart Phone Selection Using Intuitionistic Fuzzy TOPSIS”, *International Journal of Computational Intelligence Systems*, 9(4), 709–725. <https://doi.org/10.1080/18756891.2016.1204119>
- [46] Hai V. P.; Hoang, T. L.; Hung, Q. N.; Phung, K. T. (2023), “Proposed Intelligent Decision Support System Using Hedge Algebra Integrated with Picture Fuzzy Relations for Improvement of Decision-Making in Medical Diagnoses”, *Springer*, 25, 3260–3270. <https://doi.org/10.1007/s40815-023-01548-4>

- [47] Hai, V. P.; P, Moore; Khang, D.T. (2014), “Context Matching with Reasoning and Decision Support using Hedge Algebra with Kansei Evaluation”, *Proc of the fifth symposium on Information and Communication Technology (SoICT 2014)*; ACM, ACM, New York, NY, USA: Hanoi, Vietnam, 202-210. doi: 10.1145/2676585.2676598
- [48] Han, W.; Zhen-Song, C.; Yaya, L.; Luis M.; Zeng qiang, W. (2026). Enhancing risk management efficiency for building information modelling application via basic uncertain linguistic large-scale group decision making. *Applied Soft Computing*, 186, Part B, 114086 <https://doi.org/10.1016/j.asoc.2025.114086>
- [49] Herrera, F., & Martínez, L. (2000), “A 2-tuple fuzzy linguistic representation model for computing with words”, *IEEE Transactions on Fuzzy Systems*, 8(6), 746–752. <https://doi.org/10.1109/91.890332>
- [50] Hieu, D. N.; Ho, C. N; Phong, D. P; Lan, N. V.; Hiep, H. P. (2022). Scalable human knowledge about numeric time series variation and its role in improving forecasting results. *Journal of Computer Science and Cybernetics*, 38, 03–130. <https://doi.org/10.15625/1813-9663/38/2/16125>
- [51] Hieu, D.N; Phong, D.P. (2021). A Novel High-order Linguistic Time Series Forecasting Model with the Growth of Declared Word-set. *International Journal of Advanced Computer Science and Applications*, 12(6), 6371. doi: 10.14569/IJACSA.2021.0120609.
- [52] Ho, C. N., & Long, V. N. (2007), “Fuzziness measure on complete hedges algebras and quantifying semantics of terms in linear hedge algebras. *Fuzzy Sets and Systems*, 158(4), 452–471. <https://doi.org/10.1016/j.fss.2006.10.023>
- [53] Ho, C. N., & Wechler, W. (1990), “Hedge algebras: An algebraic approach to structures of sets of linguistic domains of linguistic truth variables”, *Fuzzy Sets and Systems*, 35(3), 281–293. [https://doi.org/10.1016/0165-0114\(90\)90002-N](https://doi.org/10.1016/0165-0114(90)90002-N)
- [54] Ho, C. N., Huynh, V. N., & Pedrycz, W. (2014), “A construction of sound semantic linguistic scales using 4-tuple representation of term semantics”, *International Journal of Approximate Reasoning*, 55(3), 763–786. <https://doi.org/10.1016/j.ijar.2013.10.012>
- [55] Ho, C. N., Long, V. N., & Thong, H. V. (2015), “A discussion on interpretability of linguistic rule base systems and its application to solve

- regression problems”, *Knowledge-Based Systems*, 88, 107–133.
<https://doi.org/10.1016/j.knosys.2015.08.002>
- [56] Huang, Q., & Tang, Y. (2025), “Enterprise digital transformation strategy: The impact of digital platforms”, *Entropy*, 27, 295.
<https://doi.org/10.3390/e27030295>
- [57] Huang, J., Zhang, Y., Xu, M., Lv, Y., Zhang, J., & Shafieezadeh, M. M. (2025), “Hybrid FAHP-FTOPSIS methodology for objective football player selection and ranking”, *Scientific Reports*, 15, 27913.
<https://doi.org/10.1038/s41598-025-13973-6>
- [58] Hue, T. T., Tuan, A. N., Van, H. L., Lien, T. L., Huong, D. D., Anh, T. L., Huy, X. N., & Dat, Q. L. (2022), “Prioritization of factors impacting lecturer research productivity using an improved fuzzy analytic hierarchy process approach”, *Sustainability*, 14, 6134. <https://doi.org/10.3390/su14106134>
- [59] Hwang, C.-L., & Yoon, K. (1981), “Multiple attribute decision making: Methods and applications, a state-of-the-art survey”, *Springer*, 1, XI, 269.
<https://doi.org/10.1007/978-3-642-48318-9>
- [60] Iman, B.; Kashif, H. K.; Shujaat, N. A. & E. O. (2025), “Applying dynamic TOPSIS: a multi-criteria decision-making approach to economic corridors under uncertainty—the case of IMEEC”, *Cogent Economics & Finance*, 13, 2558028, <https://doi.org/10.1080/23322039.2025.2558028>
- [61] Jiang, Y. (2018). A general type-2 fuzzy model for computing with words. *Intl. J. Intell. Syst*, 33, 713–758. <https://doi.org/10.1002/int.21952>.
- [62] Jin, W., Wang, K. & Zhu, E. (2025), “A framework based on multiple reliability criteria and Markov chain for optimal selection of cloud services”, *Computing*, 107, 230. <https://doi-org.dbvista.idm.oclc.org/10.1007/s00607-025-01575-z>
- [63] Jin, F., Guo, S., Cai, Y., Liu, J., & Zhou, L. (2023). 2-tuple linguistic decision-making with consistency adjustment strategy and data envelopment analysis. *Engineering Applications of Artificial Intelligence*, 118, 105671. <https://doi.org/10.1016/j.engappai.2022.105671>
- [64] Ju, Y. B., & Wang, A. H. (2012), “Emergency alternative evaluation under group decision makers: A method of incorporating DS/AHP with extended TOPSIS”, *Expert Systems with Applications*, 39(1), 1315–1323.
<https://doi.org/10.1016/j.eswa.2011.08.012>

- [65] Kao, L. J., Chiu, C. C., Lin, H. T., Hung, Y. W., & Lu, C. C. (2022), “Evaluating the digital transformation performance of retail by the DEA approach”. *Axioms*, 11, 284. <https://doi.org/10.3390/axioms11060284>
- [66] Khodadad, O.; Hadi-Vencheh, A.; Karbassi, Y. (2025), “An Enhanced Hierarchical Fuzzy TOPSIS-ANP Method for Supplier Selection in an Uncertain Environment”, *Mathematics*, 13(21), 3417. <https://doi.org/10.3390/math13213417>
- [67] Kendall, M. G. (1938), “A new measure of rank correlation”, *Biometrika*, 30(1–2), 81–93. <https://doi.org/10.2307/2332226>
- [68] Kim, J. H., & Ahn, B. S. (2019). Extended VIKOR method using incomplete criteria weights. *Expert Systems with Applications*, 126, 124–132. <https://doi.org/10.1016/j.eswa.2019.02.019>
- [69] Lan, T. H. L., Hien, T. T. D., Thong, T. T., Smarandache, F., & Giang, L. N. (2023). An ANP-TOPSIS model for tourist destination choice problems under temporal neutrosophic environment. *Applied Soft Computing*, 136, 110146. <https://doi.org/10.1016/j.asoc.2023.110146>
- [70] Latifi, F., Nassiri, R., Mohsenzadeh, M., & Mostafaei, H. (2025), “A Markov chain-based multi-criteria framework for dynamic cloud service selection using user feedback”, *The Journal of Supercomputing*, 81, Article 89. <https://doi.org/10.1007/s11227-024-06508-9>
- [71] Lazarashouri, H.; Najafi, S.E. (2024), “Enhancing emergency department efficiency through simulation and fuzzy multi-criteria decision-making integration”, *J. Oper. Strateg. Anal*, 2, 56–71. <https://doi.org/10.56578/josa020106>.
- [72] Le, H. B., An, N. T., & Quoc, H. C. (2023). “Active control based on hedge-algebras theory of seismic-excited buildings with upgraded tuned liquid column damper”, *Journal of Engineering Mechanics*, 149(1), 04022091. <https://doi.org/10.1061/JENMDT.EMENG-6821>
- [73] Li, T.; Wen, J.; Zeng, D.; Liu, K. (2022), “Has enterprise digital transformation improved the efficiency of enterprise technological innovation? A case study on Chinese listed companies”, *Math. Biosci. Eng. Aug.*, 19, 12632–12654. <https://doi.org/10.3934/mbe.2022590>.
- [74] Li, Y., Kou, G., Peng, Y., et al. (2024), “Z-number linguistic term set for multi-criteria group decision-making and its application in predicting the

