



Fuzzy linguistic summarization of time series with interval type-2 fuzzy c-means: BIST100 sample stock application

İlyas Özdoğan^{1*}, Fatih Emre Boran², Oktay Yıldız¹

¹Turkish Standards Institution, 06100, Ankara, Türkiye

²Energy Systems Engineering, Faculty of Technology, Gazi University, 06500, Ankara, Türkiye

³Computer Engineering, Faculty of Engineering, Gazi University, 06570, Ankara, Türkiye

Highlights:

- Probabilistic and possibilistic approaches to linguistic summarization
- Effect of IT2FCM partition on fuzzy linguistic summarization methods
- Comparison of uniform and IT2FCM partitioning

Keywords:

- Linguistic Summarization
- Interval Type-2 Fuzzy C-Means
- Data Mining
- Time Series

Article Info:

Research Article

Received: 11.03.2023

Accepted: 16.01.2025

DOI:

10.17341/gazimmfd.1263678

Correspondence:

Author: İlyas Özdoğan

e-mail:

ozdoganilyas@gmail.com

phone: +90 312 416 6200

Graphical/Tabular Abstract

In this research, the truth degree of fuzzy linguistic summaries is evaluated through probabilistic and possibilistic approaches using the daily closing price data of three stocks over the last 10 years. Fuzzy sets are generated using both uniform and IT2FCM partition techniques, and Figure A presents the IT2FCM partition fuzzy sets for THYAO stock.

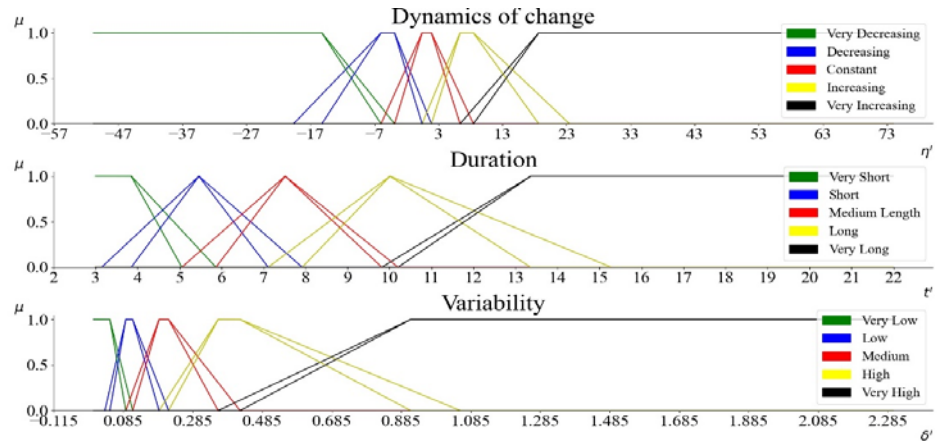


Figure A. Fuzzy sets obtained by IT2FCM partitioning of stock THYAO

Purpose:

The purpose of this study is to investigate the impact of numerical values of triangular and trapezoidal functions determined by IT2FCM partition and uniform partition on possibilistic and probabilistic approaches.

Theory and Methods:

To calculate truth degrees, which is a crucial step in linguistic summarization, possibilistic and probabilistic approaches are used, along with the macd-histogram indicator to generate local trends. The fuzzy sets are created using triangle and trapezoidal functions, and numerical values for these functions are determined by IT2FCM. The variances of the truth degrees are calculated to compare uniform partition and IT2FCM partition.

Results:

The results show that IT2FCM partitioning generates 5 linguistic summaries for THYAO, all of which are the same using both approaches. For PETKM, 5 linguistic summaries are generated using the probabilistic approach and 3 using the possibilistic approach, with 3 of them being the same. For GARAN, 6 linguistic summaries are obtained using the probabilistic approach and 2 using the possibilistic approach, with 2 of them being the same.

Conclusion:

The study investigates the stability of the truth degrees using the variance calculation method of interval valued data (Hu and Hu (2020)). Methods with small variance values are predicted to yield more stable results. The calculations show that IT2FCM partition with the possibilistic approach is more stable than the other methods.



Aralıklı tip-2 bulanık c-ortalama ile zaman serilerinin bulanık dilsel özetlemesi: BIST100 örnek hisse uygulaması

İlyas Özdoğan^{1*}, Fatih Emre Boran², Oktay Yıldız¹

¹Türk Standardları Enstitüsü, 06100, Ankara, Türkiye

²Gazi Üniversitesi, Teknoloji Fakültesi, Enerji Sistemleri Mühendisliği Bölümü, 06500, Ankara, Türkiye

³Gazi Üniversitesi, Mühendislik Fakültesi, Bilgisayar Mühendisliği Bölümü, 06570, Ankara, Türkiye

Ö N E Ç İ K A N L A R

- Dilsel özetlemede olasılık ve olabilirlik temelli yaklaşımlar
- AT2BCO bölümlenimin bulanık dilsel özetleme yöntemlerine etkisi
- Eşit ve AT2BCO bölümlenimin karşılaştırılması

Makale Bilgileri

Araştırma Makalesi

Geliş: 11.03.2023

Kabul: 16.01.2025

DOI:

10.17341/gazimmfd.1263678

Anahtar Kelimeler:

Dilsel özetleme,
aralıklı tip-2 bulanık C-
ortalama,
veri madenciliği,
zaman serileri

ÖZ

Son yıllarda büyük ilgi gören dilsel özetleme, büyük verilerden insanların anlayabileceği çıkarımlar sağlamaktadır. Dilsel özetlemenin önemli bir aşaması, verilerin ne kadar doğru yansıtıldığını gösteren doğruluk derecesinin belirlenmesidir. Bu doğruluk derecelerinin hesaplanmasında literatürde eşit bölümlenme yöntemi kullanılmaktadır fakat bu yöntem verilerin yoğunlaştığı aralıkları göz ardı etmektedir. Bu çalışmada, verilerin yoğun olarak yer aldığı aralıkları dikkate alarak bulanık kümeleri dağıtan Aralıklı Tip-2 Bulanık C Ortalama (AT2BCO) bölümlenme yöntemi önerilmektedir. Önerilen yöntem ile aralıklı tip-2 bulanık kümeler oluşturulmuş, bulanık kardinalite tabanlı olasılık ve olabilirlik temelli yaklaşımlar kullanılarak dilsel özetlerin doğruluk dereceleri hesaplanmıştır. Önerilen yaklaşım adım adım açıklanmış ve literatürde kullanılan eşit bölümlenme yöntemi ile sonuçları karşılaştırılmıştır. Önerilen AT2BCO bölümlenme yöntemi, Borsa İstanbul'da (BIST) işlem gören üç hisse senedine ait son on yıllık dönemi kapsayan finansal zaman serilerine uygulanmıştır. Elde edilen doğruluk derecelerinin varyansı, eşit bölme çalışmalarındaki varyanstan daha düşüktür, bu da önerilen yönteminin daha kararlı sonuçlar verdiğini gösterir.

Fuzzy linguistic summarization of time series with interval type-2 fuzzy c-means: BIST100 sample stock application

H I G H L I G H T S

- Probabilistic and possibilistic approaches to linguistic summarization
- Effect of IT2FCM partition on fuzzy linguistic summarization methods
- Comparison of uniform and IT2FCM partitioning

Article Info

Research Article

Received: 11.03.2023

Accepted: 16.01.2025

DOI:

10.17341/gazimmfd.1333162

Keywords:

Linguistic summarization,
interval type-2 fuzzy C-
means,
data mining,
time series

ABSTRACT

Linguistic summarization, which has garnered significant attention in recent years, facilitates the derivation of human-understandable insights from vast amounts of data. One critical aspect of linguistic summarization involves determining the truth degree that reflects the extent to which the underlying data has been appropriately represented. In the extant literature, the numerical values of fuzzy sets employed to calculate truth degrees have been generated using the uniform partitioning method, which neglects the intervals of data concentration. To address this limitation, the present study proposes the Interval Type-2 Fuzzy C-Means (IT2FCM) partitioning method, which distributes fuzzy sets, taking into account the intervals of data density. By using the proposed method, interval type-2 fuzzy sets are created, and the truth degrees of linguistic summaries are calculated using fuzzy cardinality-based probability and possibility based approaches. The proposed approach is explicated step by step, and its outcomes are compared to those of the uniform partitioning method used in prior research. To evaluate the efficacy of the proposed IT2FCM partitioning approach, we apply it to financial time series covering the last decade of three stocks traded on the Borsa İstanbul (BIST). The variance of the obtained truth degrees is lower than that in equal partitioning studies, which indicates that the proposed method produces more stable results.

1. Giriş (Introduction)

Son yıllarda, bilgisayar teknolojilerindeki hızlı gelişmeler ve yapay zeka tekniklerinin uygulama alanlarının genişlemesi, bilgisayarların çevrelerindeki dünyayı daha iyi anlama, algılama ve yorumlama becerilerinin artmasına olanak sağlamıştır. Bu durum, yapay zeka teknolojilerinin birçok bilim alanında kullanılabilmesine zemin hazırlamıştır.

Yapay zeka alanındaki son araştırmalar, destek vektör makineleri (SVM) ve derin öğrenme gibi tekniklerin kullanımıyla görüntülerdeki nesnelerin tespitinden kontrol ve tahmin problemlerine kadar geniş bir yelpazede çözümler sunabileceğini göstermektedir [1, 2]. Bu teknikler ayrıca doğal dil işleme[3-5], sağlık [6-8], otomatik üretim kontrolü ve akıllı üretim kontrolü gibi alanlarda da uygulanmaktadır. [9-12].

