

Title	不完全な言語的情報を基に多属性による意思決定モデルに関する研究
Author(s)	郭, 文涛
Citation	
Issue Date	2015-03
Type	Thesis or Dissertation
Text version	ETD
URL	http://hdl.handle.net/10119/12754
Rights	
Description	Supervisor:Huynh Nam Van, 知識科学研究科, 博士

Doctoral Dissertation

A Study on Evaluation Models for Multiple Attribute
Decision Making with Incomplete Linguistic Information

Wentao Guo

Supervisor: Associate Professor Van Nam Huynh

School of Knowledge Science

Japan Advanced Institute of Science and Technology

March 2015

Abstract

In practice, most of the multiple attribute decision making (MADM) problems involve both kinds of qualitative and quantitative attributes, which may be represented by a hierarchy. While quantitative attributes can be measured by means of numeric scales in the form of numbers, intervals or fuzzy numbers, qualitative attributes which are often associated with imprecise, vagueness and uncertain information perhaps can only be assessed by linguistic information. In such situations, how to represent and aggregate linguistic information essentially plays an important role in decision analysis. In the literature, one of reasonable ways is the use of “fuzzy linguistic approach which provides tools to model and represent qualitative attributes by means of linguistic values of linguistic variables” (Zadeh, 1975). The use of linguistic information implies the necessity of operating with the mechanism for “computing with words (CW)” (Zadeh, 1996) so as to fusion linguistic information and then provide an evaluation for decision making.

In this research, we first briefly recall some key concepts of CW. Then, through a further study on CW and fuzzy linguistic approach, we analyze the relationship between MADM with linguistic information and CW, and the mechanism that how fuzzy linguistic approach is used to deal with linguistic information in the decision making process. Further, according to three categories of linguistic computational models based on fuzzy linguistic approach in the literature, we review the main features of several classic linguistic computational models in detail, including “linguistic computational model based on membership functions”, “linguistic computational model based on ordinal scales”, “linguistic computational model based on 2-tuple representation” and “linguistic computational model based on proportional 2-tuple representation”. Meanwhile, the limitations and restrictions of these previous models have been found during the review process, such as loss of information during the evaluation process, with too much requirements when applied to MADM problems, without considering uncertain subjective judgments represented by linguistic distributions over the linguistic term set, without taking into account incomplete linguistic information and so on.

Inspired by providing more efficient measures to represent and aggregate linguistic information, three evaluation models, i.e., proportional 3-tuple fuzzy linguistic representation model, proportional fuzzy linguistic distribution model, interval fuzzy linguistic distribution model, are developed in this research aiming at overcoming the main limitations and restrictions of previous models, and meanwhile, providing some new ways to deal with more general cases of linguistic assessments. Some related concepts, such as preference-preserving proportional 3-tuple transformation, which is used for the transformation and unification of linguistic assessments represented by proportional 3-tuple between two different linguistic term sets, expected utility in proportional or interval fuzzy linguistic distribution, which is employed for obtaining an ranking order among different alternatives provided to decision makers as a reference for their final decisions are proposed in this research. Further, some corresponding aggregation operators are developed for the three evaluation models respectively according to their own representation forms of linguistic information. Besides, three practical application examples taken from the literature as well as a simple illustration example are used respectively in order to compare the results with previous models, and also for the purpose of illuminating the features and capabilities of the proposed models.

After illustration by examples, it is shown that the proposed evaluation models in this research not only overcome

the limitations and restrictions of previous models, but are also inherent with some special features, such as no loss of information during the evaluation process, ease operation in the complicated linguistic context, flexible operation space for evaluators under uncertainty, taking the ignoring information into account and so forth. These features of the three evaluation models can help decision makers to easily deal with MADM problems with incomplete linguistic information, largely improve the precision, reasonability and reliability of final results, and finally, provide a more comprehensive guidance for decision makers.

Finally, four interesting aspects for future work are explored, which can be as the directions for continuing this research in order to extend the applicability of these three evaluation models proposed in this research. Meanwhile, the contributions of this research to Knowledge Science are summarized.

Keywords: Computing with words, decision making, incomplete assessments, linguistic modeling, multiple attribute.

Acknowledgments

First of all, I would like to express my great thanks to my supervisor Associate Professor Van Nam Huynh of Japan Advanced Institute of Science and Technology (JAIST) for his constant encouragement and kind guidance during this work. In the past three years of JAIST, Prof. Huynh not only gave me so many good instructions on my research, but also helped me a lot on my life. The knowledge and experience that Prof. Huynh conveyed to me will accompany me for my whole life.

I also want to devote my sincere thanks and appreciation to Prof. Nakamori (JAIST), Prof. Kosaka (JAIST), Prof. Yoshida (JAIST), and Prof. Tetsuya Murai (Hokkaido University), who gave me many important suggestions for improving my dissertation.

I also would like to give my thanks to my colleagues and friends in Nakamori Lab and Huynh lab during the last three years. They helped me a lot on my research and life. I am very glad to study and play with you. They are Jin Jingzhong, Jin Guangzhong, Ju Hongli, Guo Wei, Sun Jing and Zhang Wei.

The last but not least, I would like to express my sincere gratitude to my parents and other relatives for their endless loves and everlasting supports. Without you, I cannot finish my research for the PH.D. smoothly.

Thanks very much for all of you. Wish you happy and healthy.

Contents

Abstract	i
Acknowledgments	iii
1 Introduction	1
1.1 Research Background and Research Purpose	1
1.2 Research Motivation	2
1.3 Research Objectives	3
1.4 Related Knowledge	3
1.4.1 Decision making	3
1.4.2 Multiple attribute decision making problems	4
1.4.3 Computing with words	5
1.4.4 Fuzzy linguistic approach	6
1.4.5 Linguistic decision making resolution scheme	9
1.5 Organization of the Dissertation	10
2 A Survey on Decision Making with Fuzzy Linguistic Models	12
2.1 Linguistic Computational Model Based on Membership Functions	12
2.2 Linguistic Symbolic Computational Models Based on Ordinal Scales	14

2.3	Linguistic Symbolic Computational Models Based on 2-Tuple Representation	15
2.3.1	Representation model	16
2.3.2	Computational model	17
2.3.3	The use of the 2-tuple linguistic representation model	18
2.4	Linguistic Symbolic Computational Models Based on Proportional 2-Tuple Representation	19
2.4.1	Representation model	20
2.4.2	Computational model	20
2.5	Conclusion	23
3	A Proportional 3-Tuple Fuzzy Linguistic Representation Model	25
3.1	Proportional 3-Tuple	26
3.2	Canonical Characteristic Values	27
3.3	Computation Operator of Proportional 3-Tuple	28
3.4	Preference-Preserving Proportional 3-Tuple Transformation	30
3.5	The Procedure of Proportional 3-Tuple Fuzzy Linguistic Representation Model . . .	32
3.6	An Illustration Example	35
3.6.1	The description of new product project screening problem	35
3.6.2	Selecting evaluation criteria	36
3.6.3	Selecting linguistic term sets and associated semantics	36
3.6.4	Assessing merit/risk ratings and weights of criteria	38
3.6.5	The unification of original linguistic assessments represented by proportional 3-tuples	41
3.6.6	Computing the evaluation result via proportional 3-tuple fuzzy linguistic representation model	43

3.6.7	Comparative study	44
3.6.8	New product project screening problem with revised linguistic assessments .	44
3.6.9	The unification of the revised linguistic assessments represented by proportional 3-tuples	46
3.6.10	The evaluation result of revised linguistic assessments	47
3.7	Conclusion	48
4	A Proportional Fuzzy Linguistic Distribution Model	49
4.1	Proportional Fuzzy Linguistic Distribution	49
4.2	The Comparison of Proportional Fuzzy Linguistic Distributions	52
4.3	Computation Operator of Proportional Fuzzy Linguistic Distribution	54
4.4	Expected Utility in Proportional Fuzzy Linguistic Distribution	55
4.5	Proportional Fuzzy Linguistic Distribution Aggregation Operator	57
4.5.1	Arithmetic mean	58
4.5.2	Weighted average operator	59
4.5.3	Linguistic weighted average operator	60
4.6	An Illustration Example	63
4.6.1	Motorcycle evaluation problem	63
4.6.2	Aggregating assessments by proportional fuzzy linguistic distribution model .	65
4.6.3	Computing the expected utilities of four types of motorcycles	68
4.6.4	Motorcycle assessment problem with linguistic weights	71
4.7	Conclusion	75
5	An Interval Fuzzy Linguistic Distribution Model	77
5.1	Linguistic Assessments with Intervals	78

5.1.1	Basic interval arithmetic	78
5.1.2	Interval fuzzy linguistic distribution	80
5.2	Comparison of Interval Fuzzy Linguistic Distributions	83
5.3	Expected Utility in Interval Fuzzy Linguistic Distribution	84
5.4	Interval Fuzzy Linguistic Distribution Aggregation Operators	87
5.4.1	Arithmetic mean	87
5.4.2	Weighted average operator	88
5.4.3	Interval weighted average operator	90
5.4.4	Interval ordered weighted average operator	92
5.5	Illustration Examples	93
5.5.1	The description of the first example	93
5.5.2	Aggregating assessments of the first example via interval fuzzy linguistic distribution model	94
5.5.3	Computing the expected utilities of the first example	95
5.5.4	The description of a cargo ship selection problem	95
5.5.5	Aggregating assessments of ship selection problem via interval fuzzy linguistic distribution model	98
5.5.6	Computing the expected utilities of six cargo ships	102
5.5.7	Ranking the expected utilities using the minimax regret approach	103
5.6	Conclusion	105

6 Conclusion 107

6.1	The Main Contributions	107
6.2	Discussion and Future work	110

Bibliography	111
Publications	119

List of Figures

1.1	Computing with words scheme [59]	6
1.2	A set of seven terms with its semantics	8
1.3	The scheme of a linguistic decision making problem	10
2.1	Retranslation problem [23]	13
2.2	Example of the third linguistic symbolic computational model [23] [79]	15
3.1	The representation of the information h of a proportional 3-tuple	29
3.2	Linguistic merit rating values and associated fuzzy number semantic	39
3.3	Linguistic risk rating values and their associated fuzzy number semantics	39
3.4	Linguistic weights (success levels) and associated fuzzy number semantics	40
3.5	Transition linguistic term set and associated fuzzy number semantics	42
4.1	The relationship between a complete and an incomplete proportional fuzzy linguistic distribution	53
4.2	The relationship between two incomplete proportional fuzzy linguistic distributions	53
4.3	Two level hierarchy	56
4.4	Evaluation hierarchy for motorcycle performance assessment [87]	64
4.5	The distributed assessments on the four types of motorcycles	69
4.6	The expected utilities of four types of motorcycles	70

4.7	Linguistic weights and associated fuzzy number semantics	71
4.8	The distributed assessments on the four types of motorcycles by using linguistic weights	74
4.9	The expected utilities of four types of motorcycles by using linguistic weights . . .	75
5.1	Two level hierarchy	86
5.2	The expected utilities of two alternatives	96
5.3	The expected utilities of six cargo ships	103

List of Tables

3.1	The evaluation criteria of new product project	37
3.2	Original linguistic assessments of merit/risk ratings of criteria represented by proportional 3-tuples	40
3.3	Original linguistic assessments of weights of criteria represented by proportional 3-tuples	41
3.4	Original linguistic assessments of risk ratings of criteria represented by proportional 3-tuples in transition linguistic term set	42
3.5	Original linguistic preferences of criteria represented by proportional 3-tuples	43
3.6	Revised linguistic assessments of merit/risk ratings of criteria represented by proportional 3-tuples	45
3.7	Revised linguistic assessments of weights of criteria represented by proportional 3-tuples	46
3.8	Revised linguistic assessments of risk ratings of criteria represented by proportional 3-tuples in transition linguistic term set	46
3.9	Revised linguistic preferences of criteria represented by proportional 3-tuples	47
4.1	Generalized decision matrix for motorcycle assessment [87]	66
4.2	Generalized decision matrix for motorcycle assessment represented by proportional fuzzy linguistic distributions	67
4.3	The overall performances represented by proportional fuzzy linguistic distributions	68

4.4	Distributed assessments on four types of motorcycles	68
4.5	The expected utilities of four types of motorcycles	70
4.6	Linguistic weights represented by proportional fuzzy linguistic distributions	72
4.7	The overall performances represented by proportional fuzzy linguistic distributions by using linguistic weights	73
4.8	Distributed assessments on four types of motorcycles by using linguistic weights . .	73
4.9	The expected utilities of four types of motorcycles by using linguistic weights . . .	74
5.1	Interval fuzzy linguistic distribution assessments for two alternatives	94
5.2	The aggregation results of two alternatives	94
5.3	The distributed assessments on two alternatives	94
5.4	The expected utilities of two alternatives	95
5.5	Original assessment data for six cargo ships [74]	97
5.6	Distribution assessment matrix for the six cargo ships [74]	99
5.7	Distribution assessment matrix for the six cargo ships represented by interval fuzzy linguistic distributions	100
5.8	The overall belief degrees of six cargo ships represented by interval fuzzy linguistic distributions	101
5.9	Distributed assessments on six cargo ships	101
5.10	The expected utilities of six cargo ships	102

Chapter 1

Introduction

1.1 Research Background and Research Purpose

Computing with words (CW) was proposed by Zadeh [93] aiming at capturing the concept of automated reasoning involving linguistic terms, of which the idea was actually rooted from his previous work on linguistic variables, fuzzy constraints and fuzzy if-then rules [90], [91], [92]. Since the conception of CW, it has already attracted great attentions from the fuzzy-set research community.

So far, numerous models have been developed for reasoning and CW in the literature. Especially, CW approaches have also been applied to a wide range of decision making problems involving vague and imprecise information expressed linguistically. In decision making applications, the main problem for CW is how to represent and aggregate linguistic information for evaluation of alternatives. One of reasonable ways is the use of fuzzy linguistic approach. However, most previously developed linguistic computational models based on fuzzy linguistic approach have various limitations and restrictions, such as loss of information, without involving ignoring information and so on. These limitations and restrictions seriously affect the precision, reasonability and reliability of final results, even leading to diametrically opposite evaluation conclusions. Consequently, these irrational results can make decision makers afford huge cost sometimes, which can be avoided originally. Therefore, it is critically important of finding some effective measures to deal with or avoid the limitations mentioned above during the evaluation process.

1.2 Research Motivation

Basically, most early work on decision making with linguistic information were making use of fuzzy sets as a tool for modeling linguistic information and aggregation methods were then developed based on “Zadeh’s extension principle”, e.g., [66]. Typically also, another approach aimed to develop the “linguistic symbolic computational model based on ordinal scales” [80]. Because of the inherent operation mechanism of these two linguistic computational models, the results of a computational process usually don’t exactly match any of the initial linguistic terms and, hence, a process of linguistic approximation must be applied to convert the computational results into linguistic terms of the initial linguistic domain. This linguistic approximation process causes a loss of information and consequently leads to the lack of precision in the final results [10]. As for overcoming this limitation in the computational stage for CW, Herrera and Martínez developed a “2-tuple fuzzy linguistic representation model” based on the concept of symbolic translation so as to improve precision of the final results [28]. Although their approach has no loss of information when it was applied, they also pointed out that “it was only suitable for linguistic variables with equidistant labels” [28].

In order to overcome the limitation of “2-tuple fuzzy linguistic representation model” [28], Wang and Hao proposed “a proportional 2-tuple fuzzy linguistic representation model” for CW making use of the canonical characteristic values (*CCV*) of linguistic terms determined by their corresponding semantics [70]. Wang and Hao’s “proportional 2-tuple fuzzy linguistic representation model” interestingly provides a suitable and more flexible space in a computation stage for CW, which could allow evaluators to flexibly evaluate the performances of alternatives by not only one label but with the form of proportional 2-tuples $(\alpha A, \beta B)$, where A and B are two consecutive linguistic terms, and $\alpha, \beta \in [0, 1]$, $\alpha + \beta = 1$. However, as we can see from the definition, this model cannot deal with the decision situations where alternative performances are generally assessed by means of uncertain judgments represented by probability distributions over the linguistic term set. In addition, due to the premise that the summation of a pair of symbolic proportions must equal to 1, this model cannot handle ignoring information. In other words, it is only applicable under the context that all the linguistic assessments are complete. As a matter of fact, incomplete assessments emerge commonly when evaluators are lack of confidence, especially in the case of facing with uncertain, vague and imprecise information.

As such, it would be desirable that an appropriate extension of Wang and Hao’s “proportional 2-tuple fuzzy linguistic representation model” [70] could be developed. Therefore, modifying and

extending Wang and Hao’s “proportional 2-tuple fuzzy linguistic representation model” [70], and then developing new models for MADM problems motive our current research.

1.3 Research Objectives

According to the three categories of linguistic computational models in the literature, we will analyze and investigate the main features and limitations of these linguistic computational models in this research. Then, we will correspondingly develop three evaluation models for MADM problems, aiming at not only overcoming the limitations of the three categories of models, such as loss of information during the approximation process, without directly considering the underlying vagueness of linguistic terms, without involving the ignoring information and so on, but also being able to deal with more general cases of linguistic assessments possibly associated with uncertainty and incomplete information. Meanwhile, the three evaluation models developed in this research will also be associated with some other features, such as ease operation in the complicated linguistic context, flexible operation space for evaluators under uncertainty and so forth. These features can be regarded as some measures for evaluators to deal with MADM problems under uncertainty. Consequently, these linguistic computational models developed in this research can improve the precision, reasonability and reliability of the final results, and give decision makers a more comprehensive guidance.

1.4 Related Knowledge

The main work of this research is to develop linguistic computational models for MADM problems. Hence, it is necessary to introduce some basic knowledge regarding to decision making and CW.

1.4.1 Decision making

Decision making is a multi-discipline comprising philosophy, psychology, business, operations research, system engineering and management science. “It includes many procedures, methods, and tools for identifying, clearly representing, and formally assessing important aspects of a decision” [64]. Decision making is also a typical human mental process, resulting in the selection

of an alternative among several alternative possibilities. Because “decision making is an inherent human ability which is not necessarily rationally guided, it does not necessarily need precise and complete information about the set of feasible alternatives” [20]. Actually, we cannot acquire the exact information about each alternative that we are considering in most situation when we make a decision, and many respects of different objectives in the reality cannot be evaluated in a precise form but rather in an approximate way, which is often with vague, imprecise and uncertain knowledge. Traditionally, these decision situations are often defined under uncertain frameworks that could be managed by probabilistic models when assuming that any uncertainty can be represented by a probabilistic distribution. However, it is common that uncertainty does not always have a probabilistic nature. This fact has resulted in a reaction that many scholars and researchers apply “fuzzy sets theory” [89] to model the vagueness, imprecision and uncertainty in decision making processes [12], [21], [38], [46].

1.4.2 Multiple attribute decision making problems

In practice, most decision making problems involve multiple attribute of both quantitative and qualitative nature, “which may be represented by a hierarchy” [5], [61]. While quantitative attributes can be measured by means of numeric scales in the form of numbers, intervals or fuzzy numbers, qualitative attributes which are often associated with uncertain information cannot be assessed in the same way. In such case, we often prefer to use words in natural language instead of numerical values. These situations may appear due to different reasons [12]. For example, because of the inherent nature, some information cannot be quantified, but only by means of linguistic terms (e.g., when we look at a picture, the linguistic terms, such as “fair”, “beautiful”, “excellent” can be used). In other cases, it is quite difficult to state precise quantitative information because sometimes we cannot obtain such kind of information or the cost of obtaining it is too high and an “approximate value” can be tolerated (e.g., we can use linguistic terms like “cold”, “warm”, “hot” to describe the temperature instead of numeric values). In these situations, it is reasonable and necessary to make use of fuzzy linguistic approach which provide tools to model and represent qualitative attributes by means of linguistic values of linguistic variables [90]. The use of linguistic information implies the necessity of operating with the mechanism for CW [93] so as to fusion linguistic information and then provide an evaluation for decision making.

1.4.3 Computing with words

Computing with words, that is, “the methodology that uses words and propositions from a natural language as its main objects of computation, was firstly introduced in the seminal paper by Zadeh” [93], aiming at capturing the concept of automated reasoning involving linguistic terms, rather than numerical quantities. In fact, the idea of CW, also known as granular computing, is rooted from Zadeh’s earlier work on fuzzy constraints and linguistic variables [90], [91], [92].

About granule, Zadeh explained that “a granule is a clump of objects (or points) which are drawn together by indistinguishability, similarity, proximity or functionality” [94]. Generally speaking, granules can be seen from two perspectives. On the one hand, the granules can be explained from the crisp perspective, i.e., they can be easily differentiated because the decomposition process is made into disjoint granules. This aspect is often used in various computational methods and techniques, such as rough set theory, divide and conquer, decision trees and so on. On the other hand, the granules are related to fuzzy rather than crisp especially when we deal with the aspects associated with human reasoning and concept formation [23], [94]. Thus, “a granule can be seen as a fuzzy set of points drawn together by similarity in a lot of situations” [23]. So far, a number of researchers have extensively studied the concept of fuzzy granulation in the literature, such as [4], [17], [30], [56], [57] and so on.

In CW, there is another basic assumption, i.e., “information is conveyed by constraining the value of variables, which take possible values as linguistic ones” [23]. “A linguistic variable is variable whose values are not numbers but words or sentences in a natural or artificial language” [90]. Generally speaking, although linguistic values are less specific than numerical values, they are much closer to human cognitive processes. Thus, humans can easily express and use their linguistic knowledge to successfully solve decision making problems with uncertainty.

Formally, “a linguistic variable is characterized by a quintuple $(H, T(H), U, G, M)$ in which H is the name of the variable; $T(H)$ (or simply T) denotes the term set of H , i.e., the set of names of linguistic values of H , with each value being a fuzzy variable that is denoted generically by X and ranging across a universe of discourse U , which is associated with the base variable u ; G is a syntactic rule (which usually takes the form of a grammar) for the generation of the names of values of H ; and M is a semantic rule for associating its meaning with each H , $M(X)$, which is a fuzzy subset of U ” [90]. Since its foundations [90], [91], [92], [93], lots of researchers have carried out a number of studies on linguistic variables and the CW methodology, such as [28], [37], [39], [40], [43], [45], [52], [68], [69], [70], [78], [88] and so on.

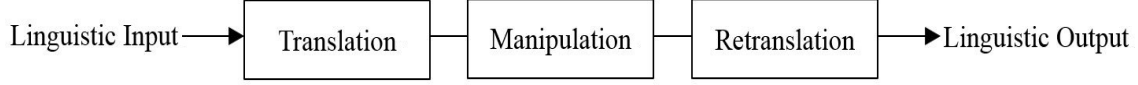


Figure 1.1. Computing with words scheme [59]

It is worth mentioning that CW has attracted great attention from the fuzzy-set research community not only since Zadeh [93] coined it, but also since the early 1980s, when different researchers such as [62], [66], [80] started to propose different computing schemes to operate with linguistic information. Such schemes are quite similar and keep a structure in which the input linguistic information should be mapped into fuzzy set models and the results should be expressed into linguistic information, which are easy to understand by human beings (see Figure 1.1) [59].

1.4.4 Fuzzy linguistic approach

CW has been and still is a key methodology in linguistic decision making problems, while fuzzy linguistic approach is a common approach to manage the uncertainty, model the linguistic information, and deal with linguistic variable. Generally speaking, it is necessary to choose appropriate linguistic term set and the associated semantics before dealing with linguistic variables. To do so, an important task is to determine the granularity of uncertainty, i.e., the level of discrimination among different counts of uncertainty, in the other words, the cardinality of the linguistic term set used to assess the linguistic variables. As mentioned in [8], “the cardinality of the term set must be small enough so as not to impose useless precision on the users, and it must be rich enough in order to allow a discrimination of the assessments in a limited number of degrees”. Usually, the values of cardinality are odd ones, such as 7 or 9. Sometimes, we also use even cardinality in order to satisfy special requirements. If odd cardinality is used in the linguistic model, the middle linguistic term often represents an assessment of “approximately 0.5”, while the other terms are around it symmetrically [7]. “These classical cardinality values seem to satisfy the Miller’s observation regarding the fact that human beings can reasonably manage to bear in mind seven or so items” [54].

Syntactically, there are two main approaches to generate a linguistic term set.

(1) Context-free grammar approach: This approach uses a context-free grammar G to define the linguistic term set. As a result, the linguistic terms are the sentences which are generated by

the grammar G [6], [9], [90], [91], [92]. “A grammar G is a 4-tuple (V_N, V_T, I, P) , where V_N is the set of nonterminal symbols, V_T is the set of terminals’ symbols, I is the starting symbol, and P is the production rules that are defined in an extended BackusCNaur form” [9]. “Among the terminal symbols of G , we can find primary terms (e.g., low, medium, high), hedges (e.g., not, much, very), relations (e.g., lower than, higher than), conjunctions (e.g., and, but), and disjunctions (e.g., or). Thus, choosing I as any nonterminal symbol and using P could be generated linguistic expressions, such as, {lower than medium, greater than high, ...}”. It is worth mentioning that using this approach may yield an infinite term set.

(2) An ordered structure approach: The linguistic term set is defined with finite and ordered structure of terms. All terms are regarded as primary ones and are distributed on a scale on which a total order has been defined [25], [82]. For instance, a linguistic term set S which includes seven linguistic terms could be given as follows:

$$S = \{ s_0 = \text{Very Low}, s_1 = \text{Low}, s_2 = \text{Fairly Low}, s_3 = \text{Medium}, \\ s_4 = \text{Fairly High}, s_5 = \text{High}, s_6 = \text{Very High} \}$$

in which the existence of the following is usually required.

- 1) A negation operator $\text{Neg}(s_i) = s_j$ so that $j = g - i$ ($g + 1$ is the granularity of the term set).
- 2) A maximization operator: $\text{Max}(s_i, s_j) = s_i$ if $s_i \geq s_j$.
- 3) A minimization operator: $\text{Min}(s_i, s_j) = s_i$ if $s_i \leq s_j$.

After determining the mechanism of generating a linguistic term set, the procedure of defining its associated semantics should be carried out. We can find three main approaches for defining the semantics of linguistic term set in the literature.

(1) Semantics based on membership functions and a semantic rule: By the use of this approach, the meaning of each linguistic term is determined by a fuzzy subset defined in the interval $[0, 1]$, which is described by membership functions [9]. Usually, this semantic approach is used when the linguistic descriptors are generated by means of a context-free grammar. “This approach consists of two elements: the primary fuzzy sets designed as associated semantics of the

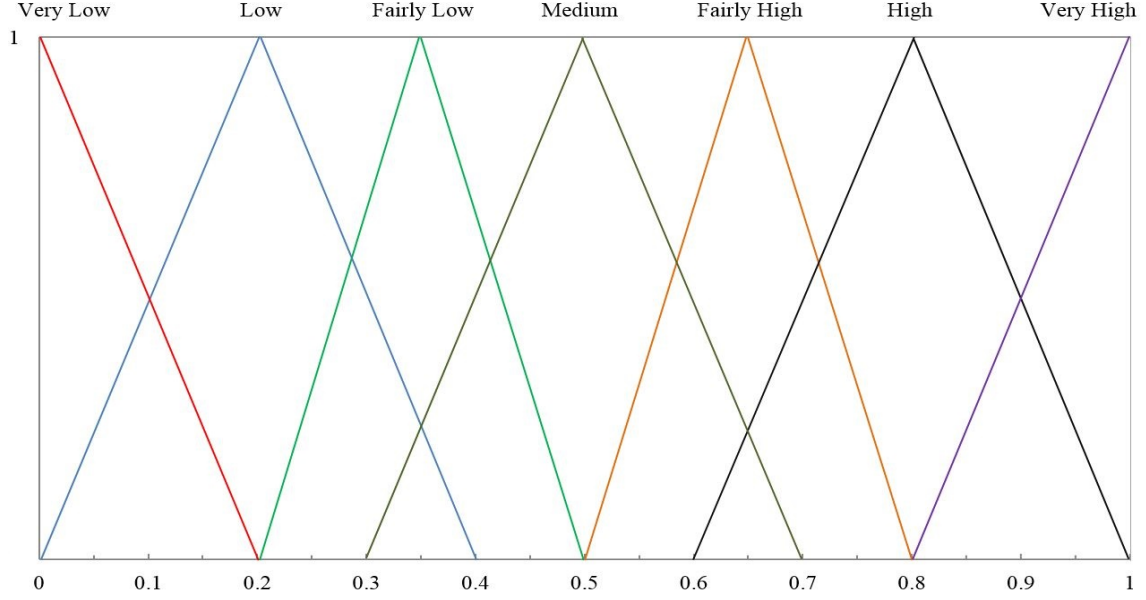


Figure 1.2. A set of seven terms with its semantics

primary linguistic terms, and a semantic rule M that provides the fuzzy sets of the non-primary linguistic terms” [90], [91], [92]. “Often, while the primary terms are labels of primary fuzzy sets which are defined subjectively and context-dependently, the semantic rule M defines linguistic hedges and connectives as mathematical operations on fuzzy sets aimed at modifying the meaning of linguistic terms applied” [32].

(2) Semantics based on an ordered structure of a finite linguistic term set: The semantics is defined over linguistic term set with finite and ordered structure of terms. Therefore, the evaluators supply their linguistic assessments by the use of an ordered linguistic term set [67], [82]. The distribution of a linguistic term set in the interval $[0, 1]$ can be distributed either symmetrically [82] or non-symmetrically [26], [67] depending on a particular situation.

(3) Mixed semantics: This is a mixed approach which uses two semantic approaches mentioned above, that is, an ordered structure of the primary linguistic terms and a fuzzy set representation of linguistic terms (see, e.g., [24], [32], [60]) for more details).

After the linguistic term set and associated semantics have been elaborately defined and established, experts or evaluators can give their linguistic assessments according to the semantics of linguistic terms. Generally speaking, it is good enough if we use linear trapezoidal membership functions to capture the vagueness and uncertainty of linguistic assessments [16]. “The parametric representation is achieved by the 4-tuple (a, b, d, c) , where b and d indicate the interval in which the membership value is 1, a and c are the left and right limits of the definition

domain of the trapezoidal membership function respectively” [7]. A special case of fuzzy numbers are triangular membership functions denoted by a 3-tuple (a, b, c) , i.e., $b = d$. Figure 1.2 shows an example which is a linguistic term set, and the semantics of their terms could be represented as

$$\begin{aligned} \text{Very High} &= (0.8, 1, 1), \text{High} = (0.6, 0.8, 1), \text{Fairly High} = (0.5, 0.65, 0.8), \\ \text{Medium} &= (0.3, 0.5, 0.7), \text{Fairly Low} = (0.2, 0.35, 0.5), \text{Low} = (0, 0.2, 0.4), \\ \text{Very Low} &= (0, 0, 0.2). \end{aligned}$$

1.4.5 Linguistic decision making resolution scheme

It is necessary to analyze the phases of a linguistic decision scheme when the linguistic information is formally modelled. In linguistic decision analysis, a common decision resolution scheme, which is shown in Figure 1.3 must comply with the following three steps [24].

- 1) Defining the linguistic term set: This step requires us to establish the linguistic expression domain that is used to supply evaluators with an instrument by which they can assess the linguistic performance values about alternatives according to the different attributes. Basically, one has to choose the granularity of the linguistic term set, its labels, and their associated semantics.
- 2) Developing the aggregation operator for linguistic information: According to different situations, develop appropriate aggregation operators for obtaining the aggregated values of the linguistic performance values provided by evaluators.
- 3) Selecting the best alternatives, including two phases:
 - Aggregation phase: Obtain the aggregated linguistic preferences of alternatives by the use of developed aggregation operator.
 - Exploitation phase: Make a ranking order among the alternatives according to the aggregated linguistic preferences and then choosing the best one(s).

Essentially, “the first two steps serve the aggregation phase in the third step, while the exploitation phase is determined depending on the choice of the semantic description of the linguistic term set” [32].

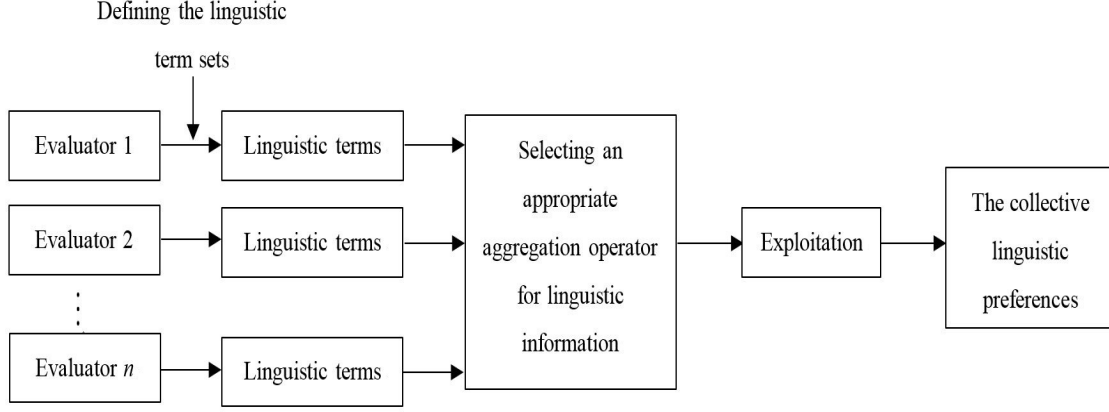


Figure 1.3. The scheme of a linguistic decision making problem

1.5 Organization of the Dissertation

This dissertation is composed of six chapters. The detailed explanation is shown as follow:

Chapter 1 first introduces research background, research purpose, research motivation, and research objective. Then, the related knowledge of this dissertation are briefly introduced for the purpose of conveniently carrying out subsequent chapters.

Chapter 2 recalls some main fuzzy linguistic approaches applied to the applications of multiple attribute decision making problems.

Chapter 3 proposes a proportional 3-tuple fuzzy linguistic representation model for multiple attribute decision making with incomplete linguistic information. Besides, an important notion, called preference-preserving proportional 3-tuple transformation, will be proposed serving for transforming linguistic assessments between two different linguistic term sets without loss of information. An illustration example taken from previous literature will be used in order to illustrate the proposed model.

Chapter 4 develops a proportional fuzzy linguistic distribution model for multiple attribute decision making with incomplete linguistic information. Several aggregation operators and expected utility in proportional fuzzy linguistic distribution will be proposed. An illustration example taken from previous literature will be employed so as to illustrate the proposed model.

Chapter 5 develops an interval fuzzy linguistic distribution model for multiple attribute decision making with incomplete linguistic information. Several interval aggregation operators and expected utility in interval fuzzy linguistic distribution will be proposed. Two illustration examples will

be used for the purpose of illustrating the proposed model from both a easily comprehensible perspective and a practical perspective.

Chapter 6 presents conclusions of this research. Meanwhile, the contributions of this research and future work will be discussed.

Chapter 2

A Survey on Decision Making with Fuzzy Linguistic Models

In decision making applications, the main problem for CW is how to represent and aggregate linguistic information for evaluation of alternatives. For this respect, lots of linguistic computational models have been proposed in the literature. In this chapter, we make a review of several main linguistic computational models based on fuzzy linguistic approach.

2.1 Linguistic Computational Model Based on Membership Functions

This kind of linguistic computational model is based on fuzzy linguistic approach and uses the extension principle to make the computations directly on the membership functions of the linguistic terms [7], [13]. “The use of extended arithmetic based on the extension principle increases the vagueness of the results. Therefore, the results obtained by the fuzzy linguistic operators based on the extension principle are fuzzy numbers that usually do not match with any linguistic term in the initial term set” [10]. According to different purposes, the results can be present either by means of the fuzzy numbers themselves (ranking purposes) [2], [22], or by means of linguistic labels which are computed from the fuzzy numbers that are obtained by the use of a linguistic approximation process (an interpretable and linguistic result purpose) [13], [51], [83].

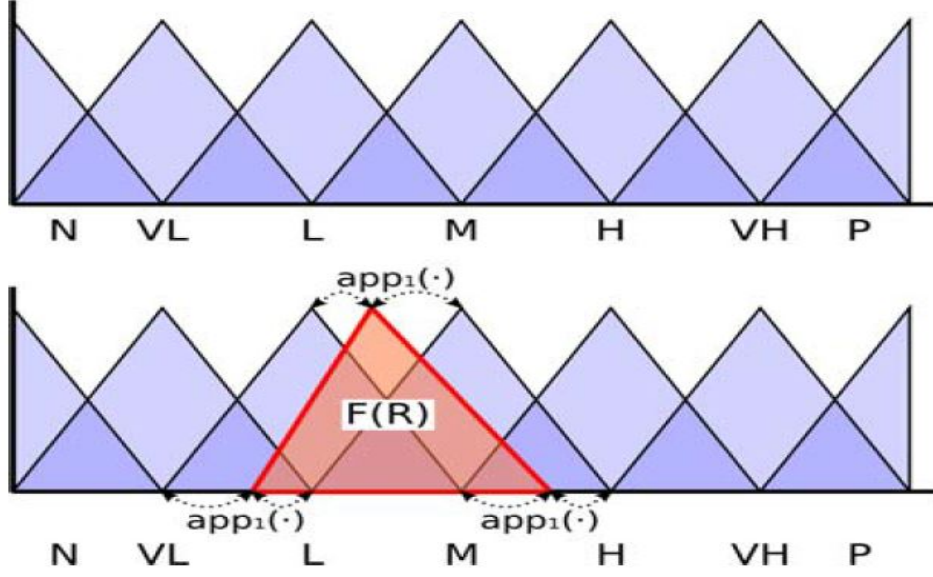


Figure 2.1. Retranslation problem [23]

If latter purpose is required, then an approximation function $app_1(\cdot)$ is applied to associate the fuzzy result $F(\mathcal{R})$ with a label in S :

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

where S^n symbolizes the n Cartesian product of S , $\tilde{F}$ is an aggregation function based on the extension principle, and $F(\mathcal{R})$ is the set of fuzzy sets over the set of real numbers $\mathcal{R}$ [59].

It is worth noting that there is a loss of information when the approximation process is applied leading to the lack of accuracy of the results [10], which can be easily found in the following example of retranslation problem, shown in Figure 2.1. The obtained fuzzy number $F(\mathcal{R})$ (highlighted in the figure) does not have an associated linguistic label on a particular term set $S = \{N, VL, L, M, H, VH, P\}$. Because we need to obtain a linguistic value in the end, an approximation function $app_1(\cdot)$ that will assign one of the L or M labels to $F(\mathcal{R})$ needs to be applied [23].

2.2 Linguistic Symbolic Computational Models Based on Ordinal Scales

This kind of linguistic computational models makes direct computations on labels and uses the ordered structure of the linguistic term set to accomplish symbolic computations. Basically, there are three kinds of linguistic symbolic computational models that are based on ordinal scales [23] in the literature: “a linguistic symbolic computational model based on ordinal scales and max-min operators” [80], “a linguistic symbolic computational model based on indexes” [15], and “a linguistic symbolic computational model based on continuous term sets” [79].

In the first linguistic symbolic computational model, an ordered linguistic scale $S = \{s_0, \dots, s_g\}$ with a linear ordering is defined. The classical operators Max, Min and Neg which were introduced in Section 1.44 of Chapter 1 are used in order to be able to aggregate information expressed as linguistic labels in that ordered linguistic scale.

