

Ranking Z-numbers Using the Optimal Clustering Method (CZ-number)

Saeed Jafari^{1,2}  | Mojtaba Najafi^{3*}  | Naghi Moadabi Pirkolachahi⁴  | Najmeh Cheraghi Shirazi⁵ 

¹ Deputy of Planning and Engineering, Renewable Energy Office, Bushehr Power Distribution Company, Bushehr, Iran
Saeed.Jafari@iau.ac.ir

² Department of Electrical Engineering, Bu. C, Islamic Azad University, Bushehr, Iran.

³ Department of Electrical Engineering, Bu. C, Islamic Azad University, Bushehr, Iran.
mojtaba.najafi@iau.ac.ir

⁴ Department of Electrical Engineering, Bu. C, Islamic Azad University, Bushehr, Iran.
n.moadabi@rayavin.com

⁵ Department of Electrical Engineering, Bu. C, Islamic Azad University, Bushehr, Iran.
na.cheraghi@iau.ac.ir

Correspondence

Mojtaba Najafi, Associate Professor of Electrical Engineering, Bu. C, Islamic Azad University, Bushehr, Iran.
mojtaba.najafi@iau.ac.ir

Main Subjects:

Number clustering

Paper History:

Received: 6 June 2023

Revised: 19 August 2023

Accepted: 9 September 2023

Abstract

In recent years, modeling uncertainties through natural language has attracted growing attention, with Z-numbers, introduced by Zadeh in 2011, being a key concept. A Z-number consists of two fuzzy components: the first represents the possibility of an event's occurrence, while the second expresses the probability of that occurrence. A major challenge in applying Z-numbers lies in properly structuring these two components; inappropriate strategies can lead to inaccurate results, especially in group decision-making contexts with large data volumes. To address this, the paper proposes an optimal clustering approach for structuring the components of Z-numbers more systematically and purposefully. This method enhances the accuracy of Z-number modeling by reducing errors in component formation. The effectiveness of the proposed approach is validated through a comparison with conventional fuzzy methods, demonstrating improved performance in handling uncertainty.

Keywords: Probability-possibility, Unsupervised clustering, Z-number, CZ number, Fuzzy.

Highlights

- Using the unsupervised learning method in the optimal selection of expert opinions.
- Effectiveness of the proposed method in choosing the optimal opinions of experts in the conditions of large and complex data sets of expert opinions.
- Determining the optimal points to determine the points that make up the intervals of the Z-number.
- The use of any method or algorithm for unsupervised clustering of points forming the Z-number.

Citation: S. Jafari, M. Najafi, N. Moadabi Pirkolachahi, and N. Cheraghi Shirazi, "Ranking Z-numbers Using the Optimal Clustering Method," *Journal of Southern Communication Engineering*, vol. 15, no. 57, pp. 56–75, 2025, doi:10.30495/jce.2023.1987707.1205 [in Persian].

COPYRIGHTS

©2025 by the authors. Published by the Islamic Azad University Bushehr Branch. This article is an open-access article distributed under the terms and conditions of the Creative Commons Attribution 4.0 International (CC BY 4.0) <https://creativecommons.org/licenses/by/4.0>



1. Introduction

In the face of complex uncertainties, selecting an appropriate solution requires access to valid and reliable information. Real-world data is often associated with uncertainty, and each type of data demands a specific approach for analysis. For this purpose, various methods have been developed, including fuzzy set theory, evidence theory, distance measurement, and Z-numbers, all of which are widely applied in decision-making problems.

Among these methods, fuzzy modeling has gained popularity due to its ability to employ natural language. However, its reliance on expert judgment reduces the confidence level in the results. In 2011, Professor Lotfi Zadeh introduced the Z-number concept, which, by combining probability and possibility of occurrence, offers higher accuracy in modeling uncertainties. Nevertheless, the main challenge in this method lies in optimally forming its first and second components.

Previous studies have proposed different solutions to this problem, such as weighting methods, applying evidential reasoning theory, neural networks, defining discrete fuzzy numbers, and ranking expert opinions. However, these approaches are mostly applicable to small datasets and have limited efficiency for large-scale data.

In this research, a novel approach based on unsupervised clustering (K-means algorithm) is proposed to optimally select expert opinions for constructing the initial structure of Z-number components. This approach not only reduces dependency on a specific algorithm but also enables processing of large-scale datasets with linear complexity, while maintaining high interpretability. As a result, the proposed method can serve as an effective solution for improving the modeling of complex uncertainties across various domains.

- **Main Problem:** Analyzing data with complex uncertainties and the need for accurate and reliable methods.
- **Existing Methods:** Fuzzy theory, evidence theory, distance measurement, Z-numbers.
- **Advantage of Z-number:** Simultaneous integration of probability and possibility, higher accuracy than probabilistic methods.
- **Current Challenge:** Optimal determination of the first and second components of the Z-number.
- **Limitations of Previous Studies:** Mostly suitable for small datasets and dependent on specific methods.
- **Proposed Solution:** Using unsupervised clustering (K-means) for optimal selection of expert opinions.
- **Advantages:**
 - Reduced dependency on a specific algorithm
 - Capability to process large datasets
 - Interpretability and visualization of results
- **Application:** Improving uncertainty modeling in decision-making and complex data analysis.

2. Innovations and Contributions

- Using the unsupervised learning method in the optimal selection of expert opinions.
- Effectiveness of the proposed method in choosing the optimal opinions of experts in the conditions of large and complex data sets of expert opinions.
- Determining the optimal points to determine the points that make up the intervals of the Z-number.
- The use of any method or algorithm for unsupervised clustering of points forming the Z-number.

3. Materials and Methods

In many real-world decision-making scenarios, data often exhibit complex uncertainties that challenge the reliability of analytical outcomes. To address such challenges, researchers have developed various theoretical frameworks, among which Z-numbers—introduced by Professor Lotfi Zadeh in 2011—stand out as an effective method for modeling imprecise information. A Z-number consists of two components: (1) the *possibility* of an event, and (2) the *probability* of its occurrence. This dual-component representation enables more precise and realistic modeling compared to conventional probabilistic approaches.

Despite their advantages, Z-numbers face a critical limitation: the accurate construction of their initial possibility and probability components becomes increasingly challenging when working with large-scale and scattered expert datasets. Existing approaches, such as weighted scoring methods, evidential reasoning, and neural network-based weight estimation, often perform well only on small datasets or require strong assumptions about data relationships, which limits their applicability in complex real-world situations.

To overcome these challenges, this study introduces the **CZ-number method**, an enhanced modeling approach that integrates unsupervised clustering with the Z-number framework. Specifically, the K-means clustering algorithm is applied to expert opinions to generate well-defined clusters, from which optimal representative points are selected for both the possibility and probability components. This process ensures higher accuracy in capturing expert knowledge, even in the presence of highly dispersed or large-scale datasets.

The proposed method offers several key advantages:

- **Scalability** – capable of efficiently handling large datasets with linear time complexity.
- **Algorithmic flexibility** – can be implemented using any clustering algorithm that meets the required properties, making it non-dependent on a single technique.
- **Unsupervised learning** – eliminates the need for predefined input-output relationships, enhancing adaptability.
- **Interpretability** – provides clear data visualization and transparent component formation for better decision support.

The remainder of this paper is organized as follows:

- **Section 2** provides a brief review of the K-means clustering algorithm and the fundamentals of Z-numbers.
- **Section 3** presents the proposed CZ-number methodology in detail.
- **Section 4** demonstrates the method's effectiveness through a numerical example.
- **Section 5** concludes the study with final remarks and potential directions for future research.

Through this integration of clustering and Z-number modeling, the proposed approach contributes a robust, scalable, and interpretable solution for addressing complex uncertainty in practical decision-making environments.

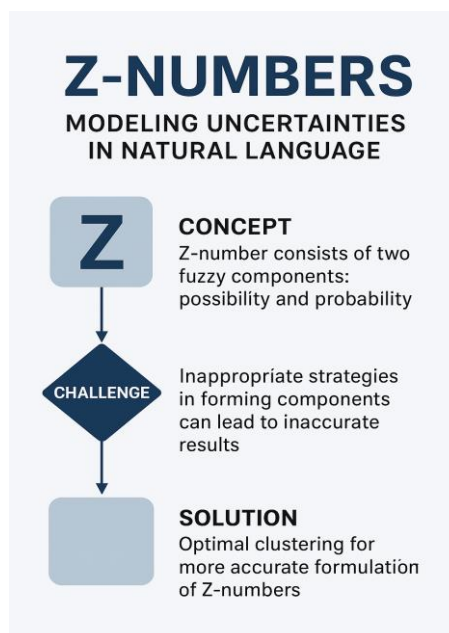


Figure I. The solution process of the proposed model

4. Results and Discussion

The main motivation for selecting PJM market price forecasts is the availability of free, real, and widely used data in the international electricity market. The PJM website provides real-time (hourly) electricity prices, which many market participants utilize to set prices and participate in the day-ahead market. In this study, it is assumed that no local data are available, and the next-hour PJM market prices are predicted using both the conventional Z-number method and the proposed approach.

Figure I shows the problem-solving process of the proposed method. In order to evaluate the obtained results, the results were compared with the results obtained by the Z-score method at 100 points [1]. Additionally, the execution time for each record was measured in MATLAB, considering 4,000 data points for each method on a Windows 7 system with a 5-core processor and 4 GB of RAM. The results indicate that the proposed method provides predictions closer to the actual values for each record. Moreover, the standard deviation of expert opinions in the proposed method is lower than that of the conventional Z-number method, reflecting a more effective utilization of expert knowledge. Although the proposed method requires slightly more computation time due to additional steps, the execution time difference is negligible. Therefore, based on the results, the proposed method produces forecasts that are more accurate and optimized compared to the conventional Z-number approach.

5. Conclusion

In this paper, a novel approach to Z-numbers is proposed, incorporating unsupervised clustering of expert data. In previous studies, a Z-number associated with a variable X is typically represented as a pair of fuzzy numbers

(A, B) , where A —the first component—is interpreted as a fuzzy restriction, and B —the second component—is considered a measure of certainty regarding the first component. From a probabilistic perspective, in our approach, both the first component A and the second component B are modeled as discrete fuzzy numbers.

The main motivation of this research is to enable straightforward modeling of any type of arithmetic operation using discrete fuzzy numbers, while providing an optimal method for determining fuzzy intervals that best capture both event characteristics and probabilities. The term *optimal* here refers to the preservation of data properties when clustering probability data, ensuring that fuzzy interval formation benefits from an efficient algorithm. In other words, these intervals can be easily managed without requiring any probability distribution to be predefined. The algorithm's outcomes can support experts in final decision-making processes involving discrete fuzzy number variables. As the Z-number framework in this study is augmented with a clustering algorithm, the computational load of the proposed method is compared against the conventional Z-number approach. This comparison was conducted using MATLAB software on a Windows 7 operating system equipped with a 5-core processor and 4 GB of RAM. The results demonstrate that while the proposed method involves additional computational steps and consequently requires marginally more processing time than the standard Z-number method, the difference in execution time remains negligible.