Zadeh [13] tarafından önerilen bulanık küme teorisi, dilsel betimlemeler biçiminde açıklanabilir ve insan tarafından yorumlanabilir sonuçlar sağlamanın çok etkili bir yolu olduğundan son zamanlarda birçok yapay zeka modelinde kullanılmaktadır. Kullanım alanlarından biri de veri madenciliğidir. Veri madenciliğinde kullanımlar yöntemler tahminsel ve tanımlayıcı olarak ikiye ayrılır. Dilsel özetleme konusu tanımlayıcı yöntemler altında yer alan açıklayıcı bir yöntemidir.

Yager [14] tarafından önerilen bulanık dilsel özetleme, bulanık kümeler kullanarak verilerden faydalı dilsel açıklama/bilgi çıkarmayı amaçlar. Bir dilsel özet, verileri üç parametre üzerinden özetler: özetleyici, niceleyici ve doğruluk derecesi. Bu fikirle, “Genç işçilerin çoğu düşük maaş alıyor” gibi dilsel özetler, çalışanların yaş ve aylık kazançlarına sahip bir veri tabanından çıkarılabilir. Bu dilsel özetin doğruluğu, doğruluk derecesi hesaplanarak bilinebilir.

Yager [14] dilsel özetleme fikrini ortaya atmasından sonra, literatürde çok sayıda çalışma gerçekleştirilmiştir. Kacprzyk ve Yager [15] verilerin dilsel olarak özetlenmesiyle ilgili temel yaklaşımı genişleterek bir yaklaşım sunmuşlardır. Dil özetleri Zadeh [16] tarafından önerilen yöntemle değerlendirilmiştir. Çalışmalarında, doğruluk derecesi, kesin olmama derecesi, örtme derecesi, uygunluk derecesi ve basitlik derecelerinin ağırlıklı ortalaması alınarak doğruluk derecesi tespit edilmiştir. Kacprzyk vd. [17] zaman serilerinden doğrusal parçalı eğilimleri çıkarmak için Sklansky, Gonzalez [18] tarafından önerilen bir yöntemi kullanmıştır. Doğrusal parçalı trendler daha sonra dinamikler, değişkenlik ve süre özellikleri açısından karakterize edilmiştir. Son olarak Yager [14] tarafından sunulan basit formlarda dilsel özetler elde edilmiştir. Kacprzyk vd. [19] doğrusal yerel trendleri çıkarmak için Sklansky ve Gonzalez [18] tarafından önerilen yöntemi değiştirmiş ve zaman serileri Yager [14] tarafından önerilen basit protoformlar biçiminde üç yöne (dinamik, değişkenlik ve süre) dayalı olarak dilsel olarak özetlemiştir. Literatürdeki önceki çalışmalardan farklı olarak Kacprzyk vd. [20] Zadeh'in [16] bulanık mantık tabanlı hesaplama yönteminde sunulan genişletilmiş protoformlar biçiminde tanımlanan dilsel özetler için bir doğruluk derecesi önermiştir. Kacprzyk vd. [21] daha önceki çalışmalarda kullanılan “Çoğu kısa trend azalıyor” gibi dilsel ifadeleri “Çoğu zaman alan kısa trendler azalıyor” gibi ifadelere genişleterek yeni bir dilsel özet formu ortaya koymuştur. Kacprzyk vd. [22, 23] göreceli kardinaliteyi hesaplamak için farklı t-normların (yani minimum, şiddetli ve Lukasiewicz operatörleri) etkilerini incelemiştir. Farklı t-normlar kullanılsa bile doğruluk derecelerinin hemen hemen aynı olduğu görülmektedir. Kacprzyk vd. [24, 25] genişletilmiş bir dilsel özetin doğruluk derecesini hesaplamak için Sugeno integralini incelemiştir. Ayrıca, Kacprzyk vd. [26], Kacprzyk vd. [24] 'nın sunduğu Sugeno integral yönteminin monotonluk özelliğini değiştirerek dilsel özetleri değerlendirmiştir. Kacprzyk vd. [27, 28],

Zadeh'in [16] yöntemine alternatif olarak doğruluk derecesini hesaplamak için Sıralı Ağırlıklı Ortalama operatörünü ve Choquet integralini uyarlamıştır.

Zaman serilerinin dilsel olarak özetlenmesinin çok sayıda uygulaması vardır. Kacprzyk ve Wilbik [29] dilsel özetleme ile benzerlik açısından iki farklı zaman serisini analiz etmişlerdir. Castillo-Ortega vd. [30] Yinelemeli Bitiş Noktası Sığdırma Algoritması aracılığıyla dilsel özetlemeye bir girdi olarak, özyineleme derinliğine ve mesafe eşiğine dayalı zaman serilerinin hiyerarşik bölümlenmelerini kullanmıştır. Bu segmentasyon, "2'den 5'e, trend sabittir" gibi dilsel özetler üretme fırsatı verebilir. Castillo-Ortega vd. [31], [30]'te sunulan bir önceki yöntemi kullanarak aynı alandaki farklı zaman serilerini karşılaştırmış ve yöntemini farklı ülkelerdeki CO₂ emisyonlarını gösteren zaman serilerine uygulamıştır. Ramos-Soto vd. [32] Galiçya'daki aylık hava raporlarını analiz ederek buluşsal bir prosedür ile heterojen veri kümelerinden orijinal raporlarla benzer raporlar üreten dilsel özetler oluşturmuştur. Hatipoğlu vd. [33], petrol fiyatının gelecekteki davranışlarının ve dalgalanmalarının tahmin edilmesi amacıyla Avrupa Brent Spot Fiyatı zaman serilerinin bulanık küme ile dilsel özetlenmesini çıkarmıştır. Sanchez-Valdes vd. [34], günlük fiziksel aktiviteleri dilsel terimlerle tanımlamak için şeffaf bir model önermiştir. Kacprzyk ve Zadrozny [35], hisse senedi fiyatı zaman serilerini ve web sunucusu günlüklerini dilbilimsel terimlerle özetlemek için bulanık mantık tabanlı bir dilsel özet yöntemi önermiştir. Kaczmarek vd. [36], zaman serilerini tahmin etmek için dilsel özetleme kullandı. Moysse ve Lesot [37], periyodikliği belirlemek için "Yaklaşık olarak Mart'tan Haziran'a kadar, dizi tam olarak 2 haftalık bir süre ile oldukça periyodiktir" gibi dilbilimsel özetler oluşturmuştur. Marin ve Sanchez [38], zaman serilerinin dilsel tanımlarını üretmek için dilsel özetleme ve doğal dil üretimine dayalı bütünsel bir yaklaşım sunmuştur. Kaczmarek ve Hryniewicz [39], zaman serilerini sınıflandırmak için DVM (Destek Vektör Makinaları)'lerle birlikte dilsel özetlemeyi önermiştir. Kaczmarek ve Hryniewicz [40], dilsel özetler yaklaşımını Bayes modeli ortalamasıyla uyumlu çoklu otoregresif modeller ile birleştirerek zaman serisi tahminini geliştirmek için bulanık nicelenmiş cümleler kullanmıştır. Zaman serilerinin dilsel özetlenmesi uygulamasına ek olarak, bulanık dilsel özetleme son zamanlarda öneri sistemi [41], aykırı değer tespiti [41], sosyal seçim [41], sağlık [42] ve ticaret ağında [43] da kullanılmıştır.

Tip-1 bulanık kümeler yalnızca kişisel belirsizliği yakalarken, aralıklı tip-2 bulanık kümeler hem kişisel hem de kişiler arası belirsizlikleri modelleme yeteneğine sahiptir. Bu faydaya rağmen, literatürde aralıklı tip-2 bulanık kümelerle dilsel özetleme için birkaç çalışma bildirilmiştir. Niewiadomski [44], Zadeh [16] ve Yager'in [14] yöntemlerini aralıklı tip-2 bulanık ortama uyarlamıştır. Niewiadomski [45], dil bilgisi üzerindeki belirsizliği temsil etmek için çok önemli olan aralıklı tip-2 bulanık kümenin farklı özelliklerini tartışmıştır. Wu ve Mendel [46] ise eğer-o zaman dilbilimsel özet formunu bulanık ortamdan aralıklı tip-2 bulanık ortama genişletmiştir. Boran ve Akay [47], aralıklı tip-2 bulanık kümelerle dilsel özetler için doğruluk derecesini hesaplamak için skaler önem düzeyi yerine bulanık önem temelli bir yöntem önermiştir. Aynı zamanda, Delgado vd. tarafından tanımlanan bir dilbilimsel özetin istenen özelliklerini genişletmiştir. Delgado vd. [48] aralıklı tip-2 bulanık ortam yönteminin tüm özelliklerini sağladığı kanıtlanmıştır. Boran vd. [49], aralıklı tip-2 bulanık kümelerle sahip dilsel özetler için doğruluk derecesini hesaplamak için olasılık yapısına sahip bir aralıklı küme atama yöntemi geliştirmiş ve tip-2 bulanık kümelerle daha doğru özetler üretmiştir. Özdoğan vd. [50] aralıklı tip-2 bulanık kümelerle dilsel özetlerin doğruluk derecesini hesaplamak için olabilirlik temelli yeni bir yöntem önermiştir ve olasılık temelli yöntemle [47] göre daha stabil sonuçlar verdiğini göstermiştir.