In the second linguistic symbolic computational model, a convex combination of linguistic labels is used during the aggregation process. “The convex combination of linguistic labels directly acts over the label indexes of the linguistic terms set $S = \{s_0, \dots, s_g\}$ in a recursive way, and produces a real value on the granularity interval of the linguistic terms set S ” [15]. Usually, this model assumes odd cardinality of the linguistic term set with linguistic labels symmetrically distributed around the middle linguistic term [23]. Because the aggregation results are numeric values, $\gamma \in [0, g]$, which usually don’t match with any linguistic labels in the initial linguistic term set, they must be approximated in each step of the process by means of an approximation function $app_2 : [0, g] \rightarrow \{0, \dots, g\}$. By this way, a numeric value can be obtained which indicates the index of the associated linguistic term, $s_{app_2(\gamma)} \in S$ [59]. Formally, it can be expressed as

$$S^n \xrightarrow{C} [0, g] \xrightarrow{app_2(\cdot)} \{0, \dots, g\} \rightarrow S,$$

where C is a symbolic linguistic aggregation operator, $app_2(\cdot)$ is an approximation function used to obtain an index $\{0, \dots, g\}$ associated with a term in $S = \{s_0, \dots, s_g\}$ from a value in $[0, g]$ [59].

In the third linguistic symbolic computational model, the linguistic term set $S = \{s_0, \dots, s_g\}$ which is discrete is extended into a linguistic term set $\bar{S} = \{s_\alpha \mid s_0 < s_\alpha \leq s_g, \alpha \in [0, g]\}$ which is continuous. In the continuous term set, if $s_\alpha \in S$, then s_α is called an original linguistic term, otherwise, s_α is called a virtual linguistic term. Generally speaking, evaluators use original

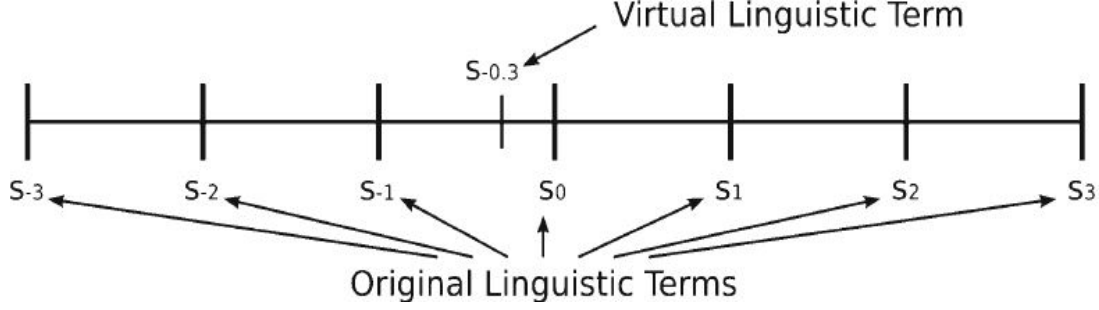


Figure 2.2. Example of the third linguistic symbolic computational model [23] [79]

linguistic terms to assess the performances of alternatives, while the virtual linguistic terms only appear in the operation process [79].

Figure 2.2 is an example that a discrete term set $S = \{s_{-3}, \dots, s_3\}$ (original linguistic terms) is extended into a continuous term set. From Figure 2.2 we can find that the virtual linguistic terms $s_{-0.3} \in [-3, 3]$ can be obtained after linguistic information aggregation in order to avoid loss of information [23], [79].

It is worth noting that since the virtual terms, which are quite different range than the original ones are created in the aggregation process, the interpretability of this computational model is limited. Therefore, as Herrera et al. points out [23], “this model also presents a retranslation problem if the results of the operations are virtual linguistic terms (and they will usually be virtual ones) and the final results must be expressed in the original linguistic term set. However, as the linguistic symbolic computational model based on linguistic terms is simple, as it avoids loss of information and as virtual linguistic terms can be used to rank alternatives and thus, to select the best of them, its use can be convenient in particular situations”. (For more information about linguistic symbolic computational models based on ordinal scales, interested readers can refer to, e.g., [23].)

2.3 Linguistic Symbolic Computational Models Based on 2-Tuple Representation

Generally speaking, when “linguistic computational models based on extension principle” [7], [13] and “linguistic symbolic computational models based on the ordered structure of linguistic term sets” [15], [80] (except “the linguistic symbolic computational model based on continuous term sets” [79]), are used in decision making for CW, the results of a computational process

usually don't exactly match any of the initial linguistic terms and, hence, a process of linguistic approximation must be applied to convert the computational results into linguistic terms of the initial linguistic domain. This linguistic approximation process causes a loss of information and consequently leads to the lack of precision in the final results [10]. In order to avoid this limitation in the computational stage for CW and improve the precision of the final results, Herrera and Martínez [28] developed the so-called “2-tuple fuzzy linguistic representation model based on the concept of symbolic translation”.

This symbolic model “extends the use of indexes modifying the fuzzy linguistic approach representation by adding a parameter to the basic linguistic representation in order to improve the accuracy of the linguistic computations after the retranslation step keeping in the CW scheme and the interpretability of the results” [49].

2.3.1 Representation model

Formally, let $S = \{s_0, s_1, \dots, s_n\}$ be a linguistic term set, and the term s_i with $i = 0, \dots, n$, represents a possible value for a linguistic variable. The total order on S is defined as: $s_i \leq s_j \Leftrightarrow i \leq j$. There is a negation operator: $\text{Neg}(s_i) = s_j$ such that $j = n - i$, where $n + 1$ is the cardinality of S . In general, using a symbolic method to aggregate linguistic information, we often get a value $\beta \in [0, n]$, and $\beta \notin \{0, \dots, n\}$. Then, an approximation function is used in order to conveniently express the index of the result in S .

To avoid any approximation process which causes a loss of information in the process of computing with words, the 2-tuple (s_i, α) that expresses the equivalent information to β is obtained with the following function:

$$\begin{aligned} \Delta : [0, n] &\rightarrow S \times [-0.5, 0.5) \\ \Delta(\beta) &= (s_i, \alpha), \text{ with } \begin{cases} s_i, & i = \text{round}(\beta) \\ \alpha = \beta - i, & \alpha \in [-0.5, 0.5) \end{cases} \end{aligned}$$

where $\text{round}(\cdot)$ is the usual round operation, s_i has the closest index label to β , and α is the value of the symbolic translation.

Inversely, a 2-tuple $(s_i, \alpha) \in S \times [-0.5, 0.5)$ can also be equivalently represented by a numerical value in $[0, n]$ by means of the following transformation:

$$\Delta^{-1} : S \times [-0.5, 0.5) \rightarrow [0, n]$$

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

2.3.2 Computational model

Based on the functions Δ and Δ^{-1} , 2-tuple fuzzy linguistic representation model is with the following operations:

(1) Comparison of 2-tuple

According to an ordinary lexicographic order, the comparison of linguistic information represented by 2-tuple is carried out as follows [23].

Let (s_k, α_1) and (s_l, α_2) be two 2-tuples with each one representing a counting of information, then

1) if $k < l$ then $(s_k, \alpha_1) < (s_l, \alpha_2)$

2) if $k = l$ then

- if $\alpha_1 = \alpha_2$ then $(s_k, \alpha_1), (s_l, \alpha_2)$ represents the same information
- if $\alpha_1 < \alpha_2$ then $(s_k, \alpha_1) < (s_l, \alpha_2)$
- if $\alpha_1 > \alpha_2$ then $(s_k, \alpha_1) > (s_l, \alpha_2)$.

(2) The negation operator of 2-tuple

The negation operator over 2-tuples is defined by

$$\text{Neg}((s_i, \alpha)) = \Delta(n - (\Delta^{-1}(s_i, \alpha)))$$

where $n + 1$ is the cardinality of S , $S = \{s_0, s_1, \dots, s_n\}$.

(3) 2-tuple aggregation operators

Because 2-tuples can be transformed into numerical values without loss of information, theoretically, conventional aggregation operators can be extended for 2-tuples.

Let $x = \{(s_1, \alpha_1), \dots, (s_n, \alpha_n)\}$ be a set of 2-tuples, and $W = \{\omega_1, \dots, \omega_n\}$ be their associated weights. Then, the 2-tuple weighted average $\bar{x}$ is computed by

$$\bar{x} = \Delta\left(\frac{\sum_{i=1}^n \Delta^{-1}(s_i, \alpha_i) \cdot \omega_i}{\sum_{i=1}^n \omega_i}\right) = \Delta\left(\frac{\sum_{i=1}^n \beta_i \cdot \omega_i}{\sum_{i=1}^n \omega_i}\right).$$

In the literature, many other 2-tuple aggregation operators have also been proposed, such as 2-tuple arithmetic mean, 2-tuple ordered weighted average operator. (For more details, see, e.g., [27], [28].)

2.3.3 The use of the 2-tuple linguistic representation model

Because the advantages of “2-tuple fuzzy linguistic representation model” [28], such as “its accuracy, its usefulness for improving linguistic solving processes in different applications, its interpretability, its ease managing of complex frameworks” [49] and so forth, it has been extensively and intensively researched and widely used as basis for different models employed in various decision making problems. Recently, several methodologies that are based on linguistic 2-tuples have been developed to deal with decision making problems under complex frameworks. In [25], Herrera et al. proposed an approach to fusion multi-granular information in decision making, in which a basic linguistic term set is selected with maximum granularity, and a transformation function that represents each linguistic performance value as a fuzzy set is defined in this basic linguistic term set. Herrera and Martínez [29] extend this methodology and using a linguistic hierarchies term sets to unify multi-granular hierarchical linguistic information using the 2-tuple linguistic model without loss of information. Huynh et al. also extended this model for MEDM in general multigranular linguistic contexts [35]. Herrera and Martínez [26] presented a 2-tuple based methodology to deal with unbalanced linguistic information, which provided an algorithm to represent the linguistic terms and a computational model to accomplish processes of CW based on the 2-tuple linguistic model. Wang and Hao [70], [71] proposed “a proportional 2-tuple fuzzy linguistic representation model” to deal with linguistic term sets that are not uniformly and symmetrically distributed. By defining the concept of the numerical scale, Dong et al. [18] proposed an integration of Herrera and Martínez’s model and Wang and Hao’s model. Dong et al. [19] proposed “an interval version of the 2-tuple fuzzy linguistic representation

model”, which generalizes the numerical scale approach to set the interval numerical scale, by considering the context where semantics of linguistic terms are defined by interval type-2 fuzzy sets.

Although 2-tuple fuzzy linguistic representation model are inspired by the symbolic models used in decision making [15], [80], [81], [82], it and its extensions have been employed in numerous applications, such as supply chain management [72], screening new product projects [33], sensory evaluation [47], [48], engineering evaluation processes [50], intelligent agent system [14], research resources management [58], risk evaluation [11] and so on. A recent overview on the 2-tuple linguistic model, its extensions, specific methodologies, and applications can be found in [49].

2.4 Linguistic Symbolic Computational Models Based on Proportional 2-Tuple Representation

In last section, we recalled Herrera and Martínez’s “2-tuple fuzzy linguistic representation model” [28], which aimed at avoiding loss of information caused by linguistic approximation process in the computational stage for CW. However, they also pointed out that “this model was only suitable for linguistic variables with equidistant labels”. In addition, as argued by Lawry [43], “although Herrera and Martínez’s symbolic approach offered a computationally much more feasible method than those approaches using the extension principle in CW, it did not directly take into account the underlying vagueness of linguistic terms”. In an attempt to improve Herrera and Martínez’s 2-tuple fuzzy linguistic representation model so as to be able to deal with unbalanced linguistic term sets while simultaneously taking the underlying semantics of terms into account, Wang and Hao [70] proposed a so-called “proportional 2-tuple fuzzy linguistic representation model for CW making use of the canonical characteristic values (*CCV*) of linguistic terms determined by their corresponding semantics”. Because our inspiration of developing a proportional 3-tuple fuzzy linguistic representation model for MADM problems comes from this symbolic linguistic computational model, we review it in this section on the one hand for appreciation, on the other hand for paving the way for next chapter.

2.4.1 Representation model

Formally, let $S = \{s_0, s_1, \dots, s_n\}$ be an ordinal term set with $s_0 < s_1 < \dots < s_n$ (“ $<$ ” represents order relation. $s_i < s_j$ if and only if $i < j$), $I = [0, 1]$ and

$$IS \equiv I \times S = \{(\alpha, s_i) : \alpha \in [0, 1] \text{ and } i = 0, 1, \dots, n\}.$$

Given a pair (s_i, s_{i+1}) of two successive ordinal terms of S , any two elements (α, s_i) , (β, s_{i+1}) of IS are called a symbolic proportion pair and α, β are called a pair of symbolic proportions of the pair (s_i, s_{i+1}) if $\alpha + \beta = 1$. A symbolic proportion pair (α, s_i) , $(1 - \alpha, s_{i+1})$ is denoted by $(\alpha s_i, (1 - \alpha) s_{i+1})$ and the set of all the symbolic proportion pairs is denoted by S^* , i.e., $S^* = \{(\alpha s_i, (1 - \alpha) s_{i+1}) : \alpha \in [0, 1] \text{ and } i = 0, 1, \dots, n - 1\}$. The set S^* is called the ordinal proportional 2-tuple set generated by S and the members of S^* are called ordinal proportional 2-tuples.

Remark: For $i = \{1, \dots, n - 1\}$, ordinal term s_i can use either $(0s_{i-1}, 1s_i)$ or $(1s_i, 0s_{i+1})$ as its representative in S^* , by abuse of notation.

Compared with 2-tuple representation, the presentation of proportional 2-tuple provides a suitable and more flexible space in a computation stage for CW. It could allow evaluators in various situations to flexibly evaluate performances of alternatives by not just one label but with the form of $(\alpha s_i, (1 - \alpha) s_{i+1})$.

2.4.2 Computational model

(1) Comparison of proportional 2-tuple

The comparison of linguistic information represented by proportional 2-tuple is carried out as follows:

Let $S = \{s_0, s_1, \dots, s_n\}$ be an ordinal term set and S^* be the ordinal proportional 2-tuple set generated by S . For any $(\alpha s_i, (1 - \alpha) s_{i+1})$, $(\beta s_j, (1 - \beta) s_{j+1}) \in S^*$, define

$$\begin{aligned} (\alpha s_i, (1 - \alpha) s_{i+1}) < (\beta s_j, (1 - \beta) s_{j+1}) &\Leftrightarrow \alpha i + (1 - \alpha)(i + 1) \\ &< \beta j + (1 - \beta)(j + 1) \Leftrightarrow i + (1 - \alpha) < j + (1 - \beta). \end{aligned}$$

Thus, for any two proportional 2-tuple $(\alpha s_i, (1 - \alpha)s_{i+1})$ and $(\beta s_j, (1 - \beta)s_{j+1})$, we obtain:

1) if $i < j$, then

- $(\alpha s_i, (1 - \alpha)s_{i+1}), (\beta s_j, (1 - \beta)s_{j+1})$ represent the same information when $i = j - 1$ and $\alpha = 0, \beta = 1$,
- $(\alpha s_i, (1 - \alpha)s_{i+1}) < (\beta s_j, (1 - \beta)s_{j+1})$ otherwise.

2) if $i = j$, then

- if $\alpha = \beta$ then $(\alpha s_i, (1 - \alpha)s_{i+1}), (\beta s_j, (1 - \beta)s_{j+1})$ represents the same information,
- if $\alpha < \beta$ then $(\alpha s_i, (1 - \alpha)s_{i+1}) > (\beta s_j, (1 - \beta)s_{j+1})$,
- if $\alpha > \beta$ then $(\alpha s_i, (1 - \alpha)s_{i+1}) < (\beta s_j, (1 - \beta)s_{j+1})$.

(2) Negation operator of a proportional 2-tuple

The negation over proportional 2-tuples is defined as

$$Neg(\alpha s_i, (1 - \alpha)s_{i+1}) = ((1 - \alpha)s_{n-i-1}, \alpha s_{n-i}),$$

where $n + 1$ is the cardinality of S , $S = \{s_0, s_1, \dots, s_n\}$.

(3) Canonical characteristic values of proportional 2-tuple

Wang and Hao introduced the so-called “canonical characteristic values (*CCV*)” to represent fuzzy number based semantics of linguistic terms and developed an efficient method for CW based on the proportional 2-tuple fuzzy linguistic representation. Specifically, “let $F(R)$ be the set of fuzzy numbers defined on R (the real numbers set). Each fuzzy number, $y_i \in F(R)$, has associated a membership function, $u_{y_i} : R \rightarrow [0, 1]$. For each fuzzy number, y_i , there is a set of characteristic values, $CV_{y_i} = \{C_i^1, C_i^2, \dots, C_i^z\}$, which are crisp values that summarize the information given by y_i , i.e., they support its meaning” [70]. Wang and Hao have enumerated some *CCV* for representing the related information of proportional 2-tuples, such as expected value, center of gravity, mean of maxima. Particularly, if the semantics of linguistic terms is defined by symmetrical triangular fuzzy numbers in $[0, 1]$, i.e., $y_i = [c - \delta, c, c + \delta]$, then the expected value (*EV*) of y_i is defined by $EV(y_i) = c$ and used as *CCV* of y_i .

With the notions of proportional 2-tuple and CCV , the computation operator used for transforming a proportional 2-tuple into a numerical value belonging to $[0, 1]$ is defined as follows.

Let $c_i \in [0, 1]$ with $c_0 < c_1 < \dots < c_n$ be the canonical characteristic values of s_i , i.e., $CCV(s_i) = c_i$ for all $i = 0, 1, \dots, n$. Then, define the function CCV on S^* by

$$\begin{aligned} CCV : S^* &\rightarrow [0, 1] \\ CCV((\alpha s_i, (1 - \alpha)s_{i+1})) &= \alpha CCV(s_i) + (1 - \alpha)CCV(s_{i+1}) \\ &= \alpha c_i + (1 - \alpha)c_{i+1} \\ &= z \in [0, 1] \end{aligned}$$

and call it the corresponding canonical characteristic value function on S^* generated by CCV on S . It has been proved by Wang and Hao [70] that the CCV is a bijection from S^* to $[c_0, c_n]$. Specifically, let us define $f : [0, n] \rightarrow [c_0, c_n]$ by

$$f(x) = c_i + \beta(c_{i+1} - c_i)$$

where $i = E(x)$, E is the integral part function and $\beta = x - i$. Then f is a bijection. Since,

$$\begin{aligned} CCV(((1 - \beta)s_i, \beta s_{i+1})) &= (1 - \beta)c_i + \beta c_{i+1} \\ &= c_i + \beta(c_{i+1} - c_i) \\ &= f(i + \beta) \\ &= f(\pi((1 - \beta)s_i, \beta s_{i+1})) \end{aligned}$$

for all $i = 0, 1, \dots, n - 1, \beta \in [0, 1]$, thus $CCV = f \circ \pi$. Here, “ π is the position index function of ordinal 2-tuples” [70], i.e.,

$$\begin{aligned} \pi : S^* &\rightarrow [0, n] \text{ by} \\ \pi((\alpha s_i, (1 - \alpha)s_{i+1})) &= i + (1 - \alpha). \end{aligned}$$

Under the identification convention, the position index function π becomes a bijection from S^* to $[c_0, c_n]$, and its inverse $\pi^{-1} : [0, n] \rightarrow S^*$ is defined by

$$\pi^{-1}(x) = ((1 - \beta)s_i, \beta s_{i+1}))$$

where $i = E(x)$, E is the integral part function and $\beta = x - i$. So, CCV is a bijection, and its inverse will be denoted by CCV^{-1} .

(4) Proportional 2-tuple aggregation operators

Based on CCV and CCV^{-1} , proportional 2-tuples can be transformed into numerical values, and vice versa, without loss of information. Thus, conventional aggregation operators can be extended easily for proportional 2-tuples.

Let $y = \{y_1, \dots, y_j, \dots, y_m\}$ be a set of proportional 2-tuples, where $y_j = (\alpha s_i, (1 - \alpha)s_{i+1})_j$. $W = \{\omega_1, \dots, \omega_m\}$ is the set of their associated weights. Then, the proportional 2-tuple weighted average $\bar{y}$ is computed by

$$\bar{y} = CCV^{-1} \left(\frac{\sum_{j=1}^m CCV((\alpha s_i, (1 - \alpha)s_{i+1})_j) \cdot \omega_j}{\sum_{j=1}^m \omega_j} \right).$$

In the literature, many other proportional 2-tuple aggregation operators have also been proposed, such as proportional 2-tuple arithmetic mean, proportional 2-tuple ordered weighted average operator. (For more details, see [70].)

2.5 Conclusion

In this chapter, we made a review of several main fuzzy linguistic computational models, including “linguistic computational model based on membership functions, linguistic symbolic computational models based on ordinal scales”, which were further classified into “a linguistic symbolic computational model based on ordinal scales and max-min operators, a linguistic symbolic computational model based on indexes, and a linguistic symbolic computational model

based on continuous term sets”. Meanwhile, we briefly introduced the characteristics of these models.

Particularly, we recalled 2-tuple fuzzy linguistic representation model in detail because of its numerous extensions and extensive applications. In addition, as the inspiration of our own model proposed in the next chapter, we made a comprehensive analysis on proportional 2-tuple fuzzy linguistic representation model.

After this review, it is not difficult to find that all of the above-mentioned linguistic computational models cannot handle ignoring information. In other words, these models are only applicable under the context that all the linguistic assessments are complete. Apparently, it is common that evaluators cannot supply complete linguistic assessments when lack of information or facing with complex nature of decision environments in real world. Therefore, it would be desirable that an appropriate model could be developed to deal with MADM problems with incomplete linguistic information. This is our motivation of extension of Wang and Hao’s “proportional 2-tuple fuzzy linguistic representation model” [70] in Chapter 3.

Chapter 3

A Proportional 3-Tuple Fuzzy Linguistic Representation Model

In this chapter, we develop a proportional 3-tuple fuzzy linguistic representation model for MADM with incomplete linguistic information. Because this model is extended from “proportional 2-tuple fuzzy linguistic representation model” [70], it inherits all the advantages of this model. For instance, when proportional 3-tuple fuzzy linguistic representation model is applied to linguistic decision making, it is not only without loss of information, but is also capable of reflecting evaluators’ confidence levels, which are represented by proportions, indicating their belief degrees that each linguistic term fits a linguistic variable. Meanwhile, there is no requirement that the linguistic labels have to be symmetrically distributed around a medium label and without the traditional requirement of having equal distance between them. Further, based on “canonical characteristic values (*CCV*) of linguistic labels, which are easy to determine by the corresponding semantics of linguistic labels, the computational technique is not only easy and efficient but also takes into account the underlying definitions of the words” [71]. Moreover, by involving a new variable representing the extent of ignoring information, the proposed model can deal with incomplete linguistic information. Thus, evaluators can avoid the dilemma that they have to supply complete linguistic assessments when they face with uncertain information. Based on these features, there are not too many restrictions and requirements for evaluators when they apply proportional 3-tuple fuzzy linguistic representation model to MADM problems, and the precision of final result will be largely improved.

3.1 Proportional 3-Tuple

Let $S = \{s_0, s_1, \dots, s_n\}$ be an ordinal term set with $s_0 < s_1 < \dots < s_n$ (“ $<$ ” represents order relation, i.e., $s_i < s_j$ if and only if $i < j$), $I = [0, 1]$ and

$$IS \equiv I \times S = \{(\alpha, s_i) : \alpha \in [0, 1] \text{ and } i = 0, 1, \dots, n\}.$$

Given a pair (s_i, s_{i+1}) of two successive ordinal terms of S , any two elements $(\alpha, s_i), (\beta, s_{i+1})$ of IS are called a symbolic proportion pair and α, β are called a pair of symbolic proportions of the pair (s_i, s_{i+1}) if $\alpha + \beta \leq 1$. A symbolic proportion pair $(\alpha, s_i), (\beta, s_{i+1})$ will be denoted by

$$\begin{cases} (\alpha s_i, \beta s_{i+1}, 0) & \text{if } \alpha + \beta = 1 \\ (\alpha s_i, \beta s_{i+1}, \varepsilon) & \text{if } \alpha + \beta < 1 \end{cases} \quad (3.1)$$

where ε represents the extent of ignoring information. The set of all the symbolic proportion sequences is denoted by S^* , i.e., $S^* = \{(\alpha s_i, \beta s_{i+1}, \varepsilon) : \alpha, \beta \in [0, 1], \varepsilon = 1 - \alpha - \beta \text{ and } i = 0, 1, \dots, n-1\}$. The set S^* is called the proportional 3-tuple set generated by S and the members of S^* are called proportional 3-tuples, which are designed for representing evaluators’ linguistic assessments. α and β indicate the confidence levels that evaluators believe a linguistic term fits a linguistic variable.

An linguistic assessment $(\alpha s_i, \beta s_{i+1}, \varepsilon)$ is called complete if $\alpha + \beta = 1$, and correspondingly incomplete if $\alpha + \beta < 1$. Because we will apply proportional 3-tuple fuzzy linguistic representation model to a new product project screening problem in Section 3.6, we use the following types of uncertain subjective judgments as examples to explain how to represent a linguistic assessments by proportional 3-tuples. Supposing an evaluator gives his/her linguistic assessments towards criteria “functional competency”, “featured differentia”, and “design quality” as follows:

- 1) The *functional competency* is *best* with a confidence degree of 1.
- 2) The *featured differentia* is evaluated to be *very good* with a confidence degree of 0.6 and to be *best* with a confidence degree of 0.3.

- 3) The *design quality* is evaluated to be *very good* with a confidence degree of 0.4, and to be *best* with a confidence degree of 0.6.

Then, the three linguistic assessments 1)–3) given above can be represented in the form of proportional 3-tuples defined by (3.1) as

$$S^*(functional\ competency) = (0s_5, 1s_6, 0)$$

$$S^*(featured\ differentia) = (0.6s_5, 0.3s_6, 0.1)$$

$$S^*(design\ quality) = (0.4s_5, 0.6s_6, 0)$$

where s_5 and s_6 are linguistic terms of the term set S_1 as shown in (3.16). It is easy to find that the second linguistic assessment is incomplete, while others are complete.

Remark: For $i = 1, 2, \dots, n - 1$, by abuse of notion, the term s_i can use either $(0s_{i-1}, \alpha s_i, \varepsilon)$ or $(\alpha s_i, 0s_{i+1}, \varepsilon)$ as its representation in S^* .

As we know, how to represent and aggregate linguistic information essentially plays an important role in linguistic decision analysis. Therefore, it would be interesting if we consider whether we can use 2-tuple mentioned in “2-tuple fuzzy linguistic representation model” [28] and proportional 2-tuple mentioned in “proportional 2-tuple fuzzy linguistic representation model” [70] to represent the three linguistic assessments. According to their definitions, we cannot represent the linguistic assessments 2) and 3) by 2-tuple because each linguistic assessment includes two linguistic terms. Similarly, we cannot represent the linguistic assessment 2) by proportional 2-tuple because this linguistic assessment is incomplete. If these two models are impossible to represent all the linguistic assessments, let alone using them to deal with the decision making problem with such kinds of linguistic assessments mentioned above. Therefore, we propose proportional 3-tuple, which has obvious advantage to represent the related linguistic assessments in order to solve such problem.

3.2 Canonical Characteristic Values

In Chapter 2, we briefly recalled canonical characteristic values (*CCV*) which was used to represent fuzzy number based semantics of linguistic terms. Considering that a symmetrical triangular fuzzy number $T = [c - \delta, c, c + \delta]$ is used in this chapter, without loss of generality, its expected value, $EV(T) = c$, will be used as a *CCV* of T in this research.

It is worth mentioning that there are two meanings about *CCV*. Besides the first meaning of *CCV*, i.e., using expected value to represent a symmetrical triangular fuzzy number, the other meaning of *CCV* is used as a function to transform a proportional 3-tuple into the form of (h, ε) , where h is a numerical value, and ε represents the extent of ignoring information. In the following section, we will introduce the *CCV* function in detail.

3.3 Computation Operator of Proportional 3-Tuple

MADM problems usually need to unify and aggregate the information. Therefore, for the unification and aggregation of linguistic information represented by proportional 3-tuples, the related computation operator on proportional 3-tuple has to be defined.

Formally, let $S = \{s_0, s_1, \dots, s_n\}$ be an ordinal term set with $s_0 < s_1 < \dots < s_n$, and S^* is the proportional 3-tuple set generated by S . Define *CCV* of proportional 3-tuple $(\alpha s_i, \beta s_{i+1}, \varepsilon)$ as follows:

$$\begin{aligned} CCV((\alpha s_i, \beta s_{i+1}, \varepsilon)) &= (\alpha CCV(s_i) + \beta CCV(s_{i+1}), \varepsilon) \\ &= ((\alpha c_i + (1 - \alpha - \varepsilon)c_{i+1}), \varepsilon) \\ &= (h, \varepsilon) \end{aligned} \tag{3.2}$$

where h is a numerical value, and $h \in [0, 1]$. ε represents the extent of ignoring information. Formula (3.2) is called the corresponding canonical characteristic value function on S^* generated by *CCV* on S . Here, $c_i \in [0, 1]$ with $c_0 < c_1 < \dots < c_n$ is the *CCV* of $s_i, i = 0, 1, \dots, n$.

Proposition: Let $S = \{s_0, s_1, \dots, s_n\}$ be an ordinal term set, S^* is the proportional 3-tuple set generated by S , and (h, ε) is the result obtained by *CCV* of proportional 3-tuple $(\alpha s_i, \beta s_{i+1}, \varepsilon) \in S^*$. Then, there is always a CCV^{-1} function such that from any given (h, ε) it returns to a proportional 3-tuple $(\alpha s_i, \beta s_{i+1}, \varepsilon) \in S^*$ and $CCV(\alpha s_i, \beta s_{i+1}, \varepsilon) = (h, \varepsilon)$.

Proof: Indeed, given (h, ε) , there exists i such that $h \in [c_i, c_{i+1}]$, as shown in Figure 3.1. If (h, ε) is the *CCV* of proportional 3-tuple $(\alpha s_i, \beta s_{i+1}, \varepsilon) \in S^*$ then we have

$$h = \alpha c_i + \beta c_{i+1}.$$

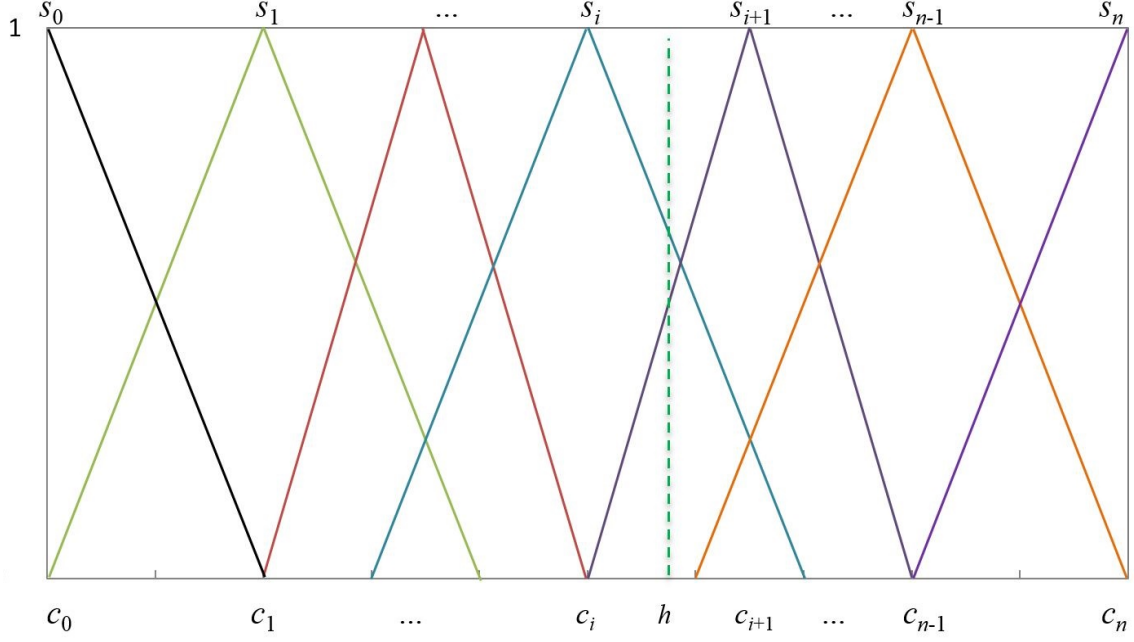


Figure 3.1. The representation of the information h of a proportional 3-tuple

Because $\beta = (1 - \alpha - \varepsilon)$, we get

$$\begin{aligned} h &= \alpha c_i + (1 - \alpha - \varepsilon) c_{i+1} \\ &= (1 - \varepsilon) c_{i+1} - \alpha (c_{i+1} - c_i) \end{aligned} \quad (3.3)$$

and hence

$$\alpha = \frac{(1 - \varepsilon) c_{i+1} - h}{c_{i+1} - c_i}. \quad (3.4)$$

This means that

$$CCV^{-1}(h, \varepsilon) = (\alpha s_i, \beta s_{i+1}, \varepsilon)$$

where α is determined by formula (3.4), and $\beta = (1 - \alpha - \varepsilon)$. This completes the proof of the Proposition.

Thus, a proportional 3-tuple can be transformed into the form of (h, ε) , and vice versa, without loss of information by the use of the functions of CCV and CCV^{-1} .

It is worth mentioning that there is always a CCV^{-1} function to transform (h, ε) back into the original proportional 3-tuple in the same ordinal term set. However, if we use CCV^{-1} function to transform (h, ε) back into a proportional 3-tuple which belongs to a different ordinal term set, the labels and proportions of the the proportional 3-tuple may be correspondingly changed. This is because different linguistic term sets may have different fuzzy number semantics. However, the information h and the extent of ignoring information ε don't change. Based on this feature, we propose a notion of preference-preserving proportional 3-tuple transformation.

3.4 Preference-Preserving Proportional 3-Tuple Transformation

The unification of the linguistic assessments represented by proportional 3-tuples of different linguistic term sets usually need to be carried out before attributes aggregation due to “inhomogeneous nature of different measurement scales/units used for different attributes in the evaluation process” [33]. Therefore, how to unify the linguistic assessments represented by proportional 3-tuples from different linguistic term sets is critically important. For this reason, we define a notion of preference-preserving proportional 3-tuple transformation serving for the unification of proportional 3-tuples between two different linguistic term sets.

Formally, let $S_1 = \{s_0^1, s_1^1, \dots, s_g^1\}$ and $S_2 = \{s_0^2, s_1^2, \dots, s_{g'}^2\}$ be two ordinal linguistic term sets, with $s_0^1 < s_1^1 < \dots < s_g^1$ and $s_0^2 < s_1^2 < \dots < s_{g'}^2$. S_1^* , S_2^* are the ordinal proportional 3-tuple sets generated by S_1 and S_2 respectively. The preference order on S_1 denoted by $<_s$ is either “in agreement with” or “reverse to” the preference order on S_2 , denoted by $\prec_s$. Suppose that we would like to transform the proportional 3-tuples in S_1^* into related proportional 3-tuples in S_2^* . Then, for the case of “in agreement with”, the greater a linguistic value in S_1 , by transformation, the greater a linguistic value will be in S_2 . However, for the case of “reverse to”, the situation is counter, i.e., the greater a linguistic value in S_1 , the smaller a linguistic value will be in S_2 .

In addition, we believe the extent of ignoring information ε of a proportional 3-tuple doesn't change after transformation. In other words, once an evaluator makes a subjective judgment towards a basic attribute, the extent of ignoring information is constant, even though the transformation process has been carried out between two different linguistic term sets. One reasonable explanation is that the subjective judgment including given information and ignoring information supplied by an evaluator has already been an established fact. The established fact

cannot change. The changing parts are labels and associated probabilities due to using different linguistic term sets with different semantics. But the given information h and the extent of ignoring information ε are constant. This prerequisite gives us a guarantee so that the preference-preserving proportional 3-tuple transformation can be carried out. Having these considerations in mind, we can define the preference-preserving proportional 3-tuple transformation. Supposing we would like to transform a proportional 3-tuple in S_1^* into the corresponding proportional 3-tuple in S_2^* , i.e.,

$$\begin{aligned} \wedge : S_1^* &\rightarrow S_2^* \\ (\alpha s_i^1, \beta s_{i+1}^1, \varepsilon) &\mapsto \wedge((\alpha s_i^1, \beta s_{i+1}^1, \varepsilon)) \\ &= (\theta s_j^2, (1 - \varepsilon - \theta) s_{j+1}^2, \varepsilon) \end{aligned} \quad (3.5)$$

with $i \in [0, g - 1]$, $j \in [0, g' - 1]$, $0 < \alpha + \beta \leq 1 - \varepsilon$, $0 \leq \theta \leq 1 - \varepsilon$.

According to formula (3.2), CCV of proportional 3-tuple $(\alpha s_i^1, \beta s_{i+1}^1, \varepsilon)$ in S_1^* is

$$\begin{aligned} CCV((\alpha s_i^1, \beta s_{i+1}^1, \varepsilon)) &= (\alpha CCV(s_i^1) + \beta CCV(s_{i+1}^1), \varepsilon) \\ &= ((\alpha c_i^1 + \beta c_{i+1}^1), \varepsilon) \\ &= (h, \varepsilon). \end{aligned}$$

(1) In the case of $< s$ is in agreement with $\prec s$, i.e., $< s \equiv \prec s$. Define

$$CCV((\theta s_j^2, (1 - \varepsilon - \theta) s_{j+1}^2)) = h \quad (3.6)$$

i.e., CCV of the two proportional 3-tuples both equal to h . Then,

$$\begin{aligned} h &= \theta c_j^2 + (1 - \varepsilon - \theta) c_{j+1}^2 \\ \theta &= \frac{(1 - \varepsilon) c_{j+1}^2 - h}{c_{j+1}^2 - c_j^2}. \end{aligned} \quad (3.7)$$

Because ε doesn't change after transformation, $h \in [0, 1]$, and $0 \leq \theta \leq 1 - \varepsilon$, we can obtain one and only one θ .

(2) In the case of $< s$ is reverse to $\prec s$, i.e., $< s^{-1} \equiv \prec s$. Define a new ordinal linguistic term set $S_t = \{s_0^t, s_1^t, \dots, s_g^t\}$, which has the same semantics with S_1 , and reversed ranking order of linguistic terms with S_1 , i.e., $s_0^t = s_g^1, s_1^t = s_{g-1}^1, \dots, s_g^t = s_0^1$. The linguistic term set S_t is called transition set. The preference order on S_t is denoted by $< s_t$, such that $< s^{-1} \equiv < s_t \equiv \prec s$. S_t^* is the ordinal proportional 3-tuple set generated by S_t . Now, the proportional 3-tuple $(\alpha s_i^1, \beta s_{i+1}^1, \varepsilon)$ in S_1^* can be easily represented by a proportional 3-tuple $(\beta s_{g-i-1}^t, \alpha s_{g-i}^t, \varepsilon)$ in S_t^* . Because of the preference order $< s_t \equiv \prec s$, we can easily transform the proportional 3-tuple $(\beta s_{g-i-1}^t, \alpha s_{g-i}^t, \varepsilon)$ in S_t^* into the corresponding proportional 3-tuple in S_2^* by formula (3.2), (3.6) and (3.7). Thus, the proportional 3-tuple can be transformed between different linguistic term sets without loss of information. Therefore, the preference-preserving proportional 3-tuple transformation can be used as a tool for unification of proportional 3-tuples between different linguistic term sets.