6. Acknowledgement

We would like to thank the esteemed reviewers of the article who guided us in increasing the quality of the article. We also express our appreciation and gratitude to the esteemed editor and the executive staff of the journal.

References

- [1] S. Jafari, M. Najafi, N. M. Pirkolachahi, and N. C. Shirazi, "Enhancing energy hub efficiency through advanced modelling and optimization techniques: A case study on micro-refinery output products and parking lot integration," *Heliyon*, vol. 10, no. 18, 2023, doi: [10.1016/j.heliyon.2024.e36233](https://doi.org/10.1016/j.heliyon.2024.e36233).

Declaration of Competing Interest: The Authors do not have a conflict of interest. The content of the paper is approved by the authors.

Author Contributions: **Saeed Jafari:** Software, methodology, writing original draft preparation; **Mojtaba Najafi, Najmeh Cheraghi Shirazi:** Resources, methodology, manuscript editing; **Naghi Moadabi Pirkolachahi:** Resources, manuscript editing.

Open Access: Journal of Southern Communication Engineering is an open-access journal. All papers are immediately available to read and reuse upon publication.

مقاله پژوهشی

رتبه‌بندی عدد زد با استفاده از روش خوشه‌بندی بهینه (CZ-number)

سعید جعفری^{۱،۲} | مجتبی نجفی^۳ | نقی مؤدبی پیرکلاچاهی^۴ | نجمه چراغی شیرازی^۵

چکیده:

استفاده از مفاهیم جدید در مدل‌سازی عدم قطعیت‌ها توسط کلمات زبان طبیعی، در سال‌های اخیر مورد توجه قرار گرفته است. مفهوم عدد زد یکی از این مفاهیم است که توسط دکتر زاده در سال ۲۰۱۱ مطرح گردید. هدف از عدد زد مدل‌سازی جملات غیردقیق زبان طبیعی توسط یک جفت اعداد فازی (A,B) است که اولین مؤلفه این عدد نشان‌دهنده امکان رخداد و مؤلفه دوم نشان‌دهنده احتمال رخداد عدد اول است. امروزه یکی از چالش‌های مهم عدد زد در بین محققان، نحوه تشکیل اولیه ساختار مؤلفه اول و مؤلفه دوم است، در صورتی که محققان راهکارهای نامناسبی در تشکیل این مؤلفه‌ها مورد استفاده قرار دهند، این عامل باعث می‌گردد که نتایج نهایی به درستی محاسبه نگردد. این ضعف در تصمیم‌گیری‌های گروهی با حجم اطلاعات بالا بیشتر نمایان می‌گردد. به منظور رفع این مشکل، نویسندگان این مقاله، خوشه‌بندی بهینه دیتاهای جمع‌آوری شده را پیشنهاد می‌نمایند. راهکار پیشنهادی باعث می‌گردد که مؤلفه‌های اول و دوم عدد زد به شکل هدفمند تشکیل گردد. به منظور صحت سنجی روش پیشنهادی، نتایج به دست آمده از روش پیشنهادی با روش فازی مقایسه می‌گردد.

کلیدواژه‌ها: امکانی-احتمالاتی، خوشه‌بندی بدون نظارت، عدد زد، عدد CZ، فازی.

^۱ معاونت برنامه‌ریزی و مهندسی، دفتر انرژی‌های تجدیدپذیر، شرکت توزیع نیروی برق استان بوشهر، بوشهر، ایران
^۲ گروه مهندسی برق، واحد بوشهر، دانشگاه آزاد اسلامی واحد بوشهر، ایران.

Saeed.Jafari@iau.ac.ir

^۳ گروه مهندسی برق، واحد بوشهر، دانشگاه آزاد اسلامی واحد بوشهر، ایران.

mojtaba.najafi@iau.ac.ir

^۴ گروه مهندسی برق، واحد بوشهر، دانشگاه آزاد اسلامی واحد بوشهر، ایران.

n.moaddabi@rayavin.com

^۵ گروه مهندسی برق، واحد بوشهر، دانشگاه آزاد اسلامی واحد بوشهر، ایران.

na.cheraghi@iau.ac.ir

نویسنده مسئول

^{*}مجتبی نجفی، دانشیار گروه مهندسی برق، دانشکده فنی و مهندسی، واحد بوشهر، دانشگاه آزاد اسلامی واحد بوشهر، ایران.

mojtaba.najafi@iau.ac.ir

موضوع اصلی:

خوشه‌بندی اعداد

تاریخچه مقاله:

تاریخ دریافت: ۱۶ خرداد ۱۴۰۲

تاریخ بازنگری: ۲۸ مرداد ۱۴۰۲

تاریخ پذیرش: ۱۸ شهریور ۱۴۰۲

تازه‌های تحقیق:

- استفاده از روش یادگیری بدون نظارت در انتخاب بهینه از نظرات متخصص.
- اثربخشی روش پیشنهادی در انتخاب نظرات بهینه متخصصان در شرایط مجموعه داده‌های بزرگ و پیچیده نظرات متخصص.
- تعیین نقاط بهینه به منظور تعیین نقاطی که فواصل Z را تشکیل می‌دهد.
- استفاده از هر روش یا الگوریتم برای خوشه‌بندی بدون نظارت از نقاط تشکیل‌دهنده عدد زد.

COPYRIGHTS

©2025 by the authors. Published by the Islamic Azad University Bushehr Branch. This article is an open-access article distributed under the terms and conditions of the Creative Commons Attribution 4.0 International (CC BY 4.0) <https://creativecommons.org/licenses/by/4.0>



۱- مقدمه

در شرایط عدم قطعیت‌های پیچیده، محققین نیازمند اطلاعات کافی و انتخاب راه‌حلی مناسب هستند. اطلاعات اولیه باید به نحوی در نظر گرفته شود که دارای اعتبار و قابل اعتماد برای محققین باشد. در دنیای واقعی، بسیاری از داده‌ها، دارای عدم قطعیت‌های پیچیده هستند که هر کدام با توجه به نوع دیتاهای در دسترس، دارای راه‌حل‌های منحصر به فرد هستند. به منظور بررسی اطلاعات نامطمئن، محققان نظریه‌ها و روش‌های مختلفی مانند نظریه مجموعه‌های فازی [۱]، نظریه شواهد [۲،۳]، اندازه‌گیری فواصل [۴] و عدد زد [۵] ارائه نموده‌اند. روش‌های ذکر شده یکی از پرکاربردترین روش‌ها در مسائل تصمیم‌گیری عملی هستند [۶-۱۳]. امروزه در بین محققان یکی از راه‌حل‌های پر کاربرد در حل عدم قطعیت‌های پیچیده، مدل‌سازی توسط زبان طبیعی (روش فازی) است. اما نکته بسیار مهم در نظریه فازی، محاسبه پارامترهای فازی توسط دانش خبرگان است و این مهم باعث می‌گردد میزان اطمینان به نظر کارشناسان کاهش یابد. در سال ۲۰۱۱، پروفسور لطفی‌زاده مفهوم جدیدی به نام عدد زد جهت مدل‌سازی متغیرهای نامشخص معرفی نمود [۱۴]. این مفهوم نسبت به روش‌های احتمالاتی دارای دقت بهتری است. در این روش (عدد زد) با در نظر گرفتن امکان رخداد و احتمال رخداد به صورت هم‌زمان، راه‌حلی بسیار مؤثر به منظور مدل‌سازی عدم قطعیت‌ها در شرایط اطلاعات نامطمئن است. از زمان معرفی، این مفهوم به طور گسترده در بین محققان جهت محاسبات عدم قطعیت‌های پیچیده مورد استفاده قرار گرفته است [۱۵]. اما امروزه یکی از مهم‌ترین چالش‌ها در عدد زد، نحوه تشکیل ساختار اولیه مؤلفه اول و مؤلفه دوم است. در مطالعات گذشته راه‌حل‌های مختلفی به منظور حل این چالش ارائه گردیده است به عنوان مثال، روش‌های امتیازدهی وزنی اطلاعات با در نظر گرفتن بهترین و بدترین حالت و وزن‌دهی به آن‌ها [۱۶]، تعیین درجه مؤلفه‌های اول و دوم بر اساس نظریه استدلال اثباتی [۱۷]، استفاده از روش وزن دهی با استفاده از تعیین ضرایب معادلات بازه فازی کلاسیک توسط شبکه عصبی [۱۸]، تعریف متغیر مختلف با هدف تعیین اعداد فازی گسسته در تشکیل بازه‌های عدد زد [۱۹]، ارزش‌گذاری و رتبه‌بندی نظرات کارشناسان [۲۰]، استفاده از مدل فازی^۱ [۲۱] به منظور تقسیم‌بندی ناحیه عدد زد و استفاده از روش چگالی هسته اقدام به منظور تشکیل فواصل مؤلفه‌های اول و دوم پیشنهاد گردیده است [۲۲]. تاکنون راه‌حل‌های ارائه شده در انتخاب نقاط تشکیل‌دهنده مؤلفه‌ها فقط در حالت دیتاهای کوچک عملی بوده است [۱۹]، در [۲۰] با نظرسنجی از کارشناسان اقدام به رتبه‌بندی بازه‌های عدد زد نموده است، در [۲۱] با در نظرگیری چندین کنترل‌کننده داخلی که هر یک از آن‌ها در یک ناحیه خاص از فضای ورودی کنترل‌کننده معتبر است، مقادیر عددی از درون‌یابی کنترل‌کننده‌های خطی به دست می‌آید که این مورد عاملی می‌گردد که نقاط به دست آمده به صورت کاملاً بهینه انتخاب نگردند. در این مقاله، یک نظریه جدید برای بهبود حل این مشکلات ارائه گردیده است. در تئوری پیشنهادی، مجموعه‌ای از نظرات کارشناسان برای تشکیل مؤلفه‌های اول و دوم عدد زد جمع‌آوری شده و سپس با انتخاب بهینه نظرات، مؤلفه‌های اول و دوم عدد زد شکل می‌گیرند. محققان با استفاده از الگوریتم‌های مبتنی بر تکرار بدون نظارت تلاش می‌نمایند مجموعه داده‌های به دست آمده (نظرات کارشناسان) را بدون همپوشانی به زیرگروه‌های مجزا تقسیم نموده، سپس یک نقطه از هر گروه را به عنوان نقطه بهینه (مشابه‌ترین ویژگی در مقایسه با تمام نقاط همان گروه) تعیین می‌نمایند. مزیت دیگر نظریه ارائه شده، یادگیری بدون نظارت در انتخاب بهینه نظرات خبرگان است و در فرآیند حل تابع هدف به منظور خوشه‌بندی، برخلاف نظریه شبکه عصبی و سایر نظریه‌های مشابه، نیازی به وجود رابطه بین داده‌های ورودی و خروجی نیست. مزیت دیگر نظریه ارائه شده نسبت به تحقیقات قبلی این است که اگر مجموع نظرات کارشناسان از نوع داده‌های بزرگ با پیچیدگی زمانی خطی باشد، انتخاب بهینه نقاط تشکیل‌دهنده مؤلفه اول و مؤلفه دوم به راحتی در این نوع مجموعه‌ها قابل حل است. ویژگی‌های نهایی نظریه پیشنهادی می‌تواند کمی‌سازی انتخابی تعداد خوشه‌ها برای تشکیل فواصل مؤلفه‌های اول و دوم مطابق با اهداف تحقیق و تفسیرپذیری و تصویرسازی داده‌های انتخاب‌شده به منظور تشکیل مؤلفه‌های اول و دوم باشد. ویژگی روش پیشنهادی این است که از هر الگوریتمی که دارای این ویژگی باشد می‌توان مورد استفاده قرار گیرد، این مزیت مهمی است که محققین به الگوریتم خاصی وابسته نیستند. محققین این مقاله به منظور بررسی اثربخشی این راه‌کار یکی از الگوریتم‌های خوشه‌بندی بدون نظارت که دارای کاربرد زیادی در بین محققان است (K-means) انتخاب نموده‌اند.