- acceptance of academic papers”, *Applied Intelligence*, 54, 10962–10981.
<https://doi.org/10.1007/s10489-024-05765-8>
- [75] Li, Z.; Zhang, X.; Tao, Z.; Wang, B. (2024). Enterprise digital transformation and supply chain management. *Financ. Res. Lett.*, 60, 104883.
<https://doi.org/10.1016/j.frl.2023.104883>.
- [76] Liao, H., He, Y., Wu, X., Wu, Z., & Baušys, R. (2023), “Reimagining multi-criterion decision making by data-driven methods based on machine learning: A literature review”, *Information Fusion*, 100, 101970.
<https://doi.org/10.1016/j.inffus.2023.101970>
- [77] Liu, T.; Leng, J.; Zhu, S.; Fu, R. (2025), “Digital Transformation and Enterprise Innovation Capability: From the Perspectives of Enterprise Cooperative Culture and Innovative Culture”. *J. Theor. Appl. Electron. Commer. Res.*, 20, 136. <https://doi.org/10.3390/jtaer20020136>.
- [78] Liu, Y., Yang, Y., Sun, L., & Huang, A. (2023), “Managing multi-granular probabilistic linguistic information in large-scale group decision making: A personalized individual semantics-based consensus model”, *Expert Systems with Applications*, 235, 120645. <https://doi.org/10.1016/j.eswa.2023.120645>
- [79] Lorge, I., & Solomon, H. (1955), “Two models of group behavior in the solution of eureka-type problems”, *Psychometrika*, 20(2), 139–148.
<https://doi.org/10.1007/BF02288986>
- [80] Luo, Y.; Cui, H.; Zhong, H.; Wei, C. (2023). Business environment and enterprise digital transformation. *Financ. Res. Lett.*, 57, 104250.
<https://doi.org/10.1016/j.frl.2023.104250>.
- [81] Maghrabie, H.F., Beauregard, Y. and Schiffauerova, A. (2019), “Multi-criteria decision making problems with unknown weight information under uncertain evaluations”, *Computers & Industrial Engineering*, 133, 131–138.
<https://doi.org/10.1016/j.cie.2019.05.003>
- [82] Marković M. G., Nikola, K., Božidar K. (2025), “A Novel ANP-DEMATEL Framework for Multi-Criteria Decision-Making in Adaptive E-Learning Systems”, *Mathematics*, 13(22), 3714;
<https://doi.org/10.3390/math13223714>
- [83] Mencar, C., & Fanelli, A. M. (2008), “Interpretability constraints for fuzzy information granulation”, *Information Sciences*, 178(24), 4585–4618.
<https://doi.org/10.1016/j.ins.2008.08.015>
- [84] Miller, G. A. (1956), “The magical number seven, plus or minus two: Some limits on our capacity for processing information”, *Psychological Review*, 63(2), 81–97. <https://doi.org/10.1037/h0043158>

- [85] Nguyen, P.-H., Pham, T.-V., Nguyen, L.-A. T., Pham, H.-A. T., Nguyen, T.-H. T., & Vu, T.-G. (2024), “Assessing cybersecurity risks and prioritizing top strategies in Vietnam’s finance and banking system using strategic decision-making models-based neutrosophic sets and Z number”, *Heliyon*, 10, e37893. <https://doi.org/10.1016/j.heliyon.2024.e37893>
- [86] Nguyen, T. D., & Bui, H. L. (2023). General optimization procedure of the hedge-algebras controller for controlling dynamic systems. *Artificial Intelligence Review*, 56, 2749–2784. <https://doi.org/10.1007/s10462-022-10242-0>
- [87] Nhu, N. P. H.; Duc, D. N.; Song, T. Q. L. (2024), “Sustainable supplier selection in the apparel industry: an integrated AHP-TOPSIS model for multi-criteria decision analysis”, *Research Journal of Textile and Apparel*, 29, 4, 822-841. <https://doi.org/10.1108/RJTA-04-2024-0056>
- [88] Phong, D. P.; Lan, T. P.; Thanh, X. T. (2025). Linguistic Summarization and Outlier Detection of Blended Learning Data. *Appl. Sci.*, 15(12), 6644. <https://doi.org/10.3390/app15126644>
- [89] Phuong, A. L.; Nhan, H. T.; Uyen, T. N.; Khang, D. N. (2022), “An approach for linguistic multi-attribute decision making based on linguistic many-valued logic”, *International Journal of Advances in Intelligent Informatics*, 8, 1, 21-32. <https://doi.org/10.26555/ijain.v8i1.820>
- [90] Qian, X., Chen, M., Zhao, F., & Ling, H. (2024), “An assessment framework of global smart cities for sustainable development in a post-pandemic era”, *Cities*, 150, 104990. <https://doi.org/10.1016/j.cities.2024.104990>
- [91] Quynh, V. T. N. (2023), “An integrated dynamic generalized trapezoidal fuzzy AHP-TOPSIS approach for evaluating sustainable performance of bank”. *Advances in Decision Sciences*, 27, 68–86. <https://doi.org/10.47654/v27y2023i1p68-86>
- [92] Rani, A., Mishra, D., & Omerovic, A. (2024), “Multi-Criteria Decision-Making Methods: A Case of Software Vendor Selection”, *TEM Journal*, 13(2), 1162-1175. <https://doi.org/10.18421/TEM132-35>
- [93] Ren, F., & Hao, F. (2025), “A new fusion method of fuzzy numbers and linguistic terms based on individual semantics in mixed decision making problems”, *International Journal of Fuzzy Systems*, 27, 1887–1903. <https://doi.org/10.1007/s40815-024-01879-w>
- [94] Roy, B. (1968), “Classement et choix en présence de points de vue multiples (la méthode ELECTRE)”. *Revue française d'automatique*,