Bu çalışmanın literatüre sunduğu katkılar maddeler halinde aşağıda verilmiştir:

- Bu çalışmada, dilsel özetleri değerlendirmede kullanılan aralıklı tip-2 bulanık kümelerin oluşturulmasında eşit bölümlene yöntemi yerine Aralıklı Tip-2 Bulanık C Ortalama (AT2BCO) dağılımı önerilmiştir.
- Uygulama olarak Borsa İstanbul'da (BIST) işlem gören üç hisse senedinin son 10 yılda günlük kapanış fiyatları dilsel olarak özetlenmiştir.
- Yerel trendler, Hareketli Ortalama Yakınsama Sıpması (MACD) - histogram değerlerine göre çıkarılmıştır.
- Eşit bölümlene ve AT2BCO bölümlene ile oluşturulan bulanık kümeler ile olasılık ve olabilirlik temelli yaklaşımlarla her bir dilsel özetin doğruluk derecesi hesaplanmıştır.
- Boran ve Akay [47] ve Özdoğan vd. [50] tarafından önerilen dilsel özetlerin değerlendirilmesi ile ilgili yöntemlere entegre edilmiş ve sonuçlar karşılaştırılmıştır.

Bu çalışmanın geri kalanı şu şekilde düzenlenmiştir. Bölüm 2'de bulanık kümeler, aralıklı tip-2 bulanık kümeler, aralıklı tip-2 bulanık c ortalama ve bulanık dilsel özetlemenin kısa bir açıklaması verilmiştir. Bölüm 3'te, AT2BCO bölümlene kullanılarak dilsel özetleri değerlendirmek için önerilen yöntem anlatılmıştır. Önerilen AT2BCO bölümlene yaklaşımın hisse senedi fiyatlarının zaman serilerine uygulaması Bölüm 4'te verilmektedir. Son olarak, makalenin sonuçları Bölüm 5'te verilmektedir.

2. Yöntemler ile İlgili Ön Bilgiler (Preliminary Information on Methods)

Bu bölümde bulanık kümeler, aralıklı tip-2 bulanık kümeler, bulanık dilsel özetleme ve AT2BCO ile ilgili tanımlar kısaca açıklanmıştır.

2.1. Bulanık Kümeler (Fuzzy Sets)

X evrensel kümesine ait A bulanık alt kümesi Eş. 1'deki gibi tanımlanır. A bulanık kümesinde herhangi bir elemanın üyelik derecesi Eş. 2'deki gibi tanımlanır [13].

$$A = \{(x, \mu_A(x)) \mid x \in X\} \quad (1)$$

$$\mu_A(x) : X \rightarrow [0,1] \quad \forall x \in X \quad (2)$$

X Kümelerinin elemanları kullanarak oluşturulan A bulanık kümesinin α - kesmeleri Eş. 3'teki gibi ifade edilir.

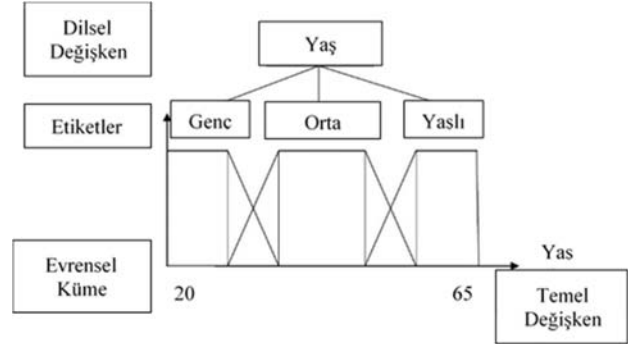
$$A_\alpha = \{x \in X \mid A(x) \geq \alpha\} \quad \alpha \in [0,1] \quad (3)$$

Kesişim, tümleyen ve kesişim kesin kümelerde tanımlandığı gibi bulanık kümeler için de tanımlanmaktadır. X evrensel kümesinde tanımlı A ve B bulanık kümelerinin birleşimi $A \cup B$ ile gösterilir. Birleşim kümesi A veya B 'de yer alan elemanları kapsar [51–53]. Bulanık kümelerde birleşim işlemi Eş. 4'de tanımlanmış olup $\oplus$ simgesi t-konorm operatörüdür:

$$\mu_{A \cup B}(x) = \oplus(\mu_A(x), \mu_B(x)) \quad (4)$$

Bulanık kümeler, herhangi bir değişkeni, sayısal değerler yerine, sözcükler veya cümleler içeren sayısal verileri dilsel olarak açıklamak için kullanılır. Bir dilsel değişken genellikle, her biri bulanık küme ile ilişkili olan bir dizi dilsel etikete ayrıştırılır. Şekil 1'de, "Yaş" dilsel bir değişkendir ve değerleri evrensel bir X kümesi üzerinde değişir. Bu dilsel değişken "Genç", "Orta Yaşlı" ve "Yaşlı" olmak üzere üç dil

etiketine ayrıştırılır ve her dil etiketi bir bulanık küme tarafından modellenir [54].



Şekil 1. Yaş dilsel değişkeni ve etiketleri
(Age linguistic variable and labels)

2.2. Aralıklı Tip-2 Bulanık Kümeler (Interval Type-2 Fuzzy Sets)

Bu alt bölümde, önerilen yöntemin anlaşılmasını kolaylaştırmak için aralıklı tip 2 bulanık küme hakkında temel tanımlar verilmiştir.

Tanım 1: İkincil üyelik fonksiyonu $\mu_{\tilde{A}'}(x, u) = 1$ 'e eşit olan $\tilde{A}'$ aralıklı tip-2 bulanık küme Eş. 5'te tanımlanmıştır [54].

$$\tilde{A}' = \int_{x \in X} \int_{u \in J_x} 1 / (x, u) \quad J_x \subseteq [0,1] \quad (5)$$

X birincil değişkendir, u , X 'in birincil üyelik fonksiyonu olan J_x etki alanına sahip ikincil değişkendir.

Tanım 2: $\tilde{A}'$ 'ya ait Belirsizliğin Kapladığı Alanın (BKA) sınırları iki tane tip-1 üyelik fonksiyonu ile gösterilir ve bunlar $\tilde{A}'$ 'nın üst ve alt üyelik fonksiyonlarıdır. BKA'nın üst sınırı (üst üyelik derecesi) $\bar{\mu}_{\tilde{A}'}(x)$, alt sınırı (alt üyelik derecesi) $\underline{\mu}_{\tilde{A}'}(x)$ ile Eş. 6 ve Eş. 7'de gösterilmiştir [54].

$$\bar{\mu}_{\tilde{A}'}(x) \equiv \overline{BKA}(\tilde{A}') \quad (6)$$

$$\underline{\mu}_{\tilde{A}'}(x) \equiv \underline{BKA}(\tilde{A}') \quad (7)$$

$J_x = [\underline{\mu}_{\tilde{A}'}(x), \bar{\mu}_{\tilde{A}'}(x)]$ $\tilde{A}'$ 'nın aralıklı tip-2 bulanık kümesidir.

Tanım 3: $\tilde{A}'$ aralıklı tip-2 bulanık kümesinin kardinalitesi, Eş. 8'de gösterildiği gibi kendisini oluşturan tip-1 bulanık kümelerin kardinalitesinin birleşiminden oluşmaktadır. [54].

$$kard(\tilde{A}') = \bigcup_{\forall A_e} kard(A_e) = [kard(\underline{A}), kard(\bar{A})] \quad (8)$$

$kard(\underline{A}) = \sum_{x \in X} \underline{\mu}_{\tilde{A}'}(x)$ ve $kard(\bar{A}) = \sum_{x \in X} \bar{\mu}_{\tilde{A}'}(x)$ 'dir. $\tilde{A}'$ 'nın ortalama kardinalitesi Eş. 9'daki gibi hesaplanır.

$$OK(\tilde{A}') = \frac{kard(\underline{A}) + kard(\bar{A})}{2} \quad (9)$$

Tanım 4: $\tilde{A}'$ aralıklı tip-2 bulanık kümesi normal olabilmesi için en az bir $x \in X$ 'in üyelik derecesi $\mu_{\tilde{A}'}(x) = 1$ olmalıdır [54].

2.3. Aralıklı Tip-2 Bulanık C-Ortalama (AT2BCO) (Interval Type-2 Fuzzy C-Means (IT2FCM))

AT2BCO (Aralıklı Tip 2 Bulanık C Ortalama), BKA'yı oluşturmak için üst ve alt üyelik değerlerine karşılık gelen iki bulanıklaştırma katsayısını (m_1, m_2) kullanan ve literatürde yer alan BCO (Bulanık C Ortalama) yönteminin bir uzantısıdır [55]. Bulanıklaştırıcıların kullanımı, minimize edilecek Eş. 10 ve Eş. 11'de tanımlanan farklı amaç fonksiyonlarına yol açar.

$$J_{m_1}(U, v) = \sum_{k=1}^N \sum_{i=1}^C (u_{ik})^{m_1 d_{ik}^2} \quad (10)$$

$$J_{m_2}(U, v) = \sum_{k=1}^N \sum_{i=1}^C (u_{ik})^{m_2 d_{ik}^2} \quad (11)$$

$d_{ik} = \|x_k - v_i\|$, x_k değişkeni ile kümenin merkezi v_i arasındaki öklid mesafesidir, C küme sayısıdır ve N eleman sayısıdır.