^۱ تاکاگی-سوگنو

در ادامه ساختار مقاله بدین شرح است: در بخش ۲، مروری بر عملکرد الگوریتم خوشه‌بندی بدون نظارت (K-means) عدد زد، در بخش ۳، روش عدد CZ^۱، در بخش ۴، مثال عددی و در پایان در بخش ۵، نتیجه‌گیری ارائه گردیده است.

۲- مروری بر عملکرد الگوریتم خوشه‌بندی بدون نظارت و عدد زد

۲-۱- الگوریتم خوشه‌بندی K-means

الگوریتم خوشه‌بندی K-means یکی از روش‌های پرکاربرد خوشه‌بندی اطلاعات است، که توسط مک کوئین در سال ۱۹۶۷ پیشنهاد گردید [۲۳]. امروزه از این روش به‌عنوان یک الگوریتم خوشه‌بندی کلاسیک در پژوهش‌های علمی و کاربردهای صنعتی مورد استفاده قرار می‌گیرد. هدف از این الگوریتم یافتن و گروه‌بندی نقاط داده در کلاس‌های مشابه است. این شباهت به‌عنوان نقطه مقابل فاصله بین داده‌ها شناخته می‌شود در واقع هر چه نقاط داده نزدیک‌تر هستند، احتمال تعلق آن‌ها به یک خوشه بیشتر است. ایده اصلی این الگوریتم، n داده را به تعداد خوشه‌های تعریف شده تقسیم‌بندی می‌نماید، به‌نحوی که مجموع مربعات نقاط داده در هر خوشه، به مرکز خوشه کمترین است [۲۴]. الگوریتم خوشه‌بندی K-means، مرکز k را محاسبه نموده و هر نقطه داده را دقیقاً به یک خوشه با هدف به حداقل رساندن واریانس درون خوشه اختصاص می‌دهد. هدف اصلی این الگوریتم که به‌عنوان مسئله بهینه‌سازی بیان می‌شود در معادله ۱ بیان گردیده است [۲۵]، مراحل پروسه حل الگوریتم خوشه‌بندی K-means به‌صورت زیر تعریف می‌گردد [۲۵].

$$P_{KM}^* = \arg \min V(P), V(P) = \sum_{l=1}^K \sum_{i=1}^{n_{pl}} d_{P_l, X_i}^2 \quad (1)$$

که در آن n_{pl} تعداد کل داده‌های اختصاص داده شده به خوشه P_l است. با توجه به معادله ۱ مراحل الگوریتم خوشه‌بندی K-means از مرحله ۱ تا ۶ تعریف می‌گردد،

مرحله اول: شاخص تکرار $\eta = 1$ و تعداد k خوشه‌ها را مقداردهی اولیه گردد [۲۵]
مرحله دوم: مقداردهی اولیه مرکز $P^{(1)}$ انجام گردد [۲۵]:

$$P^{(1)} = [P_1^{T(1)} P_2^{T(1)} P_3^{T(1)} P_K^{T(1)}]^T \in \mathbb{R}^{dk} \quad (2)$$

که در این معادله مقدار اولیه مرکز خوشه و T مخفف ماتریسی است که تمام رکوردهای مجموعه داده در آن وجود دارد و d بعد بردار و k خوشه تعریف می‌گردد.

مرحله سوم: فواصل d_{pl, X_i} برای هریک از ترکیب‌های $X_{i,i=1 \dots n}$ (رکوردهای ورودی) و $P_{l,i=1 \dots k}$ (مرکز خوشه) مطابق معادله ۳ محاسبه گردد [۲۵].

$$d_{p_l, x_i} = \|x_i - p_l\| = \sqrt{(x_{1_i} - p_{1_l})^2 + (x_{2_i} - p_{2_l})^2 + \dots + (x_{d_i} - p_{d_l})^2} \quad (3)$$

مرحله چهارم: اختصاص یک رکورد بر اساس فاصله محاسبه شده در معادله سوم به یک خوشه k [۲۵]:

$$P_l^{(\eta)} = \{x_m^{(\eta)} \mid d_{x_m, P_l} \leq d_{x_m, P_i}, \forall i = 1 \dots k\} \quad (4)$$

مرحله پنجم: مرکزهای جدید به‌دست‌آمده در تکرارهای بعدی بر حسب معادله ۵ محاسبه می‌گردد [۲۵].

$$P_l^{(\eta+1)} = \frac{1}{n_{c_j}^{(\eta)}} \sum_{x_i \in C_j^{(\eta)}} X_i, j = 1 \dots k \quad (5)$$

هنگامی که رکوردهای مجموعه داده به همه k خوشه اختصاص داده شد، تعداد کل $n_{pl}^{(\eta)}$ رکوردهای اختصاص داده شده به خوشه‌های $P_l^{(\eta)}$ به‌عنوان اصلی مجموعه‌های $P_l^{(\eta)}$ محاسبه می‌شود [۲۵].

$$n_{p_l}^{(\eta)} = |P_l^{(\eta)}|, l = 1 \dots k \quad (6)$$

¹ Clustering Z-Number

مرحله ششم: بررسی شود که آیا الگوریتم خوشه‌بندی K-means راه‌حل بهینه را پیدا کرده است. در اولین معیار توقف بررسی گردد که همه بردارهای مرکز در یک آستانه پس از دو تکرار کمتر از ε تعریف شده است. این اولین معیار توقف به‌عنوان نابرابری مرکزها بیان می‌شود، معادله ۷، [۲۵].

$$|P_1^{(\eta+1)} - P_1^{(\eta)}| < \varepsilon \quad (7)$$

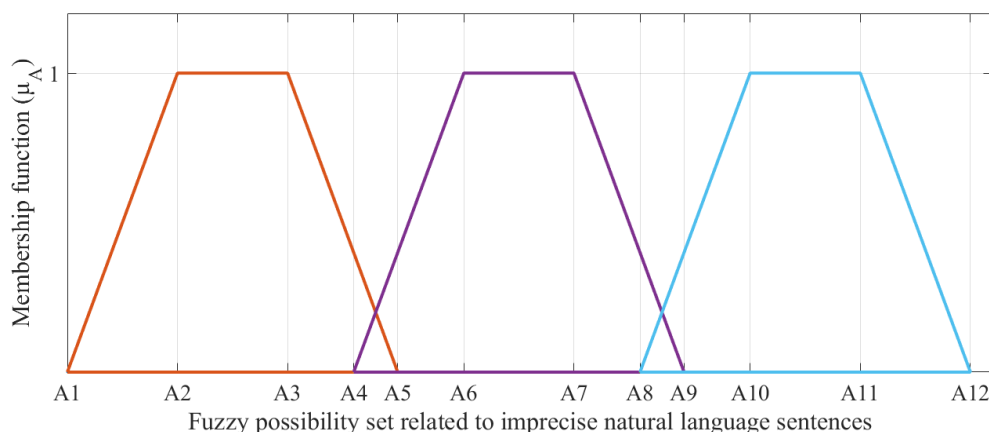
اگر شرط ۷ برآورده شود، به این معنی است که الگوریتم کامل است و معادله ۱ آخرین $P^{(\eta+1)}$ است که تاکنون به‌دست آمده است و اگر شرط ۷ برآورده نشود، می‌بایست معیار دوم توقف بررسی گردد، که از معادله ۸ استفاده می‌گردد، این معادله نشان می‌دهد که الگوریتم به حداکثر تعداد تکرار خود رسیده است [۲۵].

$$\eta = \eta_{\max} \quad (8)$$

۲-۲- عدد زد

تعریف ۱. تولید مجموعه‌های عدد زد

امروزه عدد زد یک از روش‌های مؤثر در مدل‌سازی عدم قطعیت جملات نادقیق زبان طبیعی است، در اولین مرحله، مجموعه‌های مؤلفه اول، توسط خبرگان تعیین می‌گردد. خبرگان، متغیر مؤلفه اول (امکان رخداد) را در سه مجموعه تحت عناوین کم، متوسط و زیاد تعریف می‌نمایند. شکل ۱ تابع عضویت این مجموعه‌ها را نمایش داده شده است. معادله ۹ محدوده این مجموعه‌ها را تعیین می‌نماید. [۱۴]:



شکل ۱: مجموعه امکان اعداد Z

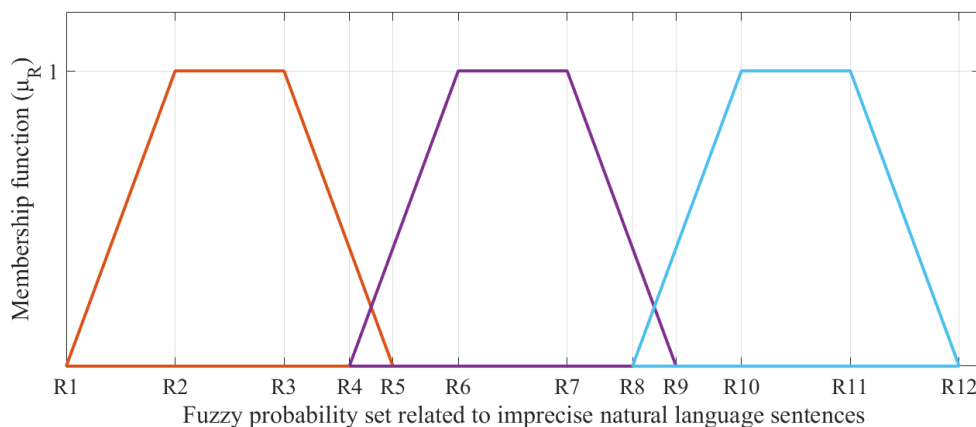
Figure 1. Possible set of Z-Numbered

$$\tilde{A}_{Poss} = \begin{cases} \tilde{A}_{Poss_1} = A_1, A_2, A_3, A_5 & Low \quad A_1 \leq A_2 \leq A_3 \leq A_5 \\ \tilde{A}_{Poss_2} = A_4, A_6, A_7, A_9 & Medium \quad A_4 \leq A_6 \leq A_7 \leq A_9 \\ \tilde{A}_{Poss_3} = A_8, A_{10}, A_{11}, A_{12} & High \quad A_8 \leq A_{10} \leq A_{11} \leq A_{12} \end{cases} \quad (9)$$

مجموعه‌های مربوط به مؤلفه دوم (احتمال رخداد) روش عدد زد، به‌صورت تجربی توسط معادله ۱۰ تعریف می‌گردند. شکل ۲ تابع عضویت این مجموعه با در نظرگیری سه مجموعه کاملاً قطعی، تقریباً ممکن و نامشخص را نمایش می‌دهد [۱۴].