- d'informatique et de recherche opérationnelle*, 8, 57-75.
<https://doi.org/10.1051/ro/196802V100571>
- [95] Saaty, T. L. (1977), "A scaling method for priorities in hierarchical structures", *Journal of Mathematical Psychology*, 15(3), 234–281.
[https://doi.org/10.1016/0022-2496\(77\)90033-5](https://doi.org/10.1016/0022-2496(77)90033-5)
- [96] Saaty, T. L. (2004), "Decision making—the analytic hierarchy and network processes (AHP/ANP)". *Journal of Systems Science and Systems Engineering*, 13(1), 1–35.
- [97] Sarsangi, V., Karimi, A., Hadavandi, E., & Hokmabadi, R. (2023), "Prioritizing risk factors of hazardous material road transportation accidents using the fuzzy AHP method", *WORK*, 75, 275–286.
<https://doi.org/10.3233/WOR-211446>
- [98] Shih, H.S.; Shyur, H.J.; Lee, E.S. (2007), "An extension of TOPSIS for group decision making", *Math. Comput. Model*, 45, 801–813.
<https://doi.org/10.1016/j.mcm.2006.03.023>.
- [99] Shu, Z., Carrasco González, R. A., García-Miguel, J. P., & Sánchez-Montañés, M. (2024), "A model integrating the 2-tuple linguistic model and the CRITIC-AHP method for hotel classification", *SN Computer Science*, 5, 9. <https://doi.org/10.1007/s42979-023-02344-5>
- [100] Spearman, C. (1904), "The proof and measurement of association between two things", *The American Journal of Psychology*, 15(1), 72–101.
<https://doi.org/10.2307/1412159>
- [101] Sobczak, E. (2025), "Digital Transformation of Enterprises and Employment in Technologically Advanced and Knowledge-Intensive Sectors in the European Union Countries". *Sustainability*, 17, 5868.
<https://doi.org/10.3390/su17135868>.
- [102] Su, Y.; Wu, J. (2024). Digital transformation and enterprise sustainable development. *Financ. Res. Lett*, 60, 104902.
<https://doi.org/10.1016/j.frl.2023.104902>.
- [103] Tai, S. N.; Thoa, T. M.; Le, H. L. Bui. (2023), "Motion Control of a Mobile Robot Using the Hedge–Algebras-Based Controller" *Journal of Robotics*, 1, 6613293. <https://doi.org/10.1155/2023/6613293>
- [104] Thong, H. V., Long, V. N., Du, D. N., Phong, D. P., & Ho, C. N. (2022), "The interpretability and scalability of linguistic-rule-based systems for

- solving regression problems”, *International Journal of Approximate Reasoning*, 149, 131–160. <https://doi.org/10.1016/j.ijar.2022.07.007>
- [105] Thong, N.T.; Smarandache, F.; Hoa, N.D.; Son, L.H.; Lan, Luong T.H.; Cu, N.G. (2020). A Novel Dynamic Multi-Criteria Decision Making Method Based on Generalized Dynamic Interval-Valued Neutrosophic. *Symmetry*, 12, 618. <https://doi.org/10.3390/sym12040618>
- [106] Van Laarhoven, P. J. M., & Pedrycz, W. (1983), “A fuzzy extension of Saaty’s priority theory”, *Fuzzy Sets and Systems*, 11(1–3), 229–241. [https://doi.org/10.1016/S0165-0114\(83\)80082-7](https://doi.org/10.1016/S0165-0114(83)80082-7)
- [107] Wadud, I.; Liang, Y.D.; Polkinghorne, M. (2025). Digital Transformation in the UK Retail Sector. *Encyclopedia*, 5, 142. <https://doi.org/10.3390/encyclopedia5030142>
- [108] Wang, Y. M., Luo, Y., & Hua, Z. S. (2008). On the extent analysis method for fuzzy AHP and its applications. *European Journal of Operational Research*, 186(2), 735–747. <https://doi.org/10.1016/j.ejor.2007.01.050>
- [109] Weihua X., Zishuo Y. (2025), “Preference ranking organization method for enrichment evaluation-based feature selection for multiple source ordered information systems”, *Engineering Applications of Artificial Intelligence*, 142, 15, 109935. <https://doi.org/10.1016/j.engappai.2024.109935>
- [110] Więckowski, J., & Sałabun, W. (2023), “Sensitivity analysis approaches in multi-criteria decision analysis: A systematic review”, *Applied Soft Computing*, 148, 110915. <https://doi.org/10.1016/j.asoc.2023.110915>
- [111] Więckowski, J., & Sałabun, W. (2025), “Supporting multi-criteria decision-making processes with unknown criteria weights”, *Engineering Applications of Artificial Intelligence*, 140, 109699. <https://doi.org/10.1016/j.engappai.2024.109699>
- [112] Yang, Y.; Meng-Qi, J.; Zhen-Song, C. (2024). Dynamic three-way multi-criteria decision making with basic uncertain linguistic information: A case study in product ranking. *Applied Soft Computing Journal*, 152, 111228. <https://doi.org/10.1016/j.asoc.2024.111228>
- [113] Yibo Wang, Xiuqin Ma, Hongwu Qin, Huanling Sun, Weiyi Wei. (2024), “Hesitant Fermatean fuzzy Bonferroni mean operators for multi-attribute decision-making”, *Complex & Intelligent Systems*, 10:1425–1457. <https://doi.org/10.1007/s40747-023-01203-3>

- [114] Yue, Z. (2011), “A method for group decision-making based on determining weights of decision makers using TOPSIS”, *Applied Mathematical Modelling*, 35(4), 1926–1936.
<https://doi.org/10.1016/j.apm.2010.11.001>
- [115] Zadeh, L. A. (1965), “Fuzzy sets”, *Information and Control*, 8, 338–353.
- [116] Zadeh, L. A. (1975a). The concept of a linguistic variable and its application to approximate reasoning—I. *Information Sciences*, 8(3), 199–249.
- Zadeh, L. A. (1975b). The concept of a linguistic variable and its application to approximate reasoning—II. *Information Sciences*, 8(4), 301–357.
- Zadeh, L. A. (1975c). The concept of a linguistic variable and its application to approximate reasoning—III. *Information Sciences*, 9(1), 43–80.
- [117] Zakeri, S., Konstantas, D., Chatterjee, P., & Zavadskas, E. K. (2025), “Soft cluster-rectangle method for eliciting criteria weights in multi-criteria decision-making”, *Scientific Reports*, 15(1), Article 284.
<https://doi.org/10.1038/s41598-024-81027-4>
- [118] Zhang, T., Shi, Z. Z., Shi, Y. R., & Chen, N. J. (2021), “Enterprise digital transformation and production efficiency: Mechanism analysis and empirical research”, *Economic Research-Ekonomska Istraživanja*, 35, 2781–2792.
<https://doi.org/10.1080/1331677X.2021.1980731>

PHỤ LỤC A. MỘT SỐ KẾT QUẢ

A1. Một số kết quả trong Chương 3

Bảng A1.1. Bộ tiêu chí đánh giá mức độ sẵn sàng chuyển đổi số của SMEs trong lĩnh vực bán lẻ tại Việt Nam

Tiêu chí chính	Tiêu chí con
Định hướng chiến lược (P1)	Lãnh đạo hiểu biết chuyển đổi số trong lĩnh vực kinh doanh của doanh nghiệp (P1.1)
	Chuyển đổi số là một nhiệm vụ trọng tâm ưu tiên mang tính chiến lược (P1.2)
	Ngân sách được phân bổ cho đổi mới và các sáng kiến chuyển đổi số nội bộ (P1.3)
	Công cụ số và phân tích dữ liệu được sử dụng để hỗ trợ các quyết định chiến lược (P1.4)
Trải nghiệm khách hàng và bán hàng đa kênh (P2)	Công cụ số giúp nâng cao hoạt động marketing, phân phối và bán hàng nhằm cải thiện trải nghiệm khách hàng và năng lực cạnh tranh (P2.1)
	Chuyển đổi số cải thiện chất lượng dịch vụ và trải nghiệm khách hàng (P2.2)
	Hệ thống quản trị quan hệ khách hàng (CRM) được tích hợp để mở rộng chức năng và năng lực (P2.3)
	Phân tích dữ liệu hỗ trợ đo lường hiệu quả hoạt động (P2.4)
	Giải pháp số cho phép dự báo bán hàng và điều chỉnh chiến lược (P2.5)
Tích hợp chuỗi cung ứng (P3)	Công cụ số hỗ trợ đánh giá hiệu suất và dự báo (P3.1)
	Chuyển đổi số nâng cao năng lực dự báo và quản trị chuỗi cung ứng (P3.2)
	Giải pháp số đồng bộ công tác lập ngân sách với kế hoạch kinh doanh (P3.3)
	Công nghệ số cho phép chuỗi cung ứng linh hoạt và đáp ứng nhanh (P3.4)
	Chuyển đổi số thúc đẩy tự động hoá trong các quy trình sản xuất (P3.5)
	Ứng dụng số nâng cao hiệu quả và tự động hoá quản lý tồn kho (P3.6)
	Hệ thống số hỗ trợ ra quyết định trong vận hành dựa trên dữ liệu (P3.7)
Quản lý tài chính, kế hoạch, nhân sự, pháp chế (P4)	Chuyển đổi số hỗ trợ nâng cao hiệu quả kinh tế và ra quyết định (P4.1)
	Chuyển đổi số tăng cường tự động hoá và tính minh bạch trong quản trị nhân sự (P4.2)
	Công cụ số hỗ trợ nhận diện và đánh giá rủi ro kinh doanh (P4.3)
Hệ thống thông tin và quản trị dữ liệu (P5)	Chuyển đổi số giúp doanh nghiệp duy trì sự phù hợp với các xu hướng công nghệ (P5.1)
	Các công nghệ số mới nổi cải thiện chi phí và hiệu quả quản trị doanh nghiệp (P5.2)
	Giải pháp số cho phép tích hợp linh hoạt với các công nghệ mới (P5.3)
	Nguồn lực được phân bổ để nâng cấp hệ thống thông tin số (P5.4)