Üst/alt üyelik dereceleri $\bar{u}_{ik}$ ve $\underline{u}_{ik}$, iki bulanıklaştırıcı m_1 ve m_2 tarafından oluşturulması ve aşağıdaki gibi belirlenmesi dışında BCO algoritmasına benzemektedir:

$$\bar{u}_{ik} = \begin{cases} \frac{1}{\sum_{j=1}^C (d_{ik} / d_{jk})^{2/(m_1-1)}} & \text{eğer } \frac{1}{\sum_{j=1}^C (d_{ik} / d_{jk})^{2/(m_1-1)}} < \frac{1}{C} \\ \frac{1}{\sum_{j=1}^C (d_{ik} / d_{jk})^{2/(m_2-1)}} & \text{eğer } \frac{1}{\sum_{j=1}^C (d_{ik} / d_{jk})^{2/(m_2-1)}} \geq \frac{1}{C} \end{cases} \quad (12)$$

$$\underline{u}_{ik} = \begin{cases} \frac{1}{\sum_{j=1}^C (d_{ik} / d_{jk})^{2/(m_1-1)}} & \text{eğer } \frac{1}{\sum_{j=1}^C (d_{ik} / d_{jk})^{2/(m_1-1)}} \geq \frac{1}{C} \\ \frac{1}{\sum_{j=1}^C (d_{ik} / d_{jk})^{2/(m_2-1)}} & \text{eğer } \frac{1}{\sum_{j=1}^C (d_{ik} / d_{jk})^{2/(m_2-1)}} < \frac{1}{C} \end{cases} \quad (13)$$

$i = \overline{1, C}, k = \overline{1, N}$ 'dir.

Her model, $\bar{u}$ ve $\underline{u}$ sınırları olarak üyelik aralığıyla birlikte geldiğinden, kümenin her bir merkezi, v^L ve v^R arasındaki aralıkla temsil edilir. Küme merkezleri, BCO algoritmasında da aynı şekilde Eş. 14 kullanılarak hesaplanır:

$$v_i = \frac{\sum_{k=1}^N (u_{ik})^m x_k}{\sum_{k=1}^N (u_{ik})^m} \quad (14)$$

$i = \overline{1, C}$ 'dir.

AT2BCO algoritması merkez noktalarını bulma adımları Şekil 2'de gösterilmiştir.

2.4. Bulanık Dilsel Özetleme (Fuzzy Linguistic Summarization)

Bu bölümde; bulanık kümeler kullanılarak verilerden faydalı dilsel açıklama ya da bilgi çıkarımını amaçlayan bulanık dilsel özetlemenin

kısa tanımı verilmiştir. X kümesine ait bir bulanık küme A ile gösterilir ve şu şekilde tanımlanır:

$$A = \{ \langle x, \mu_A(x) \rangle | x \in X \} \quad \mu_A(x) \text{ burada } x\text{'in üyelik derecesidir. } A$$

kümesinin alfa kesmeleri kesin kümedir $A_\alpha = \{ x \in X | \mu_A(x) \geq \alpha \}$.

Y nesneler kümesini $Y = \{y_1, y_2, \dots, y_M\}$, $V = \{v_1, v_2, \dots, v_K\}$ özellikler kümesini ve $X_k (k=1, 2, \dots, K)$ 'da v_k özelliğinin alanını göstermektedir., $v_k(y_m)$ de m 'ınci objenin k^{th} özelliğini gösterir. $w_g(S_g)$ öncül özetleyici (niteleyici) ve S_k ise özetleyici anlamına gelmektedir. T_1 ise her bir dilsel özetin doğruluk derecesini gösteren doğruluk derecesidir. Bahsedilen tüm özelliklerin bir arada ifade edildiği veritabanı da D ile gösterilir.

$$D = \{ \langle v_1^1, v_2^1, \dots, v_K^1, v_1^2, v_2^2, \dots, v_K^2, \dots, v_1^M, v_2^M, \dots, v_K^M \rangle \equiv \{d_1, d_2, \dots, d_M\} \}$$

Literatürde çoğunlukla kullanılan iki tür cümle yapısı bulunmaktadır, bunlar tip-I ve tip-II niceleyici cümle yapılarıdır. Tip-I niceleyici cümlelerde, niceleyici ifadeleri (tümü, yarısı gibi) Q , nesneleri oluşturan küme Y , özetleyici (yüksek not, düşük maaş gibi) ifadeleri S ve doğruluk derecesi T_1 ile gösterilir. Tip-I niceleyici cümlelerin gösterim şekli " Y 'nin Q kadarı S 'dir $[T_1]$ ". T_1 değeri Eş. 15 ve 16'da verilen formüller kullanılarak bulunur.

$$T_1 = \mu_Q \left(\frac{r}{R} \right) \quad (15)$$

$$r = \sum_{m=1}^M \mu_{S_g}(d_m) \quad (16)$$

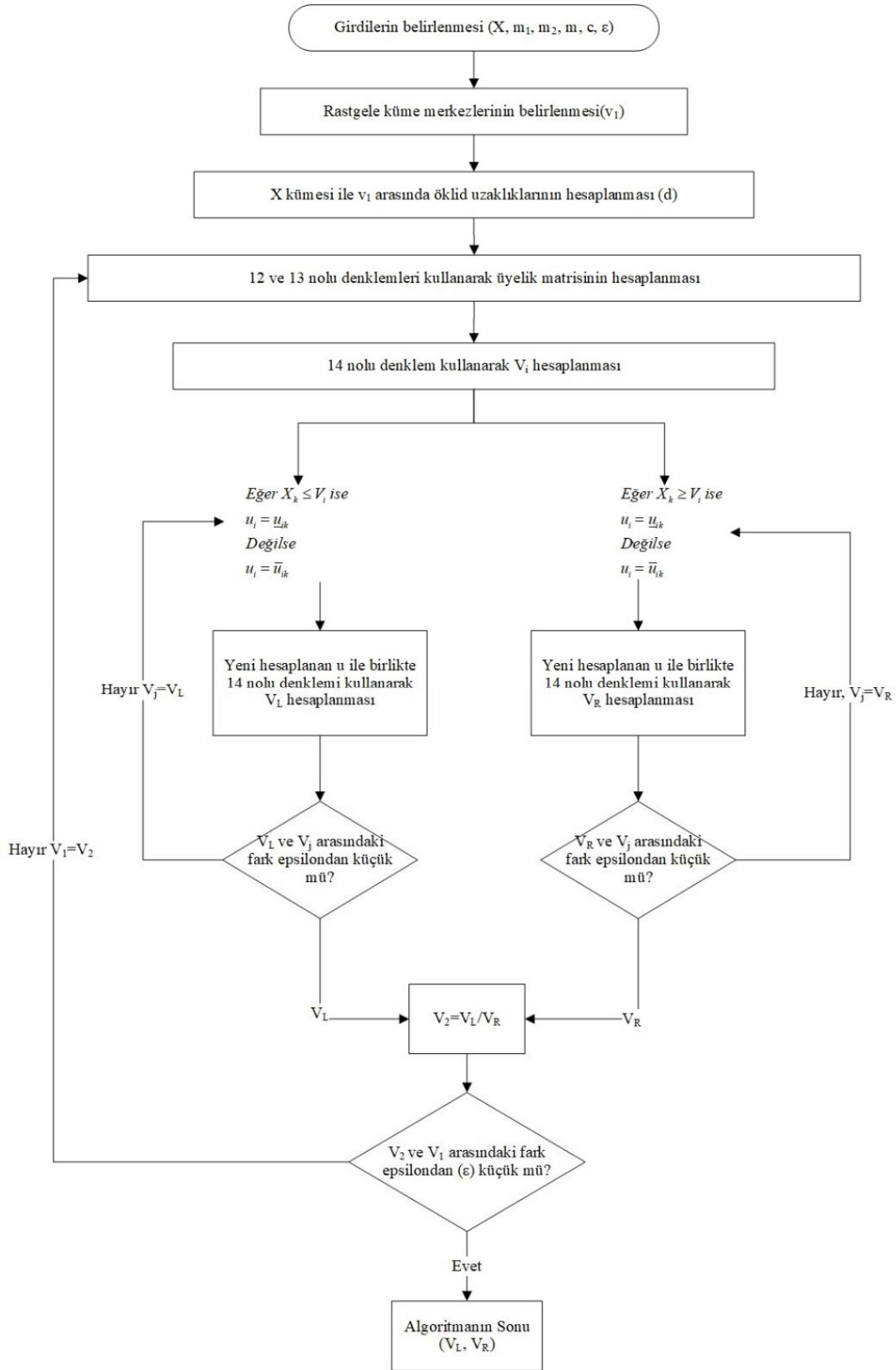
R ifadesi 1'e eşit olursa Q 'nun mutlak niceleyici olduğu anlaşılır. Eğer Q göreceli niceleyici ise R M (toplam nesne sayısı)'ye eşit olur. "Öğrencilerin çoğu matematikten düşük almıştır [0,5]" tip-I niceleyici cümlesine örnektir. Bu cümlede "öğrenciler" Y kümesini, "çoğu" Q niceleyicisini, "düşük almıştır" S özetleyicisini ve [0,5] doğruluk derecesini gösterir. Tip-II niceleyici cümle yapısı tip-I niceleyici cümle yapısından farklılık gösterir, Tip-II niceleyici cümlelerde $w_g(s_g)$ bulunmaktadır ve bu tip cümleler şu şekilde gösterilir: " $w_g(s_g)$ olan Y 'nin Q kadarı S 'dir $[T_1]$ ". Bu tip dilsel özetlerde, T_1 Eş. 17'de verilen formül ile hesaplanır.

$$T_1 = \mu_Q \left(\frac{\sum_{m=1}^M (\mu_{S_g}(d_m) \wedge \mu_{w_g}(v_g^m))}{\sum_{m=1}^M \mu_{w_g}(v_g^m)} \right) \quad (17)$$

Tip-II niceleyici yapısındaki cümleye örnek olarak "Tembel öğrencilerin tümü edebiyattan yüksek almıştır [0,05]" verilebilir. "Öğrenciler" Y kümesini, tembel" ifadesi $w_g(s_g)$ öncül özetleyiciyi, "tümü" Q niceleyiciyi, "yüksek almıştır" S özetleyiciyi, [0,05]'de doğruluk derecesini gösterir.