(۱۰)

$$\tilde{R}_{Prob} = \begin{cases} \tilde{A}_{Prob_1} = R_1, R_2, R_3, R_5 & \text{Notsure } R_1 \leq R_2 \leq R_3 \leq R_5 \\ \tilde{A}_{Prob_2} = R_4, R_6, R_7, R_9 & \text{Almostcertainly } R_4 \leq R_6 \leq R_7 \leq R_9 \\ \tilde{A}_{Prob_3} = R_8, R_{10}, R_{11}, R_{12} & \text{Quitesure } R_8 \leq R_{10} \leq R_{11} \leq R_{12} \end{cases}$$



شکل ۲: مجموعه رخداد عدد زد

Figure 2. Z-Numbered event set

پس از تعیین مجموعه‌های مؤلفه‌های اول و مؤلفه‌های دوم، معادلات احتمال وقوع حالت مجموعه‌ها، توسط معادله ۱۱ و جدول ۱ محاسبه می‌گردد [۱۴].

$$\tilde{p}_{N_i}^Z = (\tilde{A}_{Poss_i}, \tilde{R}_{Prob_i}) \quad (11)$$

جدول ۱: محدوده در نظر گرفته شده عدد زد

Table 1. Range of the considered Z-number

status	Event status	The occurrence modes of the number Z	
Low	Unknown	(Unknown, low)	$\tilde{P}_{N_1}^Z = (\tilde{A}_{Poss_1}, \tilde{R}_{Prob_1})$
	Almost certain	(Almost certain, low)	$\tilde{P}_{N_2}^Z = (\tilde{A}_{Poss_1}, \tilde{R}_{Prob_2})$
	Absolutely certain	(Absolutely certain, low)	$\tilde{P}_{N_3}^Z = (\tilde{A}_{Poss_1}, \tilde{R}_{Prob_3})$
Medium	Unknown	(Unknown, average)	$\tilde{P}_{N_4}^Z = (\tilde{A}_{Poss_2}, \tilde{R}_{Prob_1})$
	Almost certain	(Almost certain, average)	$\tilde{P}_{N_5}^Z = (\tilde{A}_{Poss_2}, \tilde{R}_{Prob_2})$
	Absolutely certain	(Absolutely certain, moderate)	$\tilde{P}_{N_6}^Z = (\tilde{A}_{Poss_3}, \tilde{R}_{Prob_3})$
High	Unknown	(Unknown, High)	$\tilde{P}_{N_7}^Z = (\tilde{A}_{Poss_3}, \tilde{R}_{Prob_1})$
	Almost certain	(Almost certain, High)	$\tilde{P}_{N_8}^Z = (\tilde{A}_{Poss_3}, \tilde{R}_{Prob_2})$
	Absolutely certain	(Absolutely certain, High)	$\tilde{P}_{N_9}^Z = (\tilde{A}_{Poss_3}, \tilde{R}_{Prob_3})$

تعریف ۲. تبدیل نتایج خروجی مجموعه عدد زد به مجموعه فازی نامنظم با استفاده از روش برش آلفا

پس از تعیین مجموعه‌های عدد زد، نتایج خروجی این مجموعه توسط رابطه ۱۲ محاسبه می‌گردد، با توجه به این که $(\tilde{A}_i, \tilde{B}_i)$ خروجی عدد زد است، می‌بایست به مجموعه فازی نامنظم تبدیل گردد که در این مطالعه روش برش آلفا مورد استفاده قرار گرفته است [۱۴].

$$(\tilde{A}_i, \tilde{B}_i) \quad (۱۲-الف)$$

$$\tilde{P}_{N_i}^Z = (\tilde{A}_{Poss_1}, \tilde{R}_{Prob_1}) = (\tilde{P}_{N_1}^Z + \tilde{P}_{N_2}^Z + \tilde{P}_{N_3}^Z + \tilde{P}_{N_4}^Z + \tilde{P}_{N_5}^Z + \tilde{P}_{N_6}^Z + \tilde{P}_{N_7}^Z + \tilde{P}_{N_8}^Z + \tilde{P}_{N_9}^Z) \quad (۱۲-ب)$$

$$-(\cap(R_4, R_5), \cap(A_4, A_5)) - (\cap(R_8, R_9), \cap(A_8, A_9))$$

$$\tilde{A}_i = A_i^\alpha \quad (۱۳)$$

در روش برش آلفا، متغیر فازی رابطه ۱۴-الف به‌عنوان یک تابع فازی چند متغیره با ورودی رابطه ۱۴-ب به‌صورت تابع رابطه ۱۴-ج بیان می‌شود. برش آلفا برای متغیر فازی ۱۴-د با تابع عضویت μ_{Ai} (مؤلفه اول A_i) به‌صورت ۱۴-ه تعریف می‌گردد [۱۴].

$$\tilde{Y} \quad (۱۴-الف)$$

$$[\tilde{X}_1, \dots, \tilde{X}_n] \quad (۱۴-ب)$$

$$Y = F(\tilde{X}_1, \dots, \tilde{X}_n) \quad (۱۴-ج)$$

$$\tilde{X}_i \quad (۱۴-د)$$

$$A_i^\alpha = \{x \mid \mu_{A_i} \geq \alpha \quad 0 \leq \alpha \leq 1\} \quad (۱۴-ه)$$

برش‌های آلفا متغیر ورودی (مؤلفه اول) مطابق معادله ۱۵ به دست می‌آیند [۱۴]:

$$q^\alpha = \left[q^\alpha \triangleq \min(f(A_i^\alpha)), q^{-\alpha} \triangleq \max(f(A_i^\alpha)) \right] \quad (۱۵)$$

در نهایت با استفاده از معادله ۱۶، با جمع برش‌های آلفای تابع عضویت، متغیر خروجی محاسبه می‌شود [۱۴].

$$A_i^\alpha = \bigcup_{\alpha=0}^{\alpha=1} q^\alpha, \alpha = \{0, \Delta\alpha, 2\Delta\alpha, \dots, 1\} \quad (۱۶)$$

پس از تبدیل مجموعه عدد زد به مجموعه فازی نامنظم توسط روش برش آلفا، می‌بایست مجموعه فازی به‌دست‌آمده فازی زدایی^۱ گردند، متداول‌ترین روش دیفازی سازی (فازی زدایی) مورد استفاده، روش مرکز ثقل^۲ است که یک مقدار قطعی را بر اساس مرکز ثقل مجموعه فازی ارائه می‌نماید. این یک روش میانگین وزنی است که در آن از تابع عضویت برای وزن‌دهی استفاده می‌گردد. بر اساس این روش، رابطه ۱۷-الف برای تبدیل یک عدد فازی به عدد قطعی فرموله شده است [۱۴]:

$$\lambda = \frac{\int_{-\infty}^{+\infty} (\tilde{R}_{Prob} \circ \mu_{\tilde{R}^t}(\tilde{R}_{Prob})) d\tilde{R}_{Prob}}{\int_{-\infty}^{+\infty} (\mu_{\tilde{R}^t}(\tilde{R}_{Prob})) d\tilde{R}_{Prob}} \quad (۱۷)$$

مقادیر محاسبه شده در رابطه ۱۷ که ضرایب وزن نامیده می‌شوند و با نماد λ نشان داده می‌شوند. این نماد در رابطه ۱۸ در تابع عضویت مؤلفه اول μ_{Ai} ضرب می‌گردد. در نتیجه، نتایج خروجی به‌صورت یک متغیر فازی نامنظم تبدیل می‌گردند. تصویر ۳ نتایج خروجی متغیر فازی نامنظم به‌دست‌آمده را نمایش می‌دهد [۱۴].

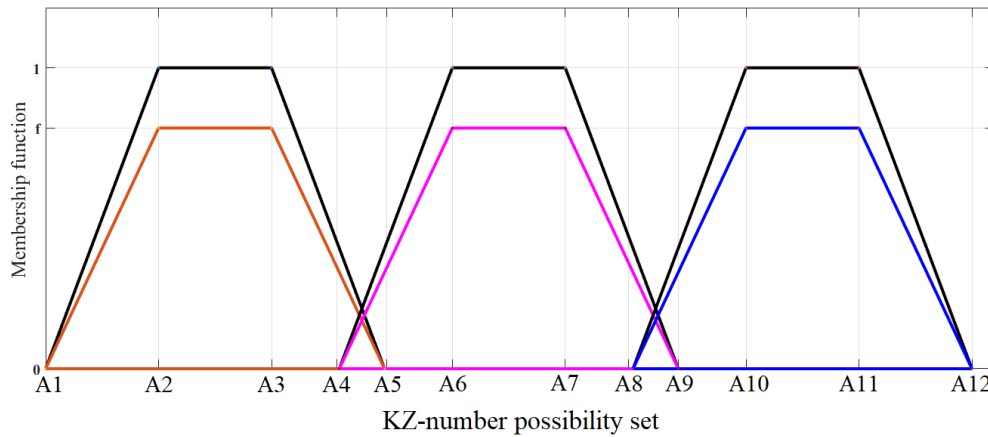
¹ De-fuzzification

² Centroid-of-Gravity

$$f = (\lambda, \mu_{A_i}) \quad (18)$$

تعریف ۳. تنظیم تابع عضویت

نتایج به دست آمده در رابطه ۱۸ دارای متغیر فازی نامنظم می باشد که می بایست به متغیر فازی منظم تبدیل گردند، به منظور این تبدیل رابطه ۱۹ استفاده گردیده است [۱۴]:



شکل ۳: مجموعه فازی نامنظم

Figure 3. Irregular fuzzy set

$$\mu_{p_{N_i}}(p) = f\left(\frac{\tilde{P}_N^Z}{\sqrt{\lambda}}\right) \quad (19)$$

در پایان مجموعه منظم به دست آمده مؤلفه اول که به رابطه ۱۲ منتقل می گردد، که رابطه جدید آن در ۲۱ ارائه گردیده است در شکل ۴ مؤلفه های منظم شده فازی را نمایش می دهد [۱۴]:

$$A_i^+ = \sqrt{\lambda} \cdot A_i \quad (20)$$

$$\tilde{P}_{N_i}^Z = (\tilde{A}_{Poss_1}^+ + \tilde{R}_{Prob_1}) = (\tilde{P}_{N_1}^Z + \tilde{P}_{N_2}^Z + \tilde{P}_{N_3}^Z + \tilde{P}_{N_4}^Z + \tilde{P}_{N_5}^Z + \tilde{P}_{N_6}^Z + \tilde{P}_{N_7}^Z + \tilde{P}_{N_8}^Z + \tilde{P}_{N_9}^Z) - (\cap(R_4, R_5), \cap(A_4, A_5)) - (\cap(R_8, R_9), \cap(A_8, A_9)) \quad (21)$$