Tiêu chí chính	Tiêu chí con
	Có các quy trình và chính sách nhằm hỗ trợ triển khai chuyển đổi số (P5.5)
Quản lý rủi ro và an ninh mạng (P6)	Doanh nghiệp nhận thức đầy đủ về rủi ro an ninh mạng và rủi ro dữ liệu trong chuyển đổi số (P6.1)
	Hệ thống giám sát, đánh giá và giảm thiểu các rủi ro liên quan đến hệ thống thông tin (P6.2)
	Hệ thống tự động phát hiện và cảnh báo các mối đe dọa an ninh mạng (P6.3)
	Có các quy trình ứng phó sự cố đối với công nghệ thông tin và an ninh mạng (P6.4)
Năng lực con người và tổ chức (P7)	Nhân sự chuyển đổi số thích ứng nhanh và tích cực trước thay đổi (P7.1)
	Tổ chức có cơ cấu linh hoạt, thích ứng với các thay đổi bên trong và bên ngoài (P7.2)
	Nhân sự có đủ kiến thức, kỹ năng và kinh nghiệm cần thiết (P7.3)
	Có chính sách thu hút và giữ chân nhân sự CNTT có năng lực, cam kết gắn bó dài hạn (P7.4)
	Chương trình đào tạo nội bộ được triển khai hiệu quả nhằm nâng cao năng lực của người lao động (P7.5)
	Chia sẻ dữ liệu theo thời gian thực hỗ trợ quản trị doanh nghiệp tích hợp và hiệu quả (P7.6)

Bảng A1.2. Ma trận so sánh mờ trung bình cho bảy tiêu chí chính

Tiêu chí chính	P1	P2	P3	P4	P5	P6	P7
P1	(0.430, 0.540, 0.654; 1)	(0.687, 0.812, 0.923; 0.9)	(0.837, 0.952, 1.069; 1)	(0.601, 0.715, 0.837; 0.9)	(0.894, 1.003, 1.069; 1)	(0.629, 0.747, 0.866; 0.9)	(0.455, 0.570, 0.690; 0.8)
P2	(1.086, 1.236, 1.465; 0.9)	(0.430, 0.540, 0.654; 1)	(0.626, 0.751, 0.869; 0.9)	(0.715, 0.847, 0.991; 0.8)	(0.712, 0.842, 0.998; 0.9)	(0.772, 0.909, 1.091; 0.9)	(0.311, 0.425, 0.546; 0.8)
P3	(0.946, 1.060, 1.208; 1)	(1.155, 1.340, 1.609; 0.9)	(0.430, 0.540, 0.654; 1)	(0.547, 0.655, 0.741; 1)	(0.542, 0.662, 0.776; 0.9)	(0.485, 0.603, 0.716; 0.8)	(0.713, 0.891, 1.111; 0.7)
P4	(1.208, 1.425, 1.708; 0.9)	(1.155, 1.375, 1.738; 0.8)	(1.396, 1.639, 2.022; 1)	(0.430, 0.540, 0.654; 1)	(0.541, 0.658, 0.780; 0.9)	(0.602, 0.719, 0.834; 0.9)	(0.369, 0.485, 0.597; 0.8)
P5	(0.946, 1.003, 1.126; 1)	(1.081, 1.279, 1.532; 0.9)	(1.299, 1.564, 1.926; 0.9)	(1.299, 1.546, 1.889; 0.9)	(0.430, 0.540, 0.654; 1)	(0.397, 0.514, 0.629; 0.8)	(0.369, 0.483, 0.602; 0.8)
P6	(1.167, 1.362, 1.627; 0.9)	(1.091, 1.317, 1.611; 0.9)	(1.495, 1.823, 2.406; 0.8)	(1.230, 1.443, 1.745; 0.9)	(1.627, 2.007, 2.668; 0.8)	(0.430, 0.540, 0.654; 1)	(0.369, 0.485, 0.597; 0.8)
P7	(1.483, 1.808, 2.319; 0.8)	(1.944, 2.587, 3.799; 0.8)	(1.663, 2.401, 4.154; 0.7)	(1.690, 2.088, 2.786; 0.8)	(1.745, 2.238, 3.118; 0.8)	(1.690, 2.088, 2.786; 0.8)	(0.430, 0.540, 0.654; 1)
Chuyên gia 1: CR = 0.04		Chuyên gia 2: CR = 0.048		Chuyên gia 3: CR = 0.037		Chuyên gia 4: CR = 0.054	

Bảng A1.3. Các giá trị mở rộng tổng hợp mờ của 7 tiêu chí chính và 34 tiêu chí con

Tiêu chí chính	Giá trị mở rộng tổng hợp mờ	Tiêu chí con	Giá trị mở rộng tổng hợp mờ
P1	(0.07, 0.10, 0.14; 0.70)	P1.1	(0.10, 0.15, 0.21; 0.80)
		P1.2	(0.14, 0.19, 0.26; 0.80)
		P1.3	(0.25, 0.33, 0.43; 0.80)
		P1.4	(0.24, 0.32, 0.43; 0.80)
P2	(0.07, 0.10, 0.15; 0.70)	P2.1	(0.11, 0.15, 0.20; 0.90)
		P2.2	(0.15, 0.19, 0.24; 0.90)
		P2.3	(0.15, 0.19, 0.25; 0.90)
		P2.4	(0.17, 0.22, 0.28; 0.90)
		P2.5	(0.19, 0.25, 0.31; 0.90)
P3	(0.07, 0.11, 0.15; 0.70)	P3.1	(0.11, 0.14, 0.19; 0.70)
		P3.2	(0.07, 0.10, 0.13; 0.70)
		P3.3	(0.09, 0.12, 0.16; 0.70)
		P3.4	(0.09, 0.12, 0.17; 0.70)
		P3.5	(0.12, 0.17, 0.23; 0.70)
		P3.6	(0.13, 0.19, 0.24; 0.70)
		P3.7	(0.13, 0.17, 0.23; 0.70)
P4	(0.09, 0.13, 0.18; 0.70)	P4.1	(0.18, 0.25, 0.32; 0.90)
		P4.2	(0.23, 0.30, 0.38; 0.90)
		P4.3	(0.37, 0.45, 0.55; 0.90)
P5	(0.09, 0.13, 0.18; 0.70)	P5.1	(0.11, 0.15, 0.20; 0.70)
		P5.2	(0.11, 0.15, 0.21; 0.70)
		P5.3	(0.13, 0.18, 0.24; 0.70)
		P5.4	(0.18, 0.24, 0.33; 0.70)
		P5.5	(0.21, 0.28, 0.37; 0.70)
P6	(0.12, 0.17, 0.24; 0.70)	P6.1	(0.14, 0.19, 0.25; 0.80)
		P6.2	(0.17, 0.22, 0.29; 0.80)
		P6.3	(0.26, 0.32, 0.41; 0.80)
		P6.4	(0.20, 0.26, 0.33; 0.80)
P7	(0.18, 0.26, 0.37; 0.70)	P7.1	(0.10, 0.14, 0.17; 0.80)
		P7.2	(0.10, 0.14, 0.19; 0.80)
		P7.3	(0.10, 0.14, 0.18; 0.80)
		P7.4	(0.11, 0.15, 0.20; 0.80)
		P7.5	(0.16, 0.21, 0.28; 0.80)
		P7.6	(0.17, 0.23, 0.30; 0.80)