3.5 The Procedure of Proportional 3-Tuple Fuzzy Linguistic Representation Model

Huynh and Nakamori (2011) proposed “a screening evaluation procedure based on 2-tuple linguistic representation” [33], which demonstrated the effectiveness for managers to make their decisions regarding to screening new product development (NPD) projects under uncertainty. Considering its advantage, we combine this evaluation framework with proportional 3-tuples, and propose the procedure of proportional 3-tuple fuzzy linguistic representation model. Since we will apply proportional 3-tuple fuzzy linguistic representation model to a new product project screening problem in Section 3.6, for convenience and consistence, but without loss generality, we will introduce the procedure of proportional 3-tuple fuzzy linguistic representation model combined with this new product project screening problem. Specifically, the procedure of proportional 3-tuple fuzzy linguistic representation model is described as follows:

(1) Proportional 3-tuple linguistic transformation and unification: “This step aims at transforming original linguistic information of a NPD project assessed by evaluators towards a set of basic attributes into a unified representation” [33], i.e., the form of proportional 3-tuples, by using the symbolic translation value s_i and associated representation method, such as 1)–3)

discussed in Section 3.1. It includes converting original linguistic assessments of merit/risk ratings and weights. After converting evaluators' linguistic assessments into corresponding proportional 3-tuples, unification operation should be carried out in order to pave the way for multiple attribute aggregation.

In the problem of screening new product projects, as shown in Section 3.6, because the preference order is counter between merit rating set (represented by S_1 in (3.16)) and risk rating set (represented by S_2 in (3.17)), the proportional 3-tuples in S_1^* and S_2^* should be unified. According to the preference-preserving proportional 3-tuple transformation, the ordinal proportional 3-tuple set S_t^* generated by the linguistic term set S_t as shown in (3.20) is chosen as a transition set used for transforming the proportional 3-tuples in S_2^* , and the ordinal proportional 3-tuple set S_p^* generated by the preference set S_p as shown in (3.21) is chosen as a unification set of proportional 3-tuples. Then, the unification process can be denoted by

$$\wedge : S_1^* \rightarrow S_p^*; S_2^* \rightarrow S_t^* \rightarrow S_p^*$$

where $\wedge$ is preference-preserving proportional 3-tuple transformation.

(2) Aggregating the average important weights and the average preferences of attributes: Because this NPD project employs multi-experts to supply their linguistic assessments including merit/risk ratings and weights to criteria, and the linguistic assessments are represented by means of proportional 3-tuples, the average mechanism for proportional 3-tuples has to be defined. According to conventional arithmetic mean, the computation and aggregation of the average weight and the average preference represented by proportional 3-tuples are defined as follows.

In terms of the weights $(\mu_p \omega_{p,j}, \rho_p \omega_{p,j+1}, \varepsilon'_p)_d$, the average weight $(\mu \omega_j, \rho \omega_{j+1}, \varepsilon'_d)_d$ is given by

$$\begin{aligned} (\mu \omega_j, \rho \omega_{j+1})_d &= CCV^{-1} \left(\sum_{p=1}^q \frac{CCV(\mu_p \omega_{p,j}, \rho_p \omega_{p,j+1})_d}{q} \right) \\ &= CCV^{-1} \left(\sum_{p=1}^q \frac{(\mu_p CCV(\omega_{p,j}) + (1 - \mu_p - \varepsilon'_p) CCV(\omega_{p,j+1}))_d}{q} \right) \end{aligned} \quad (3.8)$$

$$\varepsilon'_d = \sum_{p=1}^q \frac{(\varepsilon'_p)_d}{q} \quad (3.9)$$

with p representing the evaluator, $p \in [1, q]$, d representing the No. d criterion, and ω representing the weights of attributes.

In terms of the preferences $(\alpha_p s_{p,i}, \beta_p s_{p,i+1}, \varepsilon_p)_d$, the average preference $(\alpha s_i, \beta s_{i+1}, \varepsilon)_d$ is given by

$$\begin{aligned} (\alpha s_i, \beta s_{i+1})_d &= CCV^{-1} \left(\sum_{p=1}^q \frac{CCV(\alpha_p s_{p,i}, \beta_p s_{p,i+1})_d}{q} \right) \\ &= CCV^{-1} \left(\sum_{p=1}^q \frac{(\alpha_p CCV(s_{p,i}) + (1 - \alpha_p - \varepsilon_p) CCV(s_{p,i+1}))_d}{q} \right) \end{aligned} \quad (3.10)$$

$$\varepsilon_d = \sum_{p=1}^q \frac{(\varepsilon_p)_d}{q}. \quad (3.11)$$

(3) Computing the overall figure of merit: After obtaining the average preferences and average weights, the overall figure of merit $(\lambda r_t, \eta r_{t+1}, \varepsilon)$ typically expressing the preference regarding the NPD project under consideration is given by

$$(\lambda r_t, \eta r_{t+1}, \varepsilon) = CCV^{-1} \left(\frac{\sum_{d=1}^k CCV(\alpha s_i, \beta s_{i+1})_d \cdot CCV(\mu \omega_j, \rho \omega_{j+1})_d}{\sum_{d=1}^k CCV(\mu \omega_j, \rho \omega_{j+1})_d} \right) \quad (3.12)$$

and the extent of ignoring information can be obtained approximately by

$$\varepsilon = \frac{\sum_{d=1}^k \frac{\varepsilon_d + \varepsilon'_d}{2}}{k} \quad (3.13)$$

with r representing the overall figure of merit and $d \in [1, k]$.

(4) Proportional 3-tuple linguistic conversion: After obtaining the overall figure of merit $(\lambda r_t, \eta r_{t+1}, \varepsilon)$ in S_p^* , convert it into the corresponding proportional 3-tuple in linguistic success

levels of the set S_4^* by the use of the preference-preserve proportional 3-tuple transformation, i.e., $\wedge : S_p^* \rightarrow S_4^*$. Thus, we arrive at the final result, which will be provided to the decision maker as a reference for his/her final screening decision.

It is worth mentioning that we use the formulas (3.12) and (3.13) to compute the the overall figure of merit and related extent of ignoring information because this NPD project screening problem employs multi-expert to supply linguistic weights. Actually, formula (3.12) is a linguistic weight average operator designed for the aggregation of decision making problems with linguistic weights. If the weights are numerical values, the weighted average operator for proportional 3-tuples is give by

$$(\lambda r_t, \eta r_{t+1}, \varepsilon) = CCV^{-1} \left(\frac{\sum_{d=1}^k CCV(\alpha s_i, \beta s_{i+1})_d \cdot \omega_d}{\sum_{d=1}^k \omega_d} \right) \quad (3.14)$$

and the extent of ignoring information can be obtained approximately by

$$\varepsilon = \frac{\sum_{d=1}^k \varepsilon_d}{k} \quad (3.15)$$

with ω_d representing the weight and $d \in [1, k]$.

3.6 An Illustration Example

In this section, we will consider an illustration example of screening new product project taken from [44] in order to explain the practical application of proportional 3-tuple fuzzy linguistic representation model.

3.6.1 The description of new product project screening problem

TV is an internationally recognized CNC machine-tool company, which plans to launch a new product, called TV center-HX, in order to compete in the 21st century. However, with the

limitations imposed by both nature and the timing of new product development, there is ambiguity and uncertainty about technology and the competitive environment. So, this new product project faces with the risk of failure. In such situations, the decision makers want to make a decision whether it is appropriate to launch this new product. Because of the uncertainty, evaluators prefer to make linguistic assessments with confidence levels rather than numerical values towards the selected factors. For further detailed information related to this case, please refer to [44].

3.6.2 Selecting evaluation criteria

New product development project is very complex and cumbersome that are characterized by various of features of both quantitative and qualitative in nature. Selecting a set of criteria that can reflect a variety of features of new products and other indispensable traits is really difficult. Previous researchers have identified criteria for assessing and screening new product projects, which “provide a gauge for companies to assess design approaches and, in turn, select the most suitable design” [3], [31]. By reference to previous literatures, 13 criteria have been selected and categorized into four groups including the factors of competitive marketing advantages, superiority, technological suitability, and the unfavorable factor of risk, as shown in Table 3.1.

3.6.3 Selecting linguistic term sets and associated semantics

After selecting evaluation criteria, linguistic term sets and associated semantics should be defined, which can be used as tools for evaluators to naturally express their subjective judgments against different criteria. One of main approaches in the literature is to adopt and modify the linguistic terms and corresponding membership functions from previous studies so as to satisfy the particular requirements of respective applications. Another often used approach is to directly define a finite and ordinal linguistic term set with associated fuzzy set semantics. The latter one is Lin and Chen [44] used. Specifically, the linguistic term sets and associated fuzzy set semantics for the evaluation of TV center-HX are described as follows.

- (1) The first term set is used to linguistically evaluate the merit ratings of favorable criteria:

Table 3.1. The evaluation criteria of new product project

Criteria	
Competitive marketing advantages (C_1)	Marketing timing (C_{11})
	Price superiority (C_{12})
	Marketing competencies (C_{13})
	Marketing attractiveness (C_{14})
Superiority (C_2)	Functional competency (C_{21})
	Featured differentia (C_{22})
	Design quality (C_{31})
Technological suitability (C_3)	Material specialization (C_{32})
	Manufacturing compatibility (C_{33})
	Supply benefit (C_{34})
	Market competitiveness (C_{41})
Risk (C_4)	Technological uncertainty (C_{42})
	Monetary risk (C_{43})

$$S_1 = \{s_0^1(\text{Worst}), s_1^1(\text{Very Poor}), s_2^1(\text{Poor}), s_3^1(\text{Fair}), s_4^1(\text{Good}), s_5^1(\text{Very Good}), s_6^1(\text{Best})\} \quad (3.16)$$

and the associated fuzzy set semantics is shown in Figure 3.2. The preference order on S_1 is $s_6^1(\text{Best}) \succ s_5^1(\text{Very Good}) \succ \dots \succ s_0^1(\text{Worst})$.

(2) The second term set, which has the reversed preference order compared with other term sets, is used to linguistically evaluate the risk ratings of unfavorable criteria:

$$S_2 = \{s_0^2(\text{Low}), s_1^2(\text{Fairly Low}), s_2^2(\text{Medium}), s_3^2(\text{Fairly High}), s_4^2(\text{High}), s_5^2(\text{Very High}), s_6^2(\text{Extremely High})\} \quad (3.17)$$

and the associated fuzzy set semantics is shown in Figure 3.3. The preference order on S_2 is $s_0^2(\text{Low}) \succ s_1^2(\text{Fairly Low}) \succ \dots \succ s_6^2(\text{Extremely High})$.

(3) The third term set is used to linguistically evaluate the relative importance of criteria:

$$S_3 = \{s_0^3(\text{Very Low}), s_1^3(\text{Low}), s_2^3(\text{Fairly Low}), \\ s_3^3(\text{Fairly High}), s_4^3(\text{High}), s_5^3(\text{Very High})\} \quad (3.18)$$

and the associated fuzzy set semantics is shown in Figure 3.4.

(4) The fourth term set is used to linguistically express the success levels of the new product project:

$$S_4 = \{s_0^4(\text{Very Low}), s_1^4(\text{Low}), s_2^4(\text{Fairly Low}), \\ s_3^4(\text{Fairly High}), s_4^4(\text{High}), s_5^4(\text{Very High})\} \quad (3.19)$$

and the associated fuzzy set semantics is also shown in Figure 3.4. The preference order on S_4 is $s_5^4(\text{Very High}) \succ s_4^4(\text{High}) \succ \dots \succ s_0^4(\text{Very Low})$.

3.6.4 Assessing merit/risk ratings and weights of criteria

Once the criteria have been carefully chosen, linguistic variables and associated membership functions have been elaborately defined, four evaluators denoted by $p = \{E_1, E_2, E_3, E_4\}$ need to give linguistic assessments of merit/risk ratings and weights towards criteria. In this chapter, we first use the original linguistic assessments as in [44] in order to compare the final result with previous models. Then, we abandon the original linguistic assessments and assume more general case of linguistic assessments in order to illustrate the capability of proportional 3-tuple fuzzy linguistic representation model. For the original linguistic assessments, the corresponding proportional 3-tuples are shown in Table 3.2 and Table 3.3.

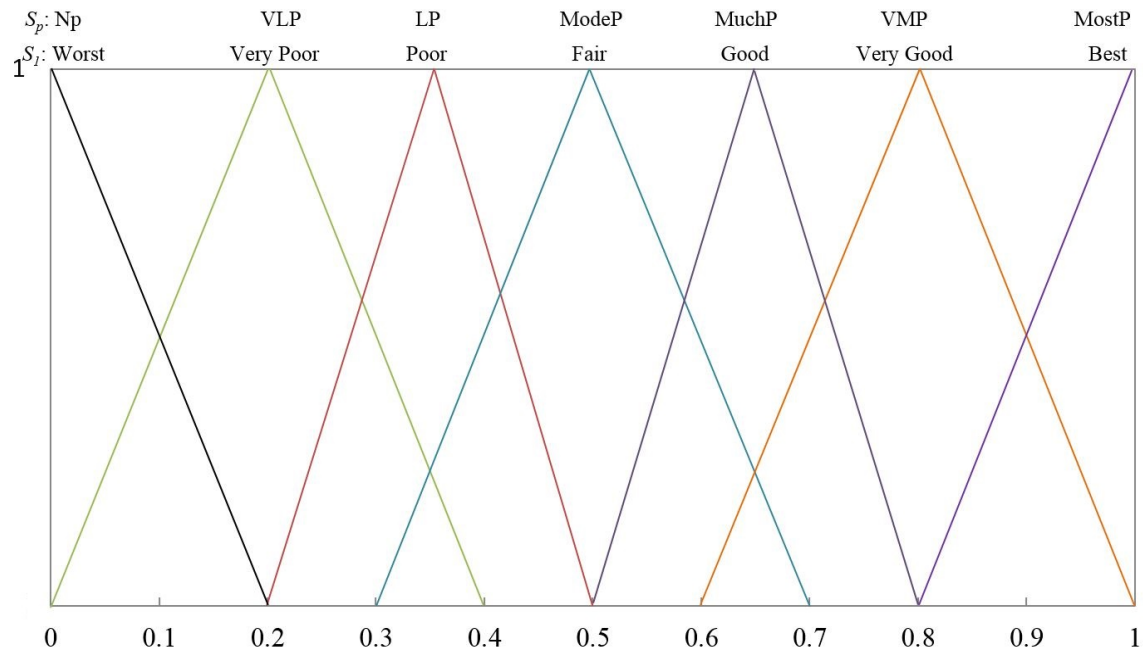


Figure 3.2. Linguistic merit rating values and associated fuzzy number semantic

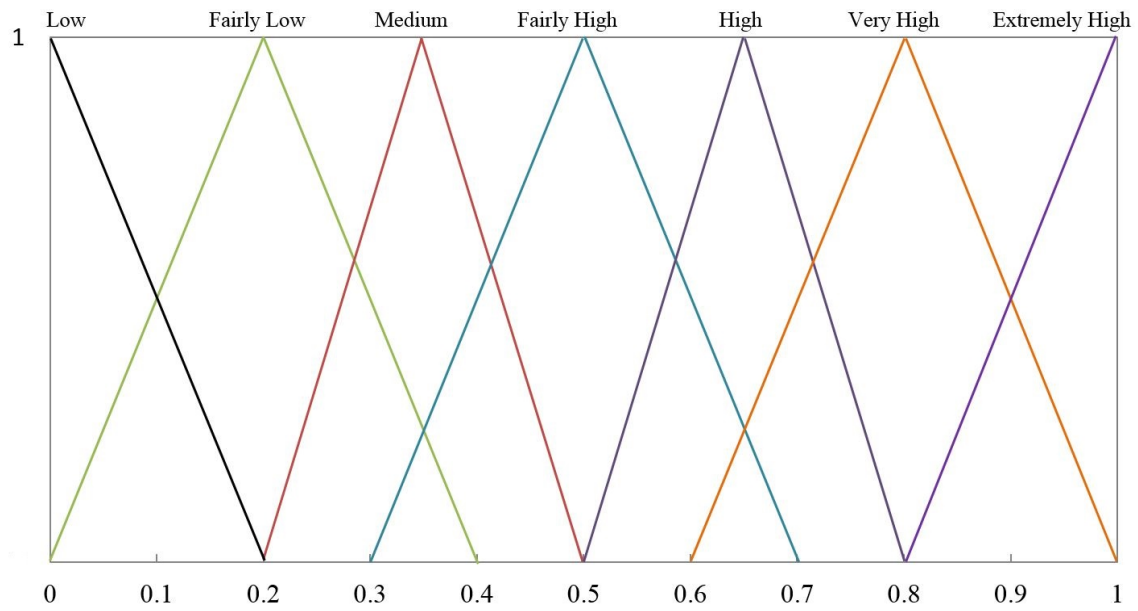


Figure 3.3. Linguistic risk rating values and their associated fuzzy number semantics

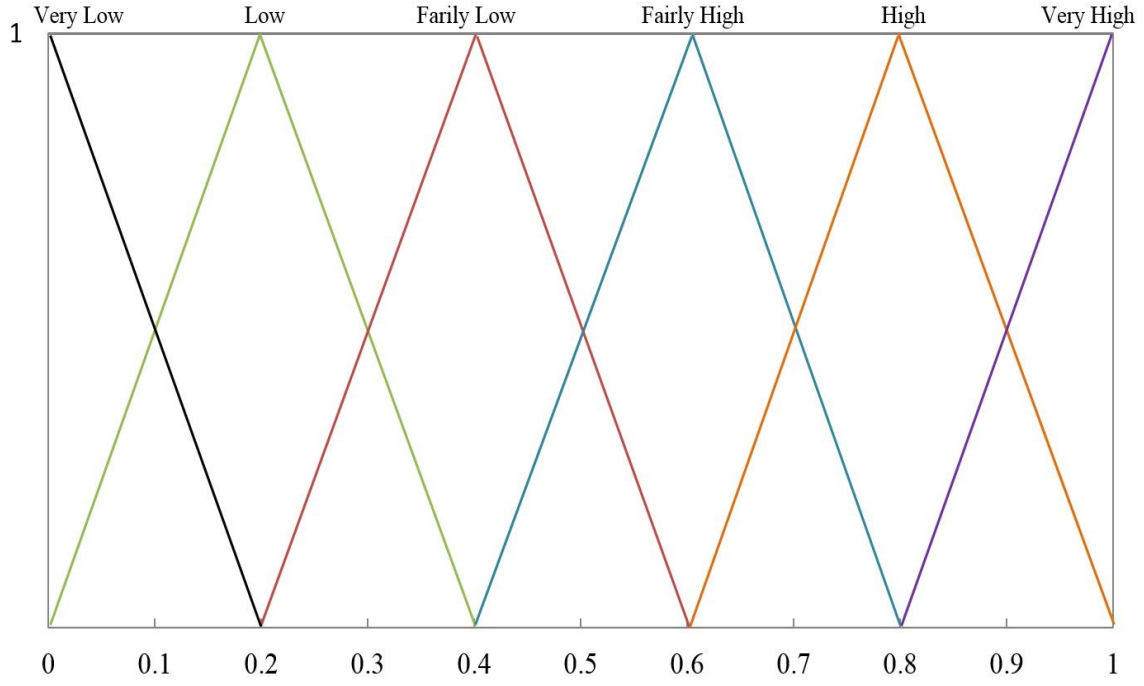


Figure 3.4. Linguistic weights (success levels) and associated fuzzy number semantics

Table 3.2. Original linguistic assessments
of merit/risk ratings of criteria represented by proportional 3-tuples

Criteria	Evaluators			
	E_1	E_2	E_3	E_4
C_{11}	$(1s_4^1, 0s_5^1, 0)$	$(0s_5^1, 1s_6^1, 0)$	$(0s_5^1, 1s_6^1, 0)$	$(0s_4^1, 1s_5^1, 0)$
C_{12}	$(0s_2^1, 1s_3^1, 0)$	$(0s_3^1, 1s_4^1, 0)$	$(1s_2^1, 0s_3^1, 0)$	$(0s_2^1, 1s_3^1, 0)$
C_{13}	$(0s_2^1, 1s_3^1, 0)$	$(1s_2^1, 0s_3^1, 0)$	$(1s_2^1, 0s_3^1, 0)$	$(0s_2^1, 1s_3^1, 0)$
C_{14}	$(1s_5^1, 0s_6^1, 0)$	$(1s_5^1, 0s_6^1, 0)$	$(0s_5^1, 1s_6^1, 0)$	$(0s_5^1, 1s_6^1, 0)$
C_{21}	$(0s_5^1, 1s_6^1, 0)$	$(0s_5^1, 1s_6^1, 0)$	$(1s_5^1, 0s_6^1, 0)$	$(0s_5^1, 1s_6^1, 0)$
C_{22}	$(1s_5^1, 0s_6^1, 0)$	$(1s_5^1, 0s_6^1, 0)$	$(0s_5^1, 1s_6^1, 0)$	$(1s_5^1, 0s_6^1, 0)$
C_{31}	$(1s_5^1, 0s_6^1, 0)$	$(1s_5^1, 0s_6^1, 0)$	$(1s_5^1, 0s_6^1, 0)$	$(0s_5^1, 1s_6^1, 0)$
C_{32}	$(1s_4^1, 0s_5^1, 0)$	$(0s_4^1, 1s_5^1, 0)$	$(1s_4^1, 0s_5^1, 0)$	$(0s_4^1, 1s_5^1, 0)$
C_{33}	$(0s_5^1, 1s_6^1, 0)$	$(1s_5^1, 0s_6^1, 0)$	$(1s_5^1, 0s_6^1, 0)$	$(1s_4^1, 0s_5^1, 0)$
C_{34}	$(1s_3^1, 0s_4^1, 0)$	$(0s_3^1, 1s_4^1, 0)$	$(1s_3^1, 0s_4^1, 0)$	$(0s_3^1, 1s_4^1, 0)$
C_{41}	$(1s_4^2, 0s_5^2, 0)$	$(0s_4^2, 1s_5^2, 0)$	$(1s_4^2, 0s_5^2, 0)$	$(0s_4^2, 1s_5^2, 0)$
C_{42}	$(1s_4^2, 0s_5^2, 0)$	$(0s_4^2, 1s_5^2, 0)$	$(0s_3^2, 1s_4^2, 0)$	$(0s_3^2, 1s_4^2, 0)$
C_{43}	$(1s_2^2, 0s_3^2, 0)$	$(0s_3^2, 1s_4^2, 0)$	$(0s_2^2, 1s_3^2, 0)$	$(1s_2^2, 0s_3^2, 0)$

Table 3.3. Original linguistic assessments
of weights of criteria represented by proportional 3-tuples

Criteria	Evaluators				Average
	E_1	E_2	E_3	E_4	$\bar{E}$
C_{11}	$(0s_4^3, 1s_5^3, 0)$	$(1s_4^3, 0s_5^3, 0)$	$(0s_4^3, 1s_5^3, 0)$	$(0s_4^3, 1s_5^3, 0)$	$(0.25s_4^3, 0.75s_5^3, 0)$
C_{12}	$(1s_2^3, 0s_3^3, 0)$	$(0s_3^3, 1s_4^3, 0)$	$(0s_3^3, 1s_4^3, 0)$	$(1s_3^3, 0s_4^3, 0)$	$(0.75s_3^3, 0.25s_4^3, 0)$
C_{13}	$(0s_4^3, 1s_5^3, 0)$	$(0s_4^3, 1s_5^3, 0)$	$(1s_4^3, 0s_5^3, 0)$	$(1s_4^3, 0s_5^3, 0)$	$(0.5s_4^3, 0.5s_5^3, 0)$
C_{14}	$(1s_4^3, 0s_5^3, 0)$	$(0s_4^3, 1s_5^3, 0)$	$(0s_4^3, 1s_5^3, 0)$	$(0s_4^3, 1s_5^3, 0)$	$(0.25s_4^3, 0.75s_5^3, 0)$
C_{21}	$(0s_4^3, 1s_5^3, 0)$	$(1s_4^3, 0s_5^3, 0)$	$(0s_4^3, 1s_5^3, 0)$	$(1s_4^3, 0s_5^3, 0)$	$(0.5s_4^3, 0.5s_5^3, 0)$
C_{22}	$(0s_2^3, 1s_3^3, 0)$	$(1s_2^3, 0s_3^3, 0)$	$(0s_2^3, 1s_3^3, 0)$	$(0s_2^3, 1s_3^3, 0)$	$(0.25s_2^3, 0.75s_3^3, 0)$
C_{31}	$(1s_4^3, 0s_5^3, 0)$	$(1s_4^3, 0s_5^3, 0)$	$(0s_4^3, 1s_5^3, 0)$	$(0s_4^3, 1s_5^3, 0)$	$(0.5s_4^3, 0.5s_5^3, 0)$
C_{32}	$(0s_3^3, 1s_4^3, 0)$	$(1s_3^3, 0s_4^3, 0)$	$(1s_3^3, 0s_4^3, 0)$	$(1s_2^3, 0s_3^3, 0)$	$(1s_3^3, 0s_4^3, 0)$
C_{33}	$(0s_3^3, 1s_4^3, 0)$	$(1s_3^3, 0s_4^3, 0)$	$(1s_2^3, 0s_3^3, 0)$	$(0s_2^3, 1s_3^3, 0)$	$(0s_2^3, 1s_3^3, 0)$
C_{34}	$(1s_3^3, 0s_4^3, 0)$	$(0s_3^3, 1s_4^3, 0)$	$(1s_3^3, 0s_4^3, 0)$	$(1s_3^3, 0s_4^3, 0)$	$(0.75s_3^3, 0.25s_4^3, 0)$
C_{41}	$(0s_4^3, 1s_5^3, 0)$	$(1s_4^3, 0s_5^3, 0)$	$(0s_4^3, 1s_5^3, 0)$	$(0s_4^3, 1s_5^3, 0)$	$(0.25s_4^3, 0.75s_5^3, 0)$
C_{42}	$(1s_4^3, 0s_5^3, 0)$	$(1s_4^3, 0s_5^3, 0)$	$(0s_4^3, 1s_5^3, 0)$	$(1s_4^3, 0s_5^3, 0)$	$(0.75s_4^3, 0.25s_5^3, 0)$
C_{43}	$(1s_3^3, 0s_4^3, 0)$	$(0s_3^3, 1s_4^3, 0)$	$(0s_2^3, 1s_3^3, 0)$	$(1s_2^3, 0s_3^3, 0)$	$(1s_3^3, 0s_4^3, 0)$

3.6.5 The unification of original linguistic assessments represented by proportional 3-tuples

As mentioned in preceding section, because we used different linguistic term sets for different criteria, the assessment results of merit and risk ratings must be unified in the evaluation process. The seven-term set S_t as shown in (3.20) is selected as transition set for proportional 3-tuples in S_2^* , and its associated fuzzy set semantics is shown in Figure 3.5. The preference order on S_t is $s_6^t(\text{Low}) \succ s_5^t(\text{Fairly Low}) \succ \cdots \succ s_0^t(\text{Extremely High})$. Thus, we can easily transform the proportional 3-tuples of criteria C_{41} , C_{42} and C_{43} in S_2^* into related proportional 3-tuples in S_t^* , which are shown in Table 3.4. “The seven-term set S_p of linguistic preferences as shown in (3.21) is selected for unifying information, and its associated fuzzy set semantics is also shown in Figure 3.2” [33]. The preference order on S_p is $s_6^p(\text{Most Preference}) \succ s_5^p(\text{Very Much Preference}) \succ \cdots \succ s_0^p(\text{No Preference})$. Because the preference orders on S_1 , S_t and S_p are the same, the overall unified information of proportional 3-tuples can be obtained via preference-preserve proportional 3-tuple transformation, and finally is showed in Table 3.5. It is worth noting that the final result of unified information doesn't depend on the granularity of S_p .

Table 3.4. Original linguistic assessments of risk ratings of criteria represented by proportional 3-tuples in transition linguistic term set

Criteria	Evaluators			
	E_1	E_2	E_3	E_4
C_{41}	$(0s_1^t, 1s_2^t, 0)$	$(1s_1^t, 0s_2^t, 0)$	$(0s_1^t, 1s_2^t, 0)$	$(1s_1^t, 0s_2^t, 0)$
C_{42}	$(0s_1^t, 1s_2^t, 0)$	$(1s_1^t, 0s_2^t, 0)$	$(1s_2^t, 0s_3^t, 0)$	$(1s_2^t, 0s_3^t, 0)$
C_{43}	$(0s_3^t, 1s_4^t, 0)$	$(1s_2^t, 0s_3^t, 0)$	$(1s_3^t, 0s_4^t, 0)$	$(0s_3^t, 1s_4^t, 0)$

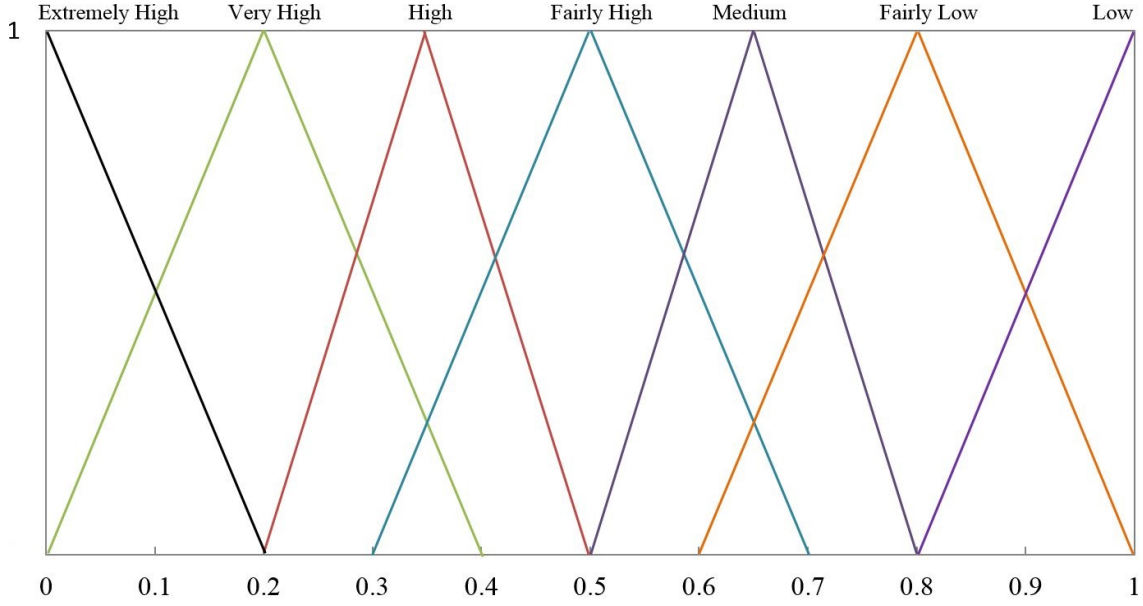


Figure 3.5. Transition linguistic term set and associated fuzzy number semantics

$$S_t = \{s_0^t \text{ (Extremely High)}, s_1^t \text{ (Very High)}, s_2^t \text{ (High)}, s_3^t \text{ (Fairly High)}, s_4^t \text{ (Medium)}, s_5^t \text{ (Fairly Low)}, s_6^t \text{ (Low)}\} \quad (3.20)$$

$$S_p = \{s_0^p \text{ (No Preference)}, s_1^p \text{ (Very Little Preference)}, s_2^p \text{ (Little Preference)}, s_3^p \text{ (Moderate Preference)}, s_4^p \text{ (Much Preference)}, s_5^p \text{ (Very Much Preference)}, s_6^p \text{ (Most Preference)}\} \quad (3.21)$$

Table 3.5. Original linguistic preferences
of criteria represented by proportional 3-tuples

Criteria	Evaluators				Average
	E_1	E_2	E_3	E_4	$\bar{E}$
C_{11}	$(1s_4^p, 0s_5^p, 0)$	$(0s_5^p, 1s_6^p, 0)$	$(0s_5^p, 1s_6^p, 0)$	$(0s_4^p, 1s_5^p, 0)$	$(0.6875s_5^p, 0.3125s_6^p, 0)$
C_{12}	$(0s_2^p, 1s_3^p, 0)$	$(0s_3^p, 1s_4^p, 0)$	$(1s_2^p, 0s_3^p, 0)$	$(0s_2^p, 1s_3^p, 0)$	$(0s_2^p, 1s_3^p, 0)$
C_{13}	$(0s_2^p, 1s_3^p, 0)$	$(1s_2^p, 0s_3^p, 0)$	$(1s_2^p, 0s_3^p, 0)$	$(0s_2^p, 1s_3^p, 0)$	$(0.5s_2^p, 0.5s_3^p, 0)$
C_{14}	$(1s_5^p, 0s_6^p, 0)$	$(1s_5^p, 0s_6^p, 0)$	$(0s_5^p, 1s_6^p, 0)$	$(0s_5^p, 1s_6^p, 0)$	$(0.5s_5^p, 0.5s_6^p, 0)$
C_{21}	$(0s_5^p, 1s_6^p, 0)$	$(0s_5^p, 1s_6^p, 0)$	$(1s_5^p, 0s_6^p, 0)$	$(0s_5^p, 1s_6^p, 0)$	$(0.25s_5^p, 0.75s_6^p, 0)$
C_{22}	$(1s_5^p, 0s_6^p, 0)$	$(1s_5^p, 0s_6^p, 0)$	$(0s_5^p, 1s_6^p, 0)$	$(1s_5^p, 0s_6^p, 0)$	$(0.75s_5^p, 0.25s_6^p, 0)$
C_{31}	$(1s_5^p, 0s_6^p, 0)$	$(1s_5^p, 0s_6^p, 0)$	$(1s_5^p, 0s_6^p, 0)$	$(0s_5^p, 1s_6^p, 0)$	$(0.75s_5^p, 0.25s_6^p, 0)$
C_{32}	$(1s_4^p, 0s_5^p, 0)$	$(0s_4^p, 1s_5^p, 0)$	$(1s_4^p, 0s_5^p, 0)$	$(0s_4^p, 1s_5^p, 0)$	$(0.5s_4^p, 0.5s_5^p, 0)$
C_{33}	$(0s_5^p, 1s_6^p, 0)$	$(1s_5^p, 0s_6^p, 0)$	$(1s_5^p, 0s_6^p, 0)$	$(1s_4^p, 0s_5^p, 0)$	$(0.9375s_5^p, 0.0625s_6^p, 0)$
C_{34}	$(1s_3^p, 0s_4^p, 0)$	$(0s_3^p, 1s_4^p, 0)$	$(1s_3^p, 0s_4^p, 0)$	$(0s_3^p, 1s_4^p, 0)$	$(0.5s_3^p, 0.5s_4^p, 0)$
C_{41}	$(0s_1^p, 1s_2^p, 0)$	$(1s_1^p, 0s_2^p, 0)$	$(0s_1^p, 1s_2^p, 0)$	$(1s_1^p, 0s_2^p, 0)$	$(0.5s_1^p, 0.5s_2^p, 0)$
C_{42}	$(0s_1^p, 1s_2^p, 0)$	$(1s_1^p, 0s_2^p, 0)$	$(1s_2^p, 0s_3^p, 0)$	$(1s_2^p, 0s_3^p, 0)$	$(0.25s_1^p, 0.75s_2^p, 0)$
C_{43}	$(0s_3^p, 1s_4^p, 0)$	$(1s_2^p, 0s_3^p, 0)$	$(1s_3^p, 0s_4^p, 0)$	$(0s_3^p, 1s_4^p, 0)$	$(0.75s_3^p, 0.25s_4^p, 0)$

3.6.6 Computing the evaluation result via proportional 3-tuple fuzzy linguistic representation model

According to the procedure of proportional 3-tuple fuzzy linguistic representation model, the aggregation of proportional 3-tuples should be carried out after information unification. Then, the average important weights and the average preferences as well as the average extent of ignoring information of criteria represented by proportional 3-tuples can be obtained via (3.8) and (3.10), (3.9) and (3.11) respectively, as shown in the last columns of Table 3.3 and Table 3.5. After that, the overall value of preference reflecting the overall figure of merit regarding the new product development project can be obtained by (3.12) and (3.13), i.e.,

$$\begin{aligned}
& (0.945s_4^p, 0.055s_5^p, 0) \\
& = (94.5\% \text{ Much Preference}, 5.5\% \text{ Very Much Preference}, 0)
\end{aligned}$$

which is then converted into the corresponding proportional 3-tuple of linguistic success levels in S_4^* , i.e.,

$$\begin{aligned}
\wedge((0.945s_4^p, 0.055s_5^p, 0)) &= (0.709s_3^4, 0.291s_4^4, 0) \\
&= (70.9\% \text{ Fairly High}, 29.1\% \text{ High}, 0).
\end{aligned}$$

Now, we obtain the final result. This proportional 3-tuple indicates that the possible success level of TV center-HX project is 70.9% fairly high and 29.1% high, which gives the decision makers a reference whether it is suitable to launch this new product project or not.

3.6.7 Comparative study

It would be interesting if we compare the final result with previous models. By using the same linguistic assessments, the final result obtained by “fuzzy-logic-based approach” [44] is a fuzzy number $(0.439, 0.666, 0.852)$, which represents its approximated linguistic expression of $s_3^4 = \text{Fairly High}$. In fact, we can easily find that the associated fuzzy number semantics of s_3^4 is $(0.4, 0.6, 0.8)$, as shown in Figure 3.4. Obviously, there is loss of information when “Fairly High” is as the final result provided to decision makers, and “Fairly High” as the final result is lack of precision. Further, the final result obtained by “2-tuple fuzzy linguistic representation model is a 2-tuple ($s_3^4 = \text{Fairly High}, 0.32$)” [33], which means the possible success level of this new product project is a little more than fairly high. Although there is no loss of information when “2-tuple fuzzy linguistic representation model” [33] was used to deal with this new product project screening problem, there is some vagueness about “0.32” in the final result so that we can only explain it as “a little more than”. Apparently, it increases the vagueness but reduces the comprehension when the 2-tuple ($s_3^4 = \text{Fairly High}, 0.32$) is as the final result provided to decision makers. In contrast, there is no loss of information in the final result obtained by proportional 3-tuple fuzzy linguistic representation model and it indicates much more information which is very comprehensible to decision makers than the obtained results by previous models. Moreover, the final result obtained by proportional 3-tuple fuzzy linguistic representation model provides more guidance to decision makers for their final screening decisions. Besides, the computation process of proportional 3-tuple fuzzy linguistic representation model is much simpler than “fuzzy-logic-based approach” [44], which needs to use other approaches to approximately construct the membership function of final result.

