



A continuous interval valued linguistic ORESTE method for multi -criteria group decision making

DOI:

[10.1016/j.knosys.2018.04.022](https://doi.org/10.1016/j.knosys.2018.04.022)

Document Version

Accepted author manuscript

[Link to publication record in Manchester Research Explorer](#)

Citation for published version (APA):

Liao, H., Wu, X., Liang, X., Yang, J.-B., Xu, D.-L., & Herrera, F. (2018). A continuous interval valued linguistic ORESTE method for multi -criteria group decision making. *Knowledge-Based Systems*, 153, 65-77. <https://doi.org/10.1016/j.knosys.2018.04.022>

Published in:

Knowledge-Based Systems

Citing this paper

Please note that where the full-text provided on Manchester Research Explorer is the Author Accepted Manuscript or Proof version this may differ from the final Published version. If citing, it is advised that you check and use the publisher's definitive version.

General rights

Copyright and moral rights for the publications made accessible in the Research Explorer are retained by the authors and/or other copyright owners and it is a condition of accessing publications that users recognise and abide by the legal requirements associated with these rights.

Takedown policy

If you believe that this document breaches copyright please refer to the University of Manchester's Takedown Procedures [<http://man.ac.uk/04Y6Bo>] or contact uml.scholarlycommunications@manchester.ac.uk providing relevant details, so we can investigate your claim.



A continuous interval-valued linguistic ORESTE method for multi-criteria group decision making

Huchang Liao^{1,2}, Xingli Wu¹, Xuedong Liang^{1,*},
Jian-Bo Yang³, Dong-Ling Xu³, Francisco Herrera^{2,4}

¹ Business School, Sichuan University, Chengdu 610064, China

² Department of Computer Science and Artificial Intelligence, University of Granada, E-18071 Granada, Spain

³ Manchester Business School, University of Manchester, M15 6PB Manchester, UK

⁴ Faculty of Computing and Information Technology, King Abdulaziz University, Jeddah, Saudi Arabia

Abstract

Considering that the uncertain linguistic variable (or interval linguistic term) has some limitations in calculation, we extend it to the continuous interval-valued linguistic term set (CIVLTS), which is equivalent to the virtual term set but has its own semantics. It has the advantages of both the uncertain linguistic variable and the virtual term set but overcomes their defenses. It not only can interpret more complex assessments by continuous terms, but also is effective in aggregating the group opinions. We propose some methods to aggregate the individual decision matrices represented by CIVLTSs to the collective matrix. The extended Gaussian-distribution-based weighting method is proposed to derive the weights for aggregating the large group opinions. Furthermore, the general ranking method ORESTE, is extended to the CIVL environment and is named as the CIVL-ORESTE method. The proposed method is excellent by no requirements of crisp criterion weights and the objective thresholds. A case study of selecting the optimal innovative sharing bike design for the "Mobike" sharing bikes is operated to show the practicability of the CIVL-ORESTE method. Finally, we compare the CIVL-ORESTE method with other ranking methods to illustrate the reliability of our method and its advantages.

Keywords: Multi-criteria group decision making; Continuous interval-valued linguistic term set; Extended Gaussian-distribution-based weighting method; ORESTE; Mobike Design

* Corresponding author. E-mail address: liaohuchang@163.com (H.C. Liao), liangxuedong@scu.edu.cn (X.D. Liang).

1. Introduction

Due to the uncertainty and complexity of objective factors, multiple criteria are defined under qualitative environment, such as quality, personality and exterior [1-4]. In addition, due to the limit of knowledge and cognition of single person, a group of individuals are invited to make judgments on alternatives to obtain the reliable evaluation information. Therefore, the multi-criteria group decision making (MCGDM) under qualitative context turns out to be a valuable research issue. This paper focuses on proposing an effective method to solve this problem, whose procedures are divided into three parts: expressing the evaluation information of each decision-maker (DM), aggregating the DMs' evaluations to group opinions, and ranking the alternatives.

In practice, the opinions of DMs are usually expressed in linguistic terms [5], which are similar to natural or artificial language and close to human's cognitive process. To avoid information loss in computational process, some enhanced models were proposed, such as the 2-tuple fuzzy linguistic representation model [6] and the virtual linguistic model [7]. These two models are finally proved to be mathematically equivalent [8]. Given that these models are both based on singleton linguistic term while human's opinions are always within an interval due to the uncertainty and vagueness in practice, Xu [7] proposed the concept of the uncertain linguistic variable whose value is expressed as the interval of linguistic terms, such as "*between good and very good*". The interval version of 2-tuple linguistic representation model were also investigated by many scholars [9]. However, people sometimes incline to express their opinions in natural language with more complex form, such as "*at least good*", "*more than high*", but both of the extended models are unable to represent these pieces of information. To solve this problem, Rodríguez et al. [10] proposed the hesitant fuzzy linguistic term set (HFLTS) whose value is a set of linguistic terms and the envelope of a HFLTS is an uncertain linguistic variable. Although the HFLTS can represent much information and has been proved to be useful in application [11-17], it also has some flaws. When people have deep understanding of an object, they may provide relatively accurate evaluations. For example, when evaluating the satisfaction degree of a product, the expert may think it is "*between a little high and high and closes to high with 60% of the proportion*". The discrete linguistic terms employed in the existed models are limited to interpret the opinion "*closes to high with 60% of the proportion*". Using

the continuous linguistic term can express this complex information and describe DMs' views more accurately than the discrete form. Thus, we extend the uncertain linguistic variable and the HFLTSSs into the continuous interval-valued linguistic term set (CIVLTS) and present the syntax and semantics.

Integrating individual opinions to collective opinion is an essential step in MCGDM. Under the linguistic environment, some literatures suggested the union-based methods to aggregate the DMs' opinions simply [10, 17, 18]. Given that the probability of each linguistic term is ignored in these union-based methods, Wu and Liao [19] proposed a group aggregation method by considering both the expert weights and the probability of the linguistic term. However, their method is not very effective to aggregate large number of experts' opinions which are expressed in continuous linguistic terms. How to determine each evaluators' weight is a challenge. The weights for group members can be intrinsically determined using their own subjective opinion values. It is appropriate to suppose that the evaluations of a large group obey Gaussian distribution [20]. In this sense, we can determine DMs' weights based on Gaussian distribution. Low weights are given to the "false" or "biased" judgments while high weights are assigned to the mid evaluations, which conforms to people's perceptions.

Ranking alternatives is a critical process to solve the MCGDM problems. There are mainly two types of ranking methods: the utility values-based methods and the outranking methods [8]. The former ranks the alternatives by aggregating the values of each alternative with respect to all criteria, such as the TOPSIS (Technique for Order Preference by Similarity to Ideal Solution) [21] and the VIKOR (VlseKriterijumska Optimizacija I Kompromisno Resenje) method [18, 22]. The obtained results with this type of methods are clear and intuitive but unable to reflect the comparability relation between two alternatives. The latter is based on pairwise comparisons, such as the ELECTRE (ELimination Et Choix Traduisant la REalité - ELimination and Choice Expressing the Reality) method [23] and the Preference Ranking Organization METHod for Enrichment Evaluations (PROMETHEE) [12, 24]. It can determine the preference (P), indifference (I) and comparability (R) relations between alternatives. However, the thresholds to distinguish the PIR relations are given by DMs subjectively, which makes the results bad robustness. The ORESTE (organisation, rangement et Synthèse de données relationnelles, in French) method [25], is an integrated ranking method which is composed by two stages. It firstly calculates the utility values to determine weak

ranking of alternatives, and then derives the PIR relations by the conflict analysis. Thus, it shows the advantages of both types of ranking methods and the thresholds are calculated objectively with less subjective factors. Furthermore, it does not require the crisp weights of criteria which are sometimes difficult or impossible to determine in linguistic environment but are indispensable in most ranking methods. However, the tedious process in the classical ORESTE method leads to information loss to some extent, and it is limited to handle the evaluations expressed in CIVLTSs.

The aim of this paper to handle the MCGDM problems in which the CIVLTSs are used to express individuals' hesitant and qualitative evaluations on alternatives and criteria importance. The aggregation methods including the extended Gauss-distribution-based weighting method (EGDBWM) is introduced to aggregate the individuals' continuous interval-valued linguistic elements (CIVLEs) to group opinions. Subsequently, we rank the alternatives by the proposed the continuous interval-valued linguistic ORESTE (CIVL-ORESTE) method based on the group opinions. The main contributions of this paper are as follows:

(1) We propose the CIVLTSs to express individuals' evaluations and collective opinions exactly. Based on the transform function, we introduce the basic operations and the comparison method of the CIVLTSs. They can overcome the defects of the operations of uncertain linguistic variables that are calculated based on the labels of linguistic terms.

(2) We divide the expert group into four types considering that different groups are suitable for different aggregation methods. The union-based method is proposed to derive the collective opinions of small size group; the average arithmetic aggregation formula-based method is used to solve the medium size group, the weighted arithmetic aggregation formula-based method is used to solve the medium size group and the EGDBWM is developed to deal with the large size group.

(3) We improve the ORESTE method by introducing the distance measure between the CIVLTSs, and derive the thresholds of the ORESTE within the context of CIVLTSs. We develop the procedure of the CIVL-ORESTE method.

(4) We provide a helpful reference for the enterprises to select the optimal innovative sharing bike design and maximize customer satisfaction based on a case study with the CIVL-ORESTE method.

The remainder of this paper is structured as follows. Section 2 introduces the CIVLTS and its semantics.

Section 3 describes some methods to aggregate individual decision matrices to collective matrix. Section 4 proposes the CIVL-ORESTE method. Section 5 introduces a case study about selecting the optimal innovative sharing bike design. Section 6 presents some conclusions.

2. Continuous interval-valued linguistic term set

This section introduces a general representation form of the uncertain linguistic variable, i.e., the CIVLTS, and then justifies its semantics in describing linguistic information.

2.1 Uncertain linguistic variable and HFLTS

To preserve more information than one linguistic term, Xu [7] introduced the concept of uncertain linguistic variable.

Definition 1 [7]. Let $S' = \{s_0, s_1, \dots, s_\tau\}$ be a linguistic term set (LTS). $s = [s_\alpha, s_\beta]$ is an uncertain linguistic variable, where $s_\alpha, s_\beta \in S'$, and s_α and s_β are the lower and the upper limits of s , respectively.

Remark 1. The subscripts of the linguistic terms s_α and s_β are integers. To avoid information loss in calculation, Xu [7] extended the discrete term set S' to a continuous term set $\bar{S} = \{s_\alpha | s_0 \leq s_\alpha \leq s_\tau, \alpha \in [0, \tau]\}$ and called it as the virtual term set. It can only appear in calculation while the original LTS is used in evaluation.

The uncertain linguistic variable can only be used to represent the linguistic expressions in the form of “between s_k and s_m ”, but individuals always have much richer expressions. In this sense, Rodríguez et al. [10] introduced the concept of HFLTSs as an ordered finite subset of the consecutive linguistic terms of S' . Afterwards, Liao et al. [11] extended and formalized the HFLTS mathematically.

Definition 2 [11]. Let $x_i \in X$ ($i=1, 2, \dots, N$) be fixed and $S = \{s_t | t = -\tau, \dots, -1, 0, 1, \dots, \tau\}$ be a LTS. A HFLTS on X , H_S , is in mathematical form of

$$H_S = \{ \langle x_i, h_S(x_i) \rangle | x_i \in X \}$$

where $h_s(x_i)$ is a set of some values in S and can be expressed as

$$h_s(x_i) = \{s_{\varphi_l}(x_i) \mid s_{\varphi_l}(x_i) \in S, l = 1, \dots, \#h_s(x_i)\}$$

with $\#h_s(x_i)$ being the number of linguistic terms in $h_s(x_i)$. $h_s(x_i)$ denotes the possible degrees of the linguistic variable x_i to S . φ_l is the subscript of $s_{\varphi_l}(x_i)$. For convenience, $h_s(x_i)$ is called the hesitant fuzzy linguistic element (HFLE).

Remark 2. The terms $s_{\varphi_l}(x_i)$ ($l = 1, \dots, \#h_s(x_i)$) in each HFLEs should be consecutive given that the linguistic terms are chosen in discrete form. Based on Remark 1, we can extend $\varphi_l \in \{-\tau, \dots, -1, 0, 1, \dots, \tau\}$ to $\varphi_l \in [-\tau, \tau]$. The integer linguistic term $s_{\varphi_l}(x_i)$ ($\varphi_l \in \{-\tau, \dots, -1, 0, 1, \dots, \tau\}$) is determined by the experts, while the virtual linguistic terms only appear in calculations.