Bảng A1.4. Véc-tơ trọng số của 7 tiêu chí chính và 34 tiêu chí con

Tiêu chí chính	Véc-tơ trọng số	Tiêu chí con	Véc-tơ trọng số
P1	0.10	P1.1	0.11
		P1.2	0.16
		P1.3	0.37
		P1.4	0.36
P2	0.10	P2.1	0.11
		P2.2	0.18
		P2.3	0.18
		P2.4	0.24
		P2.5	0.29
P3	0.11	P3.1	0.14
		P3.2	0.10
		P3.3	0.12
		P3.4	0.12
		P3.5	0.16
		P3.6	0.18
		P3.7	0.17
P4	0.12	P4.1	0.15
		P4.2	0.26
		P4.3	0.59
P5	0.12	P5.1	0.13
		P5.2	0.14
		P5.3	0.17
		P5.4	0.25
		P5.5	0.31
P6	0.16	P6.1	0.15
		P6.2	0.20
		P6.3	0.39
		P6.4	0.27
P7	0.28	P7.1	0.12
		P7.2	0.13
		P7.3	0.12
		P7.4	0.14
		P7.5	0.23
		P7.6	0.27

Bảng A1.5. Kết quả đánh giá có trọng số cho mức độ chuyển đổi số của 16 SMEs bán lẻ tại Việt Nam.

Tiêu chí chính	Tiêu chí con	Doanh nghiệp							
		E1	E2	E3	E4	E5	E6	E7	E8
P1	P1.1	0.009	0.006	0.004	0.004	0.009	0.004	0.001	0.004
	P1.2	0.005	0.009	0.005	0.009	0.009	0.002	0.005	0.002
	P1.3	0.012	0.021	0.004	0.012	0.004	0.012	0.004	0.012
	P1.4	0.012	0.012	0.012	0.004	0.004	0.004	0.012	0.004
P2	P2.1	0.006	0.006	0.004	0.004	0.004	0.004	0.009	0.004
	P2.2	0.006	0.006	0.002	0.011	0.002	0.006	0.002	0.011
	P2.3	0.002	0.006	0.006	0.015	0.006	0.002	0.002	0.006
	P2.4	0.008	0.008	0.002	0.008	0.008	0.008	0.019	0.014
	P2.5	0.010	0.003	0.003	0.003	0.024	0.010	0.017	0.010
P3	P3.1	0.005	0.005	0.001	0.001	0.005	0.005	0.005	0.005
	P3.2	0.004	0.004	0.004	0.004	0.001	0.001	0.001	0.006
	P3.3	0.001	0.004	0.001	0.004	0.004	0.004	0.004	0.010
	P3.4	0.004	0.001	0.004	0.004	0.004	0.004	0.010	0.001
	P3.5	0.002	0.006	0.002	0.010	0.002	0.002	0.002	0.006
	P3.6	0.002	0.002	0.002	0.002	0.006	0.006	0.015	0.002
	P3.7	0.006	0.002	0.002	0.002	0.006	0.002	0.014	0.002
P4	P4.1	0.006	0.006	0.006	0.006	0.002	0.006	0.002	0.006
	P4.2	0.011	0.003	0.003	0.003	0.011	0.011	0.011	0.018
	P4.3	0.024	0.024	0.007	0.024	0.024	0.007	0.024	0.024
P5	P5.1	0.006	0.002	0.002	0.002	0.002	0.006	0.002	0.009
	P5.2	0.006	0.010	0.002	0.006	0.006	0.006	0.006	0.002
	P5.3	0.007	0.007	0.007	0.016	0.007	0.007	0.002	0.002
	P5.4	0.011	0.003	0.011	0.025	0.018	0.011	0.011	0.018
	P5.5	0.013	0.021	0.021	0.030	0.021	0.013	0.013	0.013
P6	P6.1	0.008	0.014	0.008	0.008	0.008	0.014	0.002	0.014
	P6.2	0.003	0.011	0.011	0.011	0.003	0.011	0.011	0.011
	P6.3	0.022	0.006	0.006	0.006	0.022	0.006	0.022	0.006
	P6.4	0.015	0.025	0.004	0.004	0.015	0.015	0.025	0.004
P7	P7.1	0.019	0.019	0.011	0.011	0.011	0.019	0.003	0.011
	P7.2	0.012	0.020	0.020	0.012	0.012	0.020	0.003	0.012
	P7.3	0.019	0.019	0.011	0.011	0.011	0.011	0.011	0.011
	P7.4	0.013	0.022	0.013	0.022	0.013	0.022	0.004	0.004
	P7.5	0.022	0.036	0.022	0.036	0.036	0.036	0.006	0.036

	P7.6	0.042	0.042	0.007	0.007	0.025	0.042	0.007	0.025
Tiêu chí chính	Tiêu chí con	Doanh nghiệp (tiếp theo)							
		E9	E10	E11	E12	E13	E14	E15	E16
P1	P1.1	0.001	0.001	0.004	0.004	0.001	0.004	0.004	0.004
	P1.2	0.005	0.005	0.009	0.013	0.005	0.005	0.009	0.002
	P1.3	0.021	0.012	0.012	0.021	0.004	0.004	0.012	0.012
	P1.4	0.028	0.004	0.020	0.012	0.004	0.012	0.020	0.004
P2	P2.1	0.006	0.004	0.009	0.006	0.001	0.004	0.006	0.001
	P2.2	0.006	0.002	0.002	0.002	0.011	0.006	0.002	0.006
	P2.3	0.002	0.006	0.002	0.002	0.011	0.006	0.002	0.011
	P2.4	0.002	0.008	0.008	0.002	0.008	0.008	0.008	0.008
	P2.5	0.010	0.010	0.010	0.010	0.024	0.017	0.003	0.010
P3	P3.1	0.005	0.008	0.001	0.008	0.005	0.005	0.005	0.008
	P3.2	0.006	0.006	0.001	0.001	0.001	0.006	0.004	0.006
	P3.3	0.007	0.004	0.004	0.001	0.004	0.007	0.004	0.007
	P3.4	0.004	0.004	0.007	0.004	0.004	0.010	0.007	0.004
	P3.5	0.002	0.006	0.006	0.006	0.010	0.010	0.006	0.006
	P3.6	0.002	0.011	0.006	0.006	0.006	0.006	0.002	0.006
	P3.7	0.006	0.010	0.006	0.002	0.002	0.006	0.006	0.006
P4	P4.1	0.006	0.002	0.011	0.011	0.002	0.011	0.002	0.011
	P4.2	0.011	0.003	0.011	0.011	0.011	0.018	0.003	0.011
	P4.3	0.007	0.024	0.007	0.007	0.024	0.024	0.024	0.007
P5	P5.1	0.006	0.009	0.002	0.002	0.006	0.002	0.009	0.006
	P5.2	0.010	0.002	0.006	0.006	0.002	0.002	0.006	0.002
	P5.3	0.007	0.007	0.002	0.007	0.007	0.007	0.007	0.007
	P5.4	0.003	0.018	0.003	0.003	0.003	0.003	0.003	0.011
	P5.5	0.013	0.013	0.013	0.004	0.013	0.004	0.021	0.013
P6	P6.1	0.002	0.002	0.014	0.014	0.008	0.008	0.008	0.002
	P6.2	0.018	0.011	0.011	0.011	0.003	0.011	0.011	0.011
	P6.3	0.006	0.022	0.036	0.022	0.006	0.022	0.006	0.022
	P6.4	0.015	0.025	0.025	0.025	0.015	0.015	0.015	0.004
P7	P7.1	0.003	0.003	0.003	0.003	0.003	0.003	0.003	0.003
	P7.2	0.020	0.012	0.003	0.012	0.012	0.003	0.012	0.003
	P7.3	0.011	0.003	0.003	0.003	0.011	0.003	0.011	0.011
	P7.4	0.013	0.013	0.004	0.013	0.004	0.013	0.004	0.004
	P7.5	0.022	0.022	0.006	0.022	0.022	0.022	0.022	0.022
	P7.6	0.007	0.007	0.007	0.007	0.025	0.007	0.007	0.025