3.6.8 New product project screening problem with revised linguistic assessments

In order to compare the final result with previous models, we used original linguistic assessments in the preceding part. Due to the limitations of previous model, the original linguistic assessments must be complete and can only use one linguistic term to evaluate a criterion. Obviously, this is

Table 3.6. Revised linguistic assessments
of merit/risk ratings of criteria represented by proportional 3-tuples

Criteria	Evaluators			
	E_1	E_2	E_3	E_4
C_{11}	$(0.6s_4^1, 0.3s_5^1, 0.1)$	$(0.2s_5^1, 0.7s_6^1, 0.1)$	$(0.2s_5^1, 0.8s_6^1, 0)$	$(0.4s_4^1, 0.6s_5^1, 0)$
C_{12}	$(0.3s_2^1, 0.7s_3^1, 0)$	$(0.2s_3^1, 0.6s_4^1, 0.2)$	$(0.8s_2^1, 0.1s_3^1, 0.1)$	$(0.4s_2^1, 0.5s_3^1, 0.1)$
C_{13}	$(0s_2^1, 1s_3^1, 0)$	$(0.7s_2^1, 0.2s_3^1, 0.1)$	$(1s_2^1, 0s_3^1, 0)$	$(0.3s_2^1, 0.6s_3^1, 0.1)$
C_{14}	$(0.6s_5^1, 0.4s_6^1, 0)$	$(0.6s_5^1, 0.2s_6^1, 0.2)$	$(0.2s_5^1, 0.7s_6^1, 0.1)$	$(0.4s_5^1, 0.5s_6^1, 0.1)$
C_{21}	$(0.3s_5^1, 0.6s_6^1, 0.1)$	$(0.2s_5^1, 0.8s_6^1, 0)$	$(0.7s_5^1, 0.2s_6^1, 0.1)$	$(0s_5^1, 1s_6^1, 0)$
C_{22}	$(0.7s_5^1, 0.2s_6^1, 0.1)$	$(0.5s_5^1, 0.3s_6^1, 0.2)$	$(0.2s_5^1, 0.8s_6^1, 0)$	$(0.6s_5^1, 0.3s_6^1, 0.1)$
C_{31}	$(0.8s_5^1, 0.1s_6^1, 0.1)$	$(0.8s_5^1, 0.2s_6^1, 0)$	$(0.6s_5^1, 0.3s_6^1, 0.1)$	$(0.4s_5^1, 0.6s_6^1, 0)$
C_{32}	$(0.7s_4^1, 0.2s_5^1, 0.1)$	$(0.3s_4^1, 0.6s_5^1, 0.1)$	$(0.8s_4^1, 0.2s_5^1, 0)$	$(0.4s_4^1, 0.6s_5^1, 0)$
C_{33}	$(0.2s_5^1, 0.7s_6^1, 0.1)$	$(0.5s_5^1, 0.4s_6^1, 0.1)$	$(0.7s_5^1, 0.2s_6^1, 0.1)$	$(0.6s_4^1, 0.3s_5^1, 0.1)$
C_{34}	$(0.6s_3^1, 0.3s_4^1, 0.1)$	$(0.2s_3^1, 0.7s_4^1, 0.1)$	$(0.7s_3^1, 0.2s_4^1, 0.1)$	$(0.3s_3^1, 0.6s_4^1, 0.1)$
C_{41}	$(0.8s_4^2, 0.2s_5^2, 0)$	$(0.3s_4^2, 0.6s_5^2, 0.1)$	$(0.6s_4^2, 0.3s_5^2, 0.1)$	$(0.2s_4^2, 0.8s_5^2, 0)$
C_{42}	$(0.8s_4^2, 0.1s_5^2, 0.1)$	$(0.3s_4^2, 0.6s_5^2, 0.1)$	$(0.3s_3^2, 0.6s_4^2, 0.1)$	$(0.2s_3^2, 0.7s_4^2, 0.1)$
C_{43}	$(0.7s_2^2, 0.2s_3^2, 0.1)$	$(0.4s_3^2, 0.5s_4^2, 0.1)$	$(0.3s_2^2, 0.7s_3^2, 0)$	$(0.6s_2^2, 0.4s_3^2, 0)$

not reasonable. Because the nature of human judgments on uncertainty response a basic bias with probability, and sometimes evaluators cannot supply complete linguistic assessments, it is necessary and reasonable to modify the original linguistic assessments by allowing evaluators to supply more general case of linguistic assessments as discussed in Section 3.1. Besides, the revised linguistic assessments can better reflect the capability of proportional 3-tuple fuzzy linguistic representation model for dealing with MADM problems with incomplete linguistic information.

It is worth mentioning that we keep the information of original data as much as possible during the modification. For example, in [44], the evaluator E_1 supplied s_4^1 as the linguistic assessment for criterion C_{11} . We then correspondingly modify it into proportional 3-tuple as $(0.6s_4^1, 0.3s_5^1, 0.1)$, which not only keeps the most information of original assessment, i.e., s_4^1 , but also indicates the extent of ignoring information, i.e., 0.1. With this principle, we correspondingly modify evaluators' final linguistic assessment results of merit/risk ratings and important weights according to the original data in [44], which are shown in Table 3.6 and Table 3.7 respectively.

Table 3.7. Revised linguistic assessments
of weights of criteria represented by proportional 3-tuples

Criteria	Evaluators				Average
	E_1	E_2	E_3	E_4	$\bar{E}$
C_{11}	$(0.4s_4^3, 0.5s_5^3, 0.1)$	$(0.8s_4^3, 0.2s_5^3, 0)$	$(0.1s_4^3, 0.9s_5^3, 0)$	$(0.3s_4^3, 0.6s_5^3, 0.1)$	$(0.4s_4^3, 0.55s_5^3, 0.05)$
C_{12}	$(0.6s_2^3, 0.4s_3^3, 0)$	$(0.3s_3^3, 0.6s_4^3, 0.1)$	$(0.2s_3^3, 0.8s_4^3, 0)$	$(0.5s_3^3, 0.4s_4^3, 0.1)$	$(0.65s_3^3, 0.3s_4^3, 0.05)$
C_{13}	$(0.3s_4^3, 0.6s_5^3, 0.1)$	$(0.3s_4^3, 0.7s_5^3, 0)$	$(0.8s_4^3, 0.2s_5^3, 0)$	$(0.7s_4^3, 0.2s_5^3, 0.1)$	$(0.525s_4^3, 0.425s_5^3, 0.05)$
C_{14}	$(0.6s_4^3, 0.4s_5^3, 0)$	$(0.3s_4^3, 0.7s_5^3, 0)$	$(0.2s_4^3, 0.7s_5^3, 0.1)$	$(0.3s_4^3, 0.6s_5^3, 0.1)$	$(0.35s_4^3, 0.6s_5^3, 0.05)$
C_{21}	$(0.3s_4^3, 0.6s_5^3, 0.1)$	$(0.5s_4^3, 0.4s_5^3, 0.1)$	$(0s_4^3, 1s_5^3, 0)$	$(0.6s_4^3, 0.2s_5^3, 0.2)$	$(0.35s_4^3, 0.55s_5^3, 0.1)$
C_{22}	$(0.3s_2^3, 0.6s_3^3, 0.1)$	$(0.6s_2^3, 0.3s_3^3, 0.1)$	$(0.2s_2^3, 0.8s_3^3, 0)$	$(0.4s_2^3, 0.6s_3^3, 0)$	$(0.375s_2^3, 0.575s_3^3, 0.05)$
C_{31}	$(0.7s_4^3, 0.2s_5^3, 0.1)$	$(0.8s_4^3, 0.2s_5^3, 0)$	$(0.2s_4^3, 0.7s_5^3, 0.1)$	$(0.1s_4^3, 0.7s_5^3, 0.2)$	$(0.45s_4^3, 0.45s_5^3, 0.1)$
C_{32}	$(0.3s_3^3, 0.6s_4^3, 0.1)$	$(0.7s_3^3, 0.3s_4^3, 0)$	$(0.7s_3^3, 0.2s_4^3, 0.1)$	$(0.6s_3^3, 0.4s_4^3, 0)$	$(0.825s_3^3, 0.125s_4^3, 0.05)$
C_{33}	$(0.2s_3^3, 0.7s_4^3, 0.1)$	$(0.6s_3^3, 0.3s_4^3, 0.1)$	$(0.7s_2^3, 0.3s_3^3, 0)$	$(0.4s_2^3, 0.6s_3^3, 0)$	$(0.025s_2^3, 0.925s_3^3, 0.05)$
C_{34}	$(0.7s_3^3, 0.2s_4^3, 0.1)$	$(0.3s_3^3, 0.6s_4^3, 0.1)$	$(0.8s_3^3, 0.2s_4^3, 0)$	$(0.6s_3^3, 0.4s_4^3, 0)$	$(0.6s_3^3, 0.35s_4^3, 0.05)$
C_{41}	$(0.2s_4^3, 0.8s_5^3, 0)$	$(0.6s_4^3, 0.3s_5^3, 0.1)$	$(0.2s_4^3, 0.8s_5^3, 0)$	$(0.4s_4^3, 0.5s_5^3, 0.1)$	$(0.35s_4^3, 0.65s_5^3, 0.05)$
C_{42}	$(0.7s_4^3, 0.3s_5^3, 0)$	$(0.7s_4^3, 0.3s_5^3, 0)$	$(0.2s_4^3, 0.7s_5^3, 0.1)$	$(0.6s_4^3, 0.3s_5^3, 0.1)$	$(0.55s_4^3, 0.4s_5^3, 0.05)$
C_{43}	$(0.6s_3^3, 0.4s_4^3, 0)$	$(0.2s_3^3, 0.8s_4^3, 0)$	$(0.1s_2^3, 0.7s_3^3, 0.2)$	$(0.6s_2^3, 0.2s_3^3, 0.2)$	$(0.775s_3^3, 0.125s_4^3, 0.1)$

Table 3.8. Revised linguistic assessments of risk ratings of criteria
represented by proportional 3-tuples in transition linguistic term set

Criteria	Evaluators			
	E_1	E_2	E_3	E_4
C_{41}	$(0.2s_1^t, 0.8s_2^t, 0)$	$(0.6s_1^t, 0.3s_2^t, 0.1)$	$(0.3s_1^t, 0.6s_2^t, 0.1)$	$(0.8s_1^t, 0.2s_2^t, 0)$
C_{42}	$(0.1s_1^t, 0.8s_2^t, 0.1)$	$(0.6s_1^t, 0.3s_2^t, 0.1)$	$(0.6s_2^t, 0.3s_3^t, 0.1)$	$(0.7s_2^t, 0.2s_3^t, 0.1)$
C_{43}	$(0.2s_3^t, 0.7s_4^t, 0.1)$	$(0.5s_2^t, 0.4s_3^t, 0.1)$	$(0.7s_3^t, 0.3s_4^t, 0)$	$(0.4s_3^t, 0.6s_4^t, 0)$

3.6.9 The unification of the revised linguistic assessments represented by proportional 3-tuples

Similarly, we can easily transform the proportional 3-tuples of criteria C_{41} , C_{42} and C_{43} in Table 3.6 into related proportional 3-tuples in S_t^* , which are shown in Table 3.8. Then, the overall unified information of proportional 3-tuples can be obtained via preference-preserve proportional 3-tuple transformation, and is finally showed in Table 3.9.

Table 3.9. Revised linguistic preferences
of criteria represented by proportional 3-tuples

Criteria	Evaluators				Average
	E_1	E_2	E_3	E_4	$\bar{E}$
C_{11}	$(0.6s_4^p, 0.3s_5^p, 0.1)$	$(0.2s_5^p, 0.7s_6^p, 0.1)$	$(0.2s_5^p, 0.8s_6^p, 0)$	$(0.4s_4^p, 0.6s_5^p, 0)$	$(0.7625s_5^p, 0.1875s_6^p, 0.05)$
C_{12}	$(0.3s_2^p, 0.7s_3^p, 0)$	$(0.2s_3^p, 0.6s_4^p, 0.2)$	$(0.8s_2^p, 0.1s_3^p, 0.1)$	$(0.4s_2^p, 0.5s_3^p, 0.1)$	$(0.225s_2^p, 0.675s_3^p, 0.1)$
C_{13}	$(0s_2^p, 1s_3^p, 0)$	$(0.7s_2^p, 0.2s_3^p, 0.1)$	$(1s_2^p, 0s_3^p, 0)$	$(0.3s_2^p, 0.6s_3^p, 0.1)$	$(0.5s_2^p, 0.45s_3^p, 0.05)$
C_{14}	$(0.6s_5^p, 0.4s_6^p, 0)$	$(0.6s_5^p, 0.2s_6^p, 0.2)$	$(0.2s_5^p, 0.7s_6^p, 0.1)$	$(0.4s_5^p, 0.5s_6^p, 0.1)$	$(0.45s_5^p, 0.45s_6^p, 0.1)$
C_{21}	$(0.3s_5^p, 0.6s_6^p, 0.1)$	$(0.2s_5^p, 0.8s_6^p, 0)$	$(0.7s_5^p, 0.2s_6^p, 0.1)$	$(0s_5^p, 1s_6^p, 0)$	$(0.3s_5^p, 0.65s_6^p, 0.05)$
C_{22}	$(0.7s_5^p, 0.2s_6^p, 0.1)$	$(0.5s_5^p, 0.3s_6^p, 0.2)$	$(0.2s_5^p, 0.8s_6^p, 0)$	$(0.6s_5^p, 0.3s_6^p, 0.1)$	$(0.5s_5^p, 0.4s_6^p, 0.1)$
C_{31}	$(0.8s_5^p, 0.1s_6^p, 0.1)$	$(0.8s_5^p, 0.2s_6^p, 0)$	$(0.6s_5^p, 0.3s_6^p, 0.1)$	$(0.4s_5^p, 0.6s_6^p, 0)$	$(0.65s_5^p, 0.3s_6^p, 0.05)$
C_{32}	$(0.7s_4^p, 0.2s_5^p, 0.1)$	$(0.3s_4^p, 0.6s_5^p, 0.1)$	$(0.8s_4^p, 0.2s_5^p, 0)$	$(0.4s_4^p, 0.6s_5^p, 0)$	$(0.55s_4^p, 0.4s_5^p, 0.05)$
C_{33}	$(0.2s_5^p, 0.7s_6^p, 0.1)$	$(0.5s_5^p, 0.4s_6^p, 0.1)$	$(0.7s_5^p, 0.2s_6^p, 0.1)$	$(0.6s_4^p, 0.3s_5^p, 0.1)$	$(0.6875s_5^p, 0.2125s_6^p, 0.1)$
C_{34}	$(0.6s_3^p, 0.3s_4^p, 0.1)$	$(0.2s_3^p, 0.7s_4^p, 0.1)$	$(0.7s_3^p, 0.2s_4^p, 0.1)$	$(0.3s_3^p, 0.6s_4^p, 0.1)$	$(0.45s_3^p, 0.45s_4^p, 0.1)$
C_{41}	$(0.2s_1^p, 0.8s_2^p, 0)$	$(0.6s_1^p, 0.3s_2^p, 0.1)$	$(0.3s_1^p, 0.6s_2^p, 0.1)$	$(0.8s_1^p, 0.2s_2^p, 0)$	$(0.475s_1^p, 0.475s_2^p, 0.05)$
C_{42}	$(0.1s_1^p, 0.8s_2^p, 0.1)$	$(0.6s_1^p, 0.3s_2^p, 0.1)$	$(0.6s_2^p, 0.3s_3^p, 0.1)$	$(0.7s_2^p, 0.2s_3^p, 0.1)$	$(0.05s_1^p, 0.85s_2^p, 0.1)$
C_{43}	$(0.2s_3^p, 0.7s_4^p, 0.1)$	$(0.5s_2^p, 0.4s_3^p, 0.1)$	$(0.7s_3^p, 0.3s_4^p, 0)$	$(0.4s_3^p, 0.6s_4^p, 0)$	$(0.675s_3^p, 0.275s_4^p, 0.05)$

3.6.10 The evaluation result of revised linguistic assessments

After information unification, the average revised important weights and the average revised preferences as well as the average extent of ignoring information of criteria represented by proportional 3-tuples can be obtained easily and are shown in the last columns of Table 3.7 and Table 3.9. Then, the overall value of preference of the new product development project can be obtained by (3.12) and (3.13), i.e.,

$$\begin{aligned}
& (0.915s_4^p, 0.018s_5^p, 0.067) \\
& = (91.5\% \text{ Much Preference}, 1.8\% \text{ Very Much Preference}, 6.7\%)
\end{aligned}$$

which is then converted into the related proportional 3-tuple of linguistic success levels in S_4^* , i.e.,

$$\begin{aligned}
\wedge((0.915s_4^p, 0.018s_5^p, 0.067)) &= (0.687s_3^4, 0.246s_4^4, 0.067) \\
&= (68.7\% \text{ Fairly High}, 24.6\% \text{ High}, 6.7\%).
\end{aligned}$$

This is the final result of the revised linguistic assessments. This proportional 3-tuple indicates that the possible success level of TV center-HX project is 68.7% fairly high, 24.6% high, and 6.7%

ignoring information, which gives the decision makers a reference whether it is suitable to launch this new product project or not.

It is obvious proportional 3-tuple fuzzy linguistic representation model is very flexible to deal with MADM problems with incomplete linguistic information. However, the previous models cannot be used in such situations, such as “fuzzy-logic-based approach” [44], “2-tuple fuzzy linguistic representation model” [28], and “proportional 2-tuple fuzzy linguistic representation model” [70].

3.7 Conclusion

In this chapter, we introduced a proportional 3-tuple fuzzy linguistic representation model for MADM problems with incomplete linguistic information. In this model, we first defined a notion of proportional 3-tuple, by which evaluators could express their linguistic assessments with confidence levels, and meanwhile, evaluators could directly supply incomplete assessments when the complete linguistic information was not acceptable. Then, we proposed a proportional 3-tuple computation operator and a so-called preference-preserving proportional 3-tuple transformation based on *CCV* so as to unify different information represented by proportional 3-tuples.

Without loss of generality, we put forward a proportional 3-tuple fuzzy linguistic representation model combined with a new product project screening problem, which was taken from previous literature. After application, we can find the special features of this model, such as no loss of information during the evaluation process, ease operation in the complicated linguistic context, flexible operation space for evaluators under uncertainty, taking the ignoring information into account and so on. Based on these features, the proportional 3-tuple fuzzy linguistic representation model can be applied to much more complicated MADM problems, while previous models cannot deal with such problems. In addition, the final result obtained by proportional 3-tuple fuzzy linguistic representation model provides much more information which could give decision makers a more comprehensive guidance than that of previous models.

Chapter 4

A Proportional Fuzzy Linguistic Distribution Model

In this chapter, we introduce a proportional fuzzy linguistic distribution model for MADM problem with incomplete linguistic information. In this model, we use proportions as evaluators' confidence levels indicating their belief degrees that each linguistic term fits a linguistic variable. Further, "since the uncertainty may be assigned not only to any single evaluation grade but also to their rational combinations" [85], "each attribute can be directly evaluated using subjective judgments with the uncertainty being assigned to any number of adjacent single evaluation grades". Thus, when this model is employed in linguistic information, there is no obligatory requirement that evaluators have to supply the subjective judgments which are constituted by fixed number of linguistic terms. Moreover, with introducing a variable representing the extent of ignoring information, the proportional fuzzy linguistic distribution model is capable of dealing with incomplete linguistic assessments. Hence, facing with uncertain and incomplete information, evaluators can directly supply incomplete linguistic assessments. Due to these special features, evaluators can flexibly express their linguistic assessments without too many restrictions, which leave much operation space for evaluators to handle uncertainty.

4.1 Proportional Fuzzy Linguistic Distribution

Let $S = \{s_0, s_1, \dots, s_n\}$ be an ordinal term set with $s_0 < s_1 < \dots < s_n$ (" $<$ " represents order relation, i.e., $s_i < s_j$ if and only if $i < j$), $I = [0, 1]$ and

$$IS \equiv I \times S = \{(\alpha, s_i) : \alpha \in [0, 1] \text{ and } i = 0, 1, \dots, n\}.$$

Given a sequence $(s_i, s_{i+1}, \dots, s_{i+m})$ of $(m+1)$ successive ordinal terms of S , any $(m+1)$ elements $(\alpha_i, s_i), (\alpha_{i+1}, s_{i+1}), \dots, (\alpha_{i+m}, s_{i+m})$ of IS are called a symbolic proportion sequence, and it will be denoted by

$$\begin{cases} (\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, 0) & \text{if } \sum_{j=i}^{i+m} \alpha_j = 1 \\ (\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, \varepsilon) & \text{if } \sum_{j=i}^{i+m} \alpha_j < 1 \end{cases} \quad (4.1)$$

where ε represents the extent of ignoring information. The set of all the symbolic proportion sequences is denoted by S^* , i.e., $S^* = \{(\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, \varepsilon) : \alpha_i \in (0, 1], \alpha_{i+m} \in (0, 1], 0 < \sum_{j=i}^{i+m} \alpha_j \leq 1, \varepsilon = 1 - \sum_{j=i}^{i+m} \alpha_j, 0 \leq i, \text{ and } i+m \leq n\}$. The set S^* is called proportional fuzzy linguistic distribution set generated by S and the members of S^* are called proportional fuzzy linguistic distributions.

In the sequel, a proportional fuzzy linguistic distribution $(\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, \varepsilon)$ will be used to represent an evaluator's subjective judgment. Here, i is called the starting label; s_i is the No. i linguistic term; α_j is the proportional coefficient in front of the related linguistic term. It represents the confidence levels that to which degree the evaluator believes a linguistic term fits a linguistic variable. Similarly, $i+m$ is called the ending label.

An assessment $(\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, \varepsilon)$ is called complete if $\sum_{j=i}^{i+m} \alpha_j = 1$, or $\varepsilon = 0$, and incomplete if $\sum_{j=i}^{i+m} \alpha_j < 1$, or $\varepsilon > 0$. For example, we will use proportional fuzzy linguistic distribution model to evaluate the performance of four types of motorcycles [87] in this chapter. Then, for the attributes, “responsiveness”, “fuel economy”, “maneuverability”, “gearbox operation”, the following types of uncertain subjective judgments of a motorcycle, for example, “Yamaha”, are frequently used.

- 1) The *responsiveness* of engine is evaluated to be *good* with a confidence degree of 0.3, and to be *excellent* with a confidence degree of 0.6.
- 2) The *fuel economy* of engine is evaluated to be *indifferent* with a confidence degree of 1.
- 3) The *maneuverability* of handling is *excellent* with a confidence degree of 0.9.

- 4) The *gearbox operation* of transmission is evaluated to be *indifferent* with a confidence degree of 0.5, and to be *average* with a confidence degree of 0.5.

The four assessments 1)–4) given in the above can be represented in the form of proportional fuzzy linguistic distributions defined by (4.1) as

$$S^*(responsiveness) = (0.3s_3, 0.6s_4, 0.1)$$

$$S^*(fuel\ economy) = (1s_1, 0)$$

$$S^*(maneuverability) = (0.9s_4, 0.1)$$

$$S^*(gearbox\ operation) = (0.5s_1, 0.5s_2, 0)$$

where s_i with $i = 1, 2, 3$ and 4 are the linguistic terms of the set of evaluation grades S_1 as shown in (4.14), and the linguistic assessments 2) and 4) are complete, while the linguistic assessments 1) and 3) are incomplete.

It is worth mentioning that we quoted original data from [87] in order to compare the final result obtained by proportional fuzzy linguistic distribution model with that obtained by evidential reasoning approach. Hence, the above four linguistic assessments used two linguistic terms at most to describe a linguistic variable (e.g. the linguistic assessments 1) and 4)). As a matter of fact, evaluators can use the combination of any number of linguistic terms to describe a linguistic variable if they believe it is reasonable and appropriate to capture the uncertainty in some special situations. Let's still take Yamaha as an example. Supposing that the evaluator assesses the fuel economy of Yamaha based on mountain road, ordinary road and highway, the following uncertain statement perhaps can be considered to use.

The *fuel economy* of Yamaha is evaluated to be *indifferent* on mountain roads, to be *average* on ordinary roads, to be *good* on highways, and unclear on the other kinds of roads. Then the uncertain subjective judgment can be represented in the form of proportional fuzzy linguistic distribution as:

$$S^*(fuel\ economy) = (0.25s_1, 0.25s_2, 0.25s_3, 0.25).$$

4.2 The Comparison of Proportional Fuzzy Linguistic Distributions

Let $S = \{s_0, s_1, \dots, s_n\}$ be an linguistic term set and S^* be the proportional fuzzy linguistic distribution set generated by S . For any $(\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, \varepsilon)_1$, $(\beta_g s_g, \beta_{g+1} s_{g+1}, \dots, \beta_{g+f} s_{g+f}, \varepsilon)_2 \in S^*$, the comparison of proportional fuzzy linguistic distributions is described as follows.

(1) If $\varepsilon_1 = 0$ and $\varepsilon_2 = 0$, define $(\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, \varepsilon)_1 < (\beta_g s_g, \beta_{g+1} s_{g+1}, \dots, \beta_{g+f} s_{g+f}, \varepsilon)_2$

$$\begin{aligned} &\Leftrightarrow \alpha_i \cdot i + \alpha_{i+1} \cdot (i+1) + \dots + \alpha_{i+m} \cdot (i+m) \\ &< \beta_g \cdot g + \beta_{g+1} \cdot (g+1) + \dots + \beta_{g+f} \cdot (g+f) \\ &\Leftrightarrow \sum_{j=i}^{i+m} (\alpha_j \cdot j) < \sum_{k=g}^{g+f} (\beta_k \cdot k). \end{aligned} \tag{4.2}$$

(2) If $\varepsilon_1 = 0$ and $\varepsilon_2 \neq 0$, the latter will generate an interval value (φ, ψ) because it includes ignoring information. Thus, we need to allocate ε_2 in order to obtain the minimum value φ and maximum value ψ .

For the minimum value φ , we can consider an extreme situation that ε_2 is allocated to s_0 completely, i.e.,

$$\begin{aligned} \varphi &= \beta_g \cdot g + \beta_{g+1} \cdot (g+1) + \dots + \beta_{g+f} \cdot (g+f) + \varepsilon_2 \cdot 0 \\ &= \sum_{k=g}^{g+f} (\beta_k \cdot k). \end{aligned} \tag{4.3}$$

For the maximum value ψ , we can consider an extreme situation that ε_2 is allocated to s_n completely, i.e.,

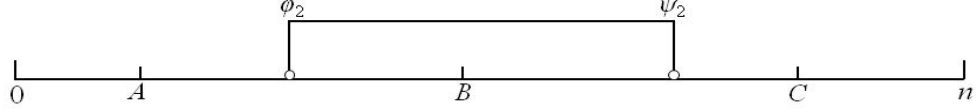


Figure 4.1. The relationship between a complete and an incomplete proportional fuzzy linguistic distribution

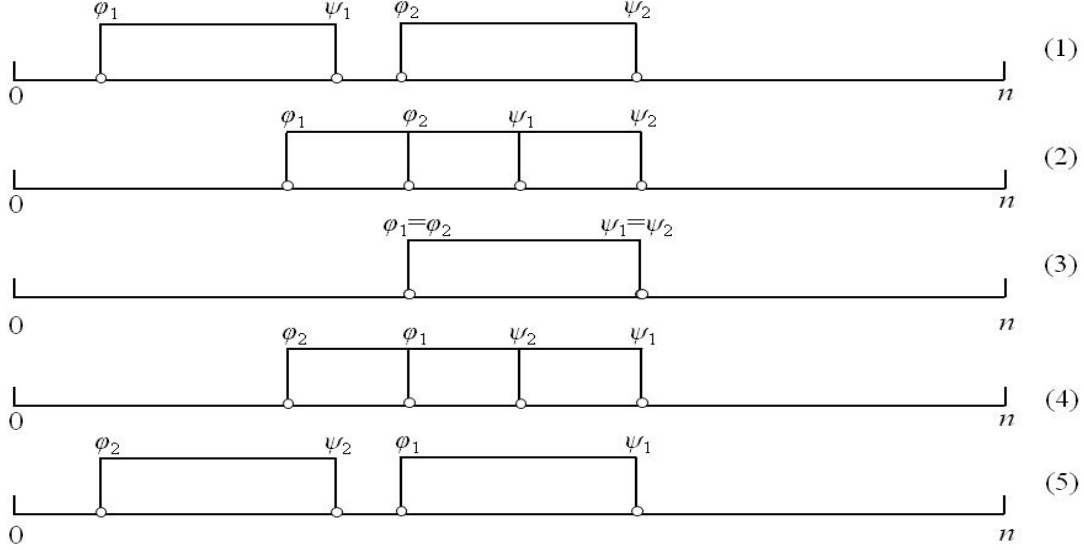


Figure 4.2. The relationship between two incomplete proportional fuzzy linguistic distributions

$$\begin{aligned}
 \psi &= \beta_g \cdot g + \beta_{g+1} \cdot (g+1) + \cdots + \beta_{g+f} \cdot (g+f) + \varepsilon_2 \cdot n \\
 &= \sum_{k=g}^{g+f} (\beta_k \cdot k) + \varepsilon_2 \cdot n.
 \end{aligned} \tag{4.4}$$

Then, the interval value can be represented as $(\sum_{k=g}^{g+f} (\beta_k \cdot k), \sum_{k=g}^{g+f} (\beta_k \cdot k) + \varepsilon_2 \cdot n)$, and the relationship between the two proportional fuzzy linguistic distributions can be described in Figure 4.1. A, B and C represent the possible relative locations of the former proportional fuzzy linguistic distribution. (φ_2, ψ_2) is the interval value generated by the latter proportional fuzzy linguistic distribution.

(3) If $\varepsilon_1 \neq 0$ and $\varepsilon_2 \neq 0$, the two proportional fuzzy linguistic distributions will respectively generate an interval value (φ_1, ψ_1) , (φ_2, ψ_2) . Similarly, the relationship between the two proportional fuzzy linguistic distributions can be described in Figure 4.2.

4.3 Computation Operator of Proportional Fuzzy Linguistic Distribution

Proportional fuzzy linguistic distribution model is a kind of symbolic model, which is very easy to operate. We can carry out the related calculations directly on the labels and proportional coefficients in most instances. This leaves us much convenience even though we use it to deal with MADM problems under complex situations. However, we have to transform a complete proportional fuzzy linguistic distribution into a numerical value first before we aggregate a set of proportional fuzzy linguistic distributions in the situation that the linguistic weights are employed during the evaluation process. In other words, the weights are represented not by numerical values but by complete proportional fuzzy linguistic distributions. Therefore, the related computation operator that transforms a linguistic weight represented by complete proportional fuzzy linguistic distribution into a numerical value has to be defined.

Formally, let $S = \{s_0, s_1, \dots, s_n\}$ be an ordinal term set with $s_0 < s_1 < \dots < s_n$, and S^* is the proportional fuzzy linguistic distribution set generated by S . Define CCV of a complete proportional fuzzy linguistic distribution $(\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, 0)$ as follows:

$$\begin{aligned}
 & CCV(\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, 0) \\
 &= (\alpha_i CCV(s_i), \alpha_{i+1} CCV(s_{i+1}), \dots, \alpha_{i+m} CCV(s_{i+m}), 0) \\
 &= (\alpha_i c_i, \alpha_{i+1} c_{i+1}, \dots, \alpha_{i+m} c_{i+m}, 0) \\
 &= z_i + z_{i+1} + \dots + z_{i+m} \\
 &= \sum_{j=i}^{i+m} z_j
 \end{aligned} \tag{4.5}$$

with $j = i, i+1, \dots, i+m$. We call it the corresponding canonical characteristic value function on S^* generated by CCV on S . $c_i < c_{i+1} < \dots < c_{i+m} \in [0, 1]$ is the CCV of $s_i, s_{i+1}, \dots, s_{i+m}$ respectively.

With the function of CCV , we can easily transform a complete linguistic weight into a numerical value without loss of information. Thus, proportional fuzzy linguistic distribution model is capable of dealing with MADM problems with linguistic weights.

4.4 Expected Utility in Proportional Fuzzy Linguistic Distribution

Because of the operation mechanism of proportional fuzzy linguistic distribution model, it provides an aggregated distribution assessment for each alternative, which is different from most of the other MADA approaches. For the aggregated distribution assessments, it is very difficult to precisely describe the ranking order among them. In such case, the notion of expected utility, which is often associated with evidential reasoning approach [87] can be employed here to compare or rank alternatives.

“Suppose a set of alternatives X with a single-valued function $u(x)$ on X , which is called expected utility” [34]. One can represent the preference relation on X , such that for any $x, y \in X$, $x \succeq y$ if and only if $u(x) \geq u(y)$. As for $u(x)$, it may be estimated using the methods mentioned in [86], [87], such as “probability assignment method” [42], [76], “constructing regression models by using partial rankings or pairwise comparisons” and so on. Then, the way used for solving the problem of choosing x can be got by maximization of $u(x)$.

Suppose a set of evaluation grades

$$S = \{s_0, s_1, \dots, s_n\}$$

which are used as an instrument supplied to evaluators for evaluating the attribute. For a proportional fuzzy linguistic distribution $(\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, \varepsilon)$, assume a utility function

$$u' : S \rightarrow [0, 1]$$

satisfying

$$u'(s_{i+1}) > u'(s_i), \text{ if } s_{i+1} \text{ is preferred to } s_i.$$

Supposing alternatives a and b have a two level hierarchy with only an attribute y on the first level, and its basic attributes $E = \{e_1, e_2, \dots, e_n\}$ which is a finite set at the bottom level, as shown

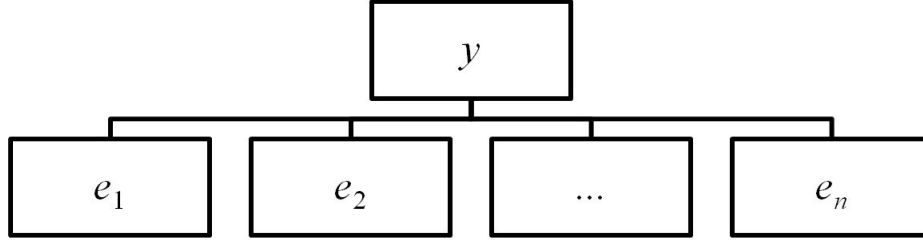


Figure 4.3. Two level hierarchy

in Figure 4.3. If all assessments for attributes are complete, i.e., $\sum_{j=i}^{i+m} \alpha_j = 1$, or $\varepsilon = 0$, then, the expected utility of the alternative a or b on the only attribute y is defined by

$$u(y) = \sum_{j=i}^{i+m} \alpha_j u'(s_j). \quad (4.6)$$

If and only if $u(y(a)) > u(y(b))$, we can say that the alternative a is preferred to another alternative b .

If there is any incomplete assessment for the basic attribute, i.e., $\sum_{j=i}^{i+m} \alpha_j < 1$, or $\varepsilon > 0$, then the assessment for y is also incomplete. In such case, the confidence interval $[\alpha_j, (\alpha_j + \varepsilon)]$ provides “the range of the likelihood to which y may be assessed to the evaluation grades” [87]. Without loss of generality, s_0 is supposed to be the least preferred grade which has the lowest utility and s_n is supposed to be the most preferred grade which has the highest utility. Then, the maximum, minimum and average expected utilities on y in proportional fuzzy linguistic distributions are given by

$$u_{\max}(y) = \sum_{j=0}^{n-1} \alpha_j u'(s_j) + (\alpha_n + \varepsilon) u'(s_n) \quad (4.7)$$

$$u_{\min}(y) = (\alpha_0 + \varepsilon) u'(s_0) + \sum_{j=1}^n \alpha_j u'(s_j) \quad (4.8)$$

$$u_{\text{avg}}(y) = \frac{u_{\max}(y) + u_{\min}(y)}{2}. \quad (4.9)$$

If all original assessments are complete, then $\varepsilon = 0$, and $u(y) = u_{\max}(y) = u_{\min}(y) = u_{\text{avg}}(y)$. If the original assessments include incomplete information, then the ranking of two alternatives a and b on y is based on their utility intervals and carried out by [34]

- $a \succ_y b$ if and only if $u_{\min}(y(a)) > u_{\max}(y(b))$
- $a \sim_y b$ if and only if $u_{\min}(y(a)) = u_{\min}(y(b))$ and $u_{\max}(y(a)) = u_{\max}(y(b))$.

Otherwise, the average expected utility can be used to generate a ranking, i.e.,

- $a \succ_y b$ on an average basis, if $u_{\text{avg}}(y(a)) > u_{\text{avg}}(y(b))$.

What we should note is that the ranking order is not reliable if the average expected utility is used. This is because there is a possibility that the special situation could happen, i.e., $u_{\text{avg}}(y(a)) > u_{\text{avg}}(y(b))$, but $u_{\max}(y(b)) > u_{\min}(y(a))$.

Further, if selecting a best alternative is not the only purpose of decision analysis, but also providing a ranking order, then, “the minimax regret approach (MRA)” [74], which is used to “calculate the maximum loss of expected utility”, can be employed to get such a ranking order. The ranking principle is that one alternative is selected as the best alternative if this alternative has the smallest maximum loss of expected utility. Then, exclude the best alternative, and calculate the maximum loss of expected utilities of other alternatives. Operate the MRA again and again until we get the ranking order of final two alternatives. Thus, the overall ranking order is also obtained. (For more details, see, e.g., [74].)

4.5 Proportional Fuzzy Linguistic Distribution Aggregation Operator

When we deal with MADA problems, we usually need to aggregate different information of attributes so as to obtain an integrated value that summarizes a set of values. The result of the aggregation of a set of proportional fuzzy linguistic distributions is also a proportional fuzzy linguistic distribution. To do so, we need to define several appropriate aggregation operators for proportional fuzzy linguistic distributions. Particularly, we defined the function CCV within the proportional fuzzy linguistic distribution framework which could transform complete proportional fuzzy linguistic distributions into numerical values in previous section. Combined with the function CCV , we are able to define a linguistic weighted aggregation operator so that the proportional fuzzy linguistic distribution model is capable of dealing with MADM problems with linguistic weights.

Let $S^* = \{(\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, \varepsilon)_1, (\alpha_g s_g, \alpha_{g+1} s_{g+1}, \dots, \alpha_{g+f} s_{g+f}, \varepsilon)_2, \dots, (\alpha_j s_j, \alpha_{j+1} s_{j+1}, \dots, \alpha_{j+q} s_{j+q}, \varepsilon)_p\}$ be a set of proportional fuzzy linguistic distributions. Here are some remarks that we should consider before introducing the aggregation operators.

- The number of proportional fuzzy linguistic distributions is p .
- It is unknown whether the starting labels of proportional fuzzy linguistic distributions are the same. For example, i perhaps doesn't equal g , but g perhaps equals j .
- It is unknown whether the ending labels of proportional fuzzy linguistic distributions are the same.
- In order to aggregate proportional fuzzy linguistic distributions more easily and directly, we will artificially make all the proportional fuzzy linguistic distributions symbolically have the same labels by adding "0" as proportional coefficients.

Taking these considerations into mind, the procedure of proportional fuzzy linguistic distribution aggregation operators is described as follows.

4.5.1 Arithmetic mean

Definition 4.1: Let $S^* = \{(\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, \varepsilon)_1, (\alpha_g s_g, \alpha_{g+1} s_{g+1}, \dots, \alpha_{g+f} s_{g+f}, \varepsilon)_2, \dots, (\alpha_j s_j, \alpha_{j+1} s_{j+1}, \dots, \alpha_{j+q} s_{j+q}, \varepsilon)_p\}$ be a set of proportional fuzzy linguistic distributions, and $(\gamma_k s_k, \gamma_{k+1} s_{k+1}, \dots, \gamma_{k+h} s_{k+h}, \bar{\varepsilon})$ be the arithmetic mean represented by proportional fuzzy linguistic distribution. Then, the procedure of calculating the arithmetic mean $(\gamma_k s_k, \gamma_{k+1} s_{k+1}, \dots, \gamma_{k+h} s_{k+h}, \bar{\varepsilon})$ is as follows.