To make judgments in human way of thinking and expressions, Rodríguez et al. [10] proposed the context-free grammar to generate linguistic expressions and gave a function E_{G_H} to translate the evaluations to HFLTSs.

Definition 3 [10]. Let $S = \{s_\alpha \mid \alpha = -\tau, \dots, -1, 0, 1, \dots, \tau\}$ be a LTS, and S_{ll} be the linguistic express domain generated by G_H (for detail of G_H , please refer to Ref. [10]). $E_{G_H} : S_{ll} \rightarrow H_S$ is a function that transforms the linguistic expressions S_{ll} to the HFLTS H_S . The linguistic expression $ll \in S_{ll}$ is converted into the HFLTS as: $E_{G_H}(s_t) = \{s_t \mid s_t \in S\}$; $E_{G_H}(\text{at most } s_k) = \{s_t \mid s_t \in S \text{ and } s_t \leq s_k\}$; $E_{G_H}(\text{lower than } s_k) = \{s_t \mid s_t \in S, s_t < s_k\}$; $E_{G_H}(\text{at least } s_k) = \{s_t \mid s_t \in S \text{ and } s_t \geq s_k\}$; $E_{G_H}(\text{greater than } s_k) = \{s_t \mid s_t \in S \text{ and } s_t > s_k\}$; $E_{G_H}(\text{between } s_k \text{ and } s_m) = \{s_t \mid s_t \in S' \text{ and } s_k \leq s_t \leq s_m\}$.

Example 1. Let $S = \{s_{-3} = \text{very bad}, s_{-2} = \text{bad}, s_{-1} = \text{a little bad}, s_0 = \text{medium}, s_1 = \text{a little good}, s_2 = \text{good}, s_3 = \text{very good}\}$ be a LTS. When evaluating the exterior of a car, someone may hold it is “between bad and a little bad”; the other insists that it is “at least a little good”. The corresponding HFLTEs are $\{s_{-2}, s_{-1}\}$,

$\{s_1, s_2, s_3\}$, respectively. Suppose that the linguistic terms are uniformly distributed. Fig. 1 shows the syntax and semantics of these two HFLTEs.

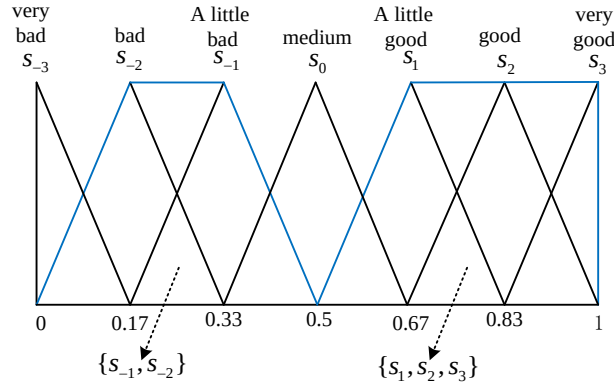


Fig. 1. The HFLTSs with their semantics

From Fig. 1, we can see that the envelope of a HFLTS is an uncertain linguistic variable.

2.2 The concept of CIVLTS

There are two problems existed in both the uncertain linguistic variable and the HFLTS about information loss. (1) The expert's preference judgments cannot be completely expressed. For example, let $S = \{s_{-3} = \text{very low}, s_{-2} = \text{low}, s_{-1} = \text{a little low}, s_0 = \text{medium}, s_1 = \text{a little high}, s_2 = \text{high}, s_3 = \text{very high}\}$ be a LTS. When evaluating the satisfaction degree of a product, the expert may think it is “20% proportion higher than a little high and 40% proportion lower than high”. Then the information can only be represented as the HFLE $\{s_1, s_2\}$ or an uncertain linguistic variable $[s_1, s_2]$. Obviously, both of them cannot reflect the precise proportion information. (2) The integrated judgments of an expert group cannot reflect the group's idea comprehensively. For example, suppose that thirty students evaluate the teaching quality of a teacher. Let S be a LTS given as above. If twenty students hold it is “high”, seven evaluate it is “very high” but two deem it is “low” and one judges it is “between low and very low”, then, the group evaluation can be represented by the HFLE $\{s_{-3}, s_{-2}, s_2, s_3\}$ using the union-based method (actually it belongs to the extended HFLTS [26] as the linguistic terms in it are not consecutive). Significantly, it cannot reflect the reality that this teacher' teaching quality is round “high”. There are few method to aggregate the evaluations expressed as uncertain linguistic variables into collective opinions.

To overcome the limitation in expressing individuals' complex and precise evaluations, we extend the uncertain linguistic variable and the HFLTS into the continuous form. Some new aggregation methods will be introduced in Section 3 to eliminate the existed defects in aggregating group opinions.

Definition 4. Let $x_i \in X$ ($i=1,2,\dots,m$) be fixed and $S = \{s_\alpha | \alpha = -\tau, \dots, -1, 0, 1, \dots, \tau\}$ be a LTS. A CIVLTS on X , $\tilde{H}_S$, is in mathematical form of

$$\tilde{H}_S = \{ \langle x_i, \tilde{h}_S(x_i) \rangle | x_i \in X \} \quad (1)$$

where $\tilde{h}_S(x_i)$ is a subset in continuous interval-valued form of S and can be expressed as

$$\tilde{h}_S(x_i) = [s_{L_i}, s_{U_i}], \quad L_i, U_i \in [-\tau, \tau] \text{ and } L_i \leq U_i \quad (2)$$

$\tilde{h}_S(x_i)$ is the continuous interval-valued linguistic element (CIVLE), denoting the possible degrees of the linguistic variable x_i to S . s_{L_i} and s_{U_i} are the lower and upper bounds of $\tilde{h}_S(x_i)$, respectively.

Remark 3. The subscripts of the CIVLE $\tilde{h}_S(x_i)$, L_i and U_i are real numbers in $[-\tau, \tau]$. This is the main difference between the CIVLTS and the uncertain linguistic variable in which the subscripts should be integers. The uncertain linguistic variable is a special case of the CIVLTS.

As we can see, the CIVLTS is mathematical equivalent to the interval-valued virtual term set. However, the interval-valued virtual term set cannot be used to represent the evaluators' assessments. The main contribution of the CIVLTS is that it has its own semantics and thus overcomes the defense of the interval-valued virtual term set. The CIVLTS can be used to represent the evaluators' judgments directly. This can be justified by Example 2.

Example 2. Suppose that someone evaluates the complexity of a certain procedure. Let $S = \{s_{-3} = \text{extremely complex}, s_{-2} = \text{very complex}, s_{-1} = \text{complex}, s_0 = \text{medium}, s_1 = \text{easy}, s_2 = \text{very easy}, s_3 = \text{extremely easy}\}$ be a LTS. If one is uncertain, he/she can say it is “complex” ($\tilde{h}_S^1$) or “between medium and complex” ($\tilde{h}_S^2$). If one can make a more accurate judgment, he/she can say it is “between medium and complex but closes to complex” ($\tilde{h}_S^3$). If one can make an extremely accurate judgment, he/she can say it is “between

medium and complex but 80% proportion closes to complex” ($\tilde{h}_s^4$). Then the corresponding CIVLEs are $\tilde{h}_s^1=[s_{-1},s_{-1}]$, $\tilde{h}_s^2=[s_{-1},s_0]$, $\tilde{h}_s^3=[s_{-1},s_{-0.5}]$, $\tilde{h}_s^4=[s_{-1},s_{-0.8}]$, respectively. Suppose that the linguistic terms are uniformly distributed. Fig. 2 shows the syntax and semantics of $\tilde{h}_s^3$.

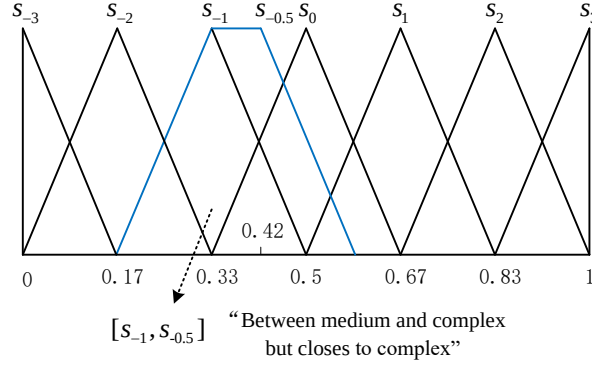


Fig. 2. The syntax and semantic of the CIVLE $\tilde{h}_s^3$

It should be noted that the absolute deviation between adjacent linguistic terms is not always equal. For example, the deviation between “medium” and “high” may be large than the deviation between “high” and “very high” in term of the quality of a product. That is to say, the symbols and semantics of the linguistic terms are disproportionate under some situations. Thus, it is irrational to calculate the linguistic terms directly by the their subscripts. Wang et al. [13] proposed some transformation functions to translate the linguistic term into its semantic .

Let s_t , $t \in [-\tau, \tau]$ be a linguistic term, and φ_t be a numeric value which denotes the semantic of s_t . The transformation function between s_t and φ_t is defined as: $g : s_t \rightarrow \varphi_t$; $g^{-1} : \varphi_t \rightarrow s_t$ ($t \in [-\tau, \tau]$). Based on Ref. [13], (1) if the absolute deviations between adjacent linguistic terms are equal, then

$$g(s_t) = (\tau + t)/2\tau \quad (3)$$

(2) if they increases with the extension from s_0 , then

$$g(s_t) = \begin{cases} \frac{\rho^\tau - \rho^{-t}}{2\rho^\tau - 2}, & \text{if } t \in [-\tau, 0] \\ \frac{\rho^\tau + \rho^t - 2}{2\rho^\tau - 2}, & \text{if } t \in (0, \tau] \end{cases} \quad (4)$$

where $\rho = \sqrt[\tau]{d}$ and d is the deviation between s_0 and s_τ which can be determined according to the actual situation.

(2) if they decreases with the extension from s_0 , then

$$g(s_t) = \begin{cases} \frac{\tau^\lambda - (-t)^\lambda}{2\tau^\lambda}, & \text{if } t \in [-\tau, 0] \\ \frac{\tau^\gamma + (t)^\gamma}{2\tau^\lambda}, & \text{if } t \in (0, \tau] \end{cases} \quad (5)$$

where $\lambda, \gamma \in (0, 1]$ can be determined based on the practical case. Especially, if $\lambda, \gamma = 1$, then $g(s_t) = (\tau + t)/2\tau$.

Definition 6. Let $S = \{s_\alpha | \alpha = -\tau, \dots, -1, 0, 1, \dots, \tau\}$ be a LTS, and $\tilde{h}_S = [s_L, s_U]$, $\tilde{h}_S^1 = [s_{L_1}, s_{U_1}]$, $\tilde{h}_S^2 = [s_{L_2}, s_{U_2}]$ be three CIVLEs, then

(1) Union: $\tilde{h}_S^1 \cup \tilde{h}_S^2 = [\min\{s_{L_1}, s_{L_2}\}, \max\{s_{U_1}, s_{U_2}\}]$;

(2) Intersection: $\tilde{h}_S^1 \cap \tilde{h}_S^2 = [\max\{s_{L_1}, s_{L_2}\}, \min\{s_{U_1}, s_{U_2}\}]$; if $\max\{s_{L_1}, s_{L_2}\} > \min\{s_{U_1}, s_{U_2}\}$, then $\tilde{h}_S^1 \cap \tilde{h}_S^2 = \emptyset$;

(3) Complement: $\tilde{h}_S^C = [s_{-\tau}, s_L] \cup [s_U, s_\tau]$;

(4) $\tilde{h}_S^1 \oplus \tilde{h}_S^2 = [s_{L_1}, s_{U_1}] \oplus [s_{L_2}, s_{U_2}] = [g^{-1}(g(s_{L_1}) + g(s_{L_2})), g^{-1}(g(s_{U_1}) + g(s_{U_2}))]$;

(5) $\lambda \tilde{h}_S = \lambda [s_L, s_U] = [g^{-1}(\lambda g(s_L)), g^{-1}(\lambda g(s_U))]$, where $\lambda \in [0, 1]$;

(6) $(\tilde{h}_S)^\lambda = [g^{-1}(g(s_L)^\lambda), g^{-1}(g(s_U)^\lambda)]$, where $\lambda \in [0, 1]$.