Bảng A1.6. Điểm số và thứ hạng cuối cùng sau khi phân tích độ nhạy ở kịch bản 1

Doanh nghiệp	S0 (Ban đầu)		S1		S2		S3		S4		S5		S6	
	Tổng điểm	Thứ hạng	Tổng điểm	Thứ hạng	Tổng điểm	Thứ hạng	Tổng điểm	Thứ hạng	Tổng điểm	Thứ hạng	Tổng điểm	Thứ hạng	Tổng điểm	Thứ hạng
E1	0.352	2	0.350	2	0.350	2	0.356	2	0.352	2	0.352	2	0.351	2
E2	0.393	1	0.389	1	0.390	1	0.400	1	0.391	1	0.392	1	0.391	1
E3	0.232	16	0.233	16	0.231	16	0.235	16	0.230	16	0.232	16	0.231	16
E4	0.338	5	0.344	4	0.336	4	0.338	5	0.337	4	0.340	4	0.337	4
E5	0.345	3	0.346	3	0.343	3	0.347	4	0.344	3	0.345	3	0.344	3
E6	0.338	4	0.335	5	0.334	5	0.347	3	0.336	5	0.337	5	0.336	5
E7	0.288	10	0.290	10	0.290	10	0.281	10	0.289	10	0.288	10	0.290	10
E8	0.325	6	0.325	6	0.323	6	0.326	6	0.325	6	0.326	6	0.324	6
E9	0.296	8	0.296	8	0.295	9	0.293	8	0.295	9	0.296	8	0.296	9
E10	0.301	7	0.305	7	0.303	7	0.298	7	0.301	7	0.302	7	0.302	7
E11	0.277	13	0.277	14	0.282	12	0.271	15	0.277	13	0.276	15	0.279	12
E12	0.283	11	0.281	11	0.285	11	0.281	11	0.283	11	0.282	11	0.284	11
E13	0.277	12	0.276	15	0.275	15	0.276	12	0.278	12	0.277	13	0.277	14
E14	0.295	9	0.294	9	0.296	8	0.290	9	0.297	8	0.295	9	0.296	8
E15	0.277	14	0.280	12	0.278	13	0.273	14	0.276	14	0.278	12	0.277	13
E16	0.276	15	0.278	13	0.275	14	0.274	13	0.276	15	0.277	14	0.276	15

A2. Một số kết quả trong Chương 4

Bảng A2.1. Bộ tiêu chí đánh giá mức độ chuyển đổi số của SMEs bán lẻ tại Việt Nam

Tiêu chí chính	Tiêu chí con
Chiến lược và tổ chức (P1)	(P1.1) Chiến lược chuyển đổi số rõ ràng, với các mục tiêu cụ thể và các chỉ số đo lường hiệu quả then chốt (KPIs) có thể định lượng
	(P1.2) Khả năng linh hoạt và thích nghi của tổ chức đối với những thay đổi công nghệ.
	(P1.3) Phân bổ nguồn lực phù hợp cho các sáng kiến số (đầu tư vào công nghệ, nguồn nhân lực, đào tạo nhân viên).

Công nghệ và năng lực số (P2)	(P2.1) Tích hợp và sử dụng các nền tảng số (thương mại điện tử, dữ liệu lớn, IoT, blockchain, v.v.)
	(P2.2) Năng lực trí tuệ nhân tạo (như học máy, hệ thống phân tích dữ liệu về khách hàng)
	(P2.3) Ứng dụng điện toán đám mây (cho lưu trữ, xử lý dữ liệu, hoặc dịch vụ phần mềm)
	(P2.4) Năng lực ra quyết định dựa trên dữ liệu (phân tích dữ liệu lớn, phân tích dự báo)
Trải nghiệm khách hàng (P3)	(P3.1) Khả năng đa kênh (Omnichannel) (tích hợp liền mạch các tương tác khách hàng trực tuyến và ngoại tuyến).
	(P3.2) Cá nhân hóa và marketing số hướng mục tiêu (đưa ra các sản phẩm phù hợp, trải nghiệm cá nhân hóa cho khách hàng).
	(P3.3) Tăng cường sự tương tác của khách hàng (sử dụng công nghệ thực tế tăng cường (AR), thực tế ảo (VR), cá nhân hóa bằng AI).
	(P3.4) Sự thuận tiện và nhanh chóng trong giao dịch số (giao hàng nhanh, hình thức nhận hàng tại cửa hàng, thanh toán qua di động)
Nhân lực và văn hóa (P4)	(P4.1) Đào tạo và phát triển kỹ năng số cho nhân viên (nâng cao trình độ kỹ thuật số, học tập liên tục).
	(P4.2) Văn hóa đổi mới và chuyển đổi số trong tổ chức (khuyến khích chấp nhận rủi ro, đổi mới số, hợp tác nội bộ).
	(P4.3) Tác động của chuyển đổi số lên cơ cấu việc làm (thay đổi trong cấu trúc lực lượng lao động, gia tăng tỷ lệ lao động kỹ năng cao và các cấu trúc phù hợp với kỹ thuật số).

Bảng A2.2. Ma trận so sánh mờ cho ba tiêu chí chính tại thời điểm τ_1

Tiêu chí	P1	P2	P3
P1	(0.430, 0.540, 0.654)	(0.783, 0.897, 1)	(0.783, 0.897, 1)
P2	(1, 1.114, 1.277)	(0.430, 0.540, 0.654)	(0.783, 0.897, 1)
P3	(1, 1.114, 1.277)	(1, 1.114, 1.277)	(0.430, 0.540, 0.654)
Chuyên gia 1: CR = 0.032		Chuyên gia 2: CR = 0.002	Chuyên gia 3: CR = 0.009

Bảng A2.3. Ma trận so sánh mờ cho tiêu chí P1 tại thời điểm τ_1

P1	P1.1	P1.2	P1.3
P1.1	(0.430, 0.540, 0.654)	(0.430, 0.540, 0.654)	(0.540, 0.654, 0.783)
P1.2	(1.528, 1.852, 2.324)	(0.430, 0.540, 0.654)	(0.540, 0.654, 0.783)
P1.3	(1.277, 1.528, 1.852)	(1.277, 1.528, 1.852)	(0.430, 0.540, 0.654)
Chuyên gia 1: CR = 0.032		Chuyên gia 2: CR = 0.021	Chuyên gia 3: CR = 0.005