- 1) Take the minimum of the starting labels of proportional fuzzy linguistic distributions in S^* , i.e., $k = \min(i, g, \dots, j)$.
- 2) Take the maximum of the ending labels of proportional fuzzy linguistic distributions in S^* , i.e., $k + h = \max(i + m, g + f, \dots, j + q)$.
- 3) Compare the proportional fuzzy linguistic distributions in S^* with arithmetic mean $(\gamma_k s_k, \gamma_{k+1} s_{k+1}, \dots, \gamma_{k+h} s_{k+h}, \bar{\varepsilon})$. For any proportional fuzzy linguistic distribution in S^* , if it is lack of corresponding linguistic terms, add "0" as symbolic proportional coefficients in front of the related linguistic terms. Thus, all the proportional fuzzy linguistic distributions

in S^* have the same starting labels and ending labels with arithmetic mean, i.e., $S^* = \{(\alpha_k s_k, \alpha_{k+1} s_{k+1}, \dots, \alpha_{k+h} s_{k+h}, \varepsilon)_1, (\alpha_k s_k, \alpha_{k+1} s_{k+1}, \dots, \alpha_{k+h} s_{k+h}, \varepsilon)_2, \dots, (\alpha_k s_k, \alpha_{k+1} s_{k+1}, \dots, \alpha_{k+h} s_{k+h}, \varepsilon)_p\}$.

4) The calculating process of arithmetic mean is given by

$$\left\{ \begin{array}{l} \gamma_k s_k = \left(\frac{\sum_{l=1}^p (\alpha_k)_l}{p} \right) s_k \\ \vdots \\ \gamma_{k+h} s_{k+h} = \left(\frac{\sum_{l=1}^p (\alpha_{k+h})_l}{p} \right) s_{k+h} \\ \bar{\varepsilon} = \frac{\sum_{l=1}^p \varepsilon_l}{p} \end{array} \right. \quad (4.10)$$

where l is the No. of proportional fuzzy linguistic distributions in S^* .

4.5.2 Weighted average operator

In the MADA problems, each attribute may have different weight, implying the importance of corresponding attribute. In this respect, weighted average operator is often used because it allows each attribute to be associated with a weight. The equivalent operator for proportional fuzzy linguistic distributions is defined as follows.

Definition 4.2: Let $S^* = \{(\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, \varepsilon)_1, (\alpha_g s_g, \alpha_{g+1} s_{g+1}, \dots, \alpha_{g+f} s_{g+f}, \varepsilon)_2, \dots, (\alpha_j s_j, \alpha_{j+1} s_{j+1}, \dots, \alpha_{j+q} s_{j+q}, \varepsilon)_p\}$ be a set of proportional fuzzy linguistic distributions, $W = \{\omega_1, \omega_2, \dots, \omega_p\}$ be their associated weights, and $(\gamma_k s_k, \gamma_{k+1} s_{k+1}, \dots, \gamma_{k+h} s_{k+h}, \bar{\varepsilon})$ be the weighted average of the set of proportional fuzzy linguistic distributions. Then, the computation and aggregation of the weighted average $(\gamma_k s_k, \gamma_{k+1} s_{k+1}, \dots, \gamma_{k+h} s_{k+h}, \bar{\varepsilon})$ is as follows.

- 1) Take the minimum of the starting labels of proportional fuzzy linguistic distributions in S^* , i.e., $k = \min(i, g, \dots, j)$.
- 2) Take the maximum of the ending labels of proportional fuzzy linguistic distributions in S^* , i.e., $k + h = \max(i + m, g + f, \dots, j + q)$.

- 3) Compare the proportional fuzzy linguistic distributions in S^* with the weighted average $(\gamma_k s_k, \gamma_{k+1} s_{k+1}, \dots, \gamma_{k+h} s_{k+h}, \bar{\varepsilon})$. For any proportional fuzzy linguistic distribution in S^* , if it is lack of corresponding linguistic terms, add “0” as symbolic proportional coefficients in front of the related linguistic terms in order to make sure all the proportional fuzzy linguistic distributions in S^* have the same starting labels and ending labels with weighted average, i.e., $S^* = \{(\alpha_k s_k, \alpha_{k+1} s_{k+1}, \dots, \alpha_{k+h} s_{k+h}, \varepsilon)_1, (\alpha_k s_k, \alpha_{k+1} s_{k+1}, \dots, \alpha_{k+h} s_{k+h}, \varepsilon)_2, \dots, (\alpha_k s_k, \alpha_{k+1} s_{k+1}, \dots, \alpha_{k+h} s_{k+h}, \varepsilon)_p\}$.
- 4) The calculating process of weighted average is given by

$$\left\{ \begin{array}{l} \gamma_k s_k = \left(\frac{\sum_{l=1}^p (\alpha_k)_l \cdot \omega_l}{\sum_{l=1}^p \omega_l} \right) s_k \\ \vdots \\ \gamma_{k+h} s_{k+h} = \left(\frac{\sum_{l=1}^p (\alpha_{k+h})_l \cdot \omega_l}{\sum_{l=1}^p \omega_l} \right) s_{k+h} \\ \bar{\varepsilon} = \frac{\sum_{l=1}^p \varepsilon_l \cdot \omega_l}{\sum_{l=1}^p \omega_l} \end{array} \right. \quad (4.11)$$

where l is the No. of proportional fuzzy linguistic distributions in S^* .

4.5.3 Linguistic weighted average operator

This is an extended version of weighted average operator, in which the weights are expressed by means of linguistic values instead of numerical values.

Definition 4.3: Let $S^* = \{(\alpha_i s_i, \alpha_{i+1} s_{i+1}, \dots, \alpha_{i+m} s_{i+m}, \varepsilon)_1, (\alpha_g s_g, \alpha_{g+1} s_{g+1}, \dots, \alpha_{g+f} s_{g+f}, \varepsilon)_2, \dots, (\alpha_j s_j, \alpha_{j+1} s_{j+1}, \dots, \alpha_{j+q} s_{j+q}, \varepsilon)_p\}$ be a set of proportional fuzzy linguistic distributions, and $W^* = \{(\beta_t \omega_t, \beta_{t+1} \omega_{t+1}, \dots, \beta_{t+d} \omega_{t+d}, 0)_1, \dots, (\beta_u \omega_u, \beta_{u+v} \omega_{u+v}, \dots, \beta_{u+v} \omega_{u+v}, 0)_p\}$ be the set of their associated linguistic weights which are represented by the form of complete proportional fuzzy linguistic distributions. Then, the procedure of calculating the linguistic weighted average $(\gamma_k s_k, \gamma_{k+1} s_{k+1}, \dots, \gamma_{k+h} s_{k+h}, \bar{\varepsilon})$ of proportional fuzzy linguistic distributions is as follows.

- 1) Take the minimum of the starting labels of proportional fuzzy linguistic distributions in S^* , i.e., $k = \min(i, g, \dots, j)$.

- 2) Take the maximum of the ending labels of proportional fuzzy linguistic distributions in S^* , i.e., $k + h = \max(i + m, g + f, \dots, j + q)$.
- 3) Compare the proportional fuzzy linguistic distributions in S^* with the linguistic weighted average $(\gamma_k s_k, \gamma_{k+1} s_{k+1}, \dots, \gamma_{k+h} s_{k+h}, \bar{\varepsilon})$. For any proportional fuzzy linguistic distribution in S^* , if it is lack of corresponding linguistic terms, add "0" as symbolic proportional coefficients in front of the related linguistic terms in order to make sure all the proportional fuzzy linguistic distributions in S^* have the same starting labels and ending labels with linguistic weighted average $(\gamma_k s_k, \gamma_{k+1} s_{k+1}, \dots, \gamma_{k+h} s_{k+h}, \bar{\varepsilon})$, i.e., $S^* = \{(\alpha_k s_k, \alpha_{k+1} s_{k+1}, \dots, \alpha_{k+h} s_{k+h}, \varepsilon)_1, (\alpha_k s_k, \alpha_{k+1} s_{k+1}, \dots, \alpha_{k+h} s_{k+h}, \varepsilon)_2, \dots, (\alpha_k s_k, \alpha_{k+1} s_{k+1}, \dots, \alpha_{k+h} s_{k+h}, \varepsilon)_p\}$.
- 4) Transform linguistic weights represented by complete proportional fuzzy linguistic distributions into numerical values by formula (4.5), i.e.,

$$\left\{ \begin{array}{l} CCV(\beta_t \omega_t, \beta_{t+1} \omega_{t+1}, \dots, \beta_{t+d} \omega_{t+d}, 0)_1 = \left(\sum_{b=t}^{t+d} Z_b \right)_1 = Z_1 \\ \vdots \\ CCV(\beta_u \omega_u, \beta_{u+1} \omega_{u+1}, \dots, \beta_{u+v} \omega_{u+v}, 0)_p = \left(\sum_{b=u}^{u+v} Z_b \right)_p = Z_p \end{array} \right. \quad (4.12)$$

where Z is the numerical values transformed by CCV over linguistic weights.

- 5) Thus, the calculating process of linguistic weighted average is given by

$$\left\{ \begin{array}{l} \gamma_k s_k = \left(\frac{\sum_{l=1}^p (\alpha_k)_l \cdot Z_l}{\sum_{l=1}^p Z_l} \right) s_k \\ \vdots \\ \gamma_{k+h} s_{k+h} = \left(\frac{\sum_{l=1}^p (\alpha_{k+h})_l \cdot Z_l}{\sum_{l=1}^p Z_l} \right) s_{k+h} \\ \bar{\varepsilon} = \frac{\sum_{l=1}^p \varepsilon_l \cdot Z_l}{\sum_{l=1}^p Z_l} \end{array} \right. \quad (4.13)$$

where l is the No. of proportional fuzzy linguistic distributions in S^* .

In order to clarify the procedure of calculating the linguistic weighted average of proportional fuzzy linguistic distributions, we give an example here.

Example: Suppose a set of proportional fuzzy linguistic distributions $S_1^* = \{(0.4s_2, 0.4s_3, 0.2), (0.6s_3, 0.2s_4, 0.2)\}$. Its associated linguistic weights are $S_2^* = \{(0.3s_4, 0.7s_5, 0), (0.5s_4, 0.5s_5, 0)\}$. S_1^* and S_2^* are the proportional fuzzy linguistic distribution sets generated by S_1 and S_2 respectively. S_1 and S_2 are linguistic term sets as shown in (4.14) and (4.15). The associated fuzzy number semantics of S_2 is shown in Figure 4.7. Then, the procedure of calculating the linguistic weighted average $(\gamma_k s_k, \gamma_{k+1} s_{k+1}, \dots, \gamma_{k+h} s_{k+h}, \bar{\varepsilon})$ is as follows.

- 1) For S_1^* , $k = \min(2, 3) = 2$.
- 2) For S_1^* , $k + h = \max(3, 4) = 4$.
- 3) For any proportional fuzzy linguistic distribution in S_1^* , add “0” as symbolic proportional coefficients in front of the related linguistic terms if it is lack of the corresponding linguistic terms. Then, $S_1^* = \{(0.4s_2, 0.4s_3, 0s_4, 0.2), (0s_2, 0.6s_3, 0.2s_4, 0.2)\}$.
- 4) Transform linguistic weights of S_2^* into numerical values by formula (4.12), i.e.,

$$\begin{cases} CCV(0.3s_4, 0.7s_5, 0) = 0.3 \times 0.8 + 0.7 \times 1 = 0.94 \\ CCV(0.5s_4, 0.5s_5, 0) = 0.5 \times 0.8 + 0.5 \times 1 = 0.9 \end{cases}$$

- 5) The linguistic weighted average can be obtained by

$$\begin{cases} \gamma_2 s_2 = \left(\frac{0.4 \times 0.94 + 0 \times 0.9}{0.94 + 0.9} \right) s_2 = 0.204s_2 \\ \gamma_3 s_3 = \left(\frac{0.4 \times 0.94 + 0.6 \times 0.9}{0.94 + 0.9} \right) s_3 = 0.498s_3 \\ \gamma_4 s_4 = \left(\frac{0 \times 0.94 + 0.2 \times 0.9}{0.94 + 0.9} \right) s_4 = 0.098s_4 \\ \bar{\varepsilon} = \frac{0.2 \times 0.94 + 0.2 \times 0.9}{0.94 + 0.9} = 0.2 \end{cases}$$

Therefore, the linguistic weighted average of the 2 proportional fuzzy linguistic distributions is $(0.204s_2, 0.498s_3, 0.098s_4, 0.2)$.

4.6 An Illustration Example

In this section, we apply proportional fuzzy linguistic distribution model to deal with a MADA problem with incomplete linguistic information taken from [87]. In order to compare the final result with [87], which used evidential reasoning approach to evaluate motorcycles, we first use the original data including distinct evaluation grades and weights. Then, we abandon the original weights and suppose a set of linguistic weights represented by complete proportional fuzzy linguistic distributions so that we can further analyze the capability of handling linguistic weights of this model.

4.6.1 Motorcycle evaluation problem

Simply speaking, “the problem is to evaluate the performances of four types of motorcycles, namely, Kawasaki, Yamaha, Honda, and BMW” [87]. Therefore, we have to know the overall performance of each motorcycle. “The overall performance of each motorcycle is based on evaluating three major qualitative attributes, which are quality of engine, operation, and general finish, although quantitative attributes may also be included” [84], [86]. Because the three major qualitative attributes are too general to assess easily and directly, we need to decompose them into more detailed sub-attributes so as to facilitate the assessment [87]. As a result of decomposition, an attribute hierarchy for evaluation of motorcycles is graphically depicted in Figure 4.4, where ω_i , ω_{ij} and ω_{ijk} are the weights of corresponding attributes at level 1, level 2, and level 3 respectively.

It is necessary and essential to define linguistic term set as well as associated semantics to supply evaluator for assessing the attributes of the operation of a motorcycle naturally. In this example, we use the same linguistic term set of distinct evaluation grades as in [87] for the purpose of comparing the final result, which is defined as

$$S_1 = \{s_0^1 \text{ (Poor)}, s_1^1 \text{ (Indifferent)}, s_2^1 \text{ (Average)}, \\ s_3^1 \text{ (Good)}, s_4^1 \text{ (Excellent)}\} \quad (4.14)$$

With considering the five evaluation grades, the subjective judgments for qualitative attributes can be easily expressed and summarized in Table 4.1, “where P , I , A , G , E are the abbreviations

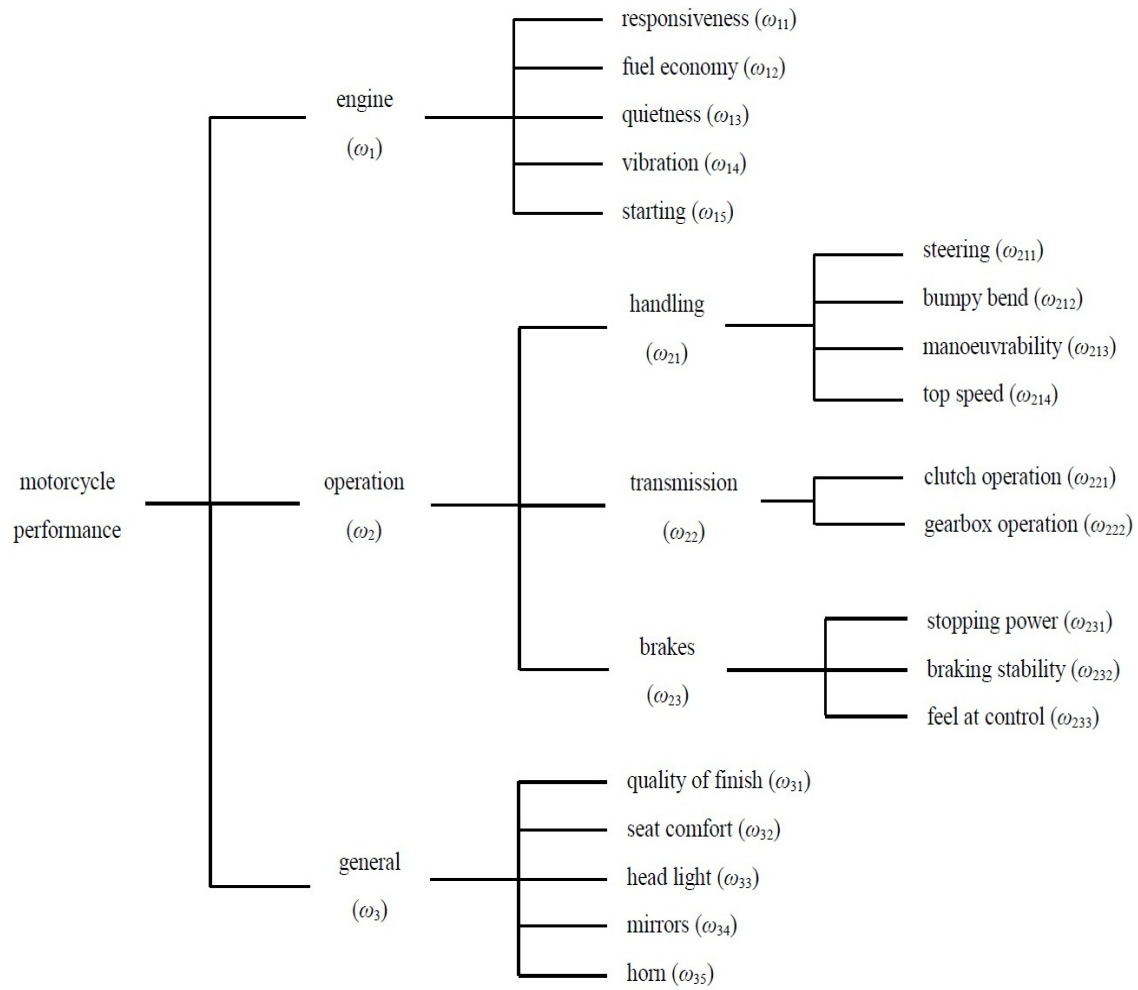


Figure 4.4. Evaluation hierarchy for motorcycle performance assessment [87]

of the evaluation grades of *Poor*, *Indifferent*, *Average*, *Good*, and *Excellent*, respectively, and a number in a bracket denotes a degree of belief to which an attribute is assessed to a grade” [87].

For the purpose of comparing the final result, but without loss of generality, we also suppose “all relevant attributes to be equal relative importance” as in [87], that is,

$$\begin{aligned}
\omega_1 &= \omega_2 = \omega_3 = 0.333 \\
\omega_{11} &= \omega_{12} = \omega_{13} = \omega_{14} = \omega_{15} = 0.2 \\
\omega_{21} &= \omega_{22} = \omega_{23} = 0.333 \\
\omega_{211} &= \omega_{212} = \omega_{213} = \omega_{214} = 0.25 \\
\omega_{221} &= \omega_{222} = 0.5 \\
\omega_{231} &= \omega_{232} = \omega_{233} = 0.333 \\
\omega_{31} &= \omega_{32} = \omega_{33} = \omega_{34} = \omega_{35} = 0.2.
\end{aligned}$$

4.6.2 Aggregating assessments by proportional fuzzy linguistic distribution model

Once the evaluator expresses all the linguistic assessments for the basic attributes, the evaluation procedure based on proportional fuzzy linguistic distribution model should be carried out. Specifically, the evaluation procedure is described as following.

(1) Proportional fuzzy linguistic distributions transformation: According to the linguistic term set of distinct evaluation grades, the original linguistic assessments shown in Table 4.1 should be converted into corresponding proportional fuzzy linguistic distributions by using symbolic translation value of $s_i, i = 0, 1, \dots, 4$ and associated representation method, such as 1)-4) discussed in Section 4.1. After transformation, the general decision matrix for motorcycle assessment represented by proportional fuzzy linguistic distributions is shown in Table 4.2. In Table 4.2, s_0, s_1, s_2, s_3 and s_4 are the representatives of *Poor*, *Indifferent*, *Average*, *Good* and *Excellent*, respectively. Besides, the numerical coefficients in front of s_0, s_1, s_2, s_3 and s_4 denote the confidence levels to which degree an attribute is assessed to a grade.

(2) Proportional fuzzy linguistic distributions computation and aggregation: According to the evaluation hierarchy of attributes, we first need to aggregate the third level attributes via formula

Table 4.1. Generalized decision matrix for motorcycle assessment [87]

General attributes	Basic attributes	Type of motorcycle (alternatives)			
		Kawasaki	Yamaha	Honda	BMW
engine ω_1	responsiveness ω_{11}	E (0.8)	G (0.3), E (0.6)	G (1.0)	I (1.0)
	fuel economy ω_{12}	A (1.0)	I (1.0)	I (0.5), A (0.5)	E (1.0)
	quietness ω_{13}	I (0.5), A (0.5)	A (1.0)	G (0.5), E (0.3)	E (1.0)
	vibration ω_{14}	G (1.0)	I (1.0)	G (0.5), E (0.5)	P (1.0)
	starting ω_{15}	G (1.0)	A (0.6), G (0.3)	G (1.0)	A (1.0)
handling ω_2	steering ω_{211}	E (0.9)	G (1.0)	A (1.0)	A (0.6)
	bumpy bends ω_{212}	A (0.5), G (0.5)	G (1.0)	G (0.8), E (0.1)	P (0.5), I (0.5)
	maneuverability ω_{213}	A (1.0)	E (0.9)	I (1.0)	P (1.0)
	top speed stability ω_{214}	E (1.0)	G (1.0)	G (1.0)	G (0.6), E (0.4)
	clutch operation ω_{221}	A (0.8)	G (1.0)	E (0.85)	I (0.2), A (0.8)
operation ω_2	gearbox operation ω_{222}	A (0.5), G (0.5)	I (0.5), A (0.5)	E (1.0)	P (1.0)
	stopping power ω_{231}	G (1.0)	A (0.3), G (0.6)	G (0.6)	E (1.0)
	braking stability ω_{232}	G (0.5), E (0.5)	G (1.0)	A (0.5), G (0.5)	E (1.0)
	feel at control ω_{233}	P (1.0)	G (0.5), E (0.5)	G (1.0)	G (0.5), E (0.5)
	quality of finish ω_{31}	P (0.5), I (0.5)	G (1.0)	E (1.0)	G (0.5), E (0.5)
general ω_3	seat comfort ω_{32}	G (1.0)	G (0.5), E (0.5)	G (0.6)	E (1.0)
	headlight ω_{33}	G (1.0)	A (1.0)	E (1.0)	G (0.5), E (0.5)
	mirrors ω_{34}	A (0.5), G (0.5)	G (0.5), E (0.5)	E (1.0)	G (1.0)
	horn ω_{35}	A (1.0)	G (1.0)	G (0.5), E (0.5)	E (1.0)

Table 4.2. Generalized decision matrix for motorcycle assessment represented by proportional fuzzy linguistic distributions

General attributes	Basic attributes	Type of motorcycle (alternatives)			
		Kawasaki	Yamaha	Honda	BMW
engine ω_1	responsiveness ω_{11}	$(0.8s_4, 0.2)$	$(0.3s_3, 0.6s_4, 0.1)$	$(1s_3, 0)$	$(1s_1, 0)$
	fuel economy ω_{12}	$(1s_2, 0)$	$(1s_1, 0)$	$(0.5s_1, 0.5s_2, 0)$	$(1s_4, 0)$
	quietness ω_{13}	$(0.5s_1, 0.5s_2, 0)$	$(1s_2, 0)$	$(0.5s_3, 0.3s_4, 0.2)$	$(1s_4, 0)$
	vibration ω_{14}	$(1s_3, 0)$	$(1s_1, 0)$	$(0.5s_3, 0.5s_4, 0)$	$(1s_0, 0)$
	starting ω_{15}	$(1s_3, 0)$	$(0.6s_2, 0.3s_3, 0.1)$	$(1s_3, 0)$	$(1s_2, 0)$
handling ω_{21}	steering ω_{211}	$(0.9s_4, 0.1)$	$(1s_3, 0)$	$(1s_2, 0)$	$(0.6s_2, 0.4)$
	bumpy bends ω_{212}	$(0.5s_2, 0.5s_3, 0)$	$(1s_3, 0)$	$(0.8s_3, 0.1s_4, 0.1)$	$(0.5s_0, 0.5s_1, 0)$
	maneuverability ω_{213}	$(1s_2, 0)$	$(0.9s_4, 0.1)$	$(1s_1, 0)$	$(1s_0, 0)$
	top speed stability ω_{214}	$(1s_4, 0)$	$(1s_3, 0)$	$(1s_3, 0)$	$(0.6s_3, 0.4s_4, 0)$
	clutch operation ω_{221}	$(0.8s_2, 0.2)$	$(1s_3, 0)$	$(0.85s_4, 0.15)$	$(0.2s_1, 0.8s_2, 0)$
operation ω_2	gearbox operation ω_{222}	$(0.5s_2, 0.5s_3, 0)$	$(0.5s_1, 0.5s_2, 0)$	$(1s_4, 0)$	$(1s_0, 0)$
	stopping power ω_{231}	$(1s_3, 0)$	$(0.3s_2, 0.6s_3, 0.1)$	$(0.6s_3, 0.4)$	$(1s_4, 0)$
	braking stability ω_{232}	$(0.5s_3, 0.5s_4, 0)$	$(1s_3, 0)$	$(0.5s_2, 0.5s_3, 0)$	$(1s_4, 0)$
	feel at control ω_{233}	$(1s_0, 0)$	$(0.5s_3, 0.5s_4, 0)$	$(1s_3, 0)$	$(0.5s_3, 0.5s_4, 0)$
	quality of finish ω_{31}	$(0.5s_0, 0.5s_1, 0)$	$(1s_3, 0)$	$(1s_4, 0)$	$(0.5s_3, 0.5s_4, 0)$
brakes ω_{23}	seat comfort ω_{32}	$(1s_3, 0)$	$(0.5s_3, 0.5s_4, 0)$	$(0.6s_3, 0.4)$	$(1s_4, 0)$
	headlight ω_{33}	$(1s_3, 0)$	$(1s_2, 0)$	$(1s_4, 0)$	$(0.5s_3, 0.5s_4, 0)$
	mirrors ω_{34}	$(0.5s_2, 0.5s_3, 0)$	$(0.5s_3, 0.5s_4, 0)$	$(1s_4, 0)$	$(1s_3, 0)$
	horn ω_{35}	$(1s_2, 0)$	$(1s_3, 0)$	$(0.5s_3, 0.5s_4, 0)$	$(1s_4, 0)$
general ω_3					

Table 4.3. The overall performances
represented by proportional fuzzy linguistic distributions

The overall performances	
Kawasaki	$(0.07s_0, 0.066s_1, 0.314s_2, 0.398s_3, 0.125s_4, 0.027)$
Yamaha	$(0s_0, 0.16s_1, 0.213s_2, 0.457s_3, 0.151s_4, 0.019)$
Honda	$(0s_0, 0.061s_1, 0.079s_2, 0.401s_3, 0.393s_4, 0.066)$
BMW	$(0.164s_0, 0.092s_1, 0.128s_2, 0.168s_3, 0.437s_4, 0.011)$

Table 4.4. Distributed assessments on four types of motorcycles

	Poor (P)	Indifferent (I)	Average (A)	Good (G)	Excellent (E)	ε
Kawasaki	0.07	0.066	0.314	0.398	0.125	0.027
Yamaha	0	0.16	0.213	0.457	0.151	0.019
Honda	0	0.061	0.079	0.401	0.393	0.066
BMW	0.164	0.092	0.128	0.168	0.437	0.011

(4.11). Then, aggregate the second level attributes and the first level attributes in proper order. After finishing the process of aggregation, we can obtain the final results of the overall performances of four types of motorcycles represented by proportional fuzzy linguistic distributions, as shown in Table 4.3.

(3) Proportional fuzzy linguistic distributions conversion: This step aims to convert the overall value of performances of four types of motorcycles represented by proportional fuzzy linguistic distributions into the corresponding linguistic terms of distinct evaluation grades, which are shown in Table 4.4. The distributed assessments on the four types of motorcycles can be shown graphically as in Figure 4.5.

4.6.3 Computing the expected utilities of four types of motorcycles

Proportional fuzzy linguistic distribution model supplies a distribution assessment for each alternative, as shown in Table 4.4 and in Figure 4.5, which is very difficult to precisely describe the related ranking order. Because the final purpose is to obtain such a ranking order among alternatives, the expected utilities of four types of motorcycles should be calculated. Therefore, the utilities of the five individual evaluation grades need to be estimated first. Assume a utility function $u' : S_1 \rightarrow [0, 1]$ and define the same utilities of five individual evaluation grades as in [87], i.e.,

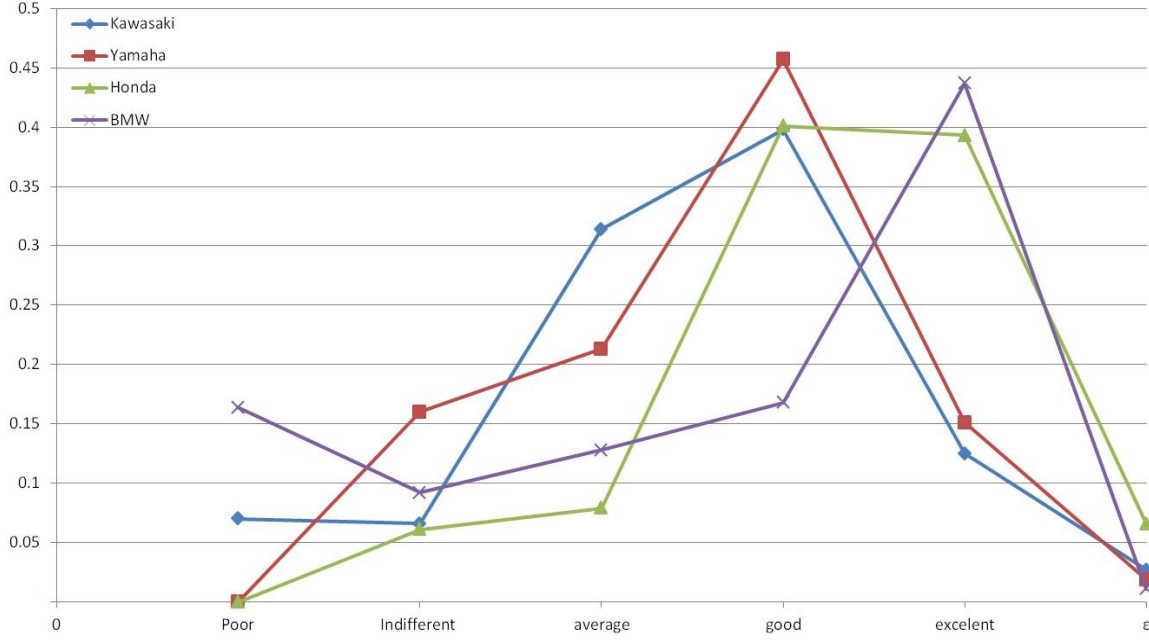


Figure 4.5. The distributed assessments on the four types of motorcycles

$$\begin{aligned}
 u'(P) &= 0, u'(I) = 0.35, u'(A) = 0.55, \\
 u'(G) &= 0.85, u'(E) = 1.
 \end{aligned}$$

We can easily obtain the expected utilities of four types of motorcycles, which are shown in Table 4.5 and in Figure 4.6 by using formula (4.7), (4.8) and (4.9). In Table 4.5, it is very obvious that the minimum expected utility of Honda is larger than the maximum expected utilities of the other three types of motorcycles. Hence, according to the ranking principle of expected utility, Honda is selected as the most preferred among the four types of motorcycles. Similarly, after excluding Honda, the minimum expected utility of Yamaha is larger than the maximum expected utilities of the other two types of motorcycles. Therefore, Yamaha is preferred to BMW and Kawasaki. In addition, although the minimum expected utility of BMW is not larger than the maximum expected utility of Kawasaki, they are very close, and the average utility of BMW is also larger than that of Kawasaki. Therefore, we believe that the overall performance of BMW is preferred to Kawasaki. Thus, based on the ranking principle of expected utility, the final ranking of the four types of motorcycles is given by

$$\text{Honda} \succ \text{Yamaha} \succ \text{BMW} \succ \text{Kawasaki}.$$

Table 4.5. The expected utilities of four types of motorcycles

	Maximum utility	Minimum utility	Average utility
Kawasaki	0.6861	0.6591	0.6726
Yamaha	0.7316	0.7126	0.7221
Honda	0.86465	0.79865	0.83165
BMW	0.6934	0.6824	0.6879

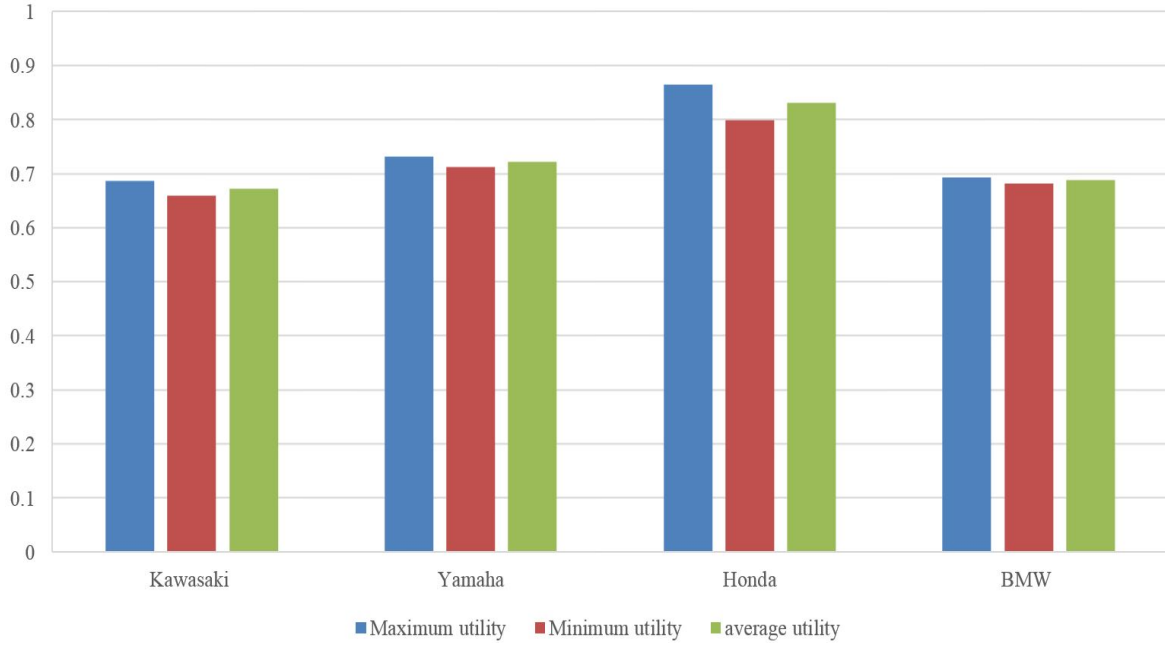


Figure 4.6. The expected utilities of four types of motorcycles

It is very clear there is not much difference between the results obtained by the proportional fuzzy linguistic distribution model and those obtained by the modified evidential reasoning approach [87]. As we can see, the expected utilities of four types of motorcycles obtained by proportional fuzzy linguistic distribution model are very similar with those obtained by the modified evidential reasoning approach [87], and the ranking order of the four types of motorcycles is the same. However, it is worth noticing that, because the proportional fuzzy linguistic distributions employ weighted average operator to aggregate multiple attribute, it clearly has a linear behavior. “While the modified evidential reasoning approach exhibits a quasi-linear behavior with equal weights and strongly nonlinear behavior with unequal weights” [34].

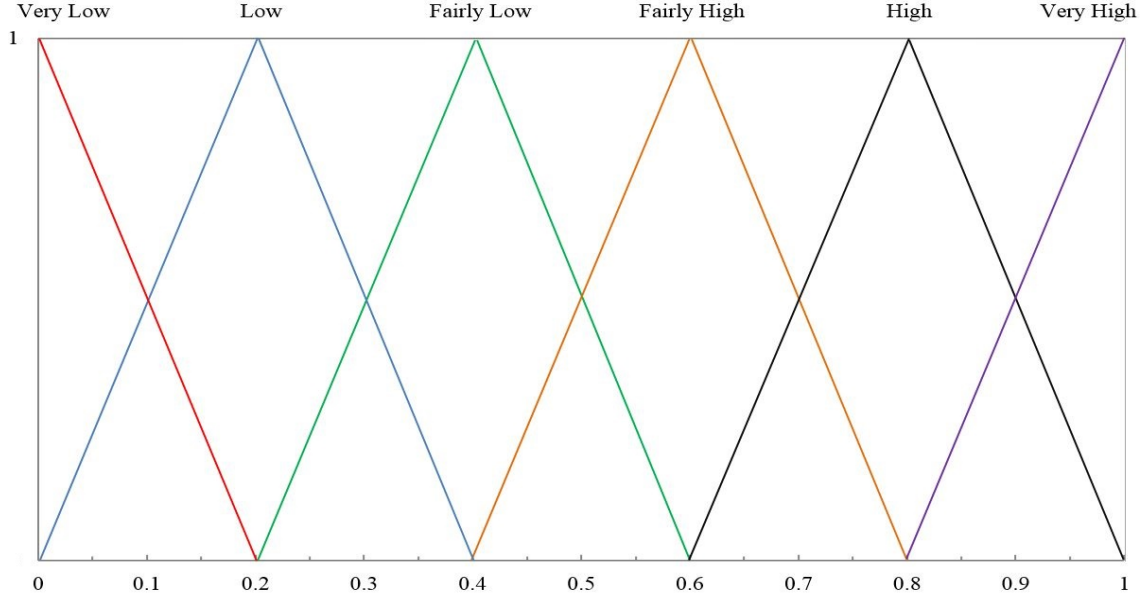


Figure 4.7. Linguistic weights and associated fuzzy number semantics

4.6.4 Motorcycle assessment problem with linguistic weights

In order to further analyze the capability of dealing with linguistic weights of proportional fuzzy linguistic distribution model, we assume a set of linguistic weights represented by means of complete proportional fuzzy linguistic distributions. To do so, a distinctive evaluation set which collectively provides a complete set of standards for evaluating the importance of each attribute need to be defined first.

Specifically, define a linguistic term set which is used to linguistically evaluate the relative importance of different attributes as S_2 ,

$$S_2 = \{s_0^2(\text{Very Low}), s_1^2(\text{Low}), s_2^2(\text{Fairly Low}), s_3^2(\text{Fairly High}), s_4^2(\text{High}), s_5^2(\text{Very High})\} \quad (4.15)$$

and the associated fuzzy set semantics is shown in Figure 4.7.