Example 3. Let $S = \{s_{-3}, \dots, s_0, \dots, s_3\}$ be a LTS, $\tilde{h}_S = [s_{1.2}, s_{1.8}]$, $\tilde{h}_S^5 = [s_{-1}, s_0]$, $\tilde{h}_S^6 = [s_0, s_1]$ be three CIVLEs and $\lambda = 0.8$. Then $\tilde{h}_S^5 \cup \tilde{h}_S^6 = [s_{-1}, s_1]$; $\tilde{h}_S^5 \cap \tilde{h}_S^6 = [s_0, s_0]$; $\tilde{h}_S^C = [s_{-3}, s_{1.2}] \cup [s_{1.8}, s_3]$; if $g(s_t)$ is given as Eq. (4) and $\rho = 2$, then $\lambda \tilde{h}_S^5 \oplus (1 - \lambda) \tilde{h}_S^6 = [s_{-0.85}, s_{0.26}]$; $(\tilde{h}_S)^\lambda = [s_{1.69}, s_{2.09}]$.

To compare the CIVLEs, we define the expect function of $\tilde{h}_s = [s_L, s_U]$ based on the transformation function as:

$$E(\tilde{h}_s) = \frac{1}{2}(g(s_L) + g(s_U)) \quad (6)$$

If $E(\tilde{h}_s^1) > E(\tilde{h}_s^2)$, then $\tilde{h}_s^1 \succ \tilde{h}_s^2$; if $E(\tilde{h}_s^1) < E(\tilde{h}_s^2)$, then $\tilde{h}_s^1 \prec \tilde{h}_s^2$; if $E(\tilde{h}_s^1) = E(\tilde{h}_s^2)$, we further define the variance function to make comparison as:

$$D(\tilde{h}_s) = \sqrt{\left(g(s_L) - E(\tilde{h}_s)\right)^2 + \left(g(s_U) - E(\tilde{h}_s)\right)^2} \quad (7)$$

When $E(\tilde{h}_s^1) = E(\tilde{h}_s^2)$, if $D(\tilde{h}_s^1) < D(\tilde{h}_s^2)$, then $\tilde{h}_s^1 \succ \tilde{h}_s^2$; if $D(\tilde{h}_s^1) > D(\tilde{h}_s^2)$, then $\tilde{h}_s^1 \prec \tilde{h}_s^2$; if $D(\tilde{h}_s^1) = D(\tilde{h}_s^2)$, then $\tilde{h}_s^1 \approx \tilde{h}_s^2$.

Example 4. Let $S = \{s_{-3}, \dots, s_0, \dots, s_3\}$ be a LTS, and $\tilde{h}_s^5 = [s_{-1}, s_0]$, $\tilde{h}_s^6 = [s_0, s_1]$ and $\tilde{h}_s^7 = [s_{0.5}, s_{0.5}]$ be three CIVLEs. If $g(s_i)$ is given as Eq. (5) and $\lambda = 0.6$, $\gamma = 0.7$, we can get $E(\tilde{h}_s^5) = 0.36$, $E(\tilde{h}_s^6) = 0.615$ and $E(\tilde{h}_s^7) = 0.64$. Thus $\tilde{h}_s^7 > \tilde{h}_s^5 > \tilde{h}_s^6$.

3. Methods to aggregate individual decision matrices to collective matrix

3.1 Description of the MCGDM problem with CIVLEs

A general MCGDM method consists of a finite set of m alternatives $A = \{a_1, \dots, a_i, \dots, a_m\}$, a set of n criteria $C = \{c_1, \dots, c_j, \dots, c_n\}$, and a set of Q DMs $E = \{e_1, \dots, e_q, \dots, e_Q\}$. The DM e_q is supposed to offer the evaluation value for alternative a_i with respect to criterion c_j in CIVLE, namely, $\tilde{h}_s^{ij(q)} = [s_L^{ij(q)}, s_U^{ij(q)}]$. Then we can construct Q judgment matrices $D^{(q)} = (\tilde{h}_s^{ij(q)})_{m \times n}$, $q = 1, \dots, Q$.

Due to the complexity of the MCGDM problem and the ambiguity of human thoughts as well as the different opinions among the DMs, in practice, it is hard to assign a crisp weight to each criterion [27]. Generally, the criterion importance degree ranges in fuzzy interval, such as “*between importance and very importance*”. Consequently, the precise criterion weights, which are given by the DMs directly or obtained

by some techniques such as the AHP [28], the entropy function [29] and the prioritized operator-based method [30], may result in information distortion and thus reduce the reliability of the final decision results. In this sense, the DM e_q is asked to give the weight of the criterion c_j in linguistic expression, which then can be transformed to the CIVLE $\tilde{h}_s^{j(q)}$. The collective criterion weight $\omega_j = \tilde{h}_s^j$ is the aggregation value of the CIVLEs $\tilde{h}_s^{j(q)}$ ($q=1, \dots, Q$) given by the DMs.

3.2 Aggregating group opinions with CIVLEs

This subsection proposes some aggregation methods to integrate the individual judgment matrices $D^{(q)}$, $q=1, \dots, Q$ to the group decision matrix $D = (\tilde{h}_s^{ij})_{m \times n}$.

Sometimes we suppose that the DMs have equal weights. However, in most cases, different DMs should have different weights because their different knowledge and experience may lead to the discrepancies in evaluation quality [31]. There are some methods to determine the weights of the DMs, such as the consistency judgement method [14] and the cluster analysis based method [15]. These methods are complicated and do not consider the different characteristics of the group members. In this section, we divide the MCGDM problems into four types according to the scale of the group. Different aggregation methods can be used with respect to different types. It should be noted that below we only give the aggregation methods over the assessments on alternatives, and the aggregation on the weights of criteria is the same.

(1) Small size group. For a group of less than three members, as it is easy to compromise with each other, computing the union is an appropriate method to integrate the DMs' evaluations, shown as Eq. (8). If there are few prejudices, we can delete them.

$$\tilde{h}_s^{ij} = \tilde{h}_s^{ij(1)} \cup \dots \cup \tilde{h}_s^{ij(q)} \cup \dots \cup \tilde{h}_s^{ij(Q)} = [\min\{s_L^{ij(1)}, \dots, s_L^{ij(q)}, \dots, s_L^{ij(Q)}\}, \max\{s_U^{ij(1)}, \dots, s_U^{ij(q)}, \dots, s_U^{ij(Q)}\}] \quad (8)$$

(2) Medium size group. For a medium scale group of three to five members, generally, there may be different opinions but all maintain referential significance. Thus, we can assign the same weight to the DMs. The average arithmetic aggregation formula is shown in Eq. (9). Note that if there are few prejudices, we can reject them; if the disparity of the evaluation quality is great, we can give different weights calculated

by Eq. (11).

$$\tilde{h}_s^{ij} = [s_L^{ij}, s_U^{ij}] = [\frac{1}{Q} \sum_{q=1}^Q s_L^{ij(q)}, \frac{1}{Q} \sum_{q=1}^Q s_U^{ij(q)}] \quad (9)$$

(3) Medium to large size group. For a group of six to thirty members with different knowledge and experience, the weight of DM e_q , $w^{(q)}$, may be assigned in advance. Then the collective opinion can be calculated by the weighted arithmetic aggregation operator shown as:

$$\tilde{h}_s^{ij} = [s_L^{ij}, s_U^{ij}] = [w^{(q)} \sum_{q=1}^Q s_L^{ij(q)}, w^{(q)} \sum_{q=1}^Q s_U^{ij(q)}] \quad (10)$$

(4) Large size group. For a large-scale group of more than thirty members, it is appropriate to suppose that the evaluations determined by the DMs obey Gaussian distribution given that most of them hold the similar opinions but a few of them insist different opinions since the evaluation on an alternative is affected by many small independent random factors. In this case, assigning the same weight to each DM is obviously unreasonable. Low weights should be given to the “false” or “biased” judgments while high weights should be assigned to the mid evaluations. The probability density function of Gaussian distribution for a random variable x is defined as $f(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-[(x-u)^2/2\sigma^2]}$, $x \in (-\infty, +\infty)$ where u is the mean value and σ is the standard deviation of x . The farther x away from u is, the smaller the value of $f(x)$ is. Inspired by this property, Xu [29] used $f(q)$ to represent the weight of each individual where q is the order of the evaluation value. However, there are some flaws in Xu’s method: 1) the discrete orders, $1, \dots, q, \dots, Q$, essentially, are disobeyed to the Gauss distribution; 2) the differences of the evaluation values were ignored in Xu’s method (which may lead to an unacceptable result that the same evaluations may get different weights while the different judgments may get the same weight); 3) it is limited to handle the linguistic evaluations.

To avoid the above flaws, we introduce an EGDBM, which utilizes the interval-valued linguistic evaluation value itself as random value, to calculate the weight of DM. Then, we can calculate the upper and lower limits of the group CIVLEs by aggregating the upper and lower bounds of the individual CIVLEs, respectively.

Let $W_L = (w_L^{(1)}, \dots, w_L^{(q)}, \dots, w_L^{(Q)})^T$ be the weight vector of the lower limits $L = (s_L^{(1)}, \dots, s_L^{(q)}, \dots, s_L^{(Q)})^T$

and $W_U = (w_U^{(1)}, \dots, w_U^{(q)}, \dots, w_U^{(Q)})^T$ be the weight vector of the upper limits $U = (s_U^{(1)}, \dots, s_U^{(q)}, \dots, s_U^{(Q)})^T$.

Based on the probability density function of Gaussian distribution, we can determine the probability density value of each lower limit as $f(s_L^{(q)}) = \frac{1}{\sqrt{2\pi}\sigma} e^{-[(L^{(q)} - u_L)^2 / 2(\sigma_L)^2]}$, $f(s_U^{(q)}) = \frac{1}{\sqrt{2\pi}\sigma} e^{-[(U^{(q)} - u_U)^2 / 2(\sigma_U)^2]}$. After normalization, the weights of the lower and upper limits are respectively calculated as

$$w_L^{(q)} = \frac{e^{-[(L^{(q)} - u_L)^2 / 2(\sigma_L)^2]}}{\sum_{q=1}^Q e^{-[(L^{(q)} - u_L)^2 / 2(\sigma_L)^2]}}, \quad w_U^{(q)} = \frac{e^{-[(U^{(q)} - u_U)^2 / 2(\sigma_U)^2]}}{\sum_{q=1}^Q e^{-[(U^{(q)} - u_U)^2 / 2(\sigma_U)^2]}}, \quad q = 1, \dots, Q \quad (11)$$

where u_L and σ_L are the mean and variance of the lower limits, u_U and σ_U are the mean and variance of the upper limits, and $L^{(q)}$ and $U^{(q)}$ are the subscripts of $s_L^{(q)}$ and $s_U^{(q)}$, respectively.

Then, we can obtain the group assessments as

$$\tilde{h}_S^{ij} = [s_L^{ij}, s_U^{ij}] = [\sum_{q=1}^Q w_L^{(q)} s_L^{ij(q)}, \sum_{q=1}^Q w_U^{(q)} s_U^{ij(q)}] \quad (12)$$

Example 5. Let $S = \{s_{-3}, \dots, s_0, \dots, s_3\}$ be a LTS. Suppose that there are thirty teachers e_q , $q = 1, \dots, 30$, who are invited to evaluate the comprehensive performance of students. Suppose that the teachers' judgments are given in CIVLEs as: $[s_{-2}, s_{-1}](1)$, $[s_{-2}, s_0](1)$, $[s_0, s_0](1)$, $[s_0, s_1](2)$, $[s_{0.5}, s_1](4)$, $[s_{0.5}, s_{1.5}](1)$, $[s_1, s_{1.5}](8)$, $[s_1, s_2](2)$, $[s_{1.2}, s_2](2)$, $[s_{1.5}, s_2](5)$, $[s_{1.5}, s_{2.5}](1)$ and $[s_2, s_{2.5}](2)$, where the number in each round bracket represents the number of teachers who provide the interval evaluation value. The lower limits of these evaluation values are $L = (s_{-2}(2), s_0(3), s_{0.5}(5), s_1(10), s_{1.2}(2), s_{1.5}(6), s_2(2))^T$. Since $u_L = 0.8$, $\sigma_L = 0.9$, by Eq. (11), we obtain $W_L = (0(2), 0.029(3), 0.041(5), 0.042(10), 0.038(2), 0.031(6), 0.017(2))^T$. Then we have $s_L = \sum_{q=1}^Q w_L^{(q)} s_L^{(q)} = s_{0.96}$. Similarly, we can obtain $W_U = (0(1), 0.007(2), 0.037(6), 0.044(9), 0.035(9), 0.018(3))^T$ and $s_U = s_{1.58}$. Hence, the overall evaluation of the group is $\tilde{h}_S = [s_{0.96}, s_{1.58}]$. We can find that $\tilde{h}_S^{(1)} = [s_{-2}, s_{-1}]$, which can be taken as a “biased” evaluation, is assigned a very small weight. Fig. 3 shows the Gaussian distribution values of the lower and upper limits of the evaluations.