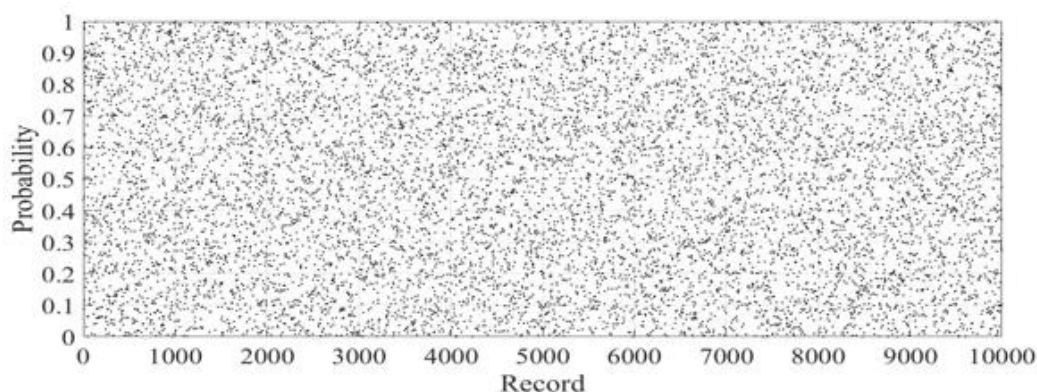
پس از محاسبه نتایج پایانی مجموعه جملات نادقیق زبان طبیعی از مجموعه عدد زد، می بایست نتایج خروجی در احتمال رخ داد که از دیتاهای منطقه مشابه مورد مطالعه استخراج گردیده ضرب گردد.

۳- روش عدد CZ

در این مطالعه یک روش جدید به منظور بهبود تشکیل ساختار اولیه عدد زد ارائه گردیده است. در این روش با استفاده از خوشه بندی بدون نظارت نظر خبرگان (در این مقاله روش K-means در نظر گرفته شده است)، در تشکیل ساختار اولیه مؤلفه اول و مؤلفه دوم عدد زد لحاظ می گردد. روش پیشنهادی باعث می گردد، در صورتی که خبرگان فازی با پراکندگی دیتاهای بی شمار روبرو شوند، ساختار ابتدایی مؤلفه اول و مؤلفه دوم روش عدد زد با دقت مناسبی تشکیل گردد.

۳-۱- معادله CZ-number

در این روش ابتدا نظرات کارشناسان جمع آوری می گردد. تصویر ۴ نمونه ای از این دیتاهای جمع آوری شده از نظرات کارشناسان نمایش داده شده است.



شکل ۴: مجموعه نظرات خبرگان در خصوص احتمال رخداد یک مسئله دارای عدم قطعیت

Figure 4. The set of expert opinions regarding the probability of an uncertain issue

در روش پیشنهادی می‌بایست دیتاهای جمع‌آوری شده نظرات کارشناسان خوشه‌بندی گردد. در این مقاله به منظور برای محاسبه این خوشه‌بندی، مراحل اول تا ششم پیشنهاد گردیده است.

$$P_{CZ}^i = PZ^i = \operatorname{argmin}_V(P), V(P) = \sum_{l=1}^K \sum_{\substack{i=1 \\ x_i \in p_l}}^{n_{p_l}} d_{p_l, x_i}^2 \quad (22)$$

در این معادله x_i مجموعه داده تعریف شده در یک خوشه با هدف اختصاص دادن هر نقطه فقط به یک خوشه، n_{p_l} تعداد کل رکوردهای اختصاص داده شده به خوشه P_l و K تعداد کل خوشه‌ها است که می‌بایست مقداردهی اولیه گردد. مرحله دوم: مقداردهی اولیه مرکز $P_{CZ}^{(1)}$ انجام گردد.

$$P_{CZ}^{(1)} = [P_{CZ_1}^{T(1)} * P_{CZ_2}^{T(1)} * P_{CZ_3}^{T(1)} * P_{CZ_n}^{T(1)}]^T \in R^{dk} \quad (23)$$

که در این معادله مقدار اولیه مرکز خوشه دیتاهای ورودی عدد زد و T مخفف ماتریسی است که تمام رکوردهای مجموعه داده در آن وجود دارد و d بعد بردار و n خوشه تعریف می‌گردد.

مرحله سوم: فواصل d_{p_l, x_i} برای هریک از ترکیب‌های $j, i=1, \dots, n$ (x_i رکوردهای ورودی) و $P_{CZ_l=1 \dots n}$ (مرکز خوشه) مطابق معادله ۲ محاسبه گردد.

$$d_{p_{CZ_l}, x_i} = \|x_i - P_{CZ_l}\| = \sqrt{(x_{i_1} - P_{l_{CZ_1}})^2 + (x_{i_2} - P_{2_{CZ_1}})^2 + \dots + (x_{i_d} - P_{d_{CZ_1}})^2} \quad (24)$$

مرحله چهارم: اختصاص یک رکورد بر اساس فاصله محاسبه شده در معادله سوم به یک خوشه k .

$$P_{CZ_1}^{(\eta)} = \left\{ x_m^{(\eta)} \mid d_{x_m, P_{CZ_1}} \leq d_{x_m, P_{CZ_i}}, \forall i=1 \dots k \right\} \quad (25)$$

مرحله پنجم: مرکزهای جدید به دست آمده در تکرارهای بعدی بر حسب معادله ۵ محاسب می‌گردد.

$$P_{CZ_1}^{(\eta+1)} = \frac{1}{n_{p_l}^{(\eta)}} \sum_{x_i \in p_l^{(\eta)}} x_i, i=1 \dots k \quad (26)$$

هنگامی تمامی رکوردهای مجموعه داده به همه k خوشه اختصاص گردید، تعداد کل $nP_{CZ_1}^{(\eta)}$ رکوردهای اختصاص داده شده به خوشه‌های $P_{CZ_1}^{(\eta)}$ به عنوان اصلی مجموعه‌های $P_{CZ_1}^{(\eta)}$ محاسبه می‌گردد.

$$n_{p_{CZ_1}^{(\eta)}} = |P_{CZ_1}^{(\eta)}|, i=1 \dots k \quad (27)$$

مرحله ششم:

بررسی شود که آیا الگوریتم K-means راه حل بهینه را پیدا کرده است. در اولین معیار توقف بررسی گردد که همه بردارهای مرکز در یک آستانه پس از دو تکرار کمتر از ε تعریف شده است. این اولین معیار توقف به عنوان نابرابری مرکزها بیان می گردد، معادله ۷.

$$|p_{CZ_1}^{(\eta+1)} - p_{CZ_1}^{(\eta)}| < \varepsilon \quad (28)$$

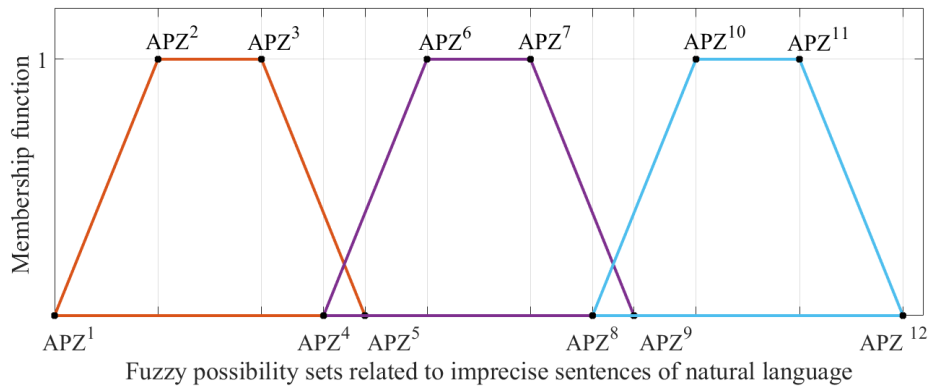
اگر شرط ۷ برآورده شود، به این معنی است که الگوریتم کامل است و معادله ۱ آخری $p_{CZ_1}^{(\eta+1)}$ است که تاکنون به دست آمده است و اگر شرط ۷ برآورده نشود، می بایست معیار دوم توقف بررسی گردد، که از معادله ۸ استفاده می گردد که این معادله نشان می دهد که الگوریتم به حداکثر تعداد تکرار خود رسیده است.

$$\eta = \eta_{\max} \quad (29)$$

مرحله هفتم تشکیل مجموعه عدد CZ:

رابطه ۳۰ مجموعه امکانی عدد CZ را تشکیل می دهد.

$$\tilde{A}_{Poss} = \begin{cases} \tilde{A}_{Poss_1} = APZ^1, APZ^2, APZ^3, APZ^5 & Low \quad APZ^1 \leq APZ^2 \leq APZ^3 \leq APZ^5 \\ \tilde{A}_{Poss_2} = APZ^4, APZ^6, APZ^7, APZ^9 & Medium \quad APZ^4 \leq APZ^6 \leq APZ^7 \leq APZ^9 \\ \tilde{A}_{Poss_3} = APZ^8, APZ^{10}, APZ^{11}, APZ^{12} & High \quad APZ^8 \leq APZ^{10} \leq APZ^{11} \leq APZ^{12} \end{cases} \quad (30)$$

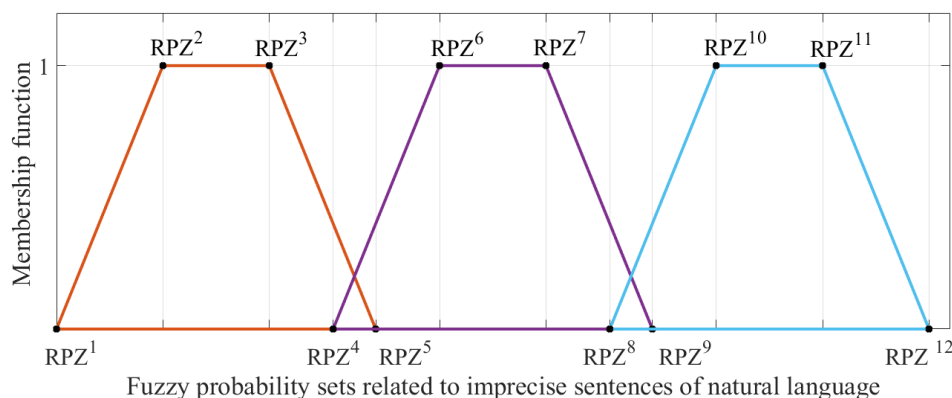


شکل ۵: مجموعه نهایی مؤلفه‌ی اول عدد CZ

Figure 5. The final set of the first component of the CZ-number

رابطه ۳۱ مجموعه احتمالی عدد CZ را تشکیل می دهد:

$$\tilde{R}_{Prob} = \begin{cases} \tilde{R}_{Prob_1} = RPZ^1, RPZ^2, RPZ^3, RPZ^5 & Low \quad RPZ^1 \leq RPZ^2 \leq RPZ^3 \leq RPZ^5 \\ \tilde{R}_{Prob_2} = RPZ^4, RPZ^6, RPZ^7, RPZ^9 & Medium \quad RPZ^4 \leq RPZ^6 \leq RPZ^7 \leq RPZ^9 \\ \tilde{R}_{Prob_3} = RPZ^8, RPZ^{10}, RPZ^{11}, RPZ^{12} & High \quad RPZ^8 \leq RPZ^{10} \leq RPZ^{11} \leq RPZ^{12} \end{cases} \quad (31)$$