With the instrument of the linguistic weight term set and associated fuzzy number semantics, the evaluator can express the relative importance of different attributes of each motorcycle by the form of uncertain statements as discussed in Section 4.1, which are then represented by proportional

Table 4.6. Linguistic weights represented
by proportional fuzzy linguistic distributions

Attributes	The third level	Linguistic weights
handling ω_{21}	steering ω_{211}	$(0.4s_3, 0.6s_4, 0)$
	bumpy bends ω_{212}	$(1s_4, 0)$
	maneuverability ω_{213}	$(0.5s_4, 0.5s_5, 0)$
	top speed stability ω_{214}	$(1s_3, 0)$
transmission ω_{22}	clutch operation ω_{221}	$(0.5s_3, 0.5s_4, 0)$
	gearbox operation ω_{222}	$(0.5s_3, 0.5s_4, 0)$
	stopping power ω_{231}	$(1s_5, 0)$
brakes ω_{23}	braking stability ω_{232}	$(1s_5, 0)$
	feel at control ω_{233}	$(0.5s_4, 0.5s_5, 0)$
Attributes	The second level	Linguistic weights
engine ω_1	responsiveness ω_{11}	$(0.4s_4, 0.6s_5, 0)$
	fuel economy ω_{12}	$(0.3s_3, 0.3s_4, 0.4s_5, 0)$
	quietness ω_{13}	$(1s_3, 0)$
	vibration	$(1s_3, 0)$
	starting ω_{15}	$(0.5s_3, 0.5s_4, 0)$
operation ω_2	handling ω_{21}	$(0.4s_4, 0.6s_5, 0)$
	transmission ω_{22}	$(1s_5, 0)$
	brakes ω_{23}	$(1s_5, 0)$
	quality of finish ω_{31}	$(1s_3, 0)$
	seat comfort ω_{32}	$(0.6s_1, 0.4s_2, 0)$
general ω_3	headlight ω_{33}	$(0.2s_3, 0.8s_4, 0)$
	mirrors ω_{34}	$(1s_2, 0)$
	horn ω_{35}	$(1s_2, 0)$
Attributes	The first level	Linguistic weights
overall performance	engine ω_1	$(0.4s_4, 0.6s_5, 0)$
	operation ω_2	$(0.4s_4, 0.6s_5, 0)$
	general ω_3	$(1s_4, 0)$

fuzzy linguistic distributions as shown in Table 4.6, where s_0, s_1, s_2, s_3, s_4 and s_5 are the expressions of *Very Low*, *Low*, *Fairly Low*, *Fairly High*, *High* and *Very High*, respectively.

Similarly, according to the aggregation procedure of proportional fuzzy linguistic distribution model mentioned above, the final results of the overall performances of four types of motorcycles represented by proportional fuzzy linguistic distributions can be obtained easily by formula (4.12)

Table 4.7. The overall performances represented by proportional fuzzy linguistic distributions by using linguistic weights

The overall performances	
Kawasaki	$(0.074s_0, 0.066s_1, 0.307s_2, 0.384s_3, 0.136s_4, 0.033)$
Yamaha	$(0s_0, 0.166s_1, 0.234s_2, 0.434s_3, 0.143s_4, 0.023)$
Honda	$(0s_0, 0.072s_1, 0.086s_2, 0.386s_3, 0.402s_4, 0.054)$
BMW	$(0.164s_0, 0.115s_1, 0.13s_2, 0.166s_3, 0.414s_4, 0.011)$

Table 4.8. Distributed assessments on four types of motorcycles by using linguistic weights

	Poor (P)	Indifferent (I)	Average (A)	Good (G)	Excellent (E)	ε
Kawasaki	0.074	0.066	0.307	0.384	0.136	0.033
Yamaha	0	0.166	0.234	0.434	0.143	0.023
Honda	0	0.072	0.086	0.386	0.402	0.054
BMW	0.164	0.115	0.13	0.166	0.414	0.011

and (4.13), as shown in Table 4.7. Then, the overall performances of four types of motorcycles are converted into the corresponding linguistic terms of distinct evaluation grades, as shown in Table 4.8 and graphically depicted in Figure 4.8. By using the same utilities of five individual evaluation grades mentioned above, the expected utilities of four types of motorcycles can be obtained by formula (4.7), (4.8) and (4.9) as shown in Table 4.9 and Figure 4.9. Thus, we can easily obtain the ranking order of the four types of motorcycles based on the ranking principle of expected utility. It is obvious that Honda is still the most preferred among the four types of motorcycles, and the ranking of the four types of motorcycles is given by

$$\text{Honda} \succ \text{Yamaha} \succ \text{BMW} \succ \text{Kawasaki}.$$

Interestingly, the same ranking result was obtained. This is probably because under the corresponding weights, the performances of some attributes of Honda are really better than those of the other three types of motorcycles. Besides, the three kinds of expected utilities of each motorcycle obtained by linguistic weighted aggregation operator are very similar with those obtained by weighted average operator. One of the possible reasons is because linguistic weighted aggregation operator is an extended version of weighted average operator.

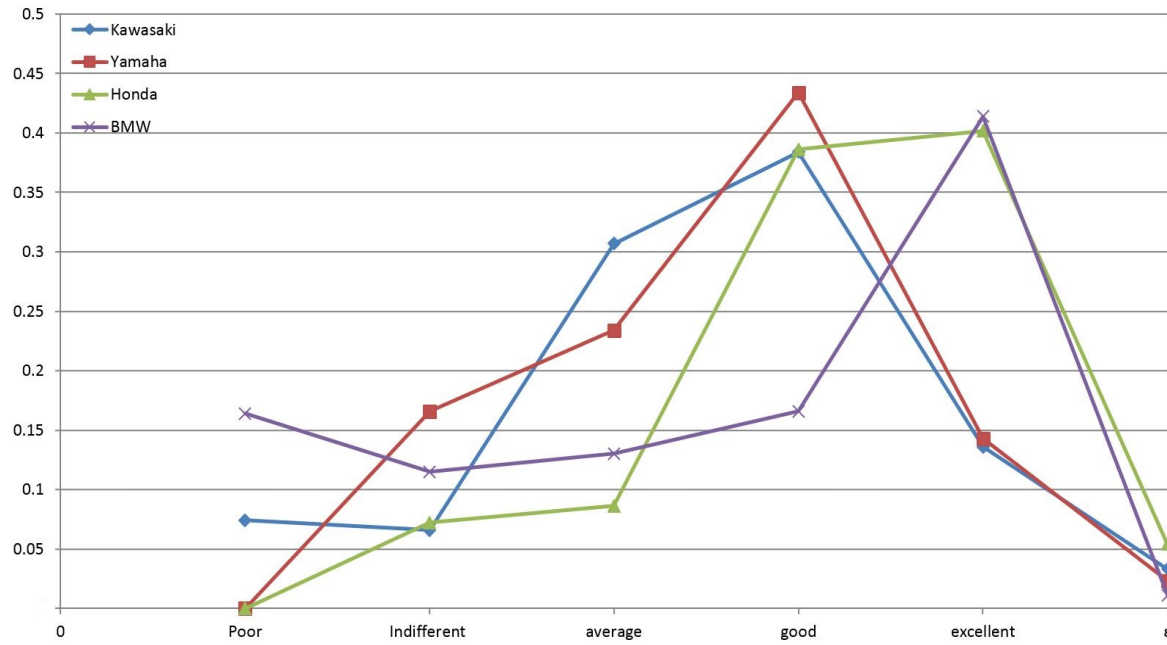


Figure 4.8. The distributed assessments on the four types of motorcycles by using linguistic weights

Table 4.9. The expected utilities of four types of motorcycles by using linguistic weights

	Maximum utility	Minimum utility	Average utility
Kawasaki	0.68735	0.65435	0.67085
Yamaha	0.7217	0.6987	0.7102
Honda	0.8566	0.8026	0.8296
BMW	0.67785	0.66685	0.67235

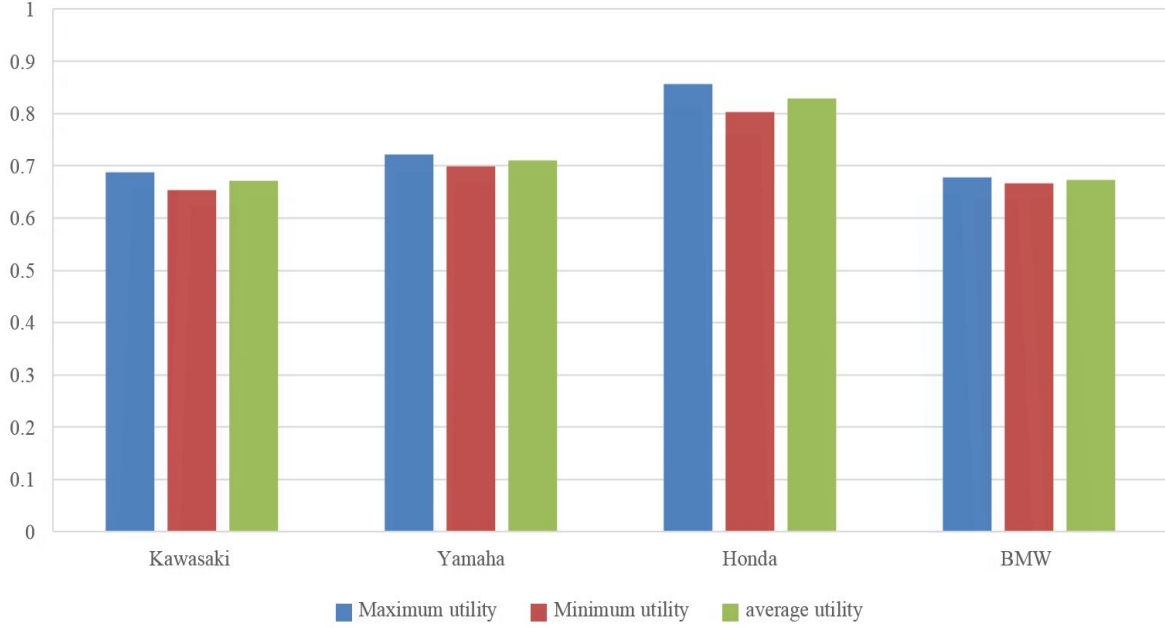


Figure 4.9. The expected utilities of four types of motorcycles by using linguistic weights

4.7 Conclusion

In this chapter, we introduced a proportional fuzzy linguistic distribution model for MADM problems with incomplete linguistic information. In this model, we used proportions as evaluators' confidence levels indicating their belief degrees that each linguistic term fits a linguistic variable. Further, by relaxing the restrictions, the rational combinations of any number of linguistic terms associated with corresponding proportions can be used as evaluator's subjective judgments. This feature can be regarded as a measure for evaluators to handle uncertainty during the evaluation process. Moreover, with introducing a new variable representing the extent of ignoring information, incomplete linguistic assessment can be involved in this model. These new features could give evaluator more flexible operation space, largely reduce the evaluator's pressure during the evaluation process, and finally, improve the reasonability and precision of the final result.

In this chapter, we also developed the computation operator for transforming a complete proportional fuzzy linguistic distribution into a numerical value so that the proportional fuzzy linguistic distribution model was capable of dealing with MADM problems with linguistic weights. Besides, we introduced arithmetic mean, weighted average operator and linguistic weighted average operator for aggregating proportional fuzzy linguistic distributions. Also, the

expected utilities in proportional fuzzy linguistic distributions were proposed for ranking the alternatives and making a final decision. Finally, a motorcycle assessment problem with numerical weights and linguistic weights was used to illustrate the proposed model. It is shown that this model has some advantages, such as accuracy, no loss of information, little restriction to evaluators in the process of evaluation, ease operating under the complex context and so forth.

Chapter 5

An Interval Fuzzy Linguistic Distribution Model

In this chapter, we propose an interval fuzzy linguistic distribution model for the purpose of supplying a new way for dealing with MADM problems with incomplete linguistic information. In this model, intervals are used as evaluators confidence levels indicating their belief degrees that a linguistic term fits a linguistic variable. Similarly with proportional fuzzy linguistic distribution model, when interval fuzzy linguistic distribution model is employed to deal with linguistic information, there is no obligatory requirement that evaluators have to supply the subjective judgments which are constituted by fixed number of linguistic terms. Each attribute can be directly evaluated by using the subjective judgment that contains any number of adjacent evaluation grades. Besides, with introducing a variable representing the extent of ignoring information, the interval fuzzy linguistic distribution model is also capable of dealing with incomplete linguistic assessments. As a matter of fact, proportion can be regarded as a special interval with left limit equaling right limit. Therefore, interval fuzzy linguistic distribution model is a generalization of proportional fuzzy linguistic distribution model. It inherits all the advantages of proportional fuzzy linguistic distribution model. Moreover, compared with proportions, intervals are able to reflect uncertain information more efficiently and comprehensively. Hence, when evaluators feel difficult to express their confidence levels exactly by proportions under uncertain and complicated situations, intervals could be considered as an alternative measure for reflecting their confidence levels.

It is worth mentioning that the final result obtained by interval fuzzy linguistic distribution

model might be lack of precision compared with proportional fuzzy linguistic distribution model. This is because intervals contain much more uncertain information than proportions. Therefore, if proportions are enough to reflect evaluators' confidence levels, proportional fuzzy linguistic distribution model should always be considered to use firstly.

5.1 Linguistic Assessments with Intervals

5.1.1 Basic interval arithmetic

Generally speaking, an interval number can be denoted by $A = [\alpha_L, \alpha_R] = \{\alpha | \alpha_L \leq \alpha \leq \alpha_R, \alpha \in R\}$, where α_L and α_R are the left limit and right limit of interval A respectively. Specially, if $\alpha_L = \alpha_R$, then, interval A reduces to a real number.

The basic interval operation rules are described as follows. Specifically, let A and B be two positive interval numbers, then,

$$A \oplus B = [\alpha_L, \alpha_R] \oplus [\beta_L, \beta_R] = [\alpha_L + \beta_L, \alpha_R + \beta_R] \quad (5.1)$$

$$A \ominus B = [\alpha_L, \alpha_R] \ominus [\beta_L, \beta_R] = [\alpha_L - \beta_L, \alpha_R - \beta_R] \quad (5.2)$$

$$A \otimes B = [\alpha_L, \alpha_R] \otimes [\beta_L, \beta_R] = [\alpha_L \beta_L, \alpha_R \beta_R] \quad (5.3)$$

$$\lambda \times A = \begin{cases} [\lambda \alpha_L, \lambda \alpha_R] & \text{if } \lambda \geq 0 \\ [\lambda \alpha_R, \lambda \alpha_L] & \text{if } \lambda \leq 0 \end{cases} \quad (5.4)$$

where λ is a scalar.

In the literature, we can find several comparison operators between two interval numbers such as [16], [36], [55], [63]. Particularly, Ishibuchi and Tanaka [36] defined a comparison operator as

$$A \leq B \Leftrightarrow \alpha_L \leq \beta_L, \alpha_R \leq \beta_R. \quad (5.5)$$

Without loss of generality, we use Ishibuchi and Tanaka's approach for comparison of interval fuzzy linguistic distribution in this chapter.

Besides, it is not difficult to find some aggregation operators with interval uncertainty in the literature, e.g., [1], [19], [53], [95]. Especially, some aggregation operators with interval uncertainty have been used in the linguistic computational model (see, e.g., [19], [53]). For the purpose of developing "interval weighted average operator (*IWA*)" and "interval ordered weighted average operator (*IOWA*)" for interval fuzzy linguistic distributions in the following section, we need to briefly recall the general *IWA* and *IOWA*.

Let $A_i = [\alpha_L, \alpha_R]_i$ be the interval attributes, with $i = 1, 2, \dots, n$, and $W_i = [\omega_L, \omega_R]_i$ be the associated interval weights. Then, the result of the general *IWA*, $Y_{IWA} = [y_L, y_R]$, is defined as follows:

$$\begin{cases} y_L = \min_{\forall \omega_i \in [\omega_L, \omega_R]_i} \frac{\sum_{i=1}^n (\alpha_L)_i \cdot \omega_i}{\sum_{i=1}^n \omega_i} \\ y_R = \max_{\forall \omega_i \in [\omega_L, \omega_R]_i} \frac{\sum_{i=1}^n (\alpha_R)_i \cdot \omega_i}{\sum_{i=1}^n \omega_i} \end{cases}. \quad (5.6)$$

Let A_{σ_i} be the i th largest elements of $\{A_1, A_2, \dots, A_n\}$. Then, the result of the general *IOWA*, $Z_{IOWA} = [z_L, z_R]$, is defined as follows:

$$\begin{cases} z_L = \min_{\forall \omega_i \in [\omega_L, \omega_R]_i} \frac{\sum_{i=1}^n (\alpha_L)_{\sigma_i} \cdot \omega_i}{\sum_{i=1}^n \omega_i} \\ z_R = \max_{\forall \omega_i \in [\omega_L, \omega_R]_i} \frac{\sum_{i=1}^n (\alpha_R)_{\sigma_i} \cdot \omega_i}{\sum_{i=1}^n \omega_i} \end{cases}. \quad (5.7)$$

Karnik and Mendel (KM) [41] developed “iterative algorithms, which are known as KM algorithms”. Wu and Mendel [77] further introduced “enhanced KM algorithms”. Based on the use of the “KM algorithms” (or “the enhanced KM algorithms”), the values of y_L, y_R, z_L and z_R can be easily obtained [19].

5.1.2 Interval fuzzy linguistic distribution

Let $S = \{s_0, s_1, \dots, s_n\}$ be an ordinal term set with $s_0 < s_1 < \dots < s_n$ (“ $<$ ” represents order relation, i.e., $s_i < s_j$ if and only if $i < j$), $I = [0, 1]$ and

$$IS \equiv I \times S = \{([\alpha_L, \alpha_R], s_i) : 0 \leq \alpha_L \leq \alpha_R \leq 1, \text{ and } i = 0, 1, \dots, n\}.$$

Given a sequence $(s_i, s_{i+1}, \dots, s_{i+m})$ of $(m + 1)$ successive ordinal terms of S , any $(m + 1)$ elements $([\alpha_L, \alpha_R]_i, s_i), ([\alpha_L, \alpha_R]_{i+1}, s_{i+1}), \dots, ([\alpha_L, \alpha_R]_{i+m}, s_{i+m})$ of IS are called a symbolic interval sequence, and it will be denoted by

$$\begin{cases} ([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, 0) & \text{if } \sum_{j=i}^{i+m} (\alpha_L)_j = 1 \\ ([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, \varepsilon) & \text{if } \sum_{j=i}^{i+m} (\alpha_L)_j < 1 \end{cases} \quad (5.8)$$

where ε represents the extent of ignoring information, and it is an interval number usually. The set of all the symbolic interval sequences is denoted by S^* , i.e., $S^* = \{([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, \varepsilon) : 0 \leq \alpha_{Li} \leq \alpha_{Ri} \leq 1, \alpha_{Ri} \neq 0, 0 \leq \alpha_{Li+m} \leq \alpha_{Ri+m} \leq 1, \alpha_{Ri+m} \neq 0, \sum_{j=i}^{i+m} (\alpha_L)_j \leq 1, \sum_{j=i}^{i+m} (\alpha_R)_j \geq 1, 0 \leq i, i + m \leq n\}$. The set S^* is called interval fuzzy linguistic distribution set generated by S and the members of S^* are called interval fuzzy linguistic distributions. If $\sum_{j=i}^{i+m} (\alpha_L)_j > 1$ or $\sum_{j=i}^{i+m} (\alpha_R)_j < 1$, then this interval linguistic distribution is said to be invalid. “Invalid interval linguistic distribution cannot be interpreted as probability and thus need to be revised or adjusted” [75].

It is worth mentioning that because an interval linguistic distribution includes interval-valued belief structure, it involves a problem whether it is normalized or non-normalized interval-valued belief structure. “For a non-normalized interval-valued belief structure, it usually means that some intervals of probability masses are too wide to be reached” [75]. Although a valid interval-valued belief structure is not necessarily to be normalized, it can avoid too wide intervals and improve

the precision in final result if we normalize all the interval-valued belief structure in advance. The following equations [65], [73], [75] can be used to verify whether an interval linguistic distribution is normalized interval-valued belief structure or not.

$$\sum_{k=i}^{i+m} (\alpha_R)_k - ((\alpha_R)_j - (\alpha_L)_j) \geq 1 \text{ and } \sum_{k=i}^{i+m} (\alpha_L)_k + ((\alpha_R)_j - (\alpha_L)_j) \leq 1 \quad (5.9)$$

for $\forall j \in \{i, \dots, i+m\}$. If $(\alpha_L)_j$ and $(\alpha_R)_j$ satisfy these requirements, it is called a normalized interval-valued belief structure. Otherwise, the following equations can be employed to normalize it so that we can screen out all infeasible belief structures [75].

$$\max \left[(\alpha_L)_j, 1 - \sum_{k \neq j} (\alpha_R)_k \right] \leq (\alpha_L)_j \text{ and } (\alpha_R)_j \leq \min \left[(\alpha_R)_j, 1 - \sum_{k \neq j} (\alpha_L)_k \right] \quad (5.10)$$

$j = i, \dots, i+m$, and $k = i, \dots, i+m$. (For more details, please see [75].)

An interval fuzzy linguistic distribution $([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, \varepsilon)$ can be used to represent an evaluator's subjective judgment. Here, i is called the starting label; s_i is the No. i linguistic term; $[\alpha_L, \alpha_R]_j$ is the interval coefficient in front of the related linguistic term. It represents the confidence levels that to which degree evaluators believe a linguistic term fits a linguistic variable. Similarly, $i+m$ is called the ending label.

A linguistic assessment $([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, \varepsilon)$ is called complete interval fuzzy linguistic distribution if $\sum_{j=i}^{i+m} (\alpha_L)_j = 1$, and incomplete interval fuzzy linguistic distribution if $\sum_{j=i}^{i+m} (\alpha_L)_j < 1$. For an incomplete linguistic assessment $([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, \varepsilon)$, the extent of ignoring information ε can be obtained by [74]

$$(\alpha_L)_\varepsilon = \max \left(0, 1 - \sum_{j=i}^{i+m} (\alpha_R)_j \right) \quad (5.11)$$

$$(\alpha_R)_\varepsilon = 1 - \sum_{j=i}^{i+m} (\alpha_L)_j. \quad (5.12)$$

For example, an evaluator assesses a cargo ship selection problem [74] and gives linguistic assessments for Ship 1 as follows:

- 1) The *load factor* of Ship 1 is evaluated to be *very good* with a confidence degree of $[0.2, 0.4]$, and to be *excellent* with a confidence degree of $[0.7, 0.8]$.
- 2) The *effective load factor* of Ship 1 is evaluated to be *very good* with a confidence degree of $[0.1, 0.2]$, and to be *excellent* with a confidence degree of $[0.8, 0.9]$.

Then for Ship 1, the two linguistic assessments 1)-2) given in above can be represented in the form of interval fuzzy linguistic distributions defined by (5.8) as

$$S^*(load\ factor) = ([0.2, 0.4]_{s_3}, [0.7, 0.8]_{s_4}, \varepsilon_1)$$

$$S^*(effective\ load\ factor) = ([0.1, 0.2]_{s_3}, [0.8, 0.9]_{s_4}, \varepsilon_2)$$

where s_i with $i = 3$ and 4 are linguistic terms of the term set S_2 as shown in (5.29). ε_1 and ε_2 can be obtained by

$$\begin{aligned} (\alpha_L)_{\varepsilon_1} &= \max\left(0, 1 - \sum_{j=3}^4 (\alpha_R)_j\right) = \max(0, 1 - 0.4 - 0.8) = 0, \\ (\alpha_R)_{\varepsilon_1} &= 1 - \sum_{j=3}^4 (\alpha_L)_j = 1 - 0.2 - 0.7 = 0.1, \end{aligned}$$

$$\begin{aligned} (\alpha_L)_{\varepsilon_2} &= \max\left(0, 1 - \sum_{j=3}^4 (\alpha_R)_j\right) = \max(0, 1 - 0.2 - 0.9) = 0, \\ (\alpha_R)_{\varepsilon_2} &= 1 - \sum_{j=3}^4 (\alpha_L)_j = 1 - 0.1 - 0.8 = 0.1. \end{aligned}$$

Thus, the two interval fuzzy linguistic distributions can be formally represented by

$$S^*(load\ factor) = ([0.2, 0.4]_{s_3}, [0.7, 0.8]_{s_4}, [0, 0.1])$$

$$S^*(effective\ load\ factor) = ([0.1, 0.2]_{s_3}, [0.8, 0.9]_{s_4}, [0, 0.1]).$$

5.2 Comparison of Interval Fuzzy Linguistic Distributions

Let $S = \{s_0, s_1, \dots, s_n\}$ be an ordinal term set and S^* be the interval fuzzy linguistic distribution set generated by S . For any $([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, \varepsilon)_1$, $([\beta_L, \beta_R]_g s_g, [\beta_L, \beta_R]_{g+1} s_{g+1}, \dots, [\beta_L, \beta_R]_{g+f} s_{g+f}, \varepsilon)_2 \in S^*$, the comparison of interval fuzzy linguistic distributions is described as follows.

(1) If $\varepsilon_1 = 0$ and $\varepsilon_2 = 0$, define $([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, 0)_1 > ([\beta_L, \beta_R]_g s_g, [\beta_L, \beta_R]_{g+1} s_{g+1}, \dots, [\beta_L, \beta_R]_{g+f} s_{g+f}, 0)_2$

$$\begin{aligned}
 &\Leftrightarrow [\alpha_L, \alpha_R]_i \cdot i + [\alpha_L, \alpha_R]_{i+1} \cdot (i+1) + \dots + [\alpha_L, \alpha_R]_{i+m} \cdot (i+m) \\
 &> [\beta_L, \beta_R]_g \cdot g + [\beta_L, \beta_R]_{g+1} \cdot (g+1) + \dots + [\beta_L, \beta_R]_{g+f} \cdot (g+f) \\
 &\Leftrightarrow \left[\sum_{j=i}^{i+m} ((\alpha_L)_j \cdot j), \sum_{j=i}^{i+m} ((\alpha_R)_j \cdot j) \right] > \left[\sum_{k=g}^{g+f} ((\beta_L)_k \cdot k), \sum_{k=g}^{g+f} ((\beta_R)_k \cdot k) \right] \\
 &\Leftrightarrow \sum_{j=i}^{i+m} ((\alpha_L)_j \cdot j) > \sum_{k=g}^{g+f} ((\beta_L)_k \cdot k), \sum_{j=i}^{i+m} ((\alpha_R)_j \cdot j) > \sum_{k=g}^{g+f} ((\beta_R)_k \cdot k). \tag{5.13}
 \end{aligned}$$

(2) If $\varepsilon_1 = 0$ and $\varepsilon_2 \neq 0$, ε_2 is an interval value $[\beta_L, \beta_R]_\varepsilon$. We need to allocate ε_2 in order to obtain the left limit φ and right limit ψ .

For the left limit φ , we can consider an extreme situation that ε_2 is allocated to s_0 completely, i.e.,

$$\begin{aligned}
 \varphi &= [\beta_L, \beta_R]_g \cdot g + [\beta_L, \beta_R]_{g+1} \cdot (g+1) + \dots + [\beta_L, \beta_R]_{g+f} \cdot (g+f) + [\beta_L, \beta_R]_\varepsilon \cdot 0 \\
 &= \sum_{k=g}^{g+f} ((\beta_L)_k \cdot k). \tag{5.14}
 \end{aligned}$$

For the right limit ψ , we can consider an extreme situation that ε_2 is allocated to s_n completely, i.e.,

$$\begin{aligned}
 \psi &= [\beta_L, \beta_R]_g \cdot g + [\beta_L, \beta_R]_{g+1} \cdot (g+1) + \dots + [\beta_L, \beta_R]_{g+f} \cdot (g+f) + [\beta_L, \beta_R]_\varepsilon \cdot n \\
 &= \sum_{k=g}^{g+f} ((\beta_R)_k \cdot k) + (\beta_R)_\varepsilon \cdot n. \tag{5.15}
 \end{aligned}$$

Then, the left limit and right limit of the second interval fuzzy linguistic distribution can be represented as $\left[\sum_{k=g}^{g+f} ((\beta_L)_k \cdot k), \sum_{k=g}^{g+f} ((\beta_R)_k \cdot k) + (\beta_R)_\varepsilon \cdot n \right]$, and $([\alpha_L, \alpha_R]_{i s_i}, [\alpha_L, \alpha_R]_{i+1 s_{i+1}}, \dots, [\alpha_L, \alpha_R]_{i+m s_{i+m}}, 0)_1 > ([\beta_L, \beta_R]_{g s_g}, [\beta_L, \beta_R]_{g+1 s_{g+1}}, \dots, [\beta_L, \beta_R]_{g+f s_{g+f}}, \varepsilon)_2$ is denoted by

$$\begin{aligned} \sum_{j=i}^{i+m} ((\alpha_L)_j \cdot j) &> \sum_{k=g}^{g+f} ((\beta_L)_k \cdot k), \\ \sum_{j=i}^{i+m} ((\alpha_R)_j \cdot j) &> \sum_{k=g}^{g+f} ((\beta_R)_k \cdot k) + (\beta_R)_\varepsilon \cdot n. \end{aligned} \quad (5.16)$$

(3) If $\varepsilon_1 \neq 0$ and $\varepsilon_2 \neq 0$, we need to respectively allocate $\varepsilon_1, \varepsilon_2$ in order to obtain the left limit φ and right limit ψ . Similarly, $([\alpha_L, \alpha_R]_{i s_i}, [\alpha_L, \alpha_R]_{i+1 s_{i+1}}, \dots, [\alpha_L, \alpha_R]_{i+m s_{i+m}}, \varepsilon)_1 > ([\beta_L, \beta_R]_{g s_g}, [\beta_L, \beta_R]_{g+1 s_{g+1}}, \dots, [\beta_L, \beta_R]_{g+f s_{g+f}}, \varepsilon)_2$ is denoted by

$$\begin{aligned} \sum_{j=i}^{i+m} ((\alpha_L)_j \cdot j) &> \sum_{k=g}^{g+f} ((\beta_L)_k \cdot k), \\ \sum_{j=i}^{i+m} ((\alpha_R)_j \cdot j) + (\alpha_R)_\varepsilon \cdot n &> \sum_{k=g}^{g+f} ((\beta_R)_k \cdot k) + (\beta_R)_\varepsilon \cdot n. \end{aligned} \quad (5.17)$$

5.3 Expected Utility in Interval Fuzzy Linguistic Distribution

Because interval fuzzy linguistic distribution model provides an aggregated distribution assessment for each alternative, it is difficult for decision makers to precisely describe the ranking order among them. In such case, the notion of expected utility can be employed to compare or rank alternatives.

Similarly with expected utility in proportional fuzzy linguistic distribution, “suppose a set of alternatives X with a single-valued function $u(x)$ on X , which is called expected utility” [34].

One can represent the preference relation on X , such that for any $x, y \in X$, $x \succeq y$ if and only if $u(x) \geq u(y)$. Then, the solution to the problem of selecting x can be got by maximization of $u(x)$.

Suppose a set of evaluation grades

$$S = \{s_0, s_1, \dots, s_n\}$$

which are used as an instrument supplied to evaluators for evaluating the attribute. For an interval fuzzy linguistic distribution $([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, \varepsilon)$, assume a utility function

$$u' : S \rightarrow [0, 1]$$

satisfying

$$u'(s_{i+1}) > u'(s_i), \text{ if } s_{i+1} \text{ is preferred to } s_i.$$

Supposing alternatives a and b have a two level hierarchy with only an attribute y on the first level, and its basic attributes $E = \{e_1, e_2, \dots, e_n\}$ which is a finite set at the bottom level, as shown in Figure 5.1. If all assessments for attributes are complete, i.e., $\sum_{j=i}^{i+m} (\alpha_L)_j = 1$, or $\varepsilon = 0$, then, the expected utility of an alternative on the only attribute y is defined by

$$u_{\min}(y) = \sum_{j=i}^{i+m} (\alpha_L)_j \cdot u'(s_j) \quad (5.18)$$

$$u_{\max}(y) = \sum_{j=i}^{i+m} (\alpha_R)_j \cdot u'(s_j) \quad (5.19)$$

$$u_{\text{avg}}(y) = \frac{u_{\max}(y) + u_{\min}(y)}{2}. \quad (5.20)$$

If there is any incomplete assessment for the basic attribute, i.e., $\sum_{j=i}^{i+m} (\alpha_L)_j < 1$, or $\varepsilon > 0$, then the assessment for y is also incomplete. In such case, we need to allocate ε in order to obtain “the range of the likelihood to which y may be assessed to the evaluation grades” [87]. Without loss of generality, s_0 is supposed to be the least preferred grade which has the lowest utility and

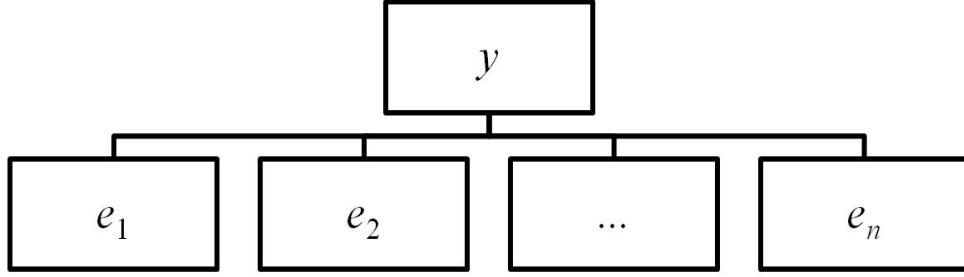


Figure 5.1. Two level hierarchy

s_n is supposed to be the most preferred grade which has the highest utility. Then, the minimum, maximum and average expected utilities on y in interval fuzzy linguistic distributions are given by

$$u_{\min}(y) = [(\alpha_L)_0 + (\alpha_L)_\varepsilon] \cdot u'(s_0) + \sum_{j=1}^n (\alpha_L)_j \cdot u'(s_j) \quad (5.21)$$

$$u_{\max}(y) = \sum_{j=0}^{n-1} (\alpha_R)_j \cdot u'(s_j) + [(\alpha_R)_n + (\alpha_R)_\varepsilon] \cdot u'(s_n) \quad (5.22)$$

$$u_{\text{avg}}(y) = \frac{u_{\max}(y) + u_{\min}(y)}{2}. \quad (5.23)$$

From formula (5.18) to formula (5.23) we can find that no matter whether the original assessments are complete or not, we can always obtain a minimum, a maximum and an average expected utility. Then, the ranking of two alternatives a and b on y is based on their utility intervals and carried out by [34]

- $a \succ_y b$ if and only if $u_{\min}(y(a)) > u_{\max}(y(b))$
- $a \sim_y b$ if and only if $u_{\min}(y(a)) = u_{\min}(y(b))$ and $u_{\max}(y(a)) = u_{\max}(y(b))$.

Otherwise, the average expected utility can be used to generate a ranking, i.e.,

- $a \succ_y b$ on an average basis, if $u_{\text{avg}}(y(a)) > u_{\text{avg}}(y(b))$.

Again, we should note that the ranking order is not reliable if the average expected utility is used. This is because there is a possibility that the special situation could happen, i.e., $u_{\text{avg}}(y(a)) > u_{\text{avg}}(y(b))$, but $u_{\max}(y(b)) > u_{\min}(y(a))$.

5.4 Interval Fuzzy Linguistic Distribution Aggregation Operators

MADM problems usually need to aggregate attributes in order to obtain an integrated value for further analysis. In this section, we introduce several aggregation operators for interval fuzzy linguistic distributions.

5.4.1 Arithmetic mean

Definition 5.1: Let $S^* = \{([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, \varepsilon)_1, ([\alpha_L, \alpha_R]_g s_g, [\alpha_L, \alpha_R]_{g+1} s_{g+1}, \dots, [\alpha_L, \alpha_R]_{g+f} s_{g+f}, \varepsilon)_2, \dots, ([\alpha_L, \alpha_R]_j s_j, [\alpha_L, \alpha_R]_{j+1} s_{j+1}, \dots, [\alpha_L, \alpha_R]_{j+q} s_{j+q}, \varepsilon)_p\}$ be a set of interval fuzzy linguistic distributions, and $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, \bar{\varepsilon})$ be the arithmetic mean represented by interval fuzzy linguistic distribution. Then, the procedure of calculating the arithmetic mean $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, \bar{\varepsilon})$ is as follows.

- 1) Take the minimum of the starting labels of interval fuzzy linguistic distributions in S^* , i.e., $k = \min(i, g, \dots, j)$.
- 2) Take the maximum of the ending labels of interval fuzzy linguistic distributions in S^* , i.e., $k + h = \max(i + m, g + f, \dots, j + q)$.
- 3) Compare the interval fuzzy linguistic distributions in S^* with arithmetic mean $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, [\gamma_L, \gamma_R]_{\bar{\varepsilon}})$. For any interval fuzzy linguistic distribution in S^* , if it is lack of corresponding linguistic terms, add “[0, 0]” as symbolic interval coefficients in front of the related linguistic terms. Thus, all the interval fuzzy linguistic distributions in S^* have the same starting labels and ending labels with arithmetic mean $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, [\gamma_L, \gamma_R]_{\bar{\varepsilon}})$, i.e., $S^* = \{([\alpha_L, \alpha_R]_k s_k, [\alpha_L, \alpha_R]_{k+1} s_{k+1}, \dots, [\alpha_L, \alpha_R]_{k+h} s_{k+h}, [\alpha_L, \alpha_R]_{\varepsilon})_1, ([\alpha_L, \alpha_R]_k s_k, [\alpha_L, \alpha_R]_{k+1} s_{k+1}, \dots, [\alpha_L, \alpha_R]_{k+h} s_{k+h}, [\alpha_L, \alpha_R]_{\varepsilon})_2, \dots, ([\alpha_L, \alpha_R]_k s_k, [\alpha_L, \alpha_R]_{k+1} s_{k+1}, \dots, [\alpha_L, \alpha_R]_{k+h} s_{k+h}, [\alpha_L, \alpha_R]_{\varepsilon})_p\}$.
- 4) The calculating process of arithmetic mean is given by

$$\left\{ \begin{array}{l} (\gamma_L)_k s_k = \left(\frac{\sum_{x=1}^p [(\alpha_L)_k]_x}{p} \right) s_k \\ (\gamma_R)_k s_k = \left(\frac{\sum_{x=1}^p [(\alpha_R)_k]_x}{p} \right) s_k \\ \vdots \\ (\gamma_L)_{k+h} s_{k+h} = \left(\frac{\sum_{x=1}^p [(\alpha_L)_{k+h}]_x}{p} \right) s_{k+h} \\ (\gamma_R)_{k+h} s_{k+h} = \left(\frac{\sum_{x=1}^p [(\alpha_R)_{k+h}]_x}{p} \right) s_{k+h} \\ (\gamma_L)_{\bar{\varepsilon}} = \frac{\sum_{x=1}^p (\alpha_L)_{\varepsilon_x}}{p} \\ (\gamma_R)_{\bar{\varepsilon}} = \frac{\sum_{x=1}^p (\alpha_R)_{\varepsilon_x}}{p} \end{array} \right. \quad (5.24)$$

where x is the No. of interval fuzzy linguistic distributions in S^* .

5.4.2 Weighted average operator

The weight average operator for interval fuzzy linguistic distributions is defined as follows.