3. Önerilen Yöntem (Proposed Method)

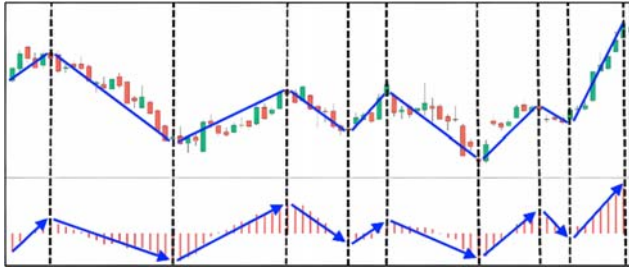
Bu bölümde, aralıklı tip-2 dilsel özetlerin değerlendirilmesi için kullanılan bulanık kümelerin Aralıklı Tip-2 Bulanık C Ortalama ile belirlenmesinin detayı adım adım anlatılmıştır.



Şekil 2. AT2BCO Merkez Noktaları Hesaplama Adımları (IT2FCM Center Points Calculations Steps)

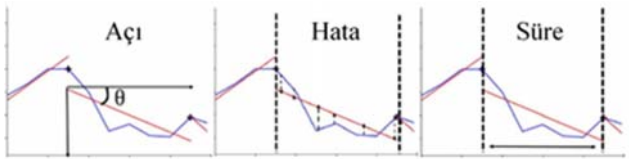
Literatürde dilsel özetleme konusunda bulanık kardinalite yöntemi ile yapılan çalışmalarda [47, 50] eşit bölümlere kullanılarak bulanık kümeler oluşturulmuştur. Eşit bölümlere yöntemi verilerin hangi aralıklarda daha çok biriktiğini göz ardı etmektedir. Bu sebeple bu çalışmada verilerin yoğun olarak yer aldığı aralıkları göz önüne alarak bulanık kümeleri dağıtan Aralıklı Tip-2 Bulanık C Ortalama (AT2BCO) bölümlere yöntemi önerilmiştir. Boran ve Akay [47] ve Özdoğan vd. [50] tarafından önerilen dilsel özetlerin değerlendirilmesi ile ilgili yöntemlere entegre edilmiştir ve sonuçlar karşılaştırılmıştır.

Adım 1: Yerel Trendlerin ve Özelliklerinin Belirlenmesi: Zaman serilerinin önemli bileşenlerinden biri, zaman serilerinin davranışını anlaşılmasını sağlayan trenddir. Eğilimler, doğrusal olarak artan, sabit veya doğrusal olarak azalan fonksiyonlar olarak tanımlanır. Bu çalışmada, ilk adım olarak zaman serilerinden elde edilen MACD-histogram değerleri kullanılarak doğrusal parçalı trendler çıkarılır [56] Doğrusal parçalı trendleri çıkarmadaki önemli adım, zaman serilerindeki dönüm noktalarını belirlemektir. MACD-histogram değeri, doğrusal parçalı trendlerin çıkarılmasında önemli bir adım olan zaman serilerinin dönüm noktalarını belirlemek için kullanılır. MACD-histogram, 12-günlük ve 26-günlük üssel hareketli ortalamalar (EMA) arasındaki farktan oluşan MACD çizgisi ile 9-günlük EMA'dan oluşan sinyal çizgisi arasındaki farkı gösterir. Üssel hareketli ortalama (EMA), daha güncel veriye daha fazla ağırlık veren bir hareketli ortalama türüdür.[50] MACD-histogram değerlerinde Şekil 3'te görüldüğü gibi devamlı artış ve azalış gözlemlendiği zaman diliminde verilerde de artan ve azalan trend oluştuğu gözlenmiştir.



Şekil 3. MACD-histogram ve hisse fiyatları
(MACD-histogram and stock prices)

Şekil 4'te görüldüğü gibi değişim özelliği olarak yerel trendlerin açısı (x eksenine ile yapılan) alınmıştır.



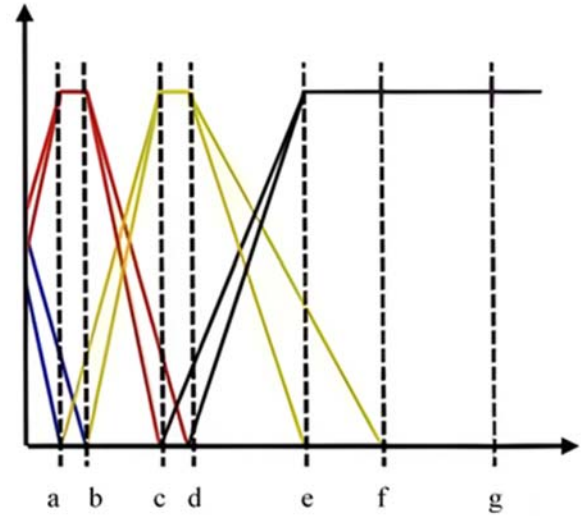
Şekil 4. Trend özellikleri (Trend features)

Trendlerin başlangıç ve bitiş zamanları arasındaki değerlerin oluşturduğu regresyon doğrusu ile değerler arasındaki farkı gösteren MAE (ortalama mutlak hata) verisi değişkenlik özelliğini; trendlerin süresi, süre özelliğini temsil eden alınmıştır. Hesaplamaların yapılması sonucu her bir trendin özelliklerine ait bilgiler veri tabanına kaydedilir. Örneğin incelenen zaman serisinden 5 trend çıktı ise trendlere ait süreler (8,10,15,22,5 gibi) veri tabanına kaydedilir.

Adım 2: Trend Özelliklerini Bulanıklaştırılması: Trendlerin her bir özelliği için belli bir sayıda etiket tanımlanarak (süre için “kısa” “orta” ve “uzun” gibi), zaman serisine ait aralıklı tip-2 fonksiyonları içeren kümeler oluşturulur.

Adım 2.1: Kümelerin merkez noktalarının bulunması: Veri tabanına kaydedilen sayısal değerler Aralıklı Tip-2 Bulanık C Ortalama ile kümeler ayrıştırılır ve her bir kümenin merkez noktaları bulunur. Örneğin süre özelliği dikkate alınarak oluşturulan kümenin (örn: {8,10,15,22,5}) kaç adet küme içereceği dikkate alınarak merkez noktaları (alt ve üst üyelik dereceleri) bulunur (küme sayısı=2 için $m1=[9,10]$ $m2=[16,18]$). Merkez noktaları üyelik dereceleri hesaplamada kullanılır.

Adım 2.2: Bulanık küme fonksiyonlarının oluşturulması: Aralıklı Tip-II Bulanık C Ortalama ile genellikle öklid uzaklığı kullanılarak üyelik derecesi hesaplanır. Geçmiş çalışmalarla[47, 50] karşılaştırma yapılabilmesi için bu çalışmada merkez noktaları esas alınıp yamuk ve üçgen fonksiyonlar oluşturulur. Fonksiyonların oluşturma mantığı Şekil 5'te gösterilmiştir. Şekil 5'te kırmızı ile belirtilen yamuk fonksiyonun ilgili kümenin merkez noktasının alt ve üst üyelik dereceleri (a ve b noktası) ve aralarındaki değerler 1'e eşitlenmiştir. Sarı ile belirtilen fonksiyonda yine ilgili kümenin merkez noktalarının alt ve üst üyelik dereceleri arasında kalan değerler (c ve d noktası) 1'e eşitlenmiştir. Şekilden de anlaşılacağı gibi kırmızı ile belirtilen fonksiyonun sıfır değeri, sarı ile belirtilen fonksiyonun 1 değeri alan noktalarına (merkez noktaları olan c ve d noktalarına) eşitlenmiştir. Siyah ile belirtilen fonksiyonda ilgili kümenin merkez noktalarının alt ve üst üyelik dereceleri (e ve f noktaları) 1'e eşitlenmiştir, g noktası ise kümenin maksimum değerini belirtir. Üçgen fonksiyonunda ise merkez noktalarının ortalaması alınarak bulunan değer 1'e eşitlenmiştir. Hazırlanan bulanık küme fonksiyonları dilsel özetlerin doğruluk derecelerini hesaplayan algoritmalarla girdi oluşturacak şekilde bulanık sonuçlar üretecektir.



Şekil 5. AT2BCO ile yamuk fonksiyonun elde edilmesi
(Obtaining the trapezoidal function with IT2FCM)

Adım 3: Dilsel Özetlerin Oluşturulması: Bu çalışmada “Kısa trendlerin çoğu çok hızlı düşmüştür [0,75]” formunda olan Tip-II niceleyici dilsel özetlerin doğruluk dereceleri hesaplanacaktır. Tip-II niceleyici dilsel özetlerde biri öncül olmak üzere iki özetleyici (kısa ve çok hızlı düşen) ve bir niceleyici (çoğu) bulunmaktadır. Eğer bir zaman serisini 3 niteleyici ve her birinde 3 etiket olan 3 özellik ile özetlemek istiyorsak 81 adet dilsel özet oluşturulması gerekir.

Adım 4: Dilsel Özetlerin Doğruluk Derecelerinin Hesaplanması: Oluşturulan dilsel özetlerin doğruluk derecelerinin (T) hesaplanması için literatürde yer alan olasılık ve olabilirlik temelli iki yöntem kullanılacaktır. [47, 50]

$\Gamma = \{\alpha_1, \alpha_2, \dots, \alpha_m\}$ benzersiz α değerleri kümesidir ve $0 = \alpha_{m+1} < \alpha_m < \dots < \alpha_2 < \alpha_1 = 1$ $\tilde{S}_g^I$ ve $\tilde{S}^I$ 'nin tüm alt ve üst üyelik derecelerini içerir.