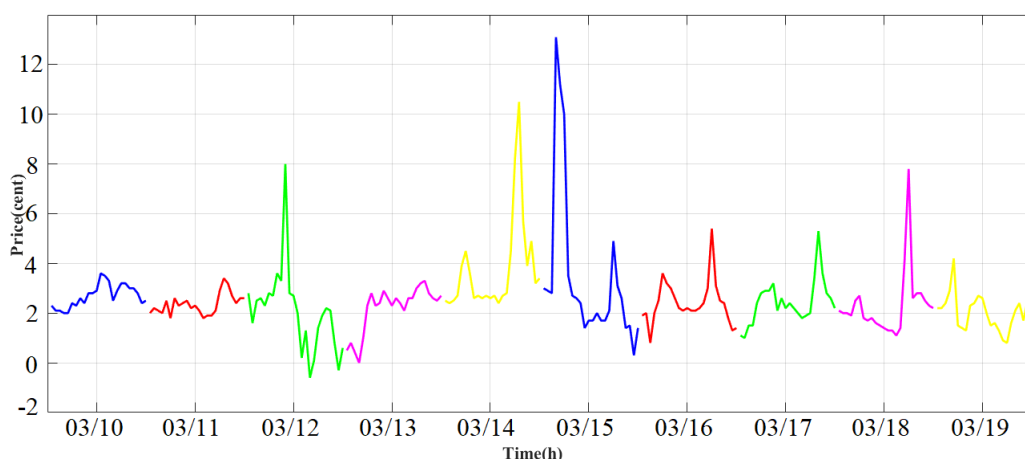
شکل ۶: مجموعه نهایی مؤلفه‌ی دوم عدد CZ
Figure 6. The final set of the second component of the CZ-number

۴- مثال عددی

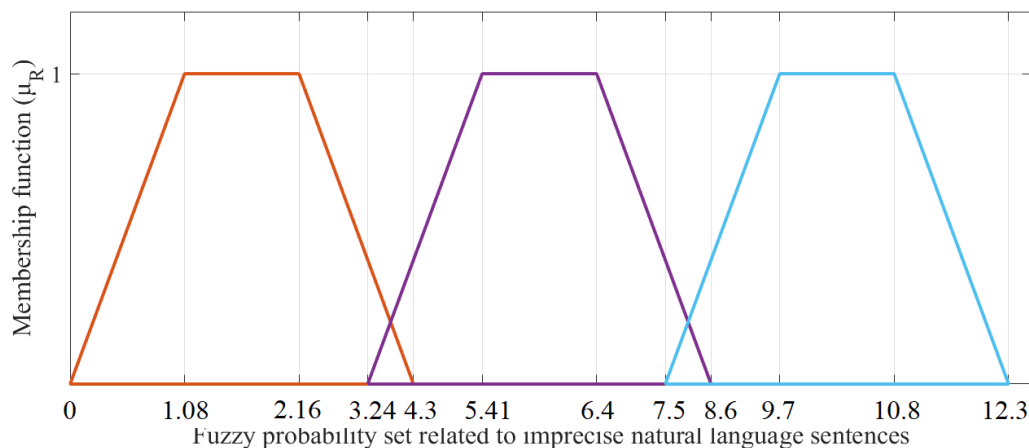
در این بخش، یک مثال عددی برای نشان دادن اثربخشی عدد CZ در نظر گرفته شده است. در این مطالعه عدم قطعیت قیمت انرژی الکتریکی بازار PJM [۲۶] در ساعت پیش رو توسط عدد CZ و عدد زد محاسبه می‌گردد و سپس نتایج به‌دست‌آمده با قیمت واقعی رخ داده مقایسه می‌گردد و در پایان نتایج با یکدیگر مقایسه می‌گردد. تصویر ۷ قیمت ساعتی انرژی الکتریکی به‌دست‌آمده [۲۶] با در نظر گرفتن دوره ۱۰ روز نشان داده شده است.

۴-۱- محاسبه عدم قطعیت قیمت انرژی الکتریکی در ساعت پیش رو توسط عدد زد

با توجه به اینکه در زمان محاسبه عدم قطعیت توسط عدد زد، خبرگان هیچ گونه اطلاعات محلی در دسترس ندارند، لذا می‌بایست از اطلاعات محلی مناطق مشابه استفاده نمایند، بدین منظور در ابتدا با استفاده از اطلاعات در دسترس قیمت بازار PJM در [۲۶] اقدام به تشکیل بازه‌های مؤلفه‌های اول و مؤلفه‌های دوم می‌نمایند. تصاویر ۱۱ و ۱۲ بازه‌های در نظر گرفته شده مؤلفه‌ی اول و مؤلفه‌ی دوم عدد زد نمایش داده شده است. به‌منظور تشکیل این بازه‌ها خبرگان کمترین و بیشترین مقدار قیمت رخ داده از بازار PJM را استخراج کرده و سپس در این بازه اقدام به تشکیل بازه‌های مربوط به مؤلفه‌ها می‌نمایند.



شکل ۷: قیمت ساعتی انرژی الکتریکی بازار PJM
Figure 7. Hourly price of electric energy in PJM market



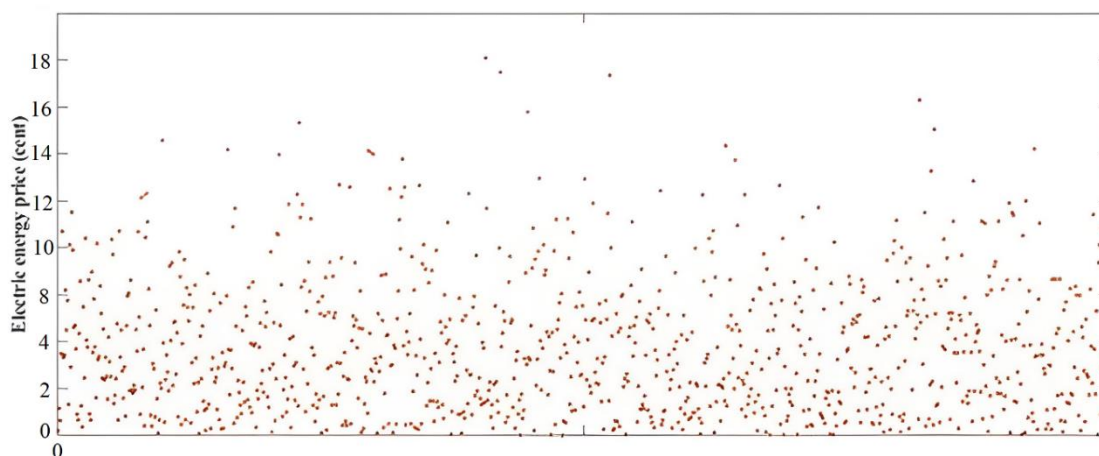
شکل ۸: مؤلفه اول عدد Z قیمت برق

Figure 8. The first component of the electricity price

پس از مشخص نمودن مجموعه‌های مؤلفه‌ی اول و دوم عدد زد، سناریوهای احتمالی مجموعه عدد زد مطابق جدول ۱ تعیین می‌گردند، در ادامه با توجه به اینکه نتایج به‌دست‌آمده از مجموعه‌های عدد زد به‌صورت اعداد زوج مرتب هستند، می‌بایست نتایج به‌دست‌آمده از سناریوها به اعداد فازی کلاسیک تبدیل گردند [۱۴] و سپس در ادامه مجموعه‌های به‌دست‌آمده با استفاده از معادلات ۱۷ و ۱۸ به اعداد واضح فازی تبدیل گردند. پس از تبدیل مجموعه عدد زد به مجموعه فازی نامنظم توسط روش برش آلفا [۱۴]، مجموعه فازی به‌دست‌آمده، فازی زدایی می‌گردند در این مطالعه از روش مرکز ثقل (معادله ۱۹) به‌منظور فازی زدایی استفاده گردیده است. در پایان توسط معادله ۲۱ قیمت انرژی الکتریکی ساعت پیش رو به دست می‌آید.

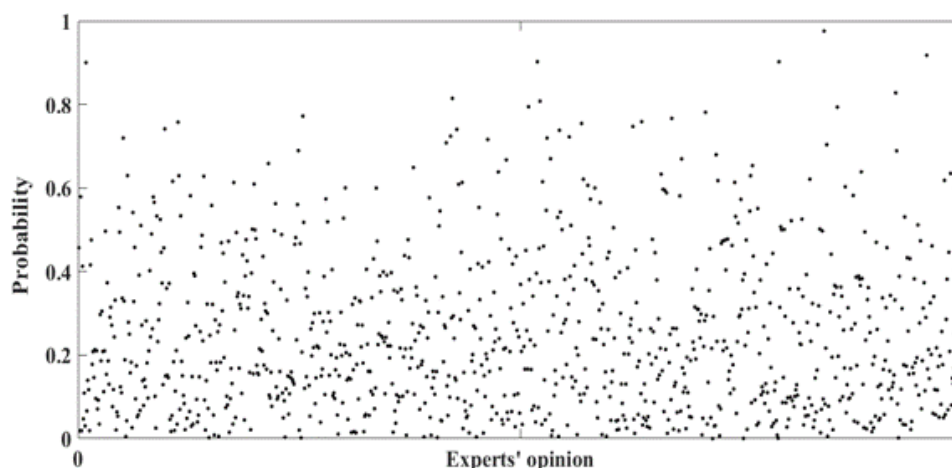
۴-۲- روش عدد CZ

همانطور که در بخش‌های قبل این مطالعه ارائه گردید، تفاوت عدد CZ با عدد زد، نحوه تشکیل مجموعه‌های احتمالاتی و امکانی اولیه توسط خبرگان است. در روش پیشنهادی، نظرات خبرگان در خصوص مجموعه امکانی قیمت انرژی الکتریکی (مؤلفه اول) و مجموعه احتمالاتی انرژی الکتریکی (مؤلفه دوم)، به‌صورت جداگانه جمع‌آوری می‌گردد، تصاویر ۹-۱۰ مجموع نظرات خبرگان، در خصوص رخداد مؤلفه‌ی اول و مؤلفه‌ی دوم نمایش داده شده است.