Definition 5.2: Let $S^* = \{([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, \varepsilon)_1, ([\alpha_L, \alpha_R]_g s_g, [\alpha_L, \alpha_R]_{g+1} s_{g+1}, \dots, [\alpha_L, \alpha_R]_{g+f} s_{g+f}, \varepsilon)_2, \dots, ([\alpha_L, \alpha_R]_j s_j, [\alpha_L, \alpha_R]_{j+1} s_{j+1}, \dots, [\alpha_L, \alpha_R]_{j+q} s_{j+q}, \varepsilon)_p\}$ be a set of interval fuzzy linguistic distributions, $W = \{\omega_1, \omega_2, \dots, \omega_p\}$ be their associated weights, and $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, \bar{\varepsilon})$ be the weighted average of the set of interval fuzzy linguistic distributions. Then, the procedure of computation and aggregation of the weighted average $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, \bar{\varepsilon})$ is as follows.

- 1) Take the minimum of the starting labels of interval fuzzy linguistic distributions in S^* , i.e., $k = \min(i, g, \dots, j)$.
- 2) Take the maximum of the ending labels of interval fuzzy linguistic distributions in S^* , i.e., $k + h = \max(i + m, g + f, \dots, j + q)$.
- 3) Compare the interval fuzzy linguistic distributions in S^* with weighted average $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, [\gamma_L, \gamma_R]_{\bar{\varepsilon}})$. For any interval fuzzy linguistic distribution in S^* , if it is lack of corresponding linguistic terms, add “[0, 0]” as symbolic

interval coefficients in front of the related linguistic terms. Thus, all the interval fuzzy linguistic distributions in S^* have the same starting labels and ending labels with weighted average $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, [\gamma_L, \gamma_R]_{\bar{\varepsilon}})$, i.e., $S^* = \{([\alpha_L, \alpha_R]_k s_k, [\alpha_L, \alpha_R]_{k+1} s_{k+1}, \dots, [\alpha_L, \alpha_R]_{k+h} s_{k+h}, [\alpha_L, \alpha_R]_{\bar{\varepsilon}})_1, ([\alpha_L, \alpha_R]_k s_k, [\alpha_L, \alpha_R]_{k+1} s_{k+1}, \dots, [\alpha_L, \alpha_R]_{k+h} s_{k+h}, [\alpha_L, \alpha_R]_{\bar{\varepsilon}})_2, \dots, ([\alpha_L, \alpha_R]_k s_k, [\alpha_L, \alpha_R]_{k+1} s_{k+1}, \dots, [\alpha_L, \alpha_R]_{k+h} s_{k+h}, [\alpha_L, \alpha_R]_{\bar{\varepsilon}})_p\}$.

4) The calculating process of weighted average is given by

$$\left\{ \begin{array}{l} (\gamma_L)_k s_k = \left(\frac{\sum_{x=1}^p [(\alpha_L)_k]_x \cdot \omega_x}{\sum_{x=1}^p \omega_x} \right) s_k \\ (\gamma_R)_k s_k = \left(\frac{\sum_{x=1}^p [(\alpha_R)_k]_x \cdot \omega_x}{\sum_{x=1}^p \omega_x} \right) s_k \\ \vdots \\ (\gamma_L)_{k+h} s_{k+h} = \left(\frac{\sum_{x=1}^p [(\alpha_L)_{k+h}]_x \cdot \omega_x}{\sum_{x=1}^p \omega_x} \right) s_{k+h} \\ (\gamma_R)_{k+h} s_{k+h} = \left(\frac{\sum_{x=1}^p [(\alpha_R)_{k+h}]_x \cdot \omega_x}{\sum_{x=1}^p \omega_x} \right) s_{k+h} \\ (\gamma_L)_{\bar{\varepsilon}} = \frac{\sum_{x=1}^p (\alpha_L)_{\bar{\varepsilon}_x} \cdot \omega_x}{\sum_{x=1}^p \omega_x} \\ (\gamma_R)_{\bar{\varepsilon}} = \frac{\sum_{x=1}^p (\alpha_R)_{\bar{\varepsilon}_x} \cdot \omega_x}{\sum_{x=1}^p \omega_x} \end{array} \right. \quad (5.25)$$

where x is the No. of interval fuzzy linguistic distributions in S^* .

In order to clearly describe the procedure of calculating the weighted average of interval fuzzy linguistic distributions, we give an example here.

Example: Suppose a set of interval fuzzy linguistic distributions $S^* = \{([0.2, 0.4] s_2, [0.7, 0.8] s_3, [0, 0.1]), ([0.1, 0.2] s_3, [0.8, 0.9] s_4, [0, 0.1])\}$. Its associated weights are $W = \{0.4, 0.6\}$. Then, the procedure of calculating the weighted average $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, [\gamma_L, \gamma_R]_{\bar{\varepsilon}})$ is as follows.

- 1) $k = \min(2, 3) = 2$.
- 2) $k + h = \max(3, 4) = 4$.
- 3) For any interval fuzzy linguistic distribution in S^* , add “[0, 0]” as symbolic interval coefficients in front of the related linguistic terms if it is lack of corresponding linguistic

terms. Then, $S^* = \{([0.2, 0.4]_{s_2}, [0.7, 0.8]_{s_3}, [0, 0]_{s_4}, [0, 0.1]), ([0, 0]_{s_2}, [0.1, 0.2]_{s_3}, [0.8, 0.9]_{s_4}, [0, 0.1])\}$.

4) The weighted average can be obtained by

$$\left\{ \begin{array}{l} (\gamma_L)_{2s_2} = \left(\frac{0.2 \times 0.4 + 0 \times 0.6}{0.4 + 0.6} \right) s_2 = (0.08)_L s_2 \\ (\gamma_R)_{2s_2} = \left(\frac{0.4 \times 0.4 + 0 \times 0.6}{0.4 + 0.6} \right) s_2 = (0.16)_R s_2 \\ (\gamma_L)_{3s_3} = \left(\frac{0.7 \times 0.4 + 0.1 \times 0.6}{0.4 + 0.6} \right) s_3 = (0.34)_L s_3 \\ (\gamma_R)_{3s_3} = \left(\frac{0.8 \times 0.4 + 0.2 \times 0.6}{0.4 + 0.6} \right) s_3 = (0.44)_R s_3 \\ (\gamma_L)_{4s_4} = \left(\frac{0 \times 0.4 + 0.8 \times 0.6}{0.4 + 0.6} \right) s_4 = (0.48)_L s_4 \\ (\gamma_R)_{4s_4} = \left(\frac{0 \times 0.4 + 0.9 \times 0.6}{0.4 + 0.6} \right) s_4 = (0.54)_R s_4 \\ (\gamma_L)_{\bar{\varepsilon}} = \frac{0 \times 0.4 + 0 \times 0.6}{0.4 + 0.6} = (0)_L \\ (\gamma_R)_{\bar{\varepsilon}} = \frac{0.1 \times 0.4 + 0.1 \times 0.6}{0.4 + 0.6} = (0.1)_R \end{array} \right.$$

Therefore, the weighted average of the two interval fuzzy linguistic distributions is $([0.08, 0.16]_{s_2}, [0.34, 0.44]_{s_3}, [0.48, 0.54]_{s_4}, [0, 0.1])$.

5.4.3 Interval weighted average operator

According to the general interval weighted average operator (5.6), the interval weighted average operator for interval fuzzy linguistic distributions is defined as follows.

Definition 5.3: Let $S^* = \{([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, \varepsilon)_1, ([\alpha_L, \alpha_R]_g s_g, [\alpha_L, \alpha_R]_{g+1} s_{g+1}, \dots, [\alpha_L, \alpha_R]_{g+f} s_{g+f}, \varepsilon)_2, \dots, ([\alpha_L, \alpha_R]_j s_j, [\alpha_L, \alpha_R]_{j+1} s_{j+1}, \dots, [\alpha_L, \alpha_R]_{j+q} s_{j+q}, \varepsilon)_p\}$ be a set of interval fuzzy linguistic distributions, $W = \{[\omega_L, \omega_R]_1, [\omega_L, \omega_R]_2, \dots, [\omega_L, \omega_R]_p\}$ be the associated interval weights, and $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, \bar{\varepsilon})$ be the interval weighted average of the set of interval fuzzy linguistic distributions. Then, the procedure of computation and aggregation of the interval weighted average $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, \bar{\varepsilon})$ is as follows.

- 1) Take the minimum of the starting labels of interval fuzzy linguistic distributions in S^* , i.e., $k = \min(i, g, \dots, j)$.

- 2) Take the maximum of the ending labels of interval fuzzy linguistic distributions in S^* , i.e., $k + h = \max(i + m, g + f, \dots, j + q)$.
- 3) Compare the interval fuzzy linguistic distributions in S^* with the interval weighted average $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, [\gamma_L, \gamma_R]_{\bar{\varepsilon}})$. For any interval fuzzy linguistic distribution in S^* , if it is lack of corresponding linguistic terms, add "[0, 0]" as symbolic interval coefficients in front of the related linguistic terms in order to make all the interval fuzzy linguistic distributions in S^* have the same starting labels and ending labels with interval weighted average $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, [\gamma_L, \gamma_R]_{\bar{\varepsilon}})$, i.e., $S^* = \{([\alpha_L, \alpha_R]_k s_k, [\alpha_L, \alpha_R]_{k+1} s_{k+1}, \dots, [\alpha_L, \alpha_R]_{k+h} s_{k+h}, [\alpha_L, \alpha_R]_{\bar{\varepsilon}})_1, ([\alpha_L, \alpha_R]_k s_k, [\alpha_L, \alpha_R]_{k+1} s_{k+1}, \dots, [\alpha_L, \alpha_R]_{k+h} s_{k+h}, [\alpha_L, \alpha_R]_{\bar{\varepsilon}})_2, \dots, ([\alpha_L, \alpha_R]_k s_k, [\alpha_L, \alpha_R]_{k+1} s_{k+1}, \dots, [\alpha_L, \alpha_R]_{k+h} s_{k+h}, [\alpha_L, \alpha_R]_{\bar{\varepsilon}})_p\}$.
- 4) The calculating process of interval weighted average is given by

$$\left\{ \begin{array}{l} (\gamma_L)_k s_k = \min_{\forall \omega_x \in [\omega_L, \omega_R]_x} \left(\frac{\sum_{x=1}^p [(\alpha_L)_k]_x \cdot \omega_x}{\sum_{x=1}^p \omega_x} \right) s_k \\ (\gamma_R)_k s_k = \max_{\forall \omega_x \in [\omega_L, \omega_R]_x} \left(\frac{\sum_{x=1}^p [(\alpha_R)_k]_x \cdot \omega_x}{\sum_{x=1}^p \omega_x} \right) s_k \\ \vdots \\ (\gamma_L)_{k+h} s_{k+h} = \min_{\forall \omega_x \in [\omega_L, \omega_R]_x} \left(\frac{\sum_{x=1}^p [(\alpha_L)_{k+h}]_x \cdot \omega_x}{\sum_{x=1}^p \omega_x} \right) s_{k+h} \\ (\gamma_R)_{k+h} s_{k+h} = \max_{\forall \omega_x \in [\omega_L, \omega_R]_x} \left(\frac{\sum_{x=1}^p [(\alpha_R)_{k+h}]_x \cdot \omega_x}{\sum_{x=1}^p \omega_x} \right) s_{k+h} \\ (\gamma_L)_{\bar{\varepsilon}} = \min_{\forall \omega_x \in [\omega_L, \omega_R]_x} \frac{\sum_{x=1}^p (\alpha_L)_{\bar{\varepsilon}_x} \cdot \omega_x}{\sum_{x=1}^p \omega_x} \\ (\gamma_R)_{\bar{\varepsilon}} = \max_{\forall \omega_x \in [\omega_L, \omega_R]_x} \frac{\sum_{x=1}^p (\alpha_R)_{\bar{\varepsilon}_x} \cdot \omega_x}{\sum_{x=1}^p \omega_x} \end{array} \right. \quad (5.26)$$

where x is the No. of interval fuzzy linguistic distributions in S^* .

- 5) Make use of KM algorithms [41] or Enhanced KM algorithms [77], the interval weighted average for interval fuzzy linguistic distributions can be easily obtained.

5.4.4 Interval ordered weighted average operator

According to the general interval ordered weighted average operator (5.7), the interval ordered weighted average operator for interval fuzzy linguistic distributions is defined as follows.

Definition 5.4: Let $S^* = \{([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, \varepsilon)_1, ([\alpha_L, \alpha_R]_g s_g, [\alpha_L, \alpha_R]_{g+1} s_{g+1}, \dots, [\alpha_L, \alpha_R]_{g+f} s_{g+f}, \varepsilon)_2, \dots, ([\alpha_L, \alpha_R]_j s_j, [\alpha_L, \alpha_R]_{j+1} s_{j+1}, \dots, [\alpha_L, \alpha_R]_{j+q} s_{j+q}, \varepsilon)_p\}$ be a set of ordered interval fuzzy linguistic distributions, with $([\alpha_L, \alpha_R]_i s_i, [\alpha_L, \alpha_R]_{i+1} s_{i+1}, \dots, [\alpha_L, \alpha_R]_{i+m} s_{i+m}, \varepsilon)_1 > ([\alpha_L, \alpha_R]_g s_g, [\alpha_L, \alpha_R]_{g+1} s_{g+1}, \dots, [\alpha_L, \alpha_R]_{g+f} s_{g+f}, \varepsilon)_2 > \dots > ([\alpha_L, \alpha_R]_j s_j, [\alpha_L, \alpha_R]_{j+1} s_{j+1}, \dots, [\alpha_L, \alpha_R]_{j+q} s_{j+q}, \varepsilon)_p$, $W = \{[\omega_L, \omega_R]_1, [\omega_L, \omega_R]_2, \dots, [\omega_L, \omega_R]_p\}$ be their associated interval weights, and $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, \bar{\varepsilon})$ be the interval ordered weighted average of the set of interval fuzzy linguistic distributions. Then, the procedure of computation and aggregation of the interval ordered weighted average $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, \bar{\varepsilon})$ is as follows.

- 1) Take the minimum of the starting labels of interval fuzzy linguistic distributions in S^* , i.e., $k = \min(i, g, \dots, j)$.
- 2) Take the maximum of the ending labels of interval fuzzy linguistic distributions in S^* , i.e., $k + h = \max(i + m, g + f, \dots, j + q)$.
- 3) Compare the interval fuzzy linguistic distributions in S^* with the interval ordered weighted average $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, [\gamma_L, \gamma_R]_{\bar{\varepsilon}})$. For any interval fuzzy linguistic distribution in S^* , if it is lack of corresponding linguistic terms, add "[0, 0]" as symbolic interval coefficients in front of the related linguistic terms in order to make all the interval fuzzy linguistic distributions in S^* have the same starting labels and ending labels with the interval ordered weighted average $([\gamma_L, \gamma_R]_k s_k, [\gamma_L, \gamma_R]_{k+1} s_{k+1}, \dots, [\gamma_L, \gamma_R]_{k+h} s_{k+h}, [\gamma_L, \gamma_R]_{\bar{\varepsilon}})$, i.e., $S^* = \{([\alpha_L, \alpha_R]_k s_k, [\alpha_L, \alpha_R]_{k+1} s_{k+1}, \dots, [\alpha_L, \alpha_R]_{k+h} s_{k+h}, [\alpha_L, \alpha_R]_{\varepsilon})_1, ([\alpha_L, \alpha_R]_k s_k, [\alpha_L, \alpha_R]_{k+1} s_{k+1}, \dots, [\alpha_L, \alpha_R]_{k+h} s_{k+h}, [\alpha_L, \alpha_R]_{\varepsilon})_2, \dots, ([\alpha_L, \alpha_R]_k s_k, [\alpha_L, \alpha_R]_{k+1} s_{k+1}, \dots, [\alpha_L, \alpha_R]_{k+h} s_{k+h}, [\alpha_L, \alpha_R]_{\varepsilon})_p\}$.
- 4) The calculating process of interval ordered weighted average is given by

$$\left\{ \begin{array}{l} (\gamma_L)_k s_k = \min_{\forall \omega_x \in [\omega_L, \omega_R]_x} \left(\frac{\sum_{x=1}^p [(\alpha_L)_k]_x \cdot \omega_x}{\sum_{x=1}^p \omega_x} \right) s_k \\ (\gamma_R)_k s_k = \max_{\forall \omega_x \in [\omega_L, \omega_R]_x} \left(\frac{\sum_{x=1}^p [(\alpha_R)_k]_x \cdot \omega_x}{\sum_{x=1}^p \omega_x} \right) s_k \\ \vdots \\ (\gamma_L)_{k+h} s_{k+h} = \min_{\forall \omega_x \in [\omega_L, \omega_R]_x} \left(\frac{\sum_{x=1}^p [(\alpha_L)_{k+h}]_x \cdot \omega_x}{\sum_{x=1}^p \omega_x} \right) s_{k+h} \\ (\gamma_R)_{k+h} s_{k+h} = \max_{\forall \omega_x \in [\omega_L, \omega_R]_x} \left(\frac{\sum_{x=1}^p [(\alpha_R)_{k+h}]_x \cdot \omega_x}{\sum_{x=1}^p \omega_x} \right) s_{k+h} \\ (\gamma_L)_{\bar{\varepsilon}} = \min_{\forall \omega_x \in [\omega_L, \omega_R]_x} \frac{\sum_{x=1}^p (\alpha_L)_{\bar{\varepsilon}_x} \cdot \omega_x}{\sum_{x=1}^p \omega_x} \\ (\gamma_R)_{\bar{\varepsilon}} = \max_{\forall \omega_x \in [\omega_L, \omega_R]_x} \frac{\sum_{x=1}^p (\alpha_R)_{\bar{\varepsilon}_x} \cdot \omega_x}{\sum_{x=1}^p \omega_x} \end{array} \right. \quad (5.27)$$

where x is the No. of interval fuzzy linguistic distributions in S^* .

- 5) Make use of KM algorithms [41] or Enhanced KM algorithms [77], the interval weighted average for interval fuzzy linguistic distributions can be easily obtained.

5.5 Illustration Examples

In this section, we apply interval fuzzy linguistic distribution model to deal with two illustration examples so as to demonstrate its capability of dealing with MADA problems. Because this model is complicated, we first use a simple example to explain it. Then, we use this model to deal with a practical application, a cargo ship selection problem taken from [74], in order to compare the final result with the extended evidential reasoning approach [74].

5.5.1 The description of the first example

Supposing an evaluator wants to choose an alternative between E_1 and E_2 , which are compared on the basis of three basic attributes, as shown in Table 5.1. The relative weights of the three basic attributes are given as: (0.4, 0.3, 0.3). The evaluator assesses these attributes according to the linguistic term set of distinct evaluation grades which is defined as follows:

Table 5.1. Interval fuzzy linguistic distribution assessments for two alternatives

Attributes	E_1	E_2	Weights
C_1	$([0.2, 0.3]_{s_2}, [0.7, 0.8]_{s_3}, [0, 0.1])$	$([0.1, 0.3]_{s_1}, [0.7, 0.8]_{s_2}, [0, 0.2])$	0.4
C_2	$([0.1, 0.2]_{s_3}, [0.6, 0.9]_{s_4}, [0, 0.3])$	$([0.5, 0.6]_{s_3}, [0.4, 0.5]_{s_4}, [0, 0.1])$	0.3
C_3	$([0.3, 0.5]_{s_0}, [0.5, 0.7]_{s_1}, [0, 0.2])$	$([0.1, 0.2]_{s_0}, [0.8, 0.9]_{s_1}, [0, 0.1])$	0.3

Table 5.2. The aggregation results of two alternatives

The overall aggregation results	
E_1	$([0.09, 0.15]_{s_0}, [0.15, 0.21]_{s_1}, [0.08, 0.12]_{s_2}, [0.31, 0.38]_{s_3}, [0.18, 0.27]_{s_4}, [0, 0.19])$
E_2	$([0.03, 0.06]_{s_0}, [0.28, 0.39]_{s_1}, [0.28, 0.32]_{s_2}, [0.15, 0.18]_{s_3}, [0.12, 0.15]_{s_4}, [0, 0.14])$

Table 5.3. The distributed assessments on two alternatives

	Poor	Indifferent	Average	Good	Excellent	ϵ
E_1	[0.09, 0.15]	[0.15, 0.21]	[0.08, 0.12]	[0.31, 0.38]	[0.18, 0.27]	[0, 0.19]
E_2	[0.03, 0.06]	[0.28, 0.39]	[0.28, 0.32]	[0.15, 0.18]	[0.12, 0.15]	[0, 0.14]

$$S_1 = \{s_0(\text{Poor}), s_1(\text{Indifferent}), s_2(\text{Average}), s_3(\text{Good}), s_4(\text{Excellent})\}. \quad (5.28)$$

Because of the uncertainty, the evaluator may use intervals as his/her confidence levels if he/she feels difficult to give precise assessments. Then, the linguistic assessments presented by interval linguistic distributions are shown in Table 5.1.

5.5.2 Aggregating assessments of the first example via interval fuzzy linguistic distribution model

According to the weighted average operator for interval fuzzy linguistic distributions, aggregate the three attributes by formula (5.25). Thus, we can obtain the aggregation results of two alternatives represented by interval fuzzy linguistic distributions, as shown in Table 5.2. Then, convert the aggregation results of two alternatives into the corresponding linguistic terms of distinct evaluation grades, which are shown in Table 5.3.

Table 5.4. The expected utilities of two alternatives

	Minimum utility	Maximum utility	Average utility
E_1	0.490	0.858	0.674
E_2	0.443	0.683	0.563

5.5.3 Computing the expected utilities of the first example

It is very difficult to precisely describe the ranking orders among the two alternatives from their distinct evaluation grades as shown in Table 5.3. In such situation, the expected utilities of two alternatives should be calculated. Define the utilities of the five individual evaluation grades as

$$\begin{aligned} u'(P) &= 0, u'(I) = 0.25, u'(A) = 0.5, \\ u'(G) &= 0.75, u'(E) = 1. \end{aligned}$$

Then, the expected utilities of two alternatives can be obtained via formula (5.21), (5.22) and (5.23), which are shown in Table 5.4, and graphically in Figure 5.2. We can find that the average expected utility of E_1 is larger than that of E_2 from Table 5.4. Therefore, according to average expected utility, E_1 is preferred to E_2 . However, as we mentioned in previous section, this is not reliable because the maximum expected utility of E_2 is larger than the minimum expected utility of E_1 .

We just used a very simple example to explain interval fuzzy linguistic distribution model from a easily comprehensible perspective. In order to test its capability of dealing with MADM problems with incomplete information from a practical perspective, we apply this model to a cargo ship selection problem taken from [74].

5.5.4 The description of a cargo ship selection problem

The problem is to “consider a cargo ship selection problem with six competing cargo ship designs, which are compared on the basis of nine basic attributes shown in Table 5.5, where *Load factor* and *Effective load factor* are two qualitative attributes and *Bale capacity*, *Deadweight*, *Speed*, *Capital investment*, *Annual M & R and manning costs*, *Sea fuel consumption* and *Off-hire* are

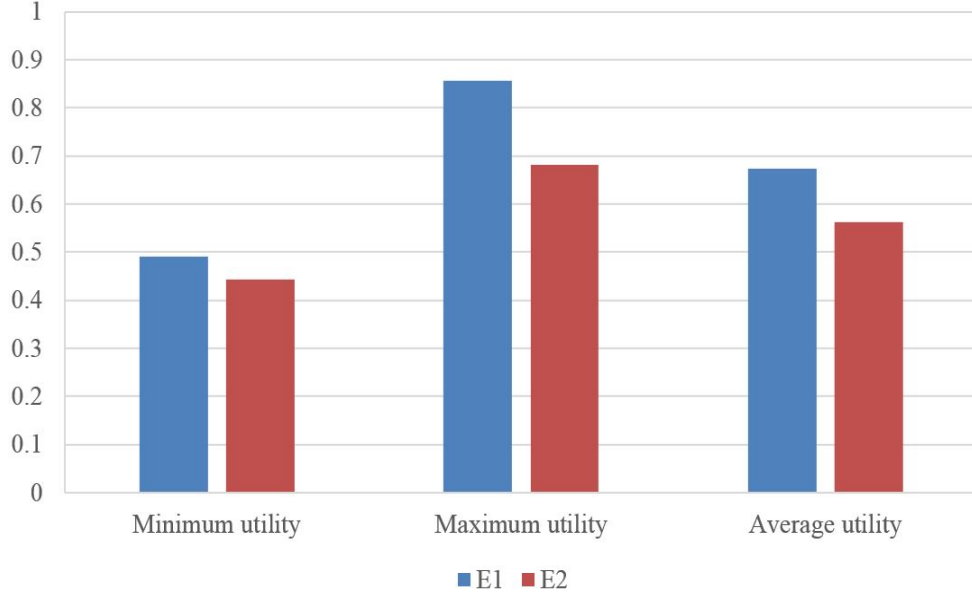


Figure 5.2. The expected utilities of two alternatives

seven quantitative attributes. Among the nine basic attributes, the first five are benefit attributes and the others are cost attributes” [74]. Because the exact values of *Sea fuel consumption* and *Off-hire* are not known, they are estimated by using interval numbers. “The relative weights of the nine basic attributes are given as: (12, 15, 10, 10, 12, 15, 10, 10, 6), which are normalized as: (0.12, 0.15, 0.1, 0.1, 0.12, 0.15, 0.1, 0.1, 0.06)” [74].

The linguistic term set of distinct evaluation grades used for assessing a cargo ship is defined as follows:

$$S_2 = \{s_0(\text{Poor}), s_1(\text{Average}), s_2(\text{Good}), s_3(\text{Very Good}), s_4(\text{Excellent})\}. \quad (5.29)$$

“Two qualitative attributes, *Load factor* and *Effective load factor*, are both assessed by the above set of assessment grades” [74]. Meanwhile, evaluator may use intervals as his/her confidence levels if he/she feels difficult to give precise assessments. Then, the linguistic assessments and original quantitative data are presented in Table 5.5, “where *P*, *A*, *G*, *V* and *E* are the abbreviations of the evaluation grades of *Poor*, *Average*, *Good*, *Very Good* and *Excellent* respectively” [74].

The quantitative data in Table 5.5 have to be modeled using proportional or interval confidence

Table 5.5. Original assessment data for six cargo ships [74]

Attributes	Ship 1	Ship 2	Ship 3	Ship 4	Ship 5	Ship 6
Load factor (12)	$\{(V, [0.2, 0.4]), (E, [0.7, 0.8])\}$	$\{(A, [0.4, 0.5]), (G, [0.55, 0.65])\}$	$\{(G, [0.3, 0.5]), (V, [0.5, 0.6])\}$	$\{(V, [0.1, 0.15]), (E, [0.8, 0.95])\}$	$\{(A, [0.3, 0.4]), (G, [0.5, 0.6])\}$	$\{(P, [0.5, 0.65]), (A, [0.4, 0.5])\}$
Bale capacity (M ³) (15)	21070	28000	31587	26718	28874	33117
Deadweight (t) (10)	17200	23200	22351	20621	22233	25500
Speed (km) (10)	14.3	14.8	17.7	15.0	18.2	17.6
Effective load factor (%) (12)	$\{(V, [0.1, 0.2]), (E, [0.8, 0.9])\}$	$\{(G, [0.2, 0.4]), (V, [0.6, 0.7])\}$	$\{(A, [0.4, 0.5]), (G, [0.5, 0.6])\}$	$\{(V, [0.15, 0.2]), (E, [0.8, 0.85])\}$	$\{(A, [0.5, 0.6]), (G, [0.4, 0.5])\}$	$\{(P, [0.6, 0.7]), (A, [0.3, 0.5])\}$
Capital investment (pound M) (15)	16.66	20.00	31.33	16.66	25.00	25.00
Annual M & R and manning costs (pound M) (10)	0.95	1.005	1.045	0.97	1.10	1.09
Sea fuel consumption (tons per day) (10)	[18, 20.1]	[23, 25.3]	[37, 39]	[20.2, 22.1]	[38, 40]	[38.5, 40.5]
Off-hire (days per year) (6)	[15, 17]	[16, 18]	[18, 20]	[18, 21]	[19, 22]	[17, 19]

levels first. In [74], the authors have modeled these quantitative data according to principle of utility equivalence (for more details, see [74]), which are shown in Table 5.6, where I_1, I_2 are two 0-1 integer variables generated during the process of modeling quantitative data, with $I_1 + I_2 = 0$, and $I_1, I_2 = 0$ or 1; H , which represents the extent of ignoring information in [74] will be replaced by ε in the sequel.

5.5.5 Aggregating assessments of ship selection problem via interval fuzzy linguistic distribution model

After modeling the quantitative data using proportional or interval confidence levels, the evaluation procedure based on interval fuzzy linguistic distribution model should be carried out and is described as follows.

(1) Interval fuzzy linguistic distributions transformation: According to the linguistic term set of distinct evaluation grades, the distribution assessment matrix shown in Table 5.6 should be converted into corresponding interval fuzzy linguistic distributions by using symbolic translation value of $s_i, i = 0, 1, \dots, 4$ and the statements with the associated representation method such as 1)-2) discussed in Section 5.1.2. Then, the decision matrix for the six cargo ships represented by interval fuzzy linguistic distributions is shown in Table 5.7, where s_0, s_1, s_2, s_3 , and s_4 are the expressions of *Poor*, *Average*, *Good*, *Very Good*, and *Excellent*, respectively, and the numerical or interval coefficients in front of s_0, s_1, s_2, s_3 , and s_4 denote the confidence levels to which degree an attribute is assessed to a grade.

(2) Interval fuzzy linguistic distributions computation and aggregation: According to formula (5.25), aggregate these nine attributes from ship 1 to ship 6. Finally, we can obtain the final results of the overall belief degrees of six cargo ships represented by interval fuzzy linguistic distributions, which are shown in Table 5.8.

(3) Interval fuzzy linguistic distributions conversion: Convert the overall belief degrees of six cargo ships represented by interval fuzzy linguistic distributions into the corresponding linguistic terms of distinct evaluation grades, which are shown in Table 5.9.

Table 5.6. Distribution assessment matrix for the six cargo ships [74]

Attributes	Ship 1	Ship 2	Ship 3	Ship 4	Ship 5	Ship 6
Load factor	$\{(V, [0.2, 0.4]), (E, [0.7, 0.8]), (H, [0, 0.1])\}$	$\{(A, [0.4, 0.5]), (G, [0.55, 0.65]), (H, [0, 0.05])\}$	$\{(G, [0.3, 0.5]), (V, [0.5, 0.6]), (H, [0, 0.2])\}$	$\{(V, [0.1, 0.15]), (E, [0.8, 0.95]), (H, [0, 0.1])\}$	$\{(A, [0.3, 0.4]), (G, [0.5, 0.6]), (H, [0, 0.2])\}$	$\{(P, [0.5, 0.65]), (A, [0.4, 0.5]), (H, [0, 0.1])\}$
Bale capacity	$\{(P, 0.8217), (A, 0.1783)\}$	$\{(A, 0.3333), (G, 0.6667)\}$	$\{(G, 0.1377), (V, 0.8623)\}$	$\{(A, 0.7607), (G, 0.2393)\}$	$\{(A, 0.042), (G, 0.958)\}$	$\{(V, 0.6277), (E, 0.3723)\}$
Deadweight	$\{(P, 0.9444), (A, 0.0556)\}$	$\{(G, 0.5556), (V, 0.4444)\}$	$\{(A, 0.0272), (G, 0.9728)\}$	$\{(A, 0.9883), (G, 0.0117)\}$	$\{(A, 0.0928), (G, 0.9072)\}$	$\{(V, 0.2778), (E, 0.7222)\}$
Speed	$\{(P, 0.85), (A, 0.15)\}$	$\{(P, 0.6), (A, 0.4)\}$	$\{(G, 0.3), (V, 0.7)\}$	$\{(P, 0.5), (A, 0.5)\}$	$\{(V, 0.8), (E, 0.2)\}$	$\{(G, 0.4), (V, 0.6)\}$
Effective load factor	$\{(V, [0.1, 0.2]), (E, [0.8, 0.9]), (H, [0, 0.1])\}$	$\{(G, [0.2, 0.4]), (V, [0.6, 0.7]), (H, [0, 0.2])\}$	$\{(A, [0.4, 0.5]), (G, [0.5, 0.6]), (H, [0, 0.1])\}$	$\{(V, [0.15, 0.2]), (E, [0.8, 0.85]), (H, [0, 0.05])\}$	$\{(A, [0.5, 0.6]), (G, [0.4, 0.5]), (H, [0, 0.1])\}$	$\{(P, [0.6, 0.7]), (A, [0.3, 0.5]), (H, [0, 0.1])\}$
Capital investment	$\{(V, 0.4882), (E, 0.5118)\}$	$\{(G, 0.4706), (V, 0.5294)\}$	$\{(P, 0.9015), (A, 0.0985)\}$	$\{(V, 0.4882), (E, 0.5118)\}$	$\{(A, 0.9412), (G, 0.0588)\}$	$\{(A, 0.9412), (G, 0.0588)\}$
Annual M & R and manning costs	$\{(E, 1.0)\}$	$\{(G, 0.375), (V, 0.625)\}$	$\{(A, 0.375), (G, 0.625)\}$	$\{(V, 0.5), (E, 0.5)\}$	$\{(P, 0.375), (A, 0.625)\}$	$\{(P, 0.25), (A, 0.75)\}$
Sea fuel consumption	$\{(V, [0, 0.4565]), (E, [0.5435, 1])\}$	$\{(G, [0.087, 0.587]), (V, [0.413, 0.913])\}$	$\{(P, [0.5652, 0.7826]), (A, [0.2174, 0.4348])\}$	$\{(V, [0.4783, 0.8913]), (E, [0.1087, 0.5217])\}$	$\{(P, [0.6739, 0.8913]), (A, [0.1087, 0.3261])\}$	$\{(P, [0.7283, 0.9457]), (A, [0.0543, 0.2717])\}$
Off-hire	$\{(G, [0, 0.5I_1]), (V, [0, I_1 + I_2]), (E, [0, I_2])\}$	$\{(G, [0, 1]), (V, [0, 1])\}$	$\{(P, [0, 0.333I_1]), (A, [0, I_1 + I_2]), (G, [0, I_2])\}$	$\{(P, [0, 0.6667I_1]), (A, [0, I_1 + I_2]), (G, [0, I_2])\}$	$\{(P, [0, 1]), (A, [0, 1])\}$	$\{(A, [0, I_1]), (G, [0, I_1 + I_2]), (V, [0, 0.5I_2])\}$

Table 5.7. Distribution assessment matrix for
the six cargo ships represented by interval fuzzy linguistic distributions

Attributes	Ship 1	Ship 2	Ship 3	Ship 4	Ship 5	Ship 6
Load factor	$([0.2, 0.4]s_3,$ $[0.7, 0.8]s_4,$ $[0, 0.1])$	$([0.4, 0.5]s_1,$ $[0.55, 0.65]s_2,$ $[0, 0.05])$	$([0.3, 0.5]s_2,$ $[0.5, 0.6]s_3,$ $[0, 0.2])$	$([0.1, 0.15]s_3,$ $[0.8, 0.95]s_4,$ $[0, 0.1])$	$([0.3, 0.4]s_1,$ $[0.5, 0.6]s_2,$ $[0, 0.2])$	$([0.5, 0.65]s_0,$ $[0.4, 0.5]s_1,$ $[0, 0.1])$
Bale capacity	$(0.8217s_0,$ $0.1783s_1, 0)$	$(0.3333s_1,$ $0.6667s_2, 0)$	$(0.1377s_2,$ $0.8623s_3, 0)$	$(0.7607s_1,$ $0.2393s_2, 0)$	$(0.042s_1,$ $0.958s_2, 0)$	$(0.6277s_3,$ $0.3723s_4, 0)$
Deadweight	$(0.9444s_0,$ $0.0556s_1, 0)$	$(0.5556s_2,$ $0.4444s_3, 0)$	$(0.0272s_1,$ $0.9728s_2, 0)$	$(0.9883s_1,$ $0.0117s_2, 0)$	$(0.0928s_1,$ $0.9072s_2, 0)$	$(0.2778s_3,$ $0.7222s_4, 0)$
Speed	$(0.85s_0, 0.15s_1, 0)$	$(0.6s_0, 0.4s_1, 0)$	$(0.3s_2, 0.7s_3, 0)$	$(0.5s_0, 0.5s_1, 0)$	$(0.8s_3, 0.2s_4, 0)$	$(0.4s_2, 0.6s_3, 0)$
Effective load factor	$([0.1, 0.2]s_3,$ $[0.8, 0.9]s_4,$ $[0, 0.1])$	$([0.2, 0.4]s_2,$ $[0.6, 0.7]s_3,$ $[0, 0.2])$	$([0.4, 0.5]s_1,$ $[0.5, 0.6]s_2,$ $[0, 0.1])$	$([0.15, 0.2]s_3,$ $[0.8, 0.85]s_4,$ $[0, 0.05])$	$([0.5, 0.6]s_1,$ $[0.4, 0.5]s_2,$ $[0, 0.1])$	$([0.6, 0.7]s_0,$ $[0.3, 0.5]s_1,$ $[0, 0.1])$
Capital investment	$(0.4882s_3,$ $0.5118s_4, 0)$	$(0.4706s_2,$ $0.5294s_3, 0)$	$(0.9015s_0,$ $0.0985s_1, 0)$	$(0.4882s_3,$ $0.5118s_4, 0)$	$(0.9412s_1,$ $0.0588s_2, 0)$	$(0.9412s_1,$ $0.0588s_2, 0)$
Annual M & R and manning costs	$(1s_3, 0)$	$(0.375s_2,$ $0.625s_3, 0)$	$(0.375s_1,$ $0.625s_2, 0)$	$(0.5s_3,$ $0.5s_4, 0)$	$(0.375s_0,$ $0.625s_1, 0)$	$(0.25s_0,$ $0.75s_1, 0)$
Sea fuel consumption	$([0, 0.4565]s_3,$ $[0.5435, 1]s_4, 0)$	$([0.087, 0.587]s_2,$ $[0.413, 0.913]s_3, 0)$	$([0.5652, 0.7826]s_0,$ $[0.2174, 0.4348]s_1, 0)$	$([0.4783, 0.8913]s_3,$ $[0.1087, 0.5217]s_4, 0)$	$([0.6739, 0.8913]s_0,$ $[0.1087, 0.3261]s_1, 0)$	$([0.7283, 0.9457]s_0,$ $[0.0543, 0.2717]s_1, 0)$
Off-hire	$([0, 0.5I_1]s_2,$ $[0, I_1 + I_2]s_3,$ $[0, I_2]s_4, 0)$	$([0, 1]s_2,$ $[0, 1]s_3, 0)$	$([0, 0.3333I_1]s_0,$ $[0, I_1 + I_2]s_1,$ $[0, I_2]s_2, 0)$	$([0, 0.6667I_1]s_0,$ $[0, I_1 + I_2]s_1,$ $[0, I_2]s_2, 0)$	$([0, 1]s_0,$ $([0, 1]s_1, 0)$	$([0, I_1]s_1,$ $[0, I_1 + I_2]s_2,$ $[0, 0.5I_2]s_3, 0)$

Table 5.8. The overall belief degrees of six cargo ships represented by interval fuzzy linguistic distributions

The overall belief degrees of six cargo ships	
Ship 1	$(0.303s_0, 0.047s_1, [0.073, 0.073 + 0.03I_1]s_2, [0.209, 0.351]s_3, [0.311, 0.381 + 0.06I_2]s_4, [0, 0.024])$
Ship 2	$(0.060s_0, [0.138, 0.150]s_1, [0.362, 0.508]s_2, [0.300, 0.422]s_3, [0, 0.03])$
Ship 3	$([0.192, 0.213 + 0.02I_1]s_0, [0.125, 0.218]s_1, [0.306, 0.342 + 0.06I_2]s_2, [0.259, 0.271]s_3, [0, 0.036])$
Ship 4	$([0.050, 0.05 + 0.04I_1]s_0, [0.263, 0.323]s_1, [0.037, 0.037 + 0.06I_2]s_2, [0.201, 0.254]s_3, [0.330, 0.395]s_4, [0, 0.018])$
Ship 5	$([0.105, 0.187]s_0, [0.326, 0.432]s_1, [0.351, 0.375]s_2, 0.080s_3, 0.020s_4, [0, 0.036])$
Ship 6	$([0.230, 0.282]s_0, [0.306, 0.363 + 0.06I_1]s_1, [0.049, 0.109]s_2, [0.182, 0.182 + 0.03I_2]s_3, 0.128s_4, [0, 0.024])$

Table 5.9. Distributed assessments on six cargo ships

	Poor	Average	Good	Very Good	Excellent	ε
Ship 1	0.303	0.047	$[0.073, 0.073 + 0.03I_1]$	$[0.209, 0.351]$	$[0.311, 0.381 + 0.06I_2]$	$[0, 0.024]$
Ship 2	0.060	$[0.138, 0.150]$	$[0.362, 0.508]$	$[0.300, 0.422]$	0	$[0, 0.03]$
Ship 3	$[0.192, 0.213 + 0.02I_1]$	$[0.125, 0.218]$	$[0.306, 0.342 + 0.06I_2]$	$[0.259, 0.271]$	0	$[0, 0.036]$
Ship 4	$[0.050, 0.05 + 0.04I_1]$	$[0.263, 0.323]$	$[0.037, 0.037 + 0.06I_2]$	$[0.201, 0.254]$	$[0.330, 0.395]$	$[0, 0.018]$
Ship 5	$[0.105, 0.187]$	$[0.326, 0.432]$	$[0.351, 0.375]$	0.080	0.020	$[0, 0.036]$
Ship 6	$[0.230, 0.282]$	$[0.306, 0.363 + 0.06I_1]$	$[0.049, 0.109]$	$[0.182, 0.182 + 0.03I_2]$	0.128	$[0, 0.024]$

Table 5.10. The expected utilities of six cargo ships

	Minimum utility	Maximum utility	Average utility
Ship 1	0.541	0.808	0.675
Ship 2	0.512	0.732	0.622
Ship 3	0.441	0.582	0.512
Ship 4	0.618	0.846	0.732
Ship 5	0.425	0.518	0.472
Ship 6	0.425	0.544	0.485

5.5.6 Computing the expected utilities of six cargo ships

It is very difficult to precisely describe the ranking orders among the six cargo ships from their distinct evaluation grades as shown in Table 5.9. In such situation, the expected utilities of six cargo ships should be calculated. In order to compare the final results with [74], we define the same utilities of the five individual evaluation grades as in [74] i.e.,

$$\begin{aligned}
 u'(P) &= 0, u'(A) = 0.4, u'(G) = 0.6, \\
 u'(V) &= 0.8, u'(E) = 1.
 \end{aligned}$$

Then, the expected utilities of six cargo ships can be obtained via formula (5.21), (5.22) and (5.23), which are shown in Table 5.10, and graphically in Figure 5.3. From Table 5.10 we can find that the minimum utility of cargo ship 4 is larger than the maximum utilities of cargo ship 3, cargo ship 5, and cargo ship 6. Hence, the cargo ship 4 is definitely better than them. However, for cargo ship 1 and cargo ship 2, it is difficult to compare them with cargo ship 4 by the same principle. Then, according to the average expected utilities, the ranking order can be obtained easily, and is given as Ship 4 $\succ$ Ship1 $\succ$ Ship2 $\succ$ Ship3 $\succ$ Ship6 $\succ$ Ship5, which is the same with the final result of [74]. However, as we mentioned in Section 5.3, this result is not reliable.