α değerleri belirlendikten sonra Tablo 1'de gösterilen kesin kümeler bulunur.

Tablo 1. Alfa kesmelerle elde edilen $\tilde{S}_g^I$ ve $\tilde{S}^I$ 'e ait kesin kümeler (Crisp sets of $\tilde{S}_g^I$ and $\tilde{S}^I$ obtained by alpha cuts)

α	$(\underline{S}_g^I)_{\alpha_i}$	$(\bar{S}_g^I)_{\alpha_i}$	$(\underline{S}^I)_{\alpha_i}$	$(\bar{S}^I)_{\alpha_i}$
α_1	$(\underline{S}_g^I)_{\alpha_1}$	$(\bar{S}_g^I)_{\alpha_1}$	$(\underline{S}^I)_{\alpha_1}$	$(\bar{S}^I)_{\alpha_1}$
α_2	$(\underline{S}_g^I)_{\alpha_2}$	$(\bar{S}_g^I)_{\alpha_2}$	$(\underline{S}^I)_{\alpha_2}$	$(\bar{S}^I)_{\alpha_2}$
α_{m-1}	$(\underline{S}_g^I)_{\alpha_{m-1}}$	$(\bar{S}_g^I)_{\alpha_{m-1}}$	$(\underline{S}^I)_{\alpha_{m-1}}$	$(\bar{S}^I)_{\alpha_{m-1}}$
α_m	$(\underline{S}_g^I)_{\alpha_m}$	$(\bar{S}_g^I)_{\alpha_m}$	$(\underline{S}^I)_{\alpha_m}$	$(\bar{S}^I)_{\alpha_m}$

Her bir alfa kesmelerde $\tilde{S}_g^I$ ve $\tilde{S}^I$ 'nin kaç tane elemana sahip olduğunu belirlenmesi gerekmektedir. S^e , $\tilde{S}^I$ 'nin içinde bir bulanık küme olsun. α değerleri ile elde edilen $(S^e)_{\alpha_i}$ kesin kümesi $(\bar{S}^I)_{\alpha_i}$ 'nin kuvvet kümesindeki alt kümelerinden biridir. Buna bağlı olarak $(S_g^e)_{\alpha_i}$ de α değerleri ile elde edilir ve $(\bar{S}_g^I)_{\alpha_i}$ 'nin kuvvet kümesindeki alt kümelerinden biridir. Sonuç olarak $|(S_g^I)_{\alpha_i} \cap (\tilde{S}^I)_{\alpha_i}| / |(S_g^I)_{\alpha_i}|$ formülünden büyük sayıda farklı değerleri oluşur. Hesaplamayı basitleştirmek için C_{α_i} 'nin olası değerler kümesi Eş. 18'de gösterilmiştir.

$$C_{\alpha_i} = \left\{ \frac{|(S_g^I)_{\alpha_i} \cap (\underline{S}^I)_{\alpha_i}|}{|(S_g^I)_{\alpha_i}|}, \frac{|(S_g^I)_{\alpha_i} \cap (\bar{S}^I)_{\alpha_i}|}{|(S_g^I)_{\alpha_i}|}, \frac{|(\bar{S}_g^I)_{\alpha_i} \cap (\underline{S}^I)_{\alpha_i}|}{|(\bar{S}_g^I)_{\alpha_i}|}, \frac{|(\bar{S}_g^I)_{\alpha_i} \cap (\bar{S}^I)_{\alpha_i}|}{|(\bar{S}_g^I)_{\alpha_i}|} \right\} \quad (18)$$

Kümenin değerleri bulunduğunda, $Q(c_{\alpha_i})$ değerleri (her α_i değeri için $c_{\alpha_i} \in C_{\alpha_i}$ olmak üzere) hesaplanabilir. Daha sonra C_{α_i} kümesindeki her bir eleman için $Q(c_{\alpha_i})$ 'nın minimum ve maksimum değerleri belirlenir. Olabilirlik temelli yaklaşım ile tip-II niceleyici cümlelerin doğruluk derecesi Eş. 19'da tanımlanan formülasyon ile hesaplanır.

$$T^L = \max_{\alpha_i \in \Gamma} \left\{ \min \left\{ \sum_{\substack{\mu_Q(\alpha) \geq \min\{\mu_Q(c_{\alpha_i})\} \\ \alpha \in [\alpha_{i+1}, \alpha_i]}} (\alpha_i - \alpha_{i+1}), \min_{c_{\alpha_i} \in C_{\alpha_i}} \{\mu_Q(c_{\alpha_i})\} \right\} \right\} \quad (19)$$

$$T^U = \max_{\alpha_i \in \Gamma} \left\{ \min \left\{ \sum_{\substack{\mu_Q(\alpha) \geq \max\{\mu_Q(c_{\alpha_i})\} \\ \alpha \in [\alpha_{i+1}, \alpha_i]}} (\alpha_i - \alpha_{i+1}), \max_{c_{\alpha_i} \in C_{\alpha_i}} \{\mu_Q(c_{\alpha_i})\} \right\} \right\} \quad (20)$$

Olasılık temelli yaklaşım ile tip-II niceleyici cümlelerin doğruluk derecesi Eş. 21'deki gibi hesaplanır.

$$(T^L, T^U) = \sum_{\alpha_i \in \Gamma} \alpha_i - \alpha_{i+1} \times \left(\min_{c_{\alpha_i} \in C_{\alpha_i}} Q(c_{\alpha_i}), \max_{c_{\alpha_i} \in C_{\alpha_i}} Q(c_{\alpha_i}) \right) \quad (21)$$

4. Deneysel Çalışma (Experimental Study)

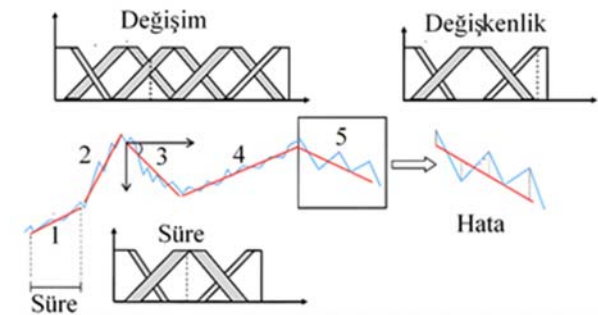
Bu bölümde Borsa İstanbul'da işlem göre üç hisse senedinin (GARAN, PETKM ve THYAO) son 10 yıllık günlük kapanış değerleri zaman serisi olarak alınmış ve bu çalışmada önerilen ve bir önceki bölümde adım adım anlatılan yöntem ile dilsel özetlerin doğruluk derecesi hesaplanmıştır. Üç hisse senedi fiyat manipülasyonlarının minimum olması için en çok işlem hacmine sahip hisselerden rastgele seçilmiştir. Bu çalışmada önerilen yöntem, Boran ve Akay [47] Özdoğan vd. [50] tarafından önerilen yöntemler ile ayrı ayrı karşılaştırılmıştır. Hisse senetlerinin zaman serileri 11 Kasım 2012 ile 11 Kasım 2022 arasında yaklaşık 2512 günü kapsamaktadır. Yapılan analize göre GARAN hisse senedi için minimum değer 4.97 Türk Lirası (TL), maksimum değeri 29.55 Türk Lirası (TL), PETKM hisse senedi için minimum değer 0.64 Türk Lirası (TL), maksimum değeri 15.25 Türk Lirası (TL), THYAO hisse senedi için minimum değer 3.63 Türk Lirası (TL), maksimum değeri 113.40 Türk Lirası (TL)'dir.

Python'da yazılmış algoritma ile MACD histogram kullanılarak GARAN, PETKM ve THYAO hisse senetlerinden sırasıyla 504, 475 ve 460 adet doğrusal yerel trend elde edilmiştir. Üç hisse senedinin 10 yıllık verilerinden oluşturulan doğrusal yerel trendlere ilişkin tanımlayıcı istatistikler Tablo 2'te verilmiştir.

Tablo 2. Yerel trendlere ilişkin istatistikler (Statistics on local trends)

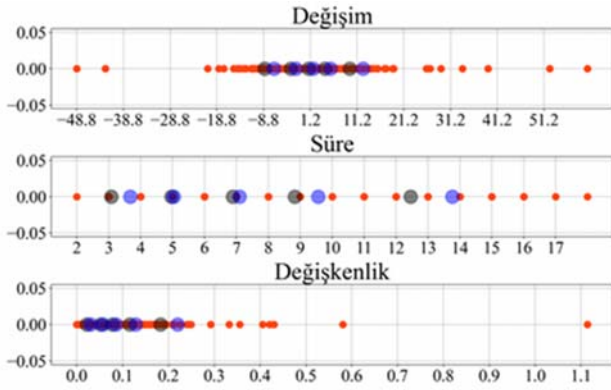
	THYAO		GARAN		PETKM	
	Min.	Maks	Min.	Maks	Min.	Maks
Trendin Süresi	3	22	2	18	3	19
x eksenine ile yaptığı açı	-50	73	-52	63	-18	38
MAE değeri	0	2.3	0	1.25	0	0.40

Yerel trendlerin açısı (x eksenine ile yapılan) değişim özelliği olarak alınmıştır. Şekil 6'da yerel trendin x eksenine ile yaptığı açı gösterilmiştir. Ortalam Mutlak Hata (MAE) değeri, regresyon doğrusu ile kapanış değerleri arasındaki farkı göstermektedir bu da değişkenlik özelliği olarak alınmıştır. Örnek bir yerel trendin hata değeri Şekil 6'da gösterilmiştir. Trendlerin süresi de süre özelliğini temsil etmesi için alınmıştır.