Bảng A2.4. Ma trận so sánh mờ cho tiêu chí P2 tại thời điểm τ_1

P2	P2.1	P2.2	P2.3	P2.4
P2.1	(0.430, 0.540, 0.654)	(0.654, 0.783, 0.897)	(0.783, 0.897, 1)	(0.897, 1, 1)
P2.2	(1.114, 1.277, 1.528)	(0.430, 0.540, 0.654)	(0.897, 1.000, 1.)	(0.654, 0.783, 0.897)
P2.3	(1, 1.114, 1.277)	(1, 1, 1.114)	(0.430, 0.540, 0.654)	(0.654, 0.783, 0.897)
P2.4	(1.114, 1.114, 1.114)	(1.528, 1.528, 1.528)	(1.528, 1.528, 1.528)	(0.430, 0.540, 0.654)
Chuyên gia 1: CR = 0.024		Chuyên gia 2: CR = 0.038	Chuyên gia 3: CR = 0.012	

Bảng A2.5. Ma trận so sánh mờ cho tiêu chí P3 tại thời điểm τ_1

P3	P3.1	P3.2	P3.3	P3.4
P3.1	(0.430, 0.540, 0.654)	(0.783, 0.897, 1.000)	(0.540, 0.654, 0.783)	(0.540, 0.654, 0.783)
P3.2	(1.000, 1.114, 1.277)	(0.430, 0.540, 0.654)	(0.783, 0.897, 1.000)	(0.654, 0.783, 0.897)
P3.3	(1.277, 1.528, 1.852)	(1.000, 1.114, 1.277)	(0.430, 0.540, 0.654)	(0.783, 0.897, 1.000)
P3.4	(1.277, 1.528, 1.852)	(1.114, 1.277, 1.528)	(1, 1.114, 1.277)	(0.430, 0.540, 0.654)
Chuyên gia 1: CR = 0.013		Chuyên gia 2: CR = 0.020	Chuyên gia 3: CR = 0.020	

Bảng A2.6. Trọng số mờ của các tiêu chí chính và tiêu chí con tại thời điểm τ_1

Tiêu chí chính	Trọng số mờ của các tiêu chí chính	Tiêu chí con	Trọng số mờ của các tiêu chí con
P1	(0.226, 0.310, 0.415)	P1.1	(0.151, 0.228, 0.227)
		P1.2	(0.231, 0.344, 0.346)
		P1.3	(0.290, 0.428, 0.427)
P2	(0.246, 0.333, 0.450)	P2.1	(0.169, 0.228, 0.222)
		P2.2	(0.185, 0.249, 0.247)
		P2.3	(0.185, 0.241, 0.242)
		P2.4	(0.217, 0.282, 0.289)
P3	(0.267, 0.358, 0.488)	P3.1	(0.137, 0.194, 0.271)
		P3.2	(0.168, 0.232, 0.317)
		P3.3	(0.198, 0.274, 0.379)
		P3.4	(0.217, 0.300, 0.422)

Bảng A2.7. Tổng hợp đánh giá của các doanh nghiệp tại thời điểm τ_1

Tiêu chí con	Doanh nghiệp	Kết nhập xếp hạng
P1.1	E1	(0.566, 0.781, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.566, 0.781, 1.000)
	E5	(0.336, 0.566, 0.781)
P1.2	E1	(0.781, 1.000, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.566, 0.781, 1.000)
	E4	(0.336, 0.566, 0.781)
	E5	(0.566, 0.781, 1.000)
P1.3	E1	(0.566, 0.781, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.566, 0.781, 1.000)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
P2.1	E1	(0.566, 0.781, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
P2.2	E1	(0.781, 1.000, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
P2.3	E1	(0.566, 0.781, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.566, 0.781, 1)
P2.4	E1	(0.336, 0.566, 0.781)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)

Tiêu chí con	Doanh nghiệp	Kết nhập xếp hạng
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
P3.1	E1	(0.336, 0.566, 0.781)
	E2	(0.781, 1, 1)
	E3	(0.1, 0.1, 0.336)
	E4	(0.1, 0.336, 0.566)
	E5	(0.1, 0.336, 0.566)
P3.2	E1	(0.781, 1.000, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
P3.3	E1	(0.781, 1.000, 1.000)
	E2	(0.566, 0.781, 1.000)
	E3	(0.566, 0.781, 1.000)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
P3.4	E1	(0.781, 1.000, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.100, 0.100, 0.336)
	E5	(0.100, 0.100, 0.336)

Bảng A2.8. Trọng số mờ của các tiêu chí chính và tiêu chí con tại thời điểm τ_2

Tiêu chí chính	Trọng số mờ của các tiêu chí chính	Tiêu chí con	Trọng số mờ của các tiêu chí con
P1	(0.187, 0.243, 0.296)	P1.1	(0.211, 0.295, 0.402)
		P1.2	(0.252, 0.348, 0.481)
		P1.3	(0.264, 0.357, 0.491)
P2	(0.186, 0.237, 0.304)	P2.1	(0.186, 0.243, 0.297)
		P2.2	(0.183, 0.241, 0.308)
		P2.3	(0.197, 0.243, 0.314)
		P2.4	(0.215, 0.273, 0.361)

Tiêu chí chính	Trọng số mờ của các tiêu chí chính	Tiêu chí con	Trọng số mờ của các tiêu chí con
P3	(0.205, 0.257, 0.323)	P3.1	(0.145, 0.200, 0.266)
		P3.2	(0.186, 0.241, 0.309)
		P3.3	(0.203, 0.268, 0.360)
		P3.4	(0.216, 0.291, 0.396)
P4	(0.210, 0.264, 0.344)	P4.1	(0.211, 0.295, 0.402)
		P4.2	(0.252, 0.348, 0.481)
		P4.3	(0.264, 0.357, 0.491)

Bảng A2.9. Tổng hợp đánh giá của các doanh nghiệp tại thời điểm τ_2

Tiêu chí con	Doanh nghiệp	Kết nhập xếp hạng
P1.1	E1	(0.566, 0.781, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.566, 0.781, 1.000)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
P1.2	E1	(0.566, 0.781, 1.000)
	E2	(0.566, 0.781, 1.000)
	E3	(0.781, 1.000, 1.000)
	E4	(0.336, 0.566, 0.781)
	E5	(0.781, 1.000, 1.000)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.566, 0.781, 1.000)
P1.3	E1	(0.566, 0.781, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.566, 0.781, 1.000)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)

Tiêu chí con	Doanh nghiệp	Kết nhập xếp hạng
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
P2.1	E1	(0.566, 0.781, 1.000)
	E2	(0.566, 0.781, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
P2.2	E1	(0.781, 1.000, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
P2.3	E1	(0.566, 0.781, 1.000)
	E2	(0.566, 0.781, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.566, 0.781, 1.000)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
P2.4	E1	(0.336, 0.566, 0.781)
	E2	(0.566, 0.781, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)

Tiêu chí con	Doanh nghiệp	Kết nhập xếp hạng
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
P3.1	E1	(0.336, 0.566, 0.781)
	E2	(0.566, 0.781, 1.000)
	E3	(0.100, 0.100, 0.336)
	E4	(0.100, 0.100, 0.336)
	E5	(0.100, 0.100, 0.336)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
P3.2	E1	(0.781, 1.000, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
P3.3	E1	(0.781, 1.000, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.781, 1.000, 1.000)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
P3.4	E1	(0.781, 1.000, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.100, 0.100, 0.336)
	E5	(0.100, 0.100, 0.336)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
	E1	(0.336, 0.566, 0.781)