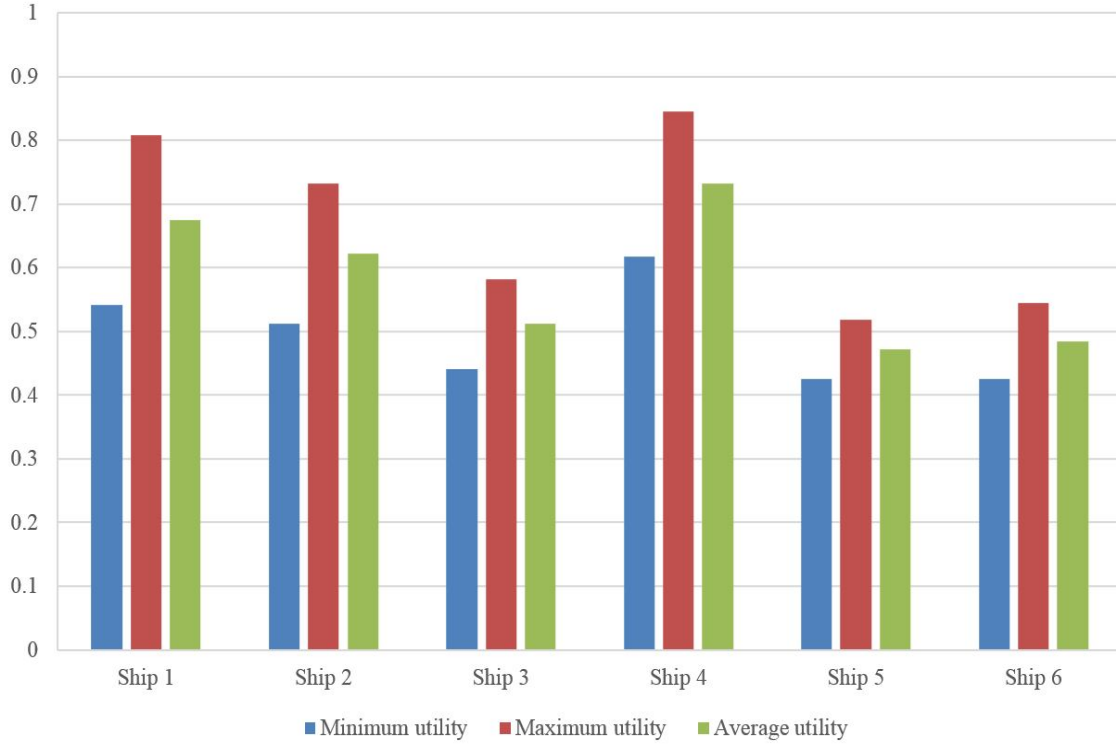


Figure 5.3. The expected utilities of six cargo ships

5.5.7 Ranking the expected utilities using the minimax regret approach

Due to the existence of overlap of expected utilities of six cargo ships, the ranking order based on average expected utilities is not reliable. In order to make a ranking order more precisely, “the minimax regret approach (MRA)” [74] can be employed here to “calculate the maximum loss of expected utility that each cargo ship may suffer”. The ranking principle is that one alternative is selected as the best alternative if this alternative has the smallest maximum loss of expected utility. Then, exclude the best alternative, and calculate the maximum loss of expected utilities of other alternatives. Operate the MRA again and again until we get the ranking order of final two alternatives. Thus, the overall ranking order can be obtained. For the cargo ship selection problem, the calculation process is given as follows.

$$\begin{aligned}
R(\text{Ship 1}) &= \max[\max(0.732, 0.582, 0.846, 0.518, 0.544) - 0.541, 0] = 0.846 - 0.541 \\
&= 0.305, \\
R(\text{Ship 2}) &= \max[\max(0.808, 0.582, 0.846, 0.518, 0.544) - 0.512, 0] = 0.846 - 0.512 \\
&= 0.334, \\
R(\text{Ship 3}) &= \max[\max(0.808, 0.732, 0.846, 0.518, 0.544) - 0.441, 0] = 0.846 - 0.441 \\
&= 0.405, \\
R(\text{Ship 4}) &= \max[\max(0.808, 0.732, 0.582, 0.518, 0.544) - 0.618, 0] = 0.808 - 0.618 \\
&= 0.19, \\
R(\text{Ship 5}) &= \max[\max(0.808, 0.732, 0.582, 0.846, 0.544) - 0.425, 0] = 0.846 - 0.425 \\
&= 0.421, \\
R(\text{Ship 6}) &= \max[\max(0.808, 0.732, 0.582, 0.846, 0.518) - 0.425, 0] = 0.846 - 0.425 \\
&= 0.421.
\end{aligned}$$

It is easily to find that Ship 4 has the smallest maximum loss of expected utility. So, Ship 4 is selected as the best cargo ship. Then, exclude Ship 4, and calculate the maximum loss of expected utility of the other 5 ships.

$$\begin{aligned}
R(\text{Ship 1}) &= \max[\max(0.732, 0.582, 0.518, 0.544) - 0.541, 0] = 0.732 - 0.541 \\
&= 0.191, \\
R(\text{Ship 2}) &= \max[\max(0.808, 0.582, 0.518, 0.544) - 0.512, 0] = 0.808 - 0.512 \\
&= 0.296, \\
R(\text{Ship 3}) &= \max[\max(0.808, 0.732, 0.518, 0.544) - 0.441, 0] = 0.808 - 0.441 \\
&= 0.367, \\
R(\text{Ship 5}) &= \max[\max(0.808, 0.732, 0.582, 0.544) - 0.425, 0] = 0.808 - 0.425 \\
&= 0.383, \\
R(\text{Ship 6}) &= \max[\max(0.808, 0.732, 0.582, 0.518) - 0.425, 0] = 0.808 - 0.425 \\
&= 0.383.
\end{aligned}$$

It is obvious that Ship 1 has the smallest maximum loss of expected utility. So, Ship 1 is selected as the best cargo ship. Then, exclude Ship 1, and calculate the maximum loss of expected utility of the other 4 ships. Based on this approach, we can finally obtain the ranking order as Ship 4 $\succ$ Ship1 $\succ$ Ship2 $\succ$ Ship3 $\succ$ Ship6 $\succ$ Ship5. The same ranking order has been obtained again, i.e., the final result obtained by interval fuzzy linguistic distribution model is the same with that obtained by evidential reasoning approach [74].

5.6 Conclusion

In this chapter, we introduced an interval fuzzy linguistic distribution model for MADM problems with incomplete linguistic information. In this model, we used intervals as evaluators' confidence levels indicating their belief degrees that each linguistic term fits a linguistic variable. Compared with proportions, the use of intervals leaves more operation space for evaluators to handle uncertain and incomplete information. In interval fuzzy linguistic distribution model, we also introduced a variable to represent the extent of ignoring information. However, different from proportional fuzzy linguistic distribution model, of which the extent of ignoring information is obvious, the extent of ignoring information of interval fuzzy linguistic distribution model needs to be calculated. This is determined by the inherent nature of interval. Besides, the expected utility in interval fuzzy linguistic distribution is also different from that in proportional fuzzy linguistic distribution model. No matter the interval fuzzy linguistic distribution is complete or incomplete, the expected utility is always an interval, while, the expected utility is a numerical value if the proportional fuzzy linguistic distribution is complete.

In this chapter, we also developed several aggregation operators for interval fuzzy linguistic distributions, such as arithmetic mean, weighted average operator, interval weighted average operator, and interval ordered weighted average operator. These aggregation operators can help decision makers to respond most MADM problems. Finally, we used two examples to illustrate the proposed model from both a easily comprehensible perspective and a practical perspective.

In the second example, we used the original data of the cargo ship selection problem in order to compare the final results with extended evidential reasoning approach [74]. However, due to the restrictions of the original data and the approach used to model quantitative data in the cargo ship selection problem, two 0-1 integer variables I_1 , I_2 were generated. This slightly affected the

explanation of distribution results of distinct evaluation grades in the cargo ship selection problem, but didn't involve the expected utilities and ranking order. This problem doesn't exist in the first example. This is another reason that we used two examples in this chapter.

In addition, there are two aspects that are worth mentioning. First, the nature of interval fuzzy linguistic distribution model is a symbolic model with the advantage of easy operating in the linguistic solving process. In order to inherit this advantage, we developed this model based on the traditional interval arithmetic rules. Therefore the proposed model is computationally simple compared with "the nonlinear optimization models used in evidential reasoning approach, which is quite computationally complicated and leaves decision makers lots of inconvenience, even they have to use software package to solve practical problems sometimes" [74], [75]. However, the intervals might become too wide to be reached after calculation based on traditional interval arithmetic rules. Especially when computing the expected utility, the situation that the maximum expected utility is larger than 1 sometimes might happen. In such case, we should artificially adjust related utilities.

The other aspect is that proportional fuzzy linguistic distribution model can be regarded as a special form of interval fuzzy linguistic distribution model actually. Hence, the latter inherits all the advantages of the former, and can be employed to deal with decision making problems under more complicated situations. However, if proportions are enough to capture the vagueness and uncertainty, it is better to use proportional fuzzy linguistic distribution model. This is because proportion is more precise than interval so that the result obtained by proportional fuzzy linguistic distribution model is also more accurate than that obtained by interval fuzzy linguistic distribution model.

Chapter 6

Conclusion

In this research, we first recalled some basic knowledge about decision making and computing with words, and discussed the relationship between multiple attribute decision making and computing with words, mainly focusing on fuzzy linguistic approach and linguistic decision making resolution scheme. Then, according to the traditional classification of linguistic computational models, we analyzed the different characteristics of “linguistic computational models based on membership functions, based on ordinal scales and based on 2-tuple representation” respectively. Meanwhile, as one of the inspirations of this research, we reviewed “proportional 2-tuple fuzzy linguistic representation model” [70] in detail in order to pave the way for proposing an extended version of linguistic computational model in Chapter 3. Finally, a proportional 3-tuple fuzzy linguistic representation model, a proportional fuzzy linguistic distribution model and an interval fuzzy linguistic distribution model were proposed in Chapter 3, Chapter 4, and Chapter 5 respectively. Four illustration examples were used to explain how these three models dealt with MADM problems with incomplete linguistic information.

6.1 The Main Contributions

The main contributions of this research can be summarized as follows:

- (1) Developed a proportional 3-tuple fuzzy linguistic representation model.

Considering that “2-tuple fuzzy linguistic representation model” [28] and “proportional 2-tuple fuzzy linguistic representation model” [70] don’t involve incomplete linguistic

assessments, while incomplete linguistic assessments emerge commonly when evaluators might not be able to compare some alternatives or evaluators might prefer to avoid introducing inconsistency in their linguistic assessments, we developed a proportional 3-tuple fuzzy linguistic representation model to solve this problem. Essentially, our idea was to introduce a new variable that represented the extent of ignoring information. Thus, the incomplete information could be considered during the calculation process. Other contributions include:

- 1) A notion of preference-preserving proportional 3-tuple transformation. It is very common that evaluators might use different linguistic term sets to express their preferences during the evaluation process. In such case, the linguistic assessments coming from different linguistic term sets have to be unified before aggregation. The notion of preference-preserving proportional 3-tuple transformation was proposed in order to unify linguistic assessments between two different linguistic term sets without loss of information.
 - 2) Aggregation operators for proportional 3-tuples. We developed arithmetic mean, weighted average operator and linguistic weighted average operator for proportional 3-tuples so that proportional 3-tuple fuzzy linguistic representation model could be applied to multi-expert decision making (MEDM) and MADM problems.
- (2) Developed a proportional fuzzy linguistic distribution model.

“Since uncertainty may be assigned not only to any single evaluation grades but also to their rational combinations, each attribute can be directly evaluated using subjective judgments with the uncertainty being assigned to any number of adjacent evaluation grades simultaneously” [85]. Therefore, we developed a proportional fuzzy linguistic distribution model to deal with linguistic distribution assessments and incomplete linguistic information. Other contributions include:

- 1) Aggregation operators for proportional fuzzy linguistic distribution. We developed arithmetic mean, weighted average operator and linguistic weighted average operator for proportional fuzzy linguistic distributions. Thus, proportional fuzzy linguistic distribution model is capable of dealing with MEDM and MADM problems with linguistic distribution assessments, with incomplete linguistic information, as well as with linguistic weights.
- 2) Expected utility in proportional fuzzy linguistic distribution. The use of proportional fuzzy linguistic distribution model with linguistic distribution assessments leaves an

aggregated distribution assessment for each alternative, which is very difficult to precisely describe the ranking order among them. Therefore, we introduced the notion of expected utility in proportional fuzzy linguistic distribution aiming at supplying a way to conveniently compare or rank alternatives.

(3) Developed an interval fuzzy linguistic distribution model.

Due to the fact that proportions might not be enough to capture the uncertainty when evaluators face with uncertain, vague and imprecise information, while interval or the combination of proportion and interval could better reflect evaluators' confidence levels, we developed an interval fuzzy linguistic distribution model to deal with MADM problems under complicated situations. Other contributions include:

- 1) Aggregation operators for interval fuzzy linguistic distribution. We developed arithmetic mean, weighted average operator, interval weighted average operator and interval ordered weighted average operator for interval fuzzy linguistic distribution model. With these aggregation operators, interval fuzzy linguistic distribution model is able to deal with MEDM and MADM problems with linguistic distribution assessments and incomplete linguistic information under complicated situations.
- 2) Expected utility in interval fuzzy linguistic distribution. For the same purpose with expected utility in proportional fuzzy linguistic distribution, we introduced the expected utility in interval fuzzy linguistic distribution. However, the difference is that no matter whether interval fuzzy linguistic distribution is complete or not, its expected utility is always an interval. Therefore, other accessorial methods may be considered to use in order to improve the reliability of final ranking order.

(4) The contribution to Knowledge Science

The three evaluation models developed in this research supply new ways of modeling evaluators' knowledge regarding the field of linguistic decision analysis, and meanwhile, they can be regarded as new tools for representing and handling tacit knowledge in decision making. Further, different aggregation operators proposed in this research can be looked as new ways for integrating personal knowledge. Moreover, the three evaluation models themselves can be as the created knowledge for decision analysis. In addition, the obtained results of this research also provide new techniques for solving multi-expert and MADM problems in practical applications.

6.2 Discussion and Future work

In this research, we developed a notion of preference-preserving proportional 3-tuple transformation in order to transform proportional 3-tuples between two different linguistic term sets without loss of information. However, we don't develop similar notions for proportional and interval fuzzy linguistic distributions. One of main problems is that linguistic assessment distributions consist of different numbers of linguistic terms. Therefore, one interesting aspect for future work is how to transform proportional and interval fuzzy linguistic distributions constituted by different numbers of linguistic terms between two different linguistic term sets without too many restrictions and without loss of information.

We developed several aggregation operators in this research. All these aggregation operators are extended from conventional linear aggregation operators. As we know, adopting linear additive method to synthesize and aggregate assessment information requires all the attributes to be additively independent. However, linear additive independence assumption may not always be acceptable in reality. Therefore, some non-linear aggregation operators may be considered to develop for the evaluation models proposed in this research.

In addition, we develop an interval fuzzy linguistic distribution model, which is based on traditional interval operation rules for the purpose of providing a computationally simple way to deal with MADM problems with incomplete linguistic information. However, when computing the expected utility, intervals might become too wide to be reached based on traditional interval arithmetic rules. Sometimes, the situation that the maximum expected utility is larger than 1 might happen. In such situation, besides artificially adjust related utilities, another interesting aspect for future work is how to define more efficient interval arithmetic rules to avoid the appearance of the situation mentioned above, and meanwhile, keep the advantage of ease operation in the complicated linguistic context.

In the future, we plan to continue our research in order to extend the applicability of these evaluation models.

Bibliography

- [1] B. S. Ahn. The uncertain OWA aggregation with weighting functions having a constant level of orness. *International journal of intelligent systems*, 21(5):469–483, 2006.
- [2] K. Anagnostopoulos, H. Doukas, and J. Psarras. A linguistic multicriteria analysis system combining fuzzy sets theory, ideal and anti-ideal points for location site selection. *Expert Systems With Applications*, 35(4):2041–2048, 2008.
- [3] T. Astebro. Key success factors for technological entrepreneurs’ R&D projects. *IEEE Transactions on Engineering Management*, 51(3):314–328, 2004.
- [4] A. Bargiela and W. Pedrycz. Granular mappings. *IEEE Transactions on Systems, Man, and Cybernetics—Part A: Systems and Humans*, 35(2):292–297, 2005.
- [5] I. Bogardi and A. Baedossy. *Essays and Surveys on Multiple Criteria Decision Making*, chapter Application of MADM to geological exploration. New York: Springer-Verlag, 1983.
- [6] P. P. Bonissone. *Approximate Reasoning in Decision Analysis*, chapter A fuzzy sets based linguistic approach: Theory and applications, pages 329–339. Amsterdam, The Netherlands: North-Holland, 1982.
- [7] P. P. Bonissone and K. S. Decker. *Uncertainty in Artificial Intelligence*, chapter Selecting uncertainty calculi and granularity: An experiment in trading-off precision and complexity, pages 217–247. Amsterdam, The Netherlands: North-Holland, 1986.
- [8] G. Bordogna, M. Fedrizzi, and G. Passi. A linguistic modeling of consensus in group decision making based on OWA operators. *IEEE Transactions on Systems, Man, and Cybernetics Part A: Systems and Humans*, SMC-27(1):126–132, 1997.
- [9] G. Bordogna and G. Pasi. A fuzzy linguistic approach generalizing boolean information retrieval: A model and its evaluation. *Journal of the American Society for Information Science*, 44:70–82, 1993.

- [10] C. Carlsson and R. Fullér. Benchmarking in linguistic importance weighted aggregations. *Fuzzy Sets and System*, 114(1):35–41, 2000.
- [11] J. Casillas, O. Cordón, M. J. del Jesús, and F. Herrera. Genetic tuning of fuzzy rule deep structures preserving interpretability and its interaction with fuzzy rule set reduction. *IEEE Transactions on Fuzzy Systems*, 13(1):13–29, 2005.
- [12] S. J. Chen and C. L. Hwang. *Fuzzy Multiple Attribute Decision Making-Methods and Applications*. Springer, Berlin, Germany, 1992.
- [13] R. Degani and G. Bortolan. The problem of linguistic approximation in clinical decision making. *International Journal of Approximate Reasoning*, 2(2):143–162, 1988.
- [14] M. Delgado, F. Herrera, E. Herrera-Viedma, M. J. Martin-Bautista, L. Martínez, and M. A. Vila. A communication model based on the 2-tuple fuzzy linguistic representation for a distributed intelligent agent system on internet. *Soft Computing*, 6:320–328, 2002.
- [15] M. Delgado, J. Verdegay, and M. Vila. On aggregation operations of linguistic labels. *International Journal of Intelligent Systems*, 8(3):351–370, 1993.
- [16] M. Delgado, M. A. Vila, and W. Voxman. On a canonical representation of fuzzy numbers. *Fuzzy Sets and Sysems*, 93:125–135, 1998.
- [17] S. Dick, A. Schenker, W. Pedrycz, and A. Kandel. Regranulation: A granular algorithm enabling communication between granular worlds. *Information Sciences*, 177(2):408–435, 2007.
- [18] Y. C. Dong, Y. F. Xu, and S. Yu. Computing the numerical scale of the linguistic term set for the 2-tuple fuzzy linguistic representation model. *IEEE Transactions on Fuzzy Systems*, 17(6):1366–1378, 2009.
- [19] Y. C. Dong, G. Q. Zhang, W. C. Hong, and S. Yu. Linguistic computational model based on 2-tuples and intervals. *IEEE Transactions on Fuzzy Systems*, 21(6):1006–1018, 2013.
- [20] T. Evangelos. *Multi-criteria Decision Making Methods: A Comparative Study*. Dordrecht: Kluwer Academic Publishers, 2000.
- [21] J. Fodor and M. Roubens. *Fuzzy preference modelling and multicriteria decision support*. Dordrecht: Kluwer Academic Publishers, 1994.

- [22] G. Fu. A fuzzy optimization method for multicriteria decision making: An application to reservoir flood control operation. *Expert Systems With Applications*, 34(1):145–149, 2008.
- [23] F. Herrera, S. Alonso, F. Chiclana, and E. Herrera-Viedma. Computing with words in decision making: Foundations, trends and prospects. *Fuzzy Optimization and Decision Making*, 8(4):337–364, 2009.
- [24] F. Herrera and E. Herrera-Viedma. Linguistic decision analysis: steps for solving decision problems under linguistic information. *Fuzzy Sets and Systems*, 115(1):67–82, 2000.
- [25] F. Herrera, E. Herrera-Viedma, and L. Martínez. A fusion approach for managing multi-granularity linguistic terms sets in decision making. *Fuzzy Sets and Systems*, 114(1):43–58, 2000.
- [26] F. Herrera, E. Herrera-Viedma, and L. Martínez. A fuzzy linguistic methodology to deal with unbalanced linguistic term sets. *IEEE Transactions on Fuzzy Systems*, 16(2):354–370, 2008.
- [27] F. Herrera, E. Herrera-Viedma, and J. L. Verdegay. Direct approach processes in group decision making using linguistic OWA operators. *Fuzzy Sets and Systems*, 79(2):175–190, 1996.
- [28] F. Herrera and L. Martínez. A 2-tuple fuzzy linguistic representation model for computing with words. *IEEE Transactions on Fuzzy Systems*, 8(6):746–752, 2000.
- [29] F. Herrera and L. Martínez. A model based on linguistic 2-tuples for dealing with multigranular hierarchical linguistic context in multi-expert decision making. *IEEE Transactions on Systems, Man, and Cybernetics—Part B: Cybernetics*, 31(2):227–234, 2001.
- [30] K. Hirota and W. Pedrycz. Fuzzy computing for data mining. *Proceedings of the IEEE*, 87(9):1575–1600, 1999.
- [31] K. M. M. Holtta and K. N. Otto. Incorporating design effort complexity measures in product architectural design and assessment. *Design Studies*, 26(5):463–485, 2005.
- [32] V. N. Huynh and Y. Nakamori. A satisfactory-oriented approach to multiexpert decision-making with linguistic assessments. *IEEE Transactions on Systems, Man, and Cybernetics—Part B: Cybernetics*, 35(2):184–196, 2005.
- [33] V. N. Huynh and Y. Nakamori. A linguistic screening evaluation model in new product development. *IEEE Transactions on Engineering Management*, 58(1):165–175, 2011.

- [34] V. N. Huynh, Y. Nakamori, T. B. Ho, and T. Murai. Multiple-attribute decision making under uncertainty: the evidential reasoning approach revisited. *IEEE Transactions on Systems, Man, and Cybernetics—Part A: Systems and Humans*, 36(4):804–822, 2006.
- [35] V. N. Huynh, C. H. Nguyen, and Y. Nakamori. MEDM in general multi-granular hierarchical linguistic contexts based on the 2-tuples linguistic model. In *IEEE International Conference on Granular Computing*, volume 2, pages 482–487. IEEE, July 2005.
- [36] H. Ishibuchi and H. Tanaka. Multiobjective programming in optimization of the interval objective function. *European Journal of Operational Research*, 48(2):219–225, 1990.
- [37] R. I. John and P. R. Innocent. Modeling uncertainty in clinical diagnosis using fuzzy logic. *IEEE Transactions on Systems, Man, and Cybernetics—Part B: Cybernetics*, 35(6):1340–1350, 2005.
- [38] J. Kacprzyk and M. Fedrizzi. *Multiperson decision making models using fuzzy sets and possibility theory*. Dordrecht: Kluwer Academic Publishers, 1990.
- [39] J. Kacprzyk and S. Zadrozny. Computing with words in decision making through individual and collective linguistic choice rules. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 9(supp01):89–102, 2001.
- [40] J. Kacprzyk and S. Zadrozny. Computing with words in intelligent database querying: Standalone and internet-based applications. *Information Sciences*, 134(1-4):71–109, 2001.
- [41] N. N. Karnik and J. M. Mendel. Centroid of a type-2 fuzzy set. *Information Sciences*, 132(1-4):195–220, 2001.
- [42] R. L. Keeney and H. Raiffa. *Decisions With Multiple Objectives*. The Press Syndicate of the University of Cambridge, 1993.
- [43] J. Lawry. An alternative approach to computing with words. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 9:3–16, 2001.
- [44] C. T. Lin and C. T. Chen. New product go/no-go evaluation at the front end: a fuzzy linguistic approach. *IEEE Transactions on Engineering Management*, 51(2):197–207, 2004.
- [45] T. Y. Lin. Granular computing: From rough sets and neighborhood systems to information granulation and computing in words. In *European Congress on Intelligent Techniques and Soft Computing*, pages 1602–1606, September 8-12 1997.

- [46] J. Lu, G. Zhang, D. Ruan, and F. Wu. *Multi-objective group decision making. Methods, software and applications with fuzzy set techniques*. London: Imperial College Press, 2007.
- [47] L. Martínez. Sensory evaluation based on linguistic decision analysis. *International Journal of Approximate Reasoning*, 44(2):148–164, 2007.
- [48] L. Martínez, M. Espinilla, and L. G. Pérez. A linguistic multigranular sensory evaluation model for olive oil. *International Journal of Computational Intelligence Systems*, 1(2):148–158, 2008.
- [49] L. Martínez and F. Herrera. An overview on the 2-tuple linguistic model for computing with words in decision making: Extensions, applications and challenges. *Information Sciences*, 207:1–18, 2012.
- [50] L. Martínez, J. Liu, D. Ruan, and J. B. Yang. Dealing with heterogeneous information in engineering evaluation processes. *Information Sciences*, 177(7):1533–1542, 2007.
- [51] L. Martínez, J. Liu, and J. B. Yang. A fuzzy model for design evaluation based on multiple criteria analysis in engineering systems. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, 14(3):317–336, 2006.
- [52] J. Mendel. Computing with words and its relationships with fuzzistics. *Information Sciences*, 177(4):988–1006, 2007.
- [53] J. M. Mendel and D. Wu. *Perceptual Computing: Aiding People in Making Subjective Judgments*. John Wiley & Sons, Inc., Hoboken, New Jersey, USA, 2010.
- [54] G. A. Miller. The magical number seven or minus two: Some limits on our capacity of processing information. *Psychological review*, 101(2):343–352, 1994.
- [55] R. E. Moore. *Method and Application of Interval Analysis*. Philadelphia, PA, USA: SIAM, 1979.
- [56] W. Pedrycz and Z. A. Sosnowski. Designing decision trees with the use of fuzzy granulation. *IEEE Transactions on Systems, Man, and Cybernetics—Part A: Systems and Humans*, 30(2):151–159, 2000.
- [57] T. Pham. Computing with words in formal methods. *International Journal of Intelligent Systems*, 15(8):801–810, 2000.

- [58] C. Porcel, A. G. López-Herrera, and E. Herrera-Viedma. A recommender system for research resources based on fuzzy linguistic modeling. *Expert Systems with Applications*, 36(3):5173–5183, 2009.
- [59] R. M. Rodríguez and L. Martínez. An analysis of symbolic linguistic computing models in decision making. *International Journal of General Systems*, 42(1):121–136, 2013.
- [60] R. M. Rodríguez, L. Martínez, and F. Herrera. Hesitant fuzzy linguistic term sets for decision making. *IEEE Transactions on Fuzzy Systems*, 20(1):109–119, 2012.
- [61] T. L. Saaty. *The Analytic Hierarchy Process*. Pittsburgh, PA: Univ. Pittsburgh, 1988.
- [62] K. Schmucker. *Fuzzy sets, natural language computations, and risk analysis*. Rockville, MD: Computer Science Press, 1984.
- [63] A. Sengupta and T. K. Pal. On comparing interval numbers. *European Journal of Operational Research*, 127(1):28–43, 2000.
- [64] D. D. Sharma. *Decision Making Style: Social and Creative Dimensions*. Global India Publications Pvt Ltd, 2009.
- [65] H. Tanaka, K. Sugihara, and Y. Maeda. Non-additive measures by interval probability functions. *Information Sciences*, 164(1-4):209–227, 2004.
- [66] R. Tong and P. Bonissone. A linguistic approach to decision making with fuzzy sets. *IEEE Transactions on Systems, Man and Cybernetics*, SMC-10(11):716–723, 1980.
- [67] V. Torra. Negation function based semantics for ordered linguistic labels. *International Journal of Intelligent Systems*, 11:975–988, 1996.
- [68] E. Trillas. On the use of words and fuzzy sets. *Information Sciences*, 176(11):1463–1487, 2006.
- [69] E. Trillas and S. Guadarrama. What about fuzzy logic’s linguistic soundness? *Fuzzy Sets and Systems*, 156(3):334–340, 2005.
- [70] J. H. Wang and J. Y. Hao. A new version of 2-tuple fuzzy linguistic representation model for computing with words. *IEEE Transactions on Fuzzy Systems*, 14(3):435–445, 2006.
- [71] J. H. Wang and J. Y. Hao. An approach to computing with words based on canonical characteristic values of linguistic labels. *IEEE Transactions on Fuzzy Systems*, 15(4):593–604, 2007.

- [72] S. Y. Wang. Applying 2-tuple multigranularity linguistic variables to determine the supply performance in dynamic environment based on product-oriented strategy. *IEEE Transactions on Fuzzy Systems*, 16(1):29–39, 2008.
- [73] Y. M. Wang and T. M. S. Elhag. On the normalization of interval and fuzzy weights. *Fuzzy Sets and Systems*, 157:2456–2471, 2006.
- [74] Y. M. Wang, J. B. Yang, D. L. Xu, and K. S. Chin. The evidential reasoning approach for multiple attribute decision analysis using interval belief degrees. *European Journal of Operational Research*, 175(1):35–66, 2006.
- [75] Y. M. Wang, J. B. Yang, D. L. Xu, and K. S. Chin. On the combination and normalization of interval-valued belief structures. *Information Sciences*, 177:1230–1247, 2007.
- [76] W. L. Winston. *Operations Research—Applications and Algorithms*. Belmont, CA: Duxbury Press, 1994.
- [77] D. Wu and J. M. Mendel. Enhanced karnikmendel algorithms. *IEEE Transactions on Fuzzy Systems*, 17(4):923–934, 2009.
- [78] D. Wu and J. M. Mendel. Perceptual reasoning for perceptual computing: A similarity based approach. *IEEE Transactions on Fuzzy Systems*, 17(6):1397–1411, 2009.
- [79] Z. Xu. A method based on linguistic aggregation operators for group decision making with linguistic preference relations. *Information Sciences*, 166(1-4):19–30, 2004.
- [80] R. R. Yager. A new methodology for ordinal multiobjective decisions based on fuzzy sets,. *Decision Sciences*, 12(4):589–600, 1981.
- [81] R. R. Yager. Non-numeric multi-criteria multi-person decision making. *Group Decision and Negotiation*, 2(1):81–93, 1993.
- [82] R. R. Yager. An approach to ordinal decision making. *International Journal of Approximate Reasoning*, 12(3-4):237–261, 1995.
- [83] R. R. Yager. On the retranslation process in Zadeh’s paradigm of computing with words. *IEEE Transactions on Systems, Man, and Cybernetics—Part B: Cybernetics*, 34(2):1184–1195, 2004.
- [84] J. B. Yang. Rule and utility based evidential reasoning approach for multiple attribute decision analysis under uncertainty. *European Journal of Operational Research*, 131(1):31–61, 2001.

- [85] J. B. Yang and P. Sen. A general multi-level evaluation process for hybrid MADM with uncertainty. *IEEE Transactions on Systems, Man, and Cybernetics*, 24(10):1458–1473, 1994.
- [86] J. B. Yang and M. G. Singh. An evidential reasoning approach for multiple attribute decision making with uncertainty. *IEEE Transactions on Systems, Man, and Cybernetics*, 24(1):1–18, 1994.
- [87] J. B. Yang and M. G. Singh. On the evidential reasoning algorithm for multiple attribute decision analysis under uncertainty. *IEEE Transactions on Systems, Man, and Cybernetics—Part A: Systems and Humans*, 32(3):289–304, 2002.
- [88] M. S. Ying. A formal model of computing with words. *IEEE Transactions on Fuzzy Systems*, 10(5):640–652, 2002.
- [89] L. A. Zadeh. Fuzzy sets. *Information and Control*, 8(3):338–353, 1965.
- [90] L. A. Zadeh. The concept of a linguistic variable and its applications to approximate reasoning part I. *Information sciences*, 8(3):199–249, 1975.
- [91] L. A. Zadeh. The concept of a linguistic variable and its applications to approximate reasoning part II. *Information sciences*, 8:301–357, 1975.
- [92] L. A. Zadeh. The concept of a linguistic variable and its applications to approximate reasoning part III. *Information sciences*, 9:43–80, 1975.
- [93] L. A. Zadeh. Fuzzy logic = computing with words. *IEEE Transactions on Fuzzy Systems*, 4(2):103–111, 1996.
- [94] L. A. Zadeh. Toward a theory of fuzzy information granulation and its centrality in human reasoning and fuzzy logic. *Fuzzy Sets and Systems*, 90(2):111–127, 1997.
- [95] L. G. Zhou, H. Y. Chen, J. M. Merigó, and A. M. Gil-Lafuente. Uncertain generalized aggregation operators. *Expert Systems with Applications*, 39(1):1105–1117, 2012.

Publications

Journals:

- [1] W. T. Guo, V. N. Huynh, and Y. Nakamori. A proportional 3-tuple fuzzy linguistic representation model for screening new product projects. *Journal of Systems Science and Systems Engineering*, Springer, accepted for publish.
- [2] W. T. Guo, V. N. Huynh, and Y. Nakamori. An interval linguistic distribution model for multiple attribute decision making problems with incomplete linguistic information. *International Journal of Knowledge and Systems Science*, IGI Global, accepted for publish.
- [3] W. T. Guo and V. N. Huynh. A model for multiple attribute decision making under uncertainty based on proportional fuzzy linguistic distributions. *Knowledge-Based Systems*, Elsevier, under revision.

Proceedings:

- [4] W. T. Guo, V. N. Huynh, and Y. Nakamori. A proportional 3-tuple fuzzy linguistic screening evaluation model in new product development. In Proceedings of *International Symposium on Knowledge and Systems Sciences, Knowledge Creation Towards Emergency Management*, JAIST Press, pages 82-88, October 25-27, 2013, Ningbo, China.
- [5] W. T. Guo and V. N. Huynh. A new version of 2-tuple fuzzy linguistic screening evaluation model in new product development. In Proceedings of *The IEEE International Conference on Industrial Engineering and Engineering Management (IEEM)*, number of paper pages: 6, December 10-13, 2013, Bangkok, Thailand.
- [6] W. T. Guo, V. N. Huynh, Y. Nakamori, and M. Kosaka. Evaluation of service based on proportional fuzzy linguistic distribution model. In Proceedings of *International Symposium on Knowledge and Systems Sciences*, JAIST Press, pages 269-274, November 1-2, 2014, Sapporo, Japan.

- [7] W. T. Guo and V. N. Huynh. A fuzzy linguistic representation model for decision making under uncertainty. In Proceedings of *The IEEE International Conference on Industrial Engineering and Engineering Management (IEEM)*, number of paper pages: 5, December 09-12, 2014, Kuala Lumpur, Malaysia.