شکل ۹: مجموعه نظرات خبرگان (مؤلفه اول)

Figure 9. Collection of the expert opinion (first component)



شکل ۱۰: مجموعه نظرات خبرگان (مؤلفه دوم)

Figure 10. Collection of the expert opinion(second component)

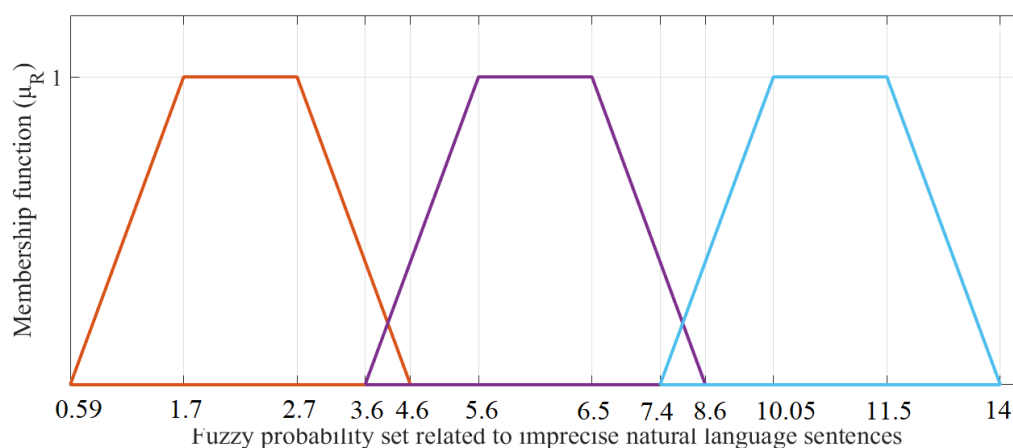
پس از جمع‌آوری نظرات خبرگان، مطابق با معادلات ۲۳-۳۱ اقدام به تشکیل بازه‌های زوج مرتب عدد CZ می‌گردد. در جدول ۲ نقاط به‌دست‌آمده توسط روش پیشنهادی در این مطالعه به‌منظور تشکیل مؤلفه‌ی اول و مؤلفه‌ی دوم عدد CZ ارائه گردیده است. تصاویر ۱۱ و ۱۲ بازه‌های تشکیل شده جزء اول و جزء دوم روش پیشنهادی را مطابق با جدول ۲ نمایش می‌دهند.

جدول ۲: نتایج عددی گروه‌بندی هدفمند بازه‌های عدد CZ

Table 2. Numerical results of targeted grouping of CZ-number intervals

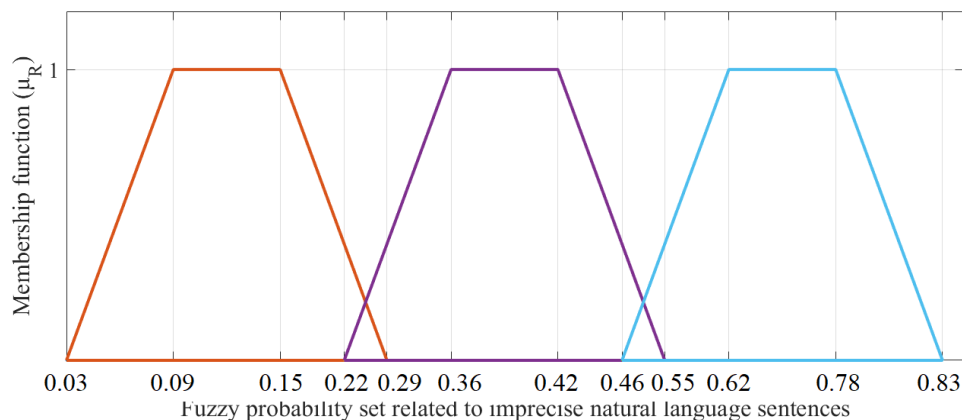
First component					second component			
Low mode	0.5974	1.7408	2.7223	4.6724	0.033	0.0934	0.1548	0.2945
Medium mode	3.6019	5.669	6.5973	8.6256	0.2264	0.3648	0.4346	0.5580
High mode	7.4462	10.0757	11.5223	14.0204	0.4966	0.6238	0.07152	0.8388

تصاویر ۱۱ و ۱۲ فواصل به‌دست‌آمده توسط مؤلفه اول و مؤلفه دوم عدد CZ، مطابق جدول ۲ نشان می‌دهد.



شکل ۱۱: مؤلفه اول عدد CZ قیمت برق

Figure 11. The first component of the CZ number of the electricity price



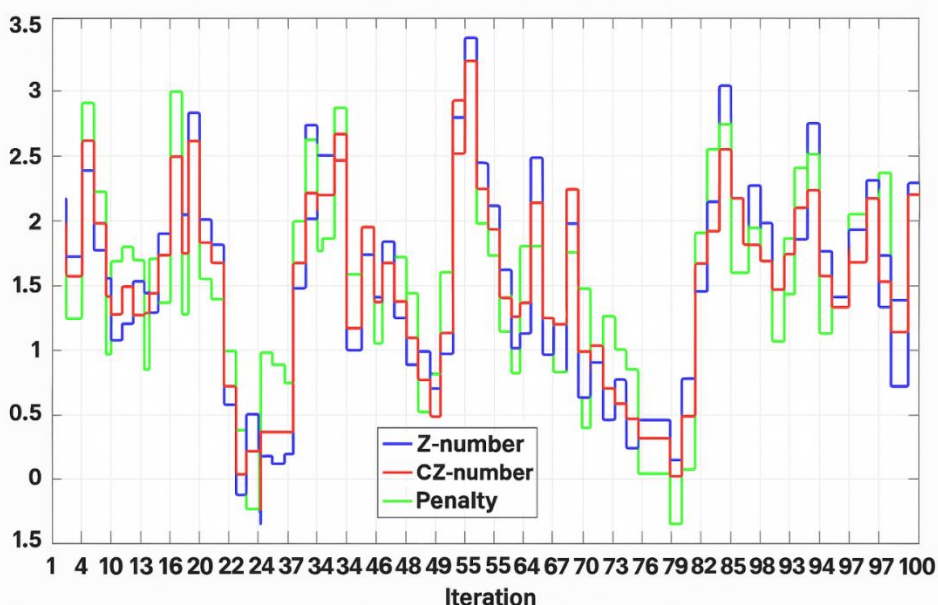
شکل ۱۲: مؤلفه دوم عدد CZ قیمت برق

Figure 12. The second component of CZ number of electricity price

۳-۴- بررسی نتایج پایانی

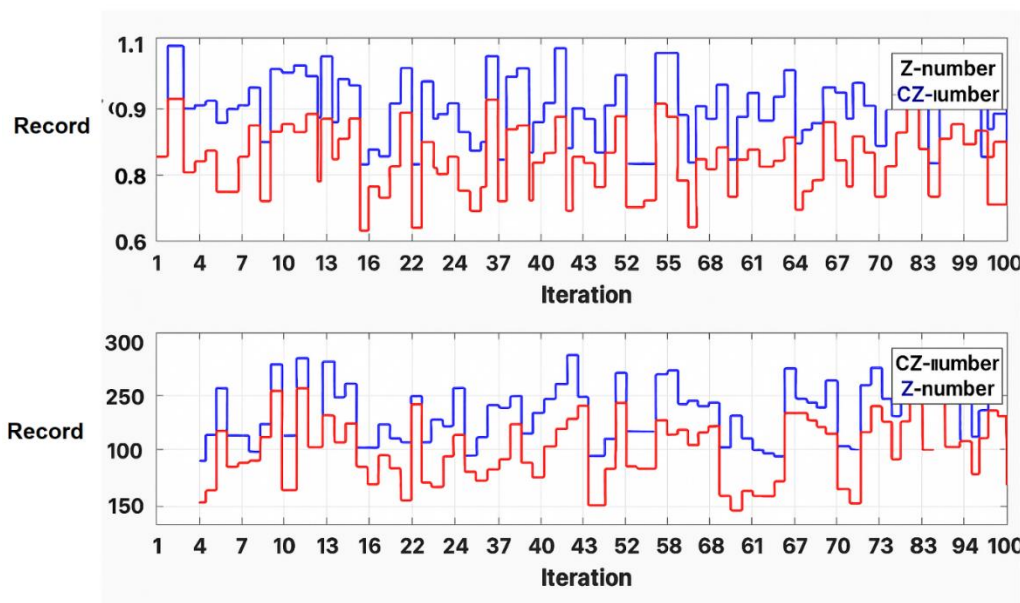
هدف از انتخاب پیش‌بینی قیمت توسط بازار PJM، در دسترس بودن یک دیتا رایگان و واقعی و کاربرد در بازار برق بین‌الملل بوده است. سایت اینترنتی PJM قیمت‌های لحظه‌ای (ساعت به ساعت) بازار برق خود را ارائه می‌دهد که بسیاری از شرکت‌های شرکت‌کننده در این بازار از این دیتاها به‌منظور تعیین قیمت و شرکت در بازار لحظه‌آتی استفاده می‌نمایند در این مقاله فرض گردیده است که هیچ‌گونه دیتا محلی در دسترس نیست و ما توسط روش عدد زد و روش پیشنهادی اقدام به پیش‌بینی بازار PJM در ساعت آتی می‌نماییم. تصاویر ۱۳ و ۱۴ مقایسه نتایج به‌دست‌آمده توسط روش عدد زد و روش پیشنهادی اعداد واقعی رخ داده شده در [۲۶] در ۱۰۰ رکورد آتی را نمایش می‌دهد.

در شکل ۱۴ انحراف معیار دیتاهای در نظر گرفته شده توسط کارشناس عدد زد و انحراف معیار دیتاهای جمع‌آوری شده نظرات کارشناسان در خصوص هر یک از رکوردهای شکل ۱۳ نمایش داده شده است، همچنین زمان اجرا هر یک از رکوردها برحسب ثانیه توسط نرم‌افزار متلب با در نظرگیری ۴۰۰۰ دیتا برای هر یک از روش‌ها در یک سیستم عامل با ویندوز ۷، پردازنده ۵ هسته‌ای و حافظه ۴ گیگابایت محاسبه و نتایج در شکل ۱۴ نمایش داده شده است.



شکل ۱۳: نتایج به‌دست‌آمده پیش‌بینی قیمت و مقایسه با عدد واقعی

Figure 13. The obtain result of price prediction and comparison with real number



شکل ۱۴: مقایسه انحراف معیار دیتاها و زمان اجرا روش پیشنهادی با روش عدد زد

Figure 14. Comparison of standard deviation of data and execution time

نتایج تصاویر ۱۳ و ۱۴ نشان می‌دهد که پیش‌بینی قیمت توسط روش پیشنهادی به عدد واقعی در هر رکورد نزدیک‌تر می‌باشد، همچنین مقدار انحراف معیار نظر کارشناسان در روش پیشنهادی نسبت به عدد زد کمتر می‌باشد و دلیل این امر استفاده بهتر از نظر کارشناسان و خبرگان است. اما با توجه به این موضوع که مراحل بیشتری در روش پیشنهادی باید زمان بیشتری سپری گردد لذا زمان اجرا تا نتیجه‌گیری هر یک از رکوردها نسبت به روش پیشنهادی مقدار بیشتری را سپری می‌نماید، اما با توجه به این که زمان اختلاف زمان اجرا ناچیز است، با توجه به نتایج به‌دست‌آمده توسط روش پیشنهادی نسبت به روش عدد زد به نتایج واقعی نزدیک‌تر است لذا محاسبات پایانی توسط روش پیشنهادی بهینه‌تر هستند.