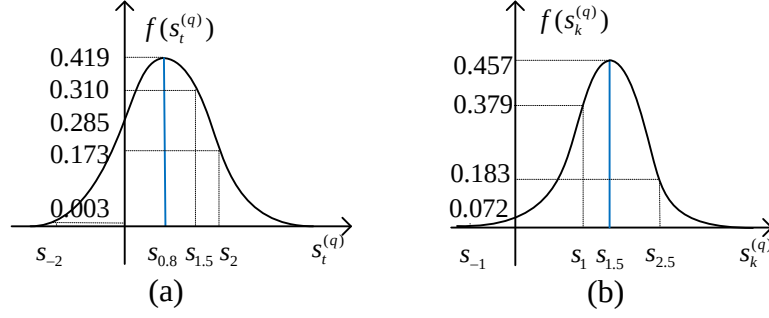


Fig. 3. The Gaussian distribution values of the lower and upper limits derived by the EGDBW method

Remark 4. If the weight of DM is assigned in advance, we can take into consideration both the assigned weight $w^{(q)}$ and the calculated weights $w_t^{(q)}$ and $w_k^{(q)}$. In this sense, the comprehensive weight of DM can be computed by:

$$\tilde{w}_L^{(q)} = \left(w^{(q)} w_L^{(q)} \right) / \sum_{q=1}^Q \left(w^{(q)} w_L^{(q)} \right), \quad \tilde{w}_U^{(q)} = \left(w^{(q)} w_U^{(q)} \right) / \sum_{q=1}^Q \left(w^{(q)} w_U^{(q)} \right) \quad (13)$$

4. Continuous interval-valued linguistic ORESTE method

In this section, we propose the CIVL-ORESTE method to rank the alternatives for the MCGDM problem in which the group decision matrix has been obtained by the aggregation methods presented in Section 3.

4.1 The classical ORESTE method

The classical ORESTE [25] consists of two stages: building the global weak ranking after computing the preference scores and building the PIR structure after an indifference and incomparability analysis (here “incomparability” is represented as “R” to distinguish it from “indifference”). In the ORESTE method, the weight of criterion is not assigned, but just given a preference structure represented by a Besson’s mean rank r_j . The merit of alternative a_i under criterion c_j is also represented as the Besson’s mean rank $r_j(a_i)$ (for more details of Besson’s mean ranks, please refer to Ref. [32, 33]). The steps of the ORESTE method are as follows.

Step 1. Let the action of alternative a_i under criterion c_j be expressed as a_{ij} . The global preference score $\tilde{D}(a_{ij})$ of a_{ij} is computed by

$$\tilde{D}(a_{ij}) = \sqrt{\varsigma r_j^2 + (1-\varsigma)r_j(a_i)^2} \quad (14)$$

where ς is the coefficient to weight the rank of the criterion and that of the alternative.

Step 2. Determine the global weak rank $r(a_{ij})$. If $\tilde{D}(a_{ij}) > \tilde{D}(a_{zv})$, $r(a_{ij}) > r(a_{zv})$; else if $\tilde{D}(a_{ij}) = \tilde{D}(a_{zv})$, $r(a_{ij}) = r(a_{zv})$, where $i, z = 1, 2, \dots, m$ and $j, v = 1, 2, \dots, n$.

Step 3. Calculate the weak rank $R(a_i)$, where

$$R(a_i) = \sum_{j=1}^n r(a_{ij}) \quad (15)$$

Step 4. Set up the PIR structure. The average preference intensity between a_i and a_z is defined as:

$$\tilde{T}(a_i, a_z) = \frac{\sum_{j=1}^n \max[r(a_{zj}) - r(a_{ij}), 0]}{(m-1)n^2} \quad (16)$$

The net preference intensity between a_i and a_z is defined as:

$$\Delta\tilde{T}(a_i, a_z) = \tilde{T}(a_i, a_z) - \tilde{T}(a_z, a_i) \quad (17)$$

The principle of the indifference and incomparability test is: If $|\Delta\tilde{T}(a_i, a_z)| \leq \tilde{\mu}$ and $\tilde{T}(a_i, a_z) \leq \tilde{\gamma}$, then $a_i I a_z$; if $\tilde{T}(a_i, a_z) / |\Delta\tilde{T}(a_i, a_z)| \geq \lambda$, then $a_i R a_z$; otherwise, if $\Delta\tilde{T}(a_i, a_z) > 0$, then $a_i P a_z$; if $\Delta\tilde{T}(a_i, a_z) < 0$, then $a_z P a_i$. The values of the thresholds $\tilde{\mu}$, $\tilde{\gamma}$ and λ are calculated as follows (for more details, refer to Ref. [33]).

$$\tilde{\mu} < 1/(m-1)n, \quad \tilde{\gamma} < \tilde{\delta}/2(m-1), \quad \lambda > (n-2)/4 \quad (18)$$

where $\tilde{\delta}$ is the minimal rank difference between alternatives a_i and a_z under criterion c_j to separate the indifference and incomparability relation. Its value is given by DM in practice.

Step 5. The results are a joint decision based on the weak rank $R(a_i)$ and the PIR structure.

Researchers subsequently analyzed its characteristics. Bourguignon and Massart [32] analyzed the necessity and significance to distinguish the indifference and incomparability relation between alternatives deeply. Pastijn and Leysen [33] carried detailed analysis and explanation on the values of thresholds in the

indifference and incomparability analysis framework. Then a sensitivity analysis for the thresholds was employed by Delhaye et al. [34], which indicated that different values have different influences on results. It has been applied in the various fields, such as agricultural investment decision [35], and Radar detection strategy selection [36], web design firm selection [37] and the firm performance efficiency order construction [38].

However, (1) the decision matrix handled by the ORESTE contains less evaluations; (2) translating the global preference scores to global weak ranks makes information loss; (3) the thresholds $\tilde{\delta}$ is hard to determine. To overcome these defects, we improve the ORESTE method, and then combine it with the CIVLTSs in the next subsection.

4.2 The CIVL-ORESTE method for MCGDM

In this part, the CIVL-ORESTE method is developed to rank the alternatives according to the collective decision matrix $D = (a_{ij})_{m \times n}$ and the weight vector $W = (\omega_1, \dots, \omega_j, \dots, \omega_n)^T$ of the criteria.

In classical ORESTE, the Besson's mean ranks $r_j(1, \dots, n)$ and $r_j(a_i)$ ($i = 1, \dots, m$) would result in information loss seriously. Example 6 can demonstrate this point.

Example 6. Suppose that three hospitals a_1, a_2, a_3 need to be assessed according to their medical levels and $S = \{s_{-3} = \text{extremely poor}, s_{-2} = \text{very poor}, s_{-1} = \text{poor}, s_0 = \text{medium}, s_1 = \text{good}, s_2 = \text{very good}, s_3 = \text{extremely good}\}$ is a given LTS. The linguistic evaluations are “ a_1 is good”, “ a_2 is between good and very good and close to good” and “ a_3 is between poor and very poor”, respectively. The corresponding CIVLEs are $\tilde{h}_s(a_1) = [s_1, s_1]$, $\tilde{h}_s(a_2) = [s_1, s_{1.5}]$ and $\tilde{h}_s(a_3) = [s_{-2}, s_{-1}]$. According to the comparison method for CIVLEs given as Eqs. (6-7), we can get the ranks $r(a_1) = 2$, $r(a_2) = 1$ and $r(a_3) = 3$. Clearly, for the medical level, a_1 is extremely close to a_2 but a_3 is far behind a_1 and a_2 . However, the ranks reflect the same degree of difference between the hospitals and thus weaken the information seriously.

Therefore, the ranks in operations are supposed to be replaced. To maintain the evaluation information

completely, the distance measure is designed to substitute for the ranks. Motivated by the Euclidean distance between HFLEs [17], we define the Euclidean distance between CIVLEs as follows.

Definition 7. Let $S = \{s_\alpha | \alpha = -\tau, \dots, -1, 0, 1, \dots, \tau\}$ be a LTS, and $\tilde{h}_S^1 = [s_L^1, s_U^1]$ and $\tilde{h}_S^2 = [s_L^2, s_U^2]$ be two CIVLEs on S . The Euclidean distance between $\tilde{h}_S^1$ and $\tilde{h}_S^2$ is

$$d(\tilde{h}_S^1, \tilde{h}_S^2) = \left[\frac{1}{2} \left(\frac{|L_1 - L_2|}{2\tau} \right)^2 + \frac{1}{2} \left(\frac{|U_1 - U_2|}{2\tau} \right)^2 \right]^{1/2} \quad (19)$$

Apparently, $0 \leq d(\tilde{h}_S^1, \tilde{h}_S^2) \leq 1$ and $d(\tilde{h}_S^1, \tilde{h}_S^2) = d(\tilde{h}_S^2, \tilde{h}_S^1)$. The smaller the distance is, the similar of $\tilde{h}_S^1$ and $\tilde{h}_S^2$ should be. Especially, if $d(\tilde{h}_S^1, \tilde{h}_S^2) = 0$, then $\tilde{h}_S^1 = \tilde{h}_S^2$.

Example 7. The Euclidean distance between $\tilde{h}_S(a_1)$, $\tilde{h}_S(a_2)$ and $\tilde{h}_S(a_3)$ in Example 6 are: $d(\tilde{h}_S(a_1), \tilde{h}_S(a_2)) = 0.0589$, $d(\tilde{h}_S(a_1), \tilde{h}_S(a_3)) = 0.5$ and $d(\tilde{h}_S(a_2), \tilde{h}_S(a_3)) = 0.4602$. It is clear that a_1 is similar to a_2 , and a_3 is highly different from a_1 and a_2 , which is fully in accordance with the reality.

We define the maximum CIVLE of a_i under criterion c_j as

$$\tilde{h}_S^{j+} = \begin{cases} \max_{i=1,2,\dots,m} \{\tilde{h}_S^{ij}\}, & \text{for the benefit criterion } c_j \\ \min_{i=1,2,\dots,m} \{\tilde{h}_S^{ij}\}, & \text{for the cost criterion } c_j \end{cases} \quad (20)$$

Additionally, the weight of the most important criterion c_j satisfies

$$\omega^+ = \max_{j=1,\dots,n} \omega_j = \max_{j=1,\dots,n} \{\tilde{h}_S^j\} \quad (21)$$

Let the distance $d(\tilde{h}_S^{ij}, \tilde{h}_S^{j+})$ be abbreviated as d_{ij} which is used to replace $r_j(a_i)$; the distance $d(\omega_j, \omega^+)$ be abbreviated as d_j to replace r_j . Then, following the classical ORESTE method, the operation processes of the CIVL-ORESTE are divided into two stages based on d_{ij} and d_j .

Stage 1. Construct a weak ranking

(1) Compute the global preference score $D(a_{ij})$. Let the coordinate of the action a_{ij} be represented

as (d_{ij}, d_j) ; let the global optimal point a_{ij+}^+ at coordinate origin be the best alternative under the most important criteria. Then, we introduce the weighted Euclidean distance between a_{ij} and a_{ij+}^+ as the global preference score $D(a_{ij})$ of a_{ij} :

$$D(a_{ij}) = d(a_{ij}, a_{ij+}^+) = \left[\xi(d_{ij})^2 + (1-\xi)(d_j)^2 \right]^{1/2} \quad (22)$$

where $\xi \in [0,1]$ denotes the relative importance between d_{ij} and d_j . This paper deems d_{ij} and d_j are equally important, and $\xi = 0.5$. Like the Euclidean distance between two CIVLEs, $D(a_{ij}) \in [0,1]$, and the smaller it is, the better a_{ij} should be.