Tiêu chí con	Doanh nghiệp	Kết nhập xếp hạng
P4.1	E2	(0.336, 0.566, 0.781)
	E3	(0.566, 0.781, 1.000)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
P4.2	E1	(0.336, 0.566, 0.781)
	E2	(0.336, 0.566, 0.781)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.566, 0.781, 1.000)
P4.3	E1	(0.336, 0.566, 0.781)
	E2	(0.336, 0.566, 0.781)
	E3	(0.566, 0.781, 1.000)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.781, 1.000, 1.000)

Bảng A2.10. Trọng số mờ của các tiêu chí chính và tiêu chí con tại thời điểm τ_3

Tiêu chí chính	Trọng số mờ của các tiêu chí	Tiêu chí con	Trọng số mờ của các tiêu chí con
P1	(0.189, 0.243, 0.294)	P1.1	(0.245, 0.321, 0.403)
		P1.2	(0.265, 0.345, 0.437)
		P1.3	(0.275, 0.333, 0.433)
P2	(0.194, 0.243, 0.302)	P2.1	(0.193, 0.249, 0.295)
		P2.2	(0.175, 0.224, 0.285)
		P2.3	(0.197, 0.242, 0.312)
		P2.4	(0.223, 0.285, 0.376)

Tiêu chí chính	Trọng số mờ của các tiêu chí	Tiêu chí con	Trọng số mờ của các tiêu chí con
P3	(0.207, 0.264, 0.332)	P3.1	(0.160, 0.218, 0.284)
		P3.2	(0.180, 0.235, 0.311)
		P3.3	(0.192, 0.255, 0.342)
		P3.4	(0.217, 0.292, 0.398)
P4	(0.212, 0.250, 0.319)	P4.1	(0.245, 0.321, 0.403)
		P4.2	(0.265, 0.345, 0.437)
		P4.3	(0.275, 0.333, 0.433)

Bảng A2.11. Tổng hợp đánh giá của các doanh nghiệp tại thời điểm τ_3

Tiêu chí con	Doanh nghiệp	Kết nhập xếp hạng
P1.1	E1	(0.566, 0.781, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.100, 0.336, 0.566)
	E4	(0.566, 0.781, 1.000)
	E5	(0.336, 0.566, 0.781)
	E6	(0.566, 0.781, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
	E9	(0.781, 1.000, 1.000)
	E10	(0.336, 0.566, 0.781)
P1.2	E1	(0.566, 0.781, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.781, 1.000, 1.000)
	E4	(0.336, 0.566, 0.781)
	E5	(0.781, 1.000, 1.000)
	E6	(0.566, 0.781, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.566, 0.781, 1.000)
	E9	(0.566, 0.781, 1.000)
	E10	(0.566, 0.781, 1.000)
P1.3	E1	(0.566, 0.781, 1.000)

Tiêu chí con	Doanh nghiệp	Kết nhập xếp hạng
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)
	E7	(0.781, 1.000, 1.000)
	E8	(0.336, 0.566, 0.781)
	E9	(0.566, 0.781, 1.000)
	E10	(0.781, 1.000, 1.000)
P2.1	E1	(0.100, 0.336, 0.566)
	E2	(0.781, 1.000, 1.000)
	E3	(0.100, 0.336, 0.566)
	E4	(0.566, 0.781, 1.000)
	E5	(0.566, 0.781, 1.000)
	E6	(0.566, 0.781, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.566, 0.781, 1.000)
	E9	(0.781, 1.000, 1.000)
	E10	(0.781, 1.000, 1.000)
P2.2	E1	(0.781, 1.000, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.100, 0.336, 0.566)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)
	E7	(0.781, 1.000, 1.000)
	E8	(0.336, 0.566, 0.781)
	E9	(0.781, 1.000, 1.000)
	E10	(0.566, 0.781, 1.000)
P2.3	E1	(0.566, 0.781, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.100, 0.336, 0.566)
	E4	(0.336, 0.566, 0.781)

Tiêu chí con	Doanh nghiệp	Kết nhập xếp hạng
	E5	(0.781, 1.000, 1.000)
	E6	(0.781, 1.000, 1.000)
	E7	(0.781, 1.000, 1.000)
	E8	(0.336, 0.566, 0.781)
	E9	(0.781, 1.000, 1.000)
	E10	(0.781, 1.000, 1.000)
P2.4	E1	(0.100, 0.336, 0.566)
	E2	(0.781, 1.000, 1.000)
	E3	(0.100, 0.336, 0.566)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)
	E7	(0.781, 1.000, 1.000)
	E8	(0.566, 0.781, 1.000)
	E9	(0.781, 1.000, 1.000)
	E10	(0.566, 0.781, 1.000)
P3.1	E1	(0.336, 0.566, 0.781)
	E2	(0.781, 1.000, 1.000)
	E3	(0.100, 0.100, 0.336)
	E4	(0.100, 0.100, 0.336)
	E5	(0.100, 0.100, 0.336)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
	E9	(0.781, 1.000, 1.000)
	E10	(0.566, 0.781, 1.000)
P3.2	E1	(0.781, 1.000, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.100, 0.336, 0.566)
	E4	(0.336, 0.566, 0.781)
	E5	(0.566, 0.781, 1.000)
	E6	(0.781, 1.000, 1.000)
	E7	(0.781, 1.000, 1.000)

Tiêu chí con	Doanh nghiệp	Kết nhập xếp hạng
	E8	(0.566, 0.781, 1.000)
	E9	(0.781, 1.000, 1.000)
	E10	(0.781, 1.000, 1.000)
P3.3	E1	(0.781, 1.000, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.566, 0.781, 1.000)
	E4	(0.336, 0.566, 0.781)
	E5	(0.566, 0.781, 1.000)
	E6	(0.781, 1.000, 1.000)
	E7	(0.781, 1.000, 1.000)
	E8	(0.566, 0.781, 1.000)
	E9	(0.781, 1.000, 1.000)
	E10	(0.781, 1.000, 1.000)
P3.4	E1	(0.566, 0.781, 1.000)
	E2	(0.781, 1.000, 1.000)
	E3	(0.100, 0.336, 0.566)
	E4	(0.100, 0.100, 0.336)
	E5	(0.100, 0.100, 0.336)
	E6	(0.781, 1.000, 1.000)
	E7	(0.781, 1.000, 1.000)
	E8	(0.566, 0.781, 1.000)
	E9	(0.781, 1.000, 1.000)
	E10	(0.781, 1.000, 1.000)
P4.1	E1	(0.100, 0.336, 0.566)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.566, 0.781, 1.000)
	E6	(0.781, 1.000, 1.000)
	E7	(0.781, 1.000, 1.000)
	E8	(0.566, 0.781, 1.000)
	E9	(0.781, 1.000, 1.000)
	E10	(0.781, 1.000, 1.000)

Tiêu chí con	Doanh nghiệp	Kết nhập xếp hạng
P4.2	E1	(0.100, 0.336, 0.566)
	E2	(0.566, 0.781, 1.000)
	E3	(0.100, 0.100, 0.336)
	E4	(0.336, 0.566, 0.781)
	E5	(0.566, 0.781, 1.000)
	E6	(0.781, 1.000, 1.000)
	E7	(0.566, 0.781, 1.000)
	E8	(0.336, 0.566, 0.781)
	E9	(0.781, 1.000, 1.000)
	E10	(0.781, 1.000, 1.000)
P4.3	E1	(0.100, 0.336, 0.566)
	E2	(0.781, 1.000, 1.000)
	E3	(0.336, 0.566, 0.781)
	E4	(0.336, 0.566, 0.781)
	E5	(0.336, 0.566, 0.781)
	E6	(0.781, 1.000, 1.000)
	E7	(0.781, 1.000, 1.000)
	E8	(0.781, 1.000, 1.000)
	E9	(0.781, 1.000, 1.000)
	E10	(0.781, 1.000, 1.000)