۵- نتیجه‌گیری

در این مقاله رویکرد جدیدی از عدد زد بر اساس خوشه‌بندی بدون نظارت داده‌های خبرگان ارائه شده است. در مطالعات گذشته، یک عدد زد مرتبط با یک متغیر، X ، معمولاً به‌عنوان یک جفت اعداد فازی (A, B) در نظر گرفته می‌شود که در آن A به‌عنوان اولین جزء اعداد فازی تفسیر می‌شود که یک محدودیت فازی است، و همچنین B به‌عنوان معیار قطعیت که به جزء دوم اعداد فازی معروف است به قطعیت جزء اول تعبیر می‌شود. از دیدگاه احتمالی در رویکرد ما، جزء اول A و جزء دوم B به‌عنوان یک عدد مدل‌سازی می‌شوند. انگیزه‌های اصلی این تحقیق این است که به‌راحتی بتوانیم هر نوع عملیات حسابی را با استفاده از اعداد فازی گسسته مدل‌سازی نماییم و بهترین مدل‌سازی را برای تعیین فواصل فازی که دارای خواص رویدادها و احتمالات در بهینه‌ترین حالت هستند ارائه نماییم. منظور از بهینه‌ترین حالت در این تحقیق این است که با در نظر گرفتن خوشه‌های داده احتمال، خاصیت داده‌ها از بین نمی‌رود و تشکیل فواصل فازی از الگوریتم کارآمدی برخوردار است. به‌عبارت‌دیگر به‌راحتی می‌توان آن‌ها را مدیریت کرد و لازم به ذکر است که برای ایجاد فاصله‌نیازی به توزیع احتمال نیست. علاوه بر این، نتایج این الگوریتم به متخصصان کمک می‌کند تا در مورد مسئله متغیرهای اعداد گسسته به تصمیم نهایی برسند. با توجه به اینکه به روش عدد زد، الگوریتم خوشه‌بندی اضافه گردیده است، به‌منظور بررسی بار محاسباتی روش پیشنهادی نسبت به عدد زد، در این مقاله از نرم‌افزار متلب با سیستم عامل ویندوز ۷ و سخت افزار سیستم با مشخصات پردازش گر ۵ هسته‌ای و حافظه ۴ گیگابایت مورد تحلیل و بررسی قرار گرفت که نتایج نشان می‌دهد با توجه به این که در روش پیشنهادی مراحل بیشتری به‌منظور حل محاسبات انجام می‌گردد، مدت زمان بیشتری به‌منظور حل محاسبات سپری می‌شود اما اختلاف زمان سپری شده نسبت به روش عدد زد ناچیز است. در پایان محققان در نظر دارند در مقالات آتی از روش پیشنهادی به‌منظور محاسبات عدم قطعیت‌های پیش آمده در سیستم قدرت استفاده نمایند.

مراجع

- [1] A. L. Zadeh, "Fuzzy sets," *Information and control*, vol. 8, no. 3, pp. 338-353, 1965, doi: [10.1016/S0019-9958\(65\)90241-X](https://doi.org/10.1016/S0019-9958(65)90241-X).
- [2] A. P. Dempster, "A generalization of Bayesian inference," *Journal of the Royal Statistical Society, Series B (Methodological)*, vol. 30, no. 2, pp. 205-232, 1968, doi: [10.1111/j.2517-6161.1968.tb00722.x](https://doi.org/10.1111/j.2517-6161.1968.tb00722.x).
- [3] G. Shafer, "A mathematical theory of evidence turns 40," *International Journal of Approximate Reasoning*, vol. 79, pp. 7-25, 2016, doi: [10.1016/j.ijar.2016.07.009](https://doi.org/10.1016/j.ijar.2016.07.009).
- [4] M. J. Mendel and J.R. Bob, "Type-2 fuzzy sets made simple," *IEEE Transactions on fuzzy systems*, vol. 10, no. 2, pp. 117-127, 2002, doi: [10.1109/91.995115](https://doi.org/10.1109/91.995115).
- [5] G. Ulutagay and V. Kreinovich, "Density-Based Fuzzy Clustering as a First Step to Learning Rules: Challenges and Solutions," in *Advance Trends in Soft Computing: Proceedings of WCSC*, 2013, San Antonio, Texas, USA, pp. 357-372, doi: [10.1007/978-3-319-03674-8](https://doi.org/10.1007/978-3-319-03674-8).
- [6] Y. Wang *et al.*, "Pairwise constraints-based semi-supervised fuzzy clustering with multi-manifold regularization," *Information Sciences*, vol. 638, pp. 118994, 2023, doi: [10.1016/j.ins.2023.118994](https://doi.org/10.1016/j.ins.2023.118994).
- [7] L. Guo *et al.*, "Pixel and region level information fusion in membership regularized fuzzy clustering for image segmentation," *Information Fusion*, vol. 92, pp. 479-497, 2023, doi: [10.1016/j.inffus.2022.12.008](https://doi.org/10.1016/j.inffus.2022.12.008).
- [8] F. Xiao, "A multiple-criteria decision-making method based on D numbers and belief entropy," *International Journal of Fuzzy Systems*, vol. 21, no. 4, pp. 1144-1153, 2019, doi: [10.1007/s40815-019-00620-2](https://doi.org/10.1007/s40815-019-00620-2).
- [9] J. Pérez-Ortega *et al.*, "POFCM: A Parallel Fuzzy Clustering Algorithm for Large Datasets," *Mathematics*, vol. 11, no. 8, p. 1920, 2023, doi: [10.3390/math11081920](https://doi.org/10.3390/math11081920).
- [10] S. Jafari, M. Najafi, N. M. Pirkolachahi, and N. C. Shirazi, "Enhancing energy hub efficiency through advanced modelling and optimization techniques: A case study on micro-refinery output products and parking lot integration," *Heliyon*, vol. 10, no. 18, 2023, doi: [10.1016/j.heliyon.2024.e36233](https://doi.org/10.1016/j.heliyon.2024.e36233).
- [11] S. Jafari, M. Najafi, N. M. Pirkolachahi, and N. C. Shirazi, "Modelling multi-stage energy optimization in the smart hub energy system with hybrid demand management and local supply by Electric Vehicles," *Research square*, 2024, doi: [10.21203/rs.3.rs-3834421/v1](https://doi.org/10.21203/rs.3.rs-3834421/v1).
- [12] R. Cerqueti and R. Mattera, "Fuzzy clustering of time series with time-varying memory," *International Journal of Approximate Reasoning*, vol. 153, pp. 193-218, 2023, doi: [10.1016/j.ijar.2022.11.021](https://doi.org/10.1016/j.ijar.2022.11.021).
- [13] A. R. Rajeswari *et al.*, "A trust-based secure neuro fuzzy clustering technique for mobile ad hoc networks," *Electronics*, vol. 12, no. 2, p. 274, 2023, doi: [10.3390/electronics12020274](https://doi.org/10.3390/electronics12020274).
- [14] R. Cheng, B. Kang and J. Zhang, "An Improved Method of Converting Z-number into Classical Fuzzy Number," *33rd Chinese Control and Decision Conference (CCDC)*, Kunming, China, 2021, pp. 3823-3828, doi: [10.1109/CCDC52312.2021.9601658](https://doi.org/10.1109/CCDC52312.2021.9601658).
- [15] A. R. Aliev, A. V. Alizadeh and O. H. Huseynov, "The arithmetic of discrete Z-numbers," *Information Sciences*, vol. 290, pp. 134-155, 2015, doi: [10.1016/j.ins.2014.08.024](https://doi.org/10.1016/j.ins.2014.08.024).
- [16] F. Liu *et al.*, "Evaluating Internet hospitals by a linguistic Z-number-based gained and lost dominance score method considering different risk preferences of experts," *Information Sciences*, vol. 630, pp. 647-668, 2023, doi: [10.1016/j.ins.2023.02.061](https://doi.org/10.1016/j.ins.2023.02.061).
- [17] F. Teng *et al.*, "Probabilistic linguistic Z number decision-making method for multiple attribute group decision-making problems with heterogeneous relationships and incomplete probability information," *International Journal of Fuzzy Systems*, vol. 24, no. 1, pp. 1-22, 2022, doi: [10.1007/s40815-021-01161-3](https://doi.org/10.1007/s40815-021-01161-3).
- [18] R. Jafari, W. Yu, and X. Li, "Numerical solution of fuzzy equations with Z-numbers using neural networks," *Intelligent Automation & Soft Computing*, vol. 24, no. 1, pp. 1-7, 2017, doi: [10.1080/10798587.2017.1327154](https://doi.org/10.1080/10798587.2017.1327154).

- [19] S. Massanet, J. V. Riera and J. Torrens, "A new vision of Zadeh's Z-numbers," *International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems*. Cham: Springer International Publishing, vol. 47, 2016, pp. 581-592, doi: [10.1007/978-3-319-40581-0_47](https://doi.org/10.1007/978-3-319-40581-0_47).
- [20] A. R. Aliev *et al.*, "Z-number-based linear programming," *International Journal of Intelligent Systems*, vol. 30, no. 5, pp. 563-589, 2015, doi: [10.1002/int.21709](https://doi.org/10.1002/int.21709).
- [21] D. Wu *et al.*, "A new medical diagnosis method based on Z-numbers," *Applied Intelligence*, vol. 48, pp. 854-867, 2018, doi: [10.1007/s10489-017-1002-4](https://doi.org/10.1007/s10489-017-1002-4).
- [22] R. Cheng, J. Zhang and B. Kang, "Ranking of Z-Numbers Based on the Developed Golden Rule Representative Value," in *IEEE Transactions on Fuzzy Systems*, vol. 30, no. 12, pp. 5196-5210, Dec. 2022, doi: [10.1109/TFUZZ.2022.3170208](https://doi.org/10.1109/TFUZZ.2022.3170208).
- [23] A. R. Aliev *et al.*, "Eigensolutions of partially reliable decision preferences described by matrices of Z-numbers," *International Journal of Information Technology & Decision Making*, vol. 19, no. 06, pp. 1429-1450, 2020, doi: [10.1142/S0219622020010063](https://doi.org/10.1142/S0219622020010063).
- [24] R. Banerjee and K. P. Sankar, "A machine-mind architecture and Z*-numbers for real-world comprehension," *Pattern Recognition and Big Data*, pp. 805-842, 2017, doi: [10.1142/9789813144552_0026](https://doi.org/10.1142/9789813144552_0026).
- [25] Q. Liu *et al.*, "On the negation of discrete z-numbers," *Information Sciences*, vol. 537, pp. 18-29, 2020, doi: [10.1016/j.ins.2020.05.106](https://doi.org/10.1016/j.ins.2020.05.106).
- [26] <https://hourlypricing.comed.com/live-prices/?date=20230310>.