(2) Compute the preference score $D(a_i)$. The preference score of alternative a_i is defined as the average of the global preference score of $a_{ij} (i=1, \dots, m)$:

$$D(a_i) = \frac{1}{n} \sum_{j=1}^n D(a_{ij}) \quad (23)$$

(3) Get the weak ranking. According to the preference score $D(a_i)$, we can obtain the weak relations between the alternatives. ① If $D(a_i) > D(a_z)$, then $r(a_i) > r(a_z)$, which is denoted as $a_i P a_z$; ② if $D(a_i) = D(a_z)$, then $r(a_i) = r(a_z)$, which is denoted as $a_i I a_z$. $r(a_i)$ is the weak rank of a_i over all alternatives (here, the “weak” ranking is named because the PI relations are only obtained by $D(a_i)$).

However, the accurate relations between alternatives are unable to be determined by the global preference score if $D(a_i)$ is large but extremely close to $D(a_z)$. The relation P assigned to a_i and a_z is unacceptable. In addition, the P and I relations cannot fully describe the relationship between the alternatives and the incomparability (R) relation must be distinguished. If $D(a_i) \approx D(a_z)$ but there are great difference between $D(a_{ij})$ and $D(a_{zj})$ under some criteria c_j , $j \in 1, 2, \dots, n$, we cannot deem $a_i I a_z$. Therefore, it is necessary to further differentiate the specific relationships between alternatives,

which is sorted out in next stage.

Stage 2. Establish the PIR structure

(1) Compute the preference intensities. Like the classical ORESTE method, the preference intensities between two alternatives are utilized to obtain the PIR relations and make the decision result acceptable. Based on the global preference scores, the preference intensity between a_i and a_z under criterion c_j is defined as:

$$T_j(a_i, a_z) = \max \left\{ \left[D(a_{zj}) - D(a_{ij}) \right], 0 \right\} \quad (24)$$

The average preference intensity between a_i and a_z is defined as:

$$T(a_i, a_z) = \frac{1}{n} \sum_{j=1}^n \max \left\{ \left[D(a_{zj}) - D(a_{ij}) \right], 0 \right\} \quad (25)$$

The net preference intensity between a_i and a_z is defined as:

$$\Delta T(a_i, a_z) = T(a_i, a_z) - T(a_z, a_i) \quad (26)$$

Obviously, $0 \leq T_j(a_i, a_z) \leq 1$, $0 \leq T(a_i, a_z) \leq 1$ and $0 \leq |\Delta T(a_i, a_z)| \leq 1$.

(2) Determine the thresholds. In CIVL-ORETSE, the PIR structure of alternatives is constructed by three thresholds: the indifference threshold (δ) to differentiate the indifference relation and the incomparability relation for each criterion, the preference threshold (μ) to separate the preference relation with the indifference relation and the incomparability relation, and the incomparability threshold (γ) to distinguish the indifference relation and the incomparability relation for all criteria. Since the ranks $r_j(a_i)$ and r_j are substituted by the distances d_{ij} and d_j , the thresholds used in the CIVL-ORESTE method are different from those in the classical ORESTE method. These thresholds are determined based on the distance between CIVLEs.

If $d(\tilde{h}_s^i, \tilde{h}_s^z) < \varepsilon$, we suppose $\tilde{h}_s^i$ is indifference to $\tilde{h}_s^z$, where ε is the CIVL indifference threshold.

In general, we deem there is absolute difference between $\tilde{h}_s^i$ and $\tilde{h}_s^z$ if $\tilde{h}_s^i = [1, 1]$ and $\tilde{h}_s^z = [1, 1.5]$. Then,

$d(\tilde{h}_s^i, \tilde{h}_s^z) = \frac{\sqrt{2}}{2} * \frac{0.5}{2\tau}$. It is a boundary to judge whether a_i is indifferent to a_z , so $\varepsilon \in [0, \frac{\sqrt{2}}{2} * \frac{0.5}{2\tau}]$.

At the first stage of the CIVL-ORESTE method, we have $|d_{ij} - d_{zj}| = |d(\tilde{h}_s^{ij}, \tilde{h}_s^{j+}) - d(\tilde{h}_s^{zj}, \tilde{h}_s^{j+})| \approx d(\tilde{h}_s^{ij}, \tilde{h}_s^{zj})$. Then the relation between $\tilde{h}_s^{ij}$ and $\tilde{h}_s^{zj}$ is indifferent if $|d_{ij} - d_{zj}| < \varepsilon$. Furthermore, to get the indifference relation between a_{ij} and a_{zj} under criterion c_j by the value of $|D(a_{ij}) - D(a_{zj})|$, we carry out the approximate calculation based on ε (To facilitate calculation, let $d_j = 0$, which does not influence the result):

$$|D(a_{ij}) - D(a_{zj})| = \left| \left[\frac{1}{2}(d_{ij})^2 + \frac{1}{2}(d_j)^2 \right]^{1/2} - \left[\frac{1}{2}(d_{zj})^2 + \frac{1}{2}(d_j)^2 \right]^{1/2} \right| = \frac{\sqrt{2}}{2} |d_{ij} - d_{zj}|$$

Definition 8. Let the preference intensity between a_i and a_z under criterion c_j be $T_j(a_i, a_z) = \max \{ [D(a_{zj}) - D(a_{ij})], 0 \}$. Suppose that a_i is indifferent to a_z under criterion c_j if $0 \leq T_j(a_i, a_z) \leq \delta_j$ with $\delta_j = \frac{\sqrt{2}}{2} \varepsilon$. δ_j is called the indifference threshold for criterion c_j .

Remark 5. If all criteria in a MCGDM problem adopt the same length of LTS, $\delta_j (j = 1, \dots, n)$ are the same, expressed as δ (in this paper, we only analyze the same δ for each criteria). For the commonly used seven LTS, $2\tau = 6$, thus $\varepsilon \in [0, 0.0589]$ and $\delta \in [0, 0.0416]$. In addition, $\delta \in [0, 0.0416]$ is only a reference value range that can range properly according to the practice problems.

In the CIVL environment, the I and R relations between a_i and a_z occur when their net preference intensities are equal or very close. The values of the thresholds μ and γ are discussed in detail by dividing the relation between a_i and a_z into three situations.

Situation 1. Let $a_i P a_z$ if $|T(a_i, a_z) - T(a_z, a_i)| \geq \frac{\delta}{n}$. Considering the extreme case that $T_j(a_i, a_z) - T_j(a_z, a_i) = 0$ ($j = 1, 2, \dots, n-1$), that is to say, as for $n-1$ criteria,

$\frac{1}{n-1} \sum_{j=1}^{n-1} \max[D(a_{ij}) - D(a_{zj}), 0] = 0$; as for the n th criterion, $|T_n(a_i, a_z)| \geq \delta$. In this case,

$|T(a_i, a_z) - T(a_z, a_i)| = \frac{\delta}{n}$, which is the minimum case for $a_i P a_z$. Therefore, let $\mu = \frac{\delta}{n}$ be the preference

threshold. Table 1 presents an example to illustrate this case (Let $n = 4$ and $\delta = 0.03$, then $\mu = 0.0075$).

Table 1. Global preference scores for the preference relation

Score	c_1	c_2	c_3	c_4
a_i	0.5	0.5	0.5	0.5
a_z	0.5	0.5	0.5	0.6

Situation 2. Let $a_i I a_z$ if $|T(a_i, a_z) - T(a_z, a_i)| < \frac{\delta}{n}$ and $T_j(a_i, a_z) < \delta$ ($j = 1, \dots, n$). In this case, if

n is odd, $T(a_i, a_z) < \frac{1}{n} \left(\frac{n}{2} \delta + \delta \right) = \frac{(n+2)\delta}{2n}$; if n is even, $T(a_i, a_z) < \frac{1}{n} \left(\frac{n}{2} \delta \right) = \frac{\delta}{2}$ (The relation a_z

to a_i should also satisfy the above conditions). We denote the indifference threshold as γ that

$\gamma = \frac{(n+2)\delta}{2n}$ if n is odd; and $\gamma = \frac{\delta}{2}$ if n is even.

Table 2 and Table 3 are the examples of the indifference relation (let $\delta = 0.03$, then $\alpha = 0.015$ if $n = 4$ and $\alpha = 0.021$ in case $n = 5$).

Table 2. Global preference scores for the indifference relation when n is even

Score	c_1	c_2	c_3	c_4
a_i	0.5	0.52	0.51	0.5
a_z	0.52	0.5	0.5	0.51

Table 3. Global preference scores for the indifference relation when n is odd

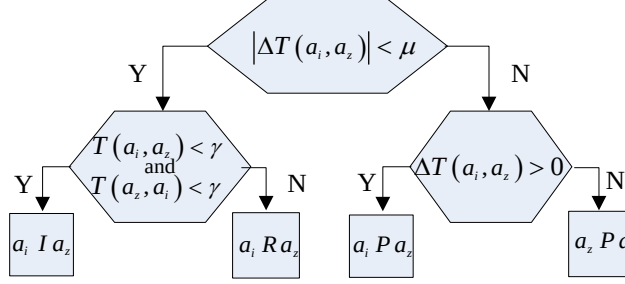
Score	c_1	c_2	c_3	c_4	c_5
a_i	0.5	0.51	0.52	0.5	0.5
a_z	0.51	0.5	0.5	0.52	0.51

Situation 3. Let $a_i R a_z$ if $|T(a_i, a_z) - T(a_z, a_i)| < \frac{\delta}{n}$ and there is at least one criterion which satisfies $T_j(a_i, a_z) > \delta$ (The relation of a_z to a_i should both satisfy the above conditions). In this case, despite the net preference score of a_i for a_z is zero or close to zero, there are great differences on preference intensity under some criteria. Thus, a_i cannot be replaced by a_z , which is essential to differ from the I relation. Table 4 is an example of the incomparability relation.

Table 4. Global preference scores for the incomparability relation

Score	c_1	c_2	c_3	c_4
a_i	0.5	0.6	0.5	0.5
a_z	0.6	0.5	0.5	0.51

(3) Conduct the indifference and incomparability analysis (Establish the PIR structure). The process of the indifference and the incomparability analyses of the CIVL-ORESTE method is shown in Fig. 4.

**Fig. 4.** The indifference and incomparability analysis of the CIVL-ORESTE method

4.3 Algorithm of the CIVL-ORESTE method

To make the CIVL-ORESTE method easy to understand and convenient for application, we summarize the algorithm as follows.

Step 1. Establish the individual decision matrix $D^{(q)} = (a_{ij}^{(q)})_{m \times n}$ and the criterion weight vector $W^{(q)} = (\omega_1^{(q)}, \dots, \omega_j^{(q)}, \dots, \omega_n^{(q)})^T$ derived from each expert e_q . The evaluations on both the merits of alternatives and the importance of criteria are expressed in linguistic expressions; then they are translated to the CIVLEs $\tilde{h}_s^{ij(q)}$ and $\tilde{h}_s^{j(q)}$. Go to the next step.

Step 2. Establish the collective decision matrix $D = (a_{ij})_{m \times n}$ and the criteria weight vector $W = (\omega_1, \dots, \omega_j, \dots, \omega_n)^T$. The CIVLEs in D and W are expressed as $\tilde{h}_s^{ij}$ and $\tilde{h}_s^j$, respectively, which are calculated by aggregation methods proposed in Section 3 based on $\tilde{h}_s^{ij(q)}$ and $\tilde{h}_s^{j(q)}$, $(q=1, \dots, Q)$. Then go to the next step.

Step 3. Calculate the CIVL distance d_{ij} and d_j . Firstly, find out the maximum CIVLE $\tilde{h}_s^{j+}(j=1, \dots, n)$ and the maximum weight ω^+ in CIVL form by Eqs. (20) and (21), respectively. Then,

compute $d(\tilde{h}_s^{ij}, \tilde{h}_s^{j+})$ as d_{ij} and $d(\omega_j, \omega^+)$ as d_j by Eq. (19). Go to the next step.

Step 4. Calculate the global preference scores $D(a_{ij})(i=1, \dots, m)(j=1, \dots, n)$ by Eq. (22). Then compute the preference scores $D(a_i)(i=1, \dots, m)$ by Eq. (23) to get the weak rankings of all alternatives. Go to the next step.

Step 5. Calculate the preference intensities: $T_j(a_i, a_k)$, $T(a_i, a_k)$ and $\Delta T(a_i, a_k)$, $\forall i, k \in 1, \dots, m$, by Eqs. (24-26), respectively. Go to the next step.