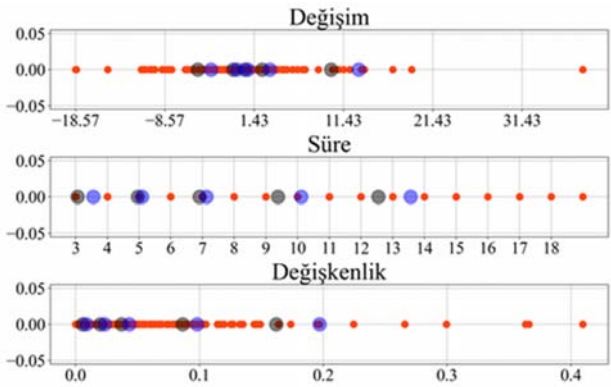


Şekil 6. Yerel trendlerin 3 özelliği (3 features of local trends)

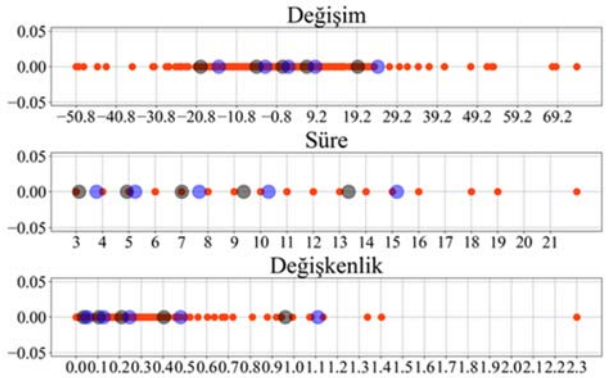
Her bir özelliğe ait 5 adet dilsel etiket tanımlanmıştır. Belirlenen özelliklere ait sayısal veriler kullanılarak aralıklı tip-2 bulanık c ortalama ile (5 adet dilsel etiket belirlendiği için 5 küme oluşturarak) kümelerin merkez noktaları hesaplanmıştır. Her bir hisse senedi için bulunan merkez noktaları mavi noktalar ile Şekil 7, 8 ve 9'da gösterilmiştir. Turuncu noktalar ise ilgili özelliğe ait verileri göstermektedir.



Şekil 7. GARAN hisse senedine ait AT2BCO ile elde edilen merkez noktalar (Center points of GARAN stock obtained with IT2FCM)

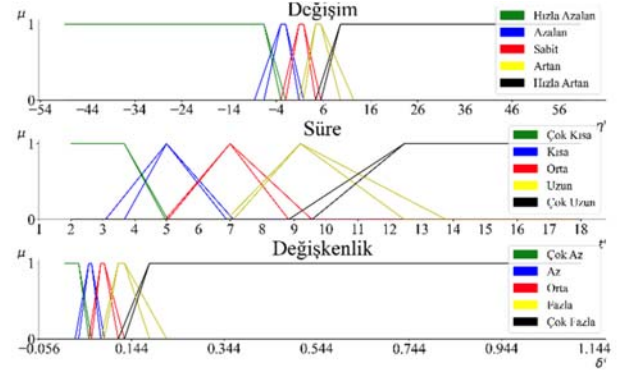


Şekil 8. PETKM hisse senedine ait AT2BCO ile elde edilen merkez noktalar (Center points of PETKM stock obtained with IT2FCM)

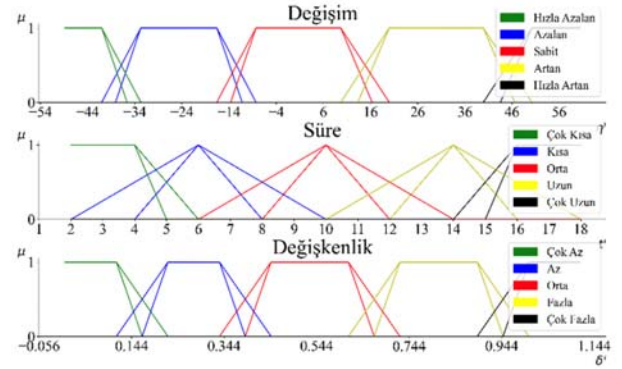


Şekil 9. THYAO hisse senedine ait AT2BCO ile elde edilen merkez noktalar (Center points of THYAO stock obtained with IT2FCM)

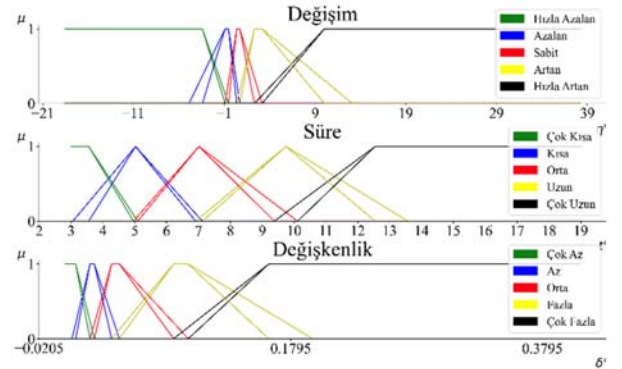
Merkez noktalarına göre oluşturulan bulanık kümeler Şekil 10, 12 ve 14'te gösterilmiştir. Boran ve Akay [47] ve Özdoğan vd. [50] tarafından yapılan çalışmaya göre oluşturulan bulanık kümeler de Şekil 11, 13 ve 15'te verilmiştir. Bu çalışmalarda bulanık kümelerin etiket değerleri maksimum ve minimum değerler dikkate eşit bölümlenme ile oluşturulmuştur. Eşit bölümlenme ve A2TBCO bölümlenme kümeleme metrikleri Xie-Beni İndeksi [57] ile hesaplanmıştır. Daha düşük Xie-Beni İndeksi değeri daha iyi bir kümeleme kalitesi gösterir. Eşit bölümlenmede küme merkezleri yamuk kümede üst noktaların orta noktaları, üçgen kümede en üst noktaları alınmıştır.



Şekil 10. GARAN hisse senedine ait AT2BCO bölümlenme ile elde edilen bulanık kümeler (Fuzzy sets of GARAN stock obtained by IT2FCM partitioning)



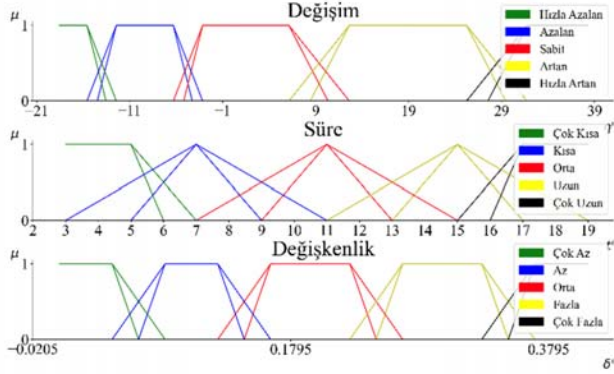
Şekil 11. GARAN hisse senedine ait eşit bölümlenme ile elde edilen bulanık kümeler (Fuzzy sets of GARAN stock obtained by uniform partitioning)



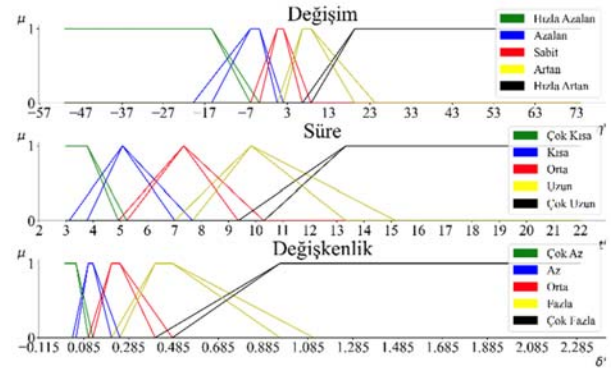
Şekil 12. PETKM hisse senedine ait AT2BCO bölümlenme ile elde edilen bulanık kümeler (Fuzzy sets of PETKM stock obtained by IT2FCM partitioning)

x veri noktası, v küme merkezi, μ üyelik derecesi, c küme sayısı kullanılarak Tablo 3'deki sonuçlar hesaplanmıştır (küme sayısı=5, bulanıklık mertebesi(m)=2 alınmıştır). Tablo 3'den anlaşılacağı üzere AT2BCO bölümlenmenin kümeleme metriği açısından daha iyi sonuç verdiği görülmektedir. Bulanık kümelerde ele alınan değişken genel olarak 3, 5 veya 7 adet etiket kullanılarak açıklanmaktadır. Bunun temel nedeni insanlar ifadelerini yansıtmak için ölçümlemede ortanca değeri nötr değer olarak 3'lü, 5'li veya 7'li ölçeği kullanmaktadır. Bulanık dilsel özetleme çalışmalarında özellik sayısının fazla olduğu durumlarda özet sayısının üstel sayıda artması nedeni ile genellikle 3'lü veya 5'li ölçek kullanılmaktadır. Zaman serileri 3

adet özellik kullanılarak özetlenmesi, kısmi trend sayısı ve trende ait özelliklerin aldığı değerlerin dağılımları göz önüne alındığında literatürde yapılan çalışmalarda genellikle 5'li ölçek kullanılmaktadır.[33, 47, 49, 50, 58]



Şekil 13. PETKM hisse senedine ait eşit bölümlenme ile elde edilen bulanık kümeler
(Fuzzy sets of PETKM stock obtained by uniform partitioning)



Şekil 14. THYAO hisse senedine ait AT2BCO bölümlenme ile elde edilen bulanık kümeler
(Fuzzy sets of THYAO stock obtained by IT2FCM partitioning)