Step 6. Determine the thresholds δ , μ and γ according to the reference values discussed above and establish the PIR structure according to Fig. 4.

Step 7. Obtain the strong rankings of all alternatives based on the weak rankings and the PIR structure.

5. A case study: "Mobike" sharing bike design selection in China

This section uses a case study concerning the selection of the innovative "Mobike" sharing bike design in Chinese market to illustrate the feasibility and effectiveness of the CIVL-ORESTE method.

5.1 Case description

Dedicated to solving the "last few kilometers of travel" problem, since the second half of 2016, sharing bikes (or bike rental) has appeared in major cities in China, and attracted great attentions. Sharing bike is a new form of sharing economy that enterprises usually cooperate with the government. It provides bicycle sharing service on campus, subway stations, residential areas and commercial areas. It adopts the Internet mobile terminal technology so that the users can use the mobile phone APP to locate bikes, and there is no limit of place and time for taking and parking bikes. Furthermore, bike rents and deposits can be paid on line. As a powerful tool for short trip (from subway stations to home or company offices, from dormitory to teaching building, riding for tourism, etc.), sharing bike has brought great convenience for people to travel and gain social recognition. It is characterized by the satisfaction of rigid demand for trip and environment protection requirements, which results in a sharp rise in demand. Due to significant market dividends from sharing bike, capitals turn into this market in such a rapid way that a growing number of

sharing bike brands are emerging. In addition to main brands such as “Mobike” and “OFO”, close to 20 brands have entered to this market, such as “Youon”, “Baicycle”, “Bluegogo”, etc. These brands are constantly expanding their market layout, trying to carve up the market to establish their positions in the entire sharing bike market. Thus, a battle for users has started.

“Mobike” was officially released in April 2016. Considering the "stocking management", “Mobike” is committed to improving the durability of bikes to reduce manual maintenance intervention. Therefore, at the beginning of designing a bicycle, too much attention is paid to improving its quality, increasing durability and reducing maintenance costs, whilst user experience is ignored. Many problems, such as unwieldy body, hard mounts, unable to adjust the height of mounts, less additional functions, etc., seriously reduce users’ satisfaction and bicycle design has been criticized by many users, which leads to reduced competitiveness seriously. To win in the fierce competition, the “Mobike” company intends to select the optimal innovative design from several new designs, which can best meet the needs of users.

Choosing the optimal innovative design for “Mobike” is a typical MCGDM problems. According to a large number of survey and analyses, we have identified users' demands for sharing bikes and propose to employ comfort c_1 , convenience c_2 , versatility c_3 , security c_4 and riding speed c_5 as evaluation criteria. The corresponding weight vector $W = (\omega_1, \omega_2, \omega_3, \omega_4, \omega_5)^T$ is expressed as CIVLEs rather than crisp data. There are five design alternatives $a_i (i = 1, \dots, 5)$ to be evaluated. Two groups are invited to make decision for this problem. Group 1, which consists of 100 users $e_q^1 (q = 1, \dots, 100)$, aims to evaluate the weights of the criteria. Group 2, which consists of 6 experts (managers), $e_q^2 (q = 1, \dots, 6)$, is to judge the merits of each alternative with respect to each criterion. In this way, we do not only obtain the real demand preferences of users but also assess the alternatives professionally by the experts (or managers).

Let $S = \{s_{-3}, \dots, s_0, \dots, s_3\}$ be a LTS. The specific meanings of the linguistic terms for the alternatives’ merits with respect to each criterion are uniformly expressed as: $s_{-3} = none, s_{-2} = very\ bad, s_{-1} = bad, s_0 = medium, s_1 = good, s_2 = very\ good, s_3 = perfect$, and as for the weights of the criteria, the specific meanings are: $s_{-3} = extremely\ unimportant, s_{-2} = very\ unimportant, s_{-1} =$

unimportant, $s_0 = \text{medium}$, $s_1 = \text{important}$, $s_2 = \text{very important}$, $s_3 = \text{extremely important}$. The evaluation results expressed in CIVLEs from both Group 1 and Group 2 are shown respectively in Table 5 and Table 6. To simplify tables and save space, we put the evaluation values of the DMs in these two table together. In these two tables, the number in a parenthesis indicates the number of DMs who give the same CIVLEs, for example, $[s_3, s_3](2)$ means two DMs give the evaluation of $[s_3, s_3]$.

Table 5. The importance of the criteria evaluated by Group 1

Importance							
c_1	$[s_3, s_3](2)$	$[s_2, s_{2.5}](20)$	$[s_2, s_2](48)$	$[s_1, s_2](15)$	$[s_0, s_1](10)$	$[s_0, s_0](3)$	$[s_{-1}, s_{-1}](2)$
c_2	$[s_3, s_3](5)$	$[s_2, s_3](18)$	$[s_2, s_2](45)$	$[s_1, s_2](22)$	$[s_0, s_1](7)$	$[s_{-1}, s_0](3)$	-
c_3	$[s_2, s_3](1)$	$[s_2, s_2](7)$	$[s_1, s_{1.5}](16)$	$[s_0, s_1](51)$	$[s_{-1}, s_0](20)$	$[s_{-2}, s_0](4)$	$[s_{-2}, s_{-1}](1)$
c_4	$[s_2, s_3](3)$	$[s_2, s_{2.5}](6)$	$[s_0, s_2](21)$	$[s_0, s_1](45)$	$[s_0, s_0](20)$	$[s_{-1}, s_{-0.5}](4)$	$[s_{-2}, s_0](1)$
c_5	$[s_2, s_3](2)$	$[s_{1.5}, s_2](3)$	$[s_0, s_1](14)$	$[s_0, s_0](41)$	$[s_{-1}, s_{-0.5}](26)$	$[s_{-2}, s_{-1}](14)$	-

Table 6. The judgments on the alternatives under criteria given by Group 2

	c_1	c_2	c_3
a_1	$[s_2, s_3](1), [s_1, s_{1.5}](3), [s_1, s_1](2)$	$[s_0, s_0](1), [s_{-0.5}, s_0](3), [s_{-1}, s_{-1}](1), [s_{-2}, s_{-1}](1)$	$[s_2, s_3](1), [s_2, s_{2.5}](4), [s_2, s_2](1)$
a_2	$[s_2, s_2](2), [s_1, s_{1.5}](2), [s_1, s_1](2)$	$[s_0, s_0](1), [s_{-0.5}, s_0](3), [s_{-1.5}, s_{-1}](2)$	$[s_{2.5}, s_3](2), [s_2, s_{2.5}](3), [s_2, s_2](1)$
a_3	$[s_2, s_3](1), [s_2, s_{2.5}](3), [s_1, s_2](2)$	$[s_1, s_2](5), [s_1, s_1](1)$	$[s_2, s_3](1), [s_2, s_2](3), [s_{1.5}, s_2](2)$
a_4	$[s_0, s_0](1), [s_{-1}, s_0](4), [s_{-1}, s_{-1}](1)$	$[s_0, s_0](1), [s_{-1}, s_0](4), [s_{-2}, s_{-1}](1)$	$[s_2, s_3](1), [s_2, s_{2.5}](4), [s_1, s_2](1)$
a_5	$[s_{-1}, s_0](4), [s_{-2}, s_{-1}](2)$	$[s_0, s_0](1), [s_{-1}, s_0](4), [s_{-1}, s_{-0.5}](1)$	$[s_{-1}, s_0](2), [s_{-1}, s_{-1}](4),$
	c_4	c_5	
a_1	$[s_{-1.5}, s_{-1}](3), [s_{-2}, s_{-1}](3)$	$[s_0, s_1](5), [s_0, s_0](1)$	
a_2	$[s_{-1}, s_{-1}](2), [s_{-2}, s_{-1}](4)$	$[s_0, s_1](2), [s_0, s_{0.5}](3), [s_{-1}, s_0](1)$	
a_3	$[s_{-1}, s_0](3), [s_{-1}, s_{-1}](3)$	$[s_1, s_1](1), [s_{0.5}, s_1](2), [s_0, s_1](2), [s_0, s_0](1)$	
a_4	$[s_{-1}, s_{-1}](3), [s_{-2}, s_{-1}](3)$	$[s_2, s_3](2), [s_2, s_{2.5}](3), [s_1, s_2](1)$	
a_5	$[s_2, s_3](2), [s_2, s_{2.5}](3), [s_1, s_2](1)$	$[s_2, s_3](2), [s_2, s_2](2), [s_1, s_2](2)$	

5.2 Solving the case by the CIVL-ORESTE method

Below we use the CIVL-ORESTE method to select the optimal innovative sharing bike design based on the evaluation information in CIVLEs given by Group 1 and Group 2.

Step 1. The EGDBW method is employed to aggregate the evaluations on the importance degrees of criteria of Group 1 due to the large number of users involved in making judgments. Based on Eqs. (11-12), the criteria weight vector is calculated as $W = ([s_{1.78}, s_{2.04}], [s_{1.77}, s_{2.1}], [s_{-0.01}, s_{0.97}], [s_{-0.01}, s_{1.07}], [s_{-0.37}, s_{-0.14}])^T$. The method that assigns the same weight is utilized to each judgment to aggregate the evaluations of Group 2 on the merits of each alternative with respect to each criterion due to the medium scale and centralized

opinions of evaluations. The group decision matrix $D = (a_{ij})_{5 \times 5}$ is obtained by Eq. (9).

$$D = \begin{matrix} & a_1 & a_2 & a_3 & a_4 & a_5 \\ \begin{matrix} a_1 \\ a_2 \\ a_3 \\ a_4 \\ a_5 \end{matrix} & \begin{bmatrix} [s_{1.17}, s_{1.58}] & [s_{-0.75}, s_{-0.33}] & [s_2, s_{2.5}] & [s_{-1.75}, s_{-1}] & [s_0, s_{0.83}] \\ [s_{1.33}, s_{1.5}] & [s_{-0.75}, s_{-0.33}] & [s_{2.17}, s_{2.58}] & [s_{-1.67}, s_{-1}] & [s_{-0.17}, s_{0.58}] \\ [s_{1.67}, s_{2.42}] & [s_1, s_{1.83}] & [s_{1.83}, s_{2.17}] & [s_{-1}, s_{-0.5}] & [s_{0.33}, s_{0.83}] \\ [s_{-0.83}, s_{-0.17}] & [s_{-1}, s_{-0.17}] & [s_{1.83}, s_{2.5}] & [s_{-1.5}, s_{-1}] & [s_{1.83}, s_{2.58}] \\ [s_{-1.33}, s_{-0.33}] & [s_{-0.83}, s_{-0.08}] & [s_{-1}, s_{-0.67}] & [s_{1.83}, s_{2.58}] & [s_{1.67}, s_{2.33}] \end{bmatrix} \end{matrix}$$

Step 2. By Eqs. (20-21), find out the maximum weight $\omega^+ = [s_{1.77}, s_{2.1}]$ and the maximum CIVLEs of the alternatives with respect to each criterion, which are $\tilde{h}_s^{1+} = [s_{1.67}, s_{2.42}]$, $\tilde{h}_s^{2+} = [s_1, s_{1.83}]$, $\tilde{h}_s^{3+} = [s_{2.17}, s_{2.58}]$, $\tilde{h}_s^{4+} = [s_{1.83}, s_{2.58}]$, $\tilde{h}_s^{5+} = [s_{1.83}, s_{2.58}]$, respectively. According to Eq. (19), we obtain the distance $d_j (j=1, \dots, 5)$ from each criterion to the most importance criterion as: $d_1 = 0.0072, d_2 = 0, d_3 = 0.2485, d_4 = 0.2424, d_5 = 0.3651$ and the distances $d_{ij} (i=1, \dots, 5) (j=1, \dots, 5)$ from each alternative to the best one under each criterion, which are shown in Table 7.

Table 7. The distances from each alternative to the best one under each criterion

Distance	c_1	c_2	c_3	c_4	c_5
a_1	0.1152	0.3276	0.0221	0.5967	0.2984
a_2	0.1156	0.3276	0	0.5900	0.3333
a_3	0	0	0.0628	0.4929	0.2716
a_4	0.4242	0.3333	0.0412	0.5762	0
a_5	0.4796	0.3117	0.5350	0	0.035

Step 3. The global preference scores are shown in Table 8 computed by Eq. (22) based on d_{ij} and d_j . The weak ranking of all alternatives is shown in Table 9 computed by Eq. (23).

Table 8. The global preference scores of CIVL-ORESTE

Global score	c_1	c_2	c_3	c_4	c_5
a_1	0.0816	0.2316	0.1764	0.4554	0.3334
a_2	0.0819	0.2316	0.1757	0.4510	0.3496
a_3	0.0051	0	0.1812	0.3884	0.3218
a_4	0.3000	0.2357	0.1781	0.4420	0.2582
a_5	0.3392	0.2204	0.4171	0.1714	0.2593

Table 9. The weak ranking of the alternatives of CIVL-ORESTE

Alternative	a_1	a_2	a_3	a_4	a_5
Score	0.25568	0.25796	0.1793	0.2828	0.28148
Weak ranking	2	3	1	5	4

Step 4-5. Calculate the preference intensities by Eqs. (24-26). Let $\delta=0.03$ in this case. Then we obtain $\gamma=\frac{(n+2)\delta}{2n}=0.021$ and $\mu=\frac{\delta}{n}=0.006$. The average preference intensities and the PIR relations of the “Mobike” innovative designs are shown in Table 10.

Table 10. The average preference intensities between pairwise alternatives

	a_1		a_2		a_3		a_4		a_5	
	$T(a_i, a_j)$	relation	$T(a_j, a_i)$	relation	$T(a_i, a_j)$	relation	$T(a_j, a_i)$	relation	$T(a_i, a_j)$	relation
a_1	0	-	0.00102	I	0.07734	>	0.01772	<	0.07386	<
a_2	0.0033	I	0	-	0.07976	>	0.02008	<	0.07622	<
a_3	0.00096	<	0.0011	<	-	-	0.01334	<	0.0559	<
a_4	0.04484	>	0.04492	>	0.11684	>	-	-	0.05718	R
a_5	0.09966	>	0.09974	>	0.15808	>	0.05586	R	-	-

Step 6. The strong ranking of all alternatives based on the weak ranking and the PIR structure is shown in Fig. 5.

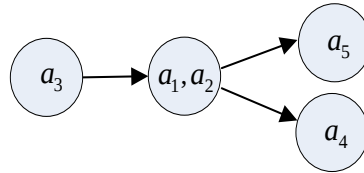


Fig. 5. The strong ranking between the designs resulted from the CIVL-ORESTE method

5.3 Solving the case by the classical ORESTE method

Below we solve the case by the classical ORESTE method. According to the alternatives' distance table (Table 8) obtained in Step 2 of the CIVL-ORETSE method, we get the ranks $r_j(1,2,...,5)$ of the criteria for the importance degrees and the ranks $r_j(a_i)$ ($i = 1,2,...,5$) of the alternatives with respect to each criterion for their merits. The average preference intensities between pairwise alternatives are shown in Table 11.

Table 11. The average preference intensities between pairwise alternatives of the ORESTE method

Intensity	a_1	a_2	a_3	a_4	a_5
a_1	-	0.095	0.26	0.175	0.275
a_2	0.04	-	0.23	0.07	0.22
a_3	0.075	0.1	-	0.105	0.045
a_4	0.185	0.185	0.3	-	0.2
a_5	0.255	0.255	0.21	0.17	-

Let $\tilde{\delta} = 2$. Then according to Eq. (18), we get $\tilde{\mu} < 1/(m-1)n = 0.05$, $\tilde{\gamma} < \tilde{\delta}/2(m-1) = 0.25$, $\lambda > (n-2)/4 = 0.75$. Let $\tilde{\mu} = 0.04$, $\tilde{\gamma} = 0.15$ and $\lambda = 2$ in this paper, we get $a_1 R a_4$, $a_1 R a_5$, $a_2 R a_5$ and $a_4 R a_5$.

The strong ranking is shown in Fig. 6.

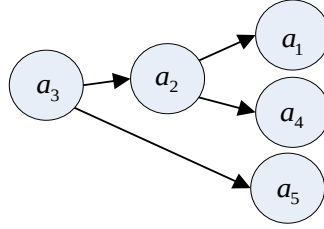


Fig. 6. The strong ranking between the alternatives resulted from the ORESTE method

Comparative analysis: There are different results derived by the CIVL-ORESTE and the ORESTE methods. The ORESTE method has been improved by three aspects: the evaluation information in decision matrix, the calculation processes and the technique to determine the thresholds, which are described in details as follows:

(1) With regard to the initial calculation data derived from evaluation information, the CIVL-ORESTE method maintains more evaluation information by employing the distance d_j and d_{ij} rather than the Besson's mean ranks r_j and $r_j(a_i)$. In the above case study, from the decision matrix, we can find that the evaluation of a_{11} is extremely close to a_{21} but a_{41} is far away to a_{21} . These similarities and differences are clearly reflected by the distances: $d_{11} = 0.1152$, $d_{21} = 0.1156$ and $d_{41} = 0.4242$, but they are obscured by the Besson's mean ranks: $r_1(a_1) = 2$, $r_1(a_2) = 3$ and $r_1(a_4) = 4$.

(2) With regard to the weak ranking, the ORESTE method converts the global preference score $\tilde{D}(a_{ij})$ into the global weak rank $r(a_{ij})$ to calculate the weak rank $R(a_i)$, which leads to information loss. For example by Eq. (14) we have $\tilde{D}(a_{32}) = 1$ corresponding to $r(a_{32}) = 1$, $\tilde{D}(a_{31}) = 1.5811$ to $r(a_{31}) = 2.5$, $\tilde{D}(a_{21}) = 2.5495$ to $r(a_{21}) = 6.5$, and $\tilde{D}(a_{22}) = 2.5739$ to $r(a_{22}) = 8.5$. It is clear that the global weak

rank weakens the score information that $\tilde{D}(a_{31}) - \tilde{D}(a_{32}) = 0.5811$ related to $r(a_{31}) - r(a_{32}) = 1.5$ but $\tilde{D}(a_{22}) - \tilde{D}(a_{21}) = 0.0244$ related to $r(a_{22}) - r(a_{21}) = 2$. However, in the CIVL-ORESTE method, the weak rank $r(a_i)$ is derived by the global preference score $D(a_{ij})$ directly.

(3) With regard to the preference intensities, they are computed by the global weak ranking in the classical ORESTE method while by the global preference scores in the CIVL-ORESTE method. From the above discussion, we can make a conclusion that the preference intensities of the ORESTE method are untrustworthy due to the less information in the global weak ranking.

(4) With regard to the thresholds, in the ORESTE method, $\tilde{\delta}$ are determined by the DM freely with less basis and the ranges of thresholds are so broad that it is difficult to choose reasonable values, which has a decisive effect on the results. For the above case, if $\tilde{\gamma} \in [0, 0.185)$, $a_1 R a_4$ and $a_4 R a_5$; if $\tilde{\gamma} \in [0.185, 0.2)$, $a_1 I a_4$ and $a_4 R a_5$; if $\tilde{\gamma} \in [0.2, 0.25)$, $a_1 I a_4$ and $a_4 I a_5$. However, in the CIVL-ORESTE method, the indifference threshold δ is derived based on the distance between two CIVLEs, and the other two thresholds μ and γ are calculated by δ , which forms a systematic process to set the values of these parameters to ensure that the rankings are generated consistently. Furthermore, they vary in smaller ranges and as the value changes, the results are stable. In the case study, if $\delta \in [0, 0.0336]$, the results obtained will be similar.

5.4 Solving the case by other ranking methods

To further illustrate the reliability of the CIVL-ORESTE method, we deal with the case by three widely used ranking methods and make some comparative analyses. Considering that the crisp criterion weights are the basis of these methods, we can determine the weights based on the evaluations of criteria by a simple formula $\tilde{\omega}_j = d(\tilde{h}_s^j, \min_j \tilde{h}_s^j) / \sum_{j=1}^n d(\tilde{h}_s^j, \min_j \tilde{h}_s^j)$, and thus obtain $W = (0.25, 0.25, 0.18, 0.18, 0.14)^T$.

(1) The results derived by the CIVL-VIKOR method

As a utility value-based ranking method, the VIKOR ranks alternatives considering both the group utility values and the individual regret values. It can avoid the defect that the selected solution may perform

badly under some criteria as in the TOPSIS method [22]. In this part, we extend the VIKOR to the CIVL context to handle the case.

The group utility values can be calculated by $GU_i = \sum_{j=1}^n \tilde{w}_j (d_{ij} / \max_i d_{ij})$ and the individual regret values can be determined by $RS_i = \max_j \tilde{w}_j (d_{ij} / \max_i d_{ij})$ where $\tilde{w}_j$ is the crisp weight of criterion c_j . Let the relative importance between GU_i and RS_i be 0.5. The results derived by the CIVL-VIKOR method are shown in Table 12.

Table 12. The results derived by the CIVL-VIKOR method

	a_1	a_2	a_3	a_4	a_5
GU_i	0.6186	0.6240	0.2839	0.6588	0.6785
RS_i	0.2457	0.2457	0.1487	0.25	0.25
Utility values	0.9029	0.9097	0	0.975	1
Rank	2	3	1	4	5

Comparative analysis: The results derived by the CIVL-VIKOR method are like the weak ranking obtained by the CIVL-ORESTE method. But the CIVL-VIKOR cannot describe a detail relation between pairwise alternatives. It deems that alternative a_1 is superior to a_2 and a_4 is superior to a_5 despite that their utility values are extremely close. Besides, the incomparability relation is ignored in the CIVL-VIKOR.

(2) The results derived by the CIVL-PROMITHEE method

The PROMITHEE method is based on pairwise comparisons between two alternatives associated to each criterion. It is characterized by six kinds of preference functions. The PIR relations of pairwise alternatives are determined by the positive outranking flows $\phi^+(a_i)$ and the negative outranking flows $\phi^-(a_i)$ in PROMITHEE I. We combine the PROMITHEE with the CIVLTS to solve the case in this part.

To display the deviations between two alternatives precisely, we conduct the preference function based on the distance measure of the CIVLTSs as

$$p_j(a_i, a_z) = \begin{cases} 0, & \text{if } \tilde{h}_s^{ij} \leq \tilde{h}_s^{zj} \\ d(\tilde{h}_s^{ij}, \tilde{h}_s^{zj}), & \text{if } \tilde{h}_s^{ij} > \tilde{h}_s^{zj} \end{cases} \quad (27)$$

where $p_j(a_i, a_z)$ is the preference value of a_i over a_z with respect to criterion c_j , and $d(\tilde{h}_s^{ij}, \tilde{h}_s^{zj})$ is the distance between $\tilde{h}_s^{ij}$ and $\tilde{h}_s^{zj}$ computed by Eq. (19). We can compare $\tilde{h}_s^{ij}$ and $\tilde{h}_s^{zj}$ by Eqs. (6-7).

Following the steps of the PROMITHEE I and PROMITHEE II [24], we obtain the results of the case as shown in Table 13.

Table 13. The results derived by the CIVL-PROMITHEE method

	a_1	a_2	a_3	a_4	a_5
ϕ^+	0.071	0.076	0.184	0.065	0.134
ϕ^-	0.085	0.086	0.05	0.12	0.191
$\phi^+ - \phi^-$	-0.014	-0.01	0.134	-0.055	-0.057
Rank	3	2	1	4	5

Comparative analysis: According to the net outranking flow $\phi^+ - \phi^-$ in PROMITHEE II, we obtain the ranking as $a_3 \succ a_2 \succ a_1 \succ a_4 \succ a_5$ which is similar to the ranking derived by the CIVL-VIKOR method and the weak ranking derived by the CIVL-ORESTE method. Furthermore, based on the principle of distinguishing the PIR relations on the basis of ϕ^+ and ϕ^- in PROMITHEE I, we obtain $a_1 R a_2$ and $a_4 R a_5$. We are easy to accept $a_4 R a_5$ since their net outranking flows are close but their positive and negative outranking flows are quite different, which implies that they have a big gap in performance under some criteria. This situation is in accordance with the collective decision matrix. However, we are hard to accept $a_1 R a_2$ since both their net outranking flows and the positive and negative outranking flows are quite close, which implies that they perform similar under all criteria. The I relation between a_i and a_z only appears when $\phi^+(a_i) = \phi^+(a_z)$ and $\phi^-(a_i) = \phi^-(a_z)$. From our cognition, we deem $a_i I a_z$ when they have a small gap in performance under each criterion rather than shown the same performance under all criteria. In CIVL-ORESTE method, we introduce some thresholds to establish the PIR relations objectively, and we obtain $a_1 R a_2$, which conforms to the fact.

(3) The results derived by the CIVL-ELECTRE method

ELECTRE is a famous outranking method, which is characterized by the concordance and discordance concepts. It determines the PIR relation by comparing pairwise alternatives under each criterion based on

some thresholds selected by DMs in advance. Liao et al. [16] extended the ELECTRE II to handle the evaluations expressed as the HFLTSSs based on the distance measure between each alternative and the positive and negative ideal solutions, respectively. Following Ref. [16], we define the concordance, indifferent and discordance sets in CIVL context as follows:

We define the minimum CIVLE of a_i under the criterion c_j as $\tilde{h}_s^{j-} = \min_{i=1,2,\dots,m} (\tilde{h}_s^{ij})$ if c_j is the benefit criterion and $\tilde{h}_s^{j-} = \max_{i=1,2,\dots,m} (\tilde{h}_s^{ij})$ if c_j is the cost criterion. The concordance set is divided into three types: the strong concordance set, the medium concordance set and the weak concordance set. The discordance set is also divided into three types: the strong discordance set, the medium discordance set and the weak discordance set. Then the indifference set is defined. Let the weights of the strong, medium, weak concordance and discordance sets, and the weight of the indifference set as: $\omega = (\omega_C, \omega_{C'}, \omega_{C''}, \omega_D, \omega_{D'}, \omega_{D''}, \omega_I)^T = (1, 0.9, 0.8, 1, 0.9, 0.8, 0.7)^T$. Following the steps of ELECTRE in Ref. [16], we only can obtain the part relations: $a_1 I a_2$, $a_3 P a_1$ and $a_3 P a_4$.

Comparative analysis: The part relations derived by the CIVL-ORESTE and the CIVL-ELECTRE method are similar. We are unable to obtain the global order of all alternatives by the ELECTRE method since it ignores the global preference values. Besides, the weights of the concordance and discordance are determined subjectively, which makes the ELECTRE method with less robustness. Meanwhile, the divisions between two alternatives are blurred by the weights of concordance and discordance.

In conclusion, compared with the ranking method mentioned above, the CIVL-ORESTE method has the following advantages:

(1) We can obtain the global orders of all alternatives and the PIR relations of pairwise alternatives by the CIVL-ORESTE method, which is convincing and easy to make final decision. The global order can only be derived by the VIKOR while the partial relation can only be obtained by the ELECTRE.

(2) The PIR relation is conducted based on some thresholds which are calculated objectively. Thus, the results are robust. The PIR relation in the PROMITHEE method is determined by the positive outranking flow and the negative outranking flow which are calculated by aggregating the preference values

over all alternatives. In this way, the PIR relation is contrary to the real case. In ELECTRE, the thresholds to distinguish the PIR relation are determined subjectively.

6. Conclusions

This paper established a CIVL-ORESTE method to solve the MCGDM problem with qualitative information. The uncertain linguistic variable is a powerful method to interpret the uncertain linguistic information, but it has some limitations in calculation and expressing the hesitant qualitative evaluations precisely. We extended it to the CIVLTS which is not only able to express complex assessments, but more flexible to aggregate group opinions. Some group aggregation methods for CIVLEs were proposed to deal with different types of groups. Especially the EGDBWM is excellent to cope with the large size group. We improved the ranking method, ORESTE, and proposed the CIVL-ORESTE method to cope with the group decision matrix. The advantages of the proposed method are concluded as follows:

- (1) The evaluation information is expressed completely. The CIVLEs can describe both the vague and accurate linguistic evaluations by the continuous interval form.
- (2) Suitable scope is broad. It can handle the experts group with any numbers, and there is no need to determine the crisp criterion weights.
- (3) The results are robust. It derives both the global order and the PIR relation of alternatives. In addition, the thresholds are determined objectively.

However, we ignore the semantics of linguistic terms regarding the asymmetrical situation when calculating the distance between two CIVLEs and aggregating individuals' CIVLEs into a collective one. This challenge will be overcome in our future study. Extending the ORESTE in wider areas when evaluations are expressed as the hesitant fuzzy number and the intuitionistic multiplicative set rather than linguistic term sets is also interesting.

Acknowledgements

The work was supported by the National Natural Science Foundation of China (No. 71771156, No. 71501135), the China Postdoctoral Science Foundation (2016T90863, 2016M602698), the Scientific

Research Foundation for Excellent Young Scholars at Sichuan University (No. 2016SCU04A23), and the International Visiting Program for Excellent Young Scholars of SCU.

References

- [1] E.K. Zavadskas, Z. Turskis, Multiple criteria decision making (MCDM) methods in economics: An overview, *Technological & Economic Development of Economy* 17(2) (2011) 397-427.
- [2] E.K. Zavadskas, Z. Turskis, S. Kildienė, State of art surveys of overviews on MCDM/MADM methods, *Technological & Economic Development of Economy* 20 (1) (2014) 165-179.
- [3] M.X. Wang, J.Q Wang, New online recommendation approach based on unbalanced linguistic label with integrated cloud, *Kybernetes* (2018).
- [4] H.G Peng, H.Y Zhang, J.Q Wang, Cloud decision support model for selecting hotels on TripAdvisor.com with probabilistic linguistic information, *International Journal of Hospitality Management* 68 (2018) 124-138.
- [5] L. A. Zadeh, The concept of a linguistic variable and its application to approximate reasoning-Part I, *Information Sciences* 8(3) (1975) 199-249.
- [6] F. Herrera, L. Martínez, A 2-tuple fuzzy linguistic representation model for computing with words, *IEEE Transactions on Fuzzy Systems* 8 (2000) 746-752.
- [7] Z.S. Xu, Uncertain linguistic aggregation operators based approach to multiple attribute group decision making under uncertain linguistic environment, *Information Sciences* 168 (1-4) (2004) 171-184.
- [8] H.C. Liao, Z.S. Xu, E. Herrera-Viedma, F. Herrera, Hesitant fuzzy linguistic term set and its application in decision making: A state-of-the art survey. *International Journal of Fuzzy Systems* (2018). DOI: 10.1007/s40815-017-0432-9.
- [9] I. Truck, Comparison and links between two 2-tuple linguistic models for decision making, *Knowledge-Based Systems* 87 (2015) 61-68.
- [10] R.M. Rodríguez, L. Martínez, F. Herrera, Hesitant fuzzy linguistic terms sets for decision making, *IEEE Transactions on Fuzzy Systems* 20 (2012) 109-119.
- [11] H.C. Liao, Z.S. Xu, X.J. Zeng, J.M. Merigó, Qualitative decision making with correlation coefficients

- of hesitant fuzzy linguistic term sets, *Knowledge-Based Systems* 76 (2015) 127-138.
- [12] R.X. Liang, J.Q. Wang, H.Y. Zhang, Projection-based PROMETHEE methods based on hesitant fuzzy linguistic term sets, *International Journal of Fuzzy Systems* (2017) 1-14.
- [13] J.Q. Wang, J.T. Wu, J. Wang, H.Y. Zhang, X.H. Chen, Interval-valued hesitant fuzzy linguistic sets and their applications in multi-criteria decision-making problems, *Information Sciences* 288 (2014) 55-72.
- [14] Z.S. Chen, K.S. Chin, Y.L. Li, Y. Yang, Proportional hesitant fuzzy linguistic term set for multiple criteria group decision making, *Information Sciences* 357 (2016) 61-87.
- [15] J.Q. Wang, J. Wang, Q.H. Chen, H.Y. Zhang, X.H. Chen, An outranking approach for multi-criteria decision-making with hesitant fuzzy linguistic term sets, *Information Sciences* 280 (2014) 338-351.
- [16] H.C. Liao, L.Y. Yang, Z.S. Xu, Two new approaches based on ELECTRE II to solve the multiple criteria decision making problems with hesitant fuzzy linguistic term sets, *Applied Soft Computing* 63 (2018) 223-234.
- [17] H.C. Liao, Z.S. Xu, X.J. Zeng, Distance and similarity measures for hesitant fuzzy linguistic term sets and their application in multi-criteria decision making, *Information Sciences* 271 (2014) 125-142.
- [18] J.D. Qin, X.W. Liu, W. Pedrycz, An extended VIKOR method based on prospect theory for multiple attribute decision making under interval type-2 fuzzy environment, *Knowledge-Based Systems* 86 (C) (2015) 116-130.
- [19] X.L. Wu, H.C. Liao, An approach to quality function deployment based on probabilistic linguistic term sets and ORESTE method for multi-expert multi-criteria decision making, *Information Fusion* 43 (2018) 13-26.
- [20] Z.S. Xu, An overview of methods for determining OWA weights, *International Journal of Intelligent Systems* 20 (8) (2005) 843-865.
- [21] L. Dymova, P. Sevastjanov, A. Tikhonenko, An interval type-2 fuzzy extension of the TOPSIS method using alpha cuts, *Knowledge-Based Systems* 83 (1) (2015) 116-127.
- [22] S. Opricovic, G.H. Tzeng, Compromise solution by MCDM methods: A comparative analysis of VIKOR and TOPSIS, *European Journal of Operational Research* 156 (2) (2004), 445-455.
- [23] H. Ma, H.B. Zhu, Z.G. Hu, K.Q. Li, W.S. Tang, Time-aware trustworthiness ranking prediction for

- cloud services using interval neutrosophic set and ELECTRE, *Knowledge-Based Systems* 138 (2017) 27-45.
- [24] J.P. Brans, B. Mareschal, Promethee methods, *Multiple Criteria Decision Analysis State of the Art Surveys* (2013) 163-186.
- [25] M. Roubens, Preference relations an actions and criteria in multicriteria decision making, *European Journal of Operational Research* 10 (1) (1982) 51-55.
- [26] H. Wang, Extended hesitant fuzzy linguistic term sets and their aggregation in group decision making, *International Journal of Computational Intelligence Systems* 8 (1) (2015) 14-33.
- [27] E. Herrera-Viedma, Fuzzy sets and fuzzy logic in multi-criteria decision making. The 50th anniversary of Prof. Lotfi Zadeh's theory: introduction, *Technological & Economic Development of Economy* 21 (5) (2015) 677-683.
- [28] C.A. Franco, On the analytic hierarchy process and decision support based on fuzzy-linguistic preference structures, *Knowledge-Based Systems* 70 (2014) 203-211.
- [29] F.F. Jin, L.D. Pei, H.Y. Chen, L.G. Zhou, Interval-valued intuitionistic fuzzy continuous weighted entropy and its application to multi-criteria fuzzy group decision making, *Knowledge-Based Systems* 59 (2) (2014) 132-141.
- [30] S.M. Yu, J Wang, J.Q. Wang, L. Li, A multi-criteria decision-making model for hotel selection with linguistic distribution assessments, *Applied Soft Computing* (2017). DOI: 10.1016/j.asoc.2017.08.009
- [31] B. Liu, Y. Shen, Y. Chen, X. Chen, Y. Wang, A two-layer weight determination method for complex multi-attribute large-group decision-making experts in a linguistic environment, *Information Fusion* 23 (2015) 156-165.
- [32] B. Bourguignon, D.L. Massart, The Oreste method for multicriteria decision making in experimental chemistry, *Chemometrics & Intelligent Laboratory Systems* 22 (2) (1994) 241-256.
- [33] H. Pastijn, J. Leysen, Constructing an outranking relation with ORESTE, *Mathematical & Computer Modelling* 12 (10-11) (1989) 1255-1268.
- [34] C. Delhaye, J. Teghem, P. Kunsch, Application of the ORESTE method to a nuclear waste management problem, *International Journal of Production Economics* 24 (1-2) (1991) 29-39.

- [35]G.V. Huylenbroeck, The conflict analysis method: bridging the gap between ELECTRE, PROMETHEE and ORESTE, European Journal of Operational Research 82 (3) (1995) 490-502.
- [36]I.D. Leeneer, H. Pastijn, Selecting land mine detection strategies by means of outranking MCDM techniques, European Journal of Operational Research 139 (2) (2002) 327-338.
- [37]E.A. Adali, A.T. Isik, Ranking web design firms with the ORESTE mthod, Ege Academic Review 17 (2) (2017) 243-253.
- [38]M.A. Yerlikaya, F. Arikan, Constructing the performance effectiveness order of sme supports programmes via Promethee and Oreste techniques, Journal of the Faculty of Engineering & Architecture of Gazi University 31 (4) (2016) 1007-1016.