Üç adet niceleyici, beş adet dilsel etiket ve 3 özellik ile her bir hisse senedi için toplam 450 adet tip-II niceleyici dilsel özet oluşturulmuştur. Oluşturulan dilsel özetler olasılık ve olabilirlik temelli yaklaşımlar ile doğruluk dereceleri hesaplanmıştır. İki yaklaşıma ait hesaplamalar Python 3.11 sürümünde, MacOS işletim

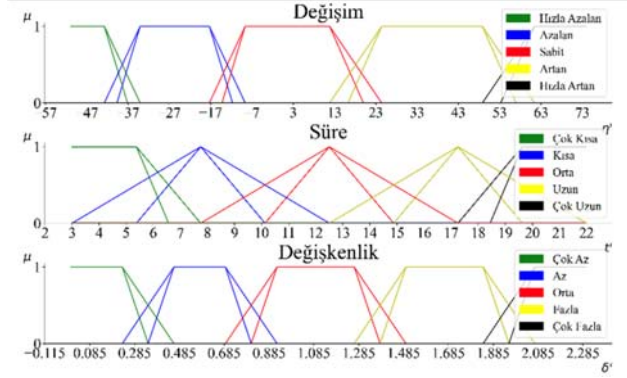
Tablo 3. Xie-Beni İndeksi Değerleri (Xie-Beni Index Values)

Eşit B.	THYAO	0,78	12,1	0,7
AT2BCO B.	THYAO	0,55	0,85	0,192
Eşit B.	PETKM	0,74	3,27	0,85
AT2BCO B.	PETKM	1,57	1,51	0,2
Eşit B.	GARAN	0,63	7,35	0,74
AT2BCO B.	GARAN	0,82	1,31	0,207

Tablo 4. PETKM hissesine ait dilsel özetlerin doğruluk dereceleri (Truth degrees of linguistic summaries of PETKM stock)

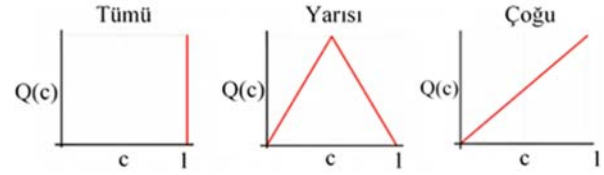
PETKM	Eşit Bölümlenme	AT2BCO Bölümlenme	Eşit Bölümlenme	AT2BCO Bölümlenme
	Olasılık Temelli Yaklaşım		Olabilirlik Temelli Yaklaşım	
Çok kısa trendlerin çoğu sabittir	[0,93 0,94]	[0,28 0,35]	[0,89 0,89]	[0,36 0,4]
Kısa trendlerin çoğu sabittir	[0,94 0,95]	[0,32 0,41]	[0,93 0,94]	[0,37 0,46]
Orta uzunluktaki trendlerin çoğu sabittir	[0,89 0,94]	[0,35 0,45]	[0,87 0,89]	[0,38 0,47]
Uzun trendlerin çoğu sabittir	[0,95 1]	[0,29 0,38]	[0,89 1]	[0,36 0,44]
Çok uzun trendlerin tümü sabittir	[1 1]	[0 0]	[1 1]	[0 0]
Çok uzun trendlerin çoğu sabittir	[1 1]	[0,28 0,39]	[1 1]	[0,37 0,44]
Hızla azalan trendlerin yarısı çok kısadır	[0,73 0,77]	[0,74 0,79]	[0,67 0,67]	[0,75 0,85]
Azalan trendlerin yarısı çok kısadır	[0,78 0,93]	[0,53 0,6]	[0,74 0,83]	[0,58 0,62]
Sabit trendlerin yarısı çok kısadır	[0,86 1]	[0,51 0,57]	[0,73 0,98]	[0,6 0,68]
Azalan trendlerin yarısının değişkenliği çok azdır	[0,77 0,89]	[0,8 0,82]	[0,69 0,82]	[0,7 0,73]
Sabit trendlerin çoğunun değişkenliği çok azdır	[0,9 0,92]	[0,66 0,68]	[0,88 0,89]	[0,65 0,66]
Çok az değişkenliğe sahip trendlerin çoğu sabittir	[0,96 0,97]	[0,44 0,54]	[0,94 0,95]	[0,47 0,54]
Orta değişkenliğe sahip trendlerin yarısı azalandır	[0,75 0,84]	[0,44 0,6]	[0,71 0,8]	[0,44 0,57]
Çok kısa trendlerin çoğunun değişkenliği çok azdır	[0,92 0,95]	[0,63 0,64]	[0,92 0,92]	[0,6 0,61]
Kısa trendlerin yarısının değişkenliği çok azdır	[0,26 0,35]	[0,77 0,79]	[0,33 0,37]	[0,69 0,74]
Kısa trendlerinin çoğunun değişkenliği çok azdır	[0,82 0,87]	[0,43 0,45]	[0,78 0,81]	[0,46 0,46]
Orta uzunluktaki trendlerin yarısının değişkenliği çok azdır	[0,5 0,61]	[0,83 0,84]	[0,49 0,52]	[0,75 0,75]
Çok uzun trendlerin çoğunun değişkenliği çok azdır	[0,85 0,92]	[0,29 0,29]	[0,71 0,85]	[0,33 0,33]
Değişkenliği çok az olan trendlerin yarısı çok kısadır	[0,8 0,93]	[0,79 0,8]	[0,68 0,92]	[0,81 0,87]

sistemi üzerinde Visual Studio Code kullanılarak geliştirilmiştir. Her bir hisse senedi için hesaplanan T değeri (doğruluk derecesi) $[0,75, 0,75]$ 'den büyük dilsel özetler Tablo 4, 5 ve 6'da verilmiştir. Her bir yöntemden $0,75$ 'den büyük dilsel özetlerin yanında diğer yöntemlerin doğruluk dereceleri de verilmiştir.



Şekil 15. THYAO hisse senedine ait eşit bölümlenme ile elde edilen bulanık kümeler
(Fuzzy sets of THYAO stock obtained by uniform partitioning)

Dilsel özetlerde kullanılan üç niceleyici “hepsi”, “yarısı” ve “çoğu” Şekil 16’da gösterilmiştir.



Şekil 16. Dilsel özetlerin oluşturulmasında kullanılan üç niceleyici
(Three quantifiers used to construct linguistic summaries)

$[0,75, 0,75]$ 'den büyük doğruluk derecesine sahip dilsel özetlere bakıldığında bazı farklılıkların ve benzerliklerin olduğu görülmektedir. AT2BCO bölümlenme yöntemi ile THYAO için her iki yaklaşımla hepsi aynı özet olmak üzere 5'er adet dilsel özetin elde edildiği görülmektedir. PETKM için olasılık temelli yaklaşım ile 5 tane, olabilirlik temelli yaklaşım ile 3 tane dilsel özet elde edildiği bunlardan 3 tanesinin aynı dilsel özet olduğu görülmektedir. GARAN için olasılık temelli yaklaşım ile 6 tane, olabilirlik temelli yaklaşım ile 2 tane dilsel özet elde edildiği, bunlardan 2 tanesinin ortak olduğu görülmektedir.

Tablo 5. GARAN hissesine ait dilsel özetlerin doğruluk dereceleri (Truth degrees of linguistic summaries of GARAN stock)

GARAN	Eşit Bölümlenme	AT2BCO Bölümlenme	Eşit Bölümlenme	AT2BCO Bölümlenme
	Olasılık Temelli Yaklaşım		Olabilirlik Temelli Yaklaşım	
Çok kısa trendlerin çoğu sabittir	[0,91 0,93]	[0,24 0,3]	[0,85 0,87]	[0,32 0,35]
Kısa trendlerin çoğu sabittir	[0,88 0,92]	[0,24 0,32]	[0,83 0,84]	[0,32 0,39]
Orta uzunluktaki trendlerin çoğu sabittir	[0,95 0,97]	[0,27 0,36]	[0,93 0,94]	[0,33 0,4]
Uzun trendlerin çoğu sabittir	[0,96 0,99]	[0,23 0,33]	[0,94 0,97]	[0,31 0,4]
Çok uzun trendlerin tümü sabittir	[1 1]	[0 0]	[1 1]	[0 0]
Çok uzun trendlerin çoğu sabittir	[1 1]	[0,38 0,47]	[1 1]	[0,44 0,54]
Hızla azalan trendlerin yarısı çok kısadır	[1 1]	[0,76 0,79]	[1 1]	[0,74 0,75]
Azalan trendlerin yarısı çok kısadır	[0,82 0,91]	[0,65 0,69]	[0,76 0,86]	[0,74 0,75]
Sabit trendlerin yarısı çok kısadır	[0,79 0,86]	[0,64 0,68]	[0,79 0,8]	[0,68 0,71]
Hızla artan trendlerin tümü çok kısadır	[0,93 1]	[0 0]	[0,93 1]	[0 0]
Hızla artan trendlerin yarısı çok kısadır	[0 0,07]	[0,87 0,9]	[0 0,07]	[0,71 0,75]
Hızla artan trendlerin çoğu çok kısadır	[0,97 1]	[0,44 0,47]	[0,93 1]	[0,47 0,49]
Çok kısa trendlerin yarısının değişkenliği azdır	[1 1]	[0,44 0,53]	[1 1]	[0,51 0,54]
Kısa trendlerin yarısının değişkenliği çok azdır	[0,47 0,72]	[0,75 0,78]	[0,5 0,67]	[0,66 0,71]
Sabit trendlerin çoğunun değişkenliği çok azdır	[0,94 0,97]	[0,34 0,36]	[0,92 0,93]	[0,37 0,41]
Çok hızlı artan trendlerin tümünün değişkenliği azdır	[0,91 0,99]	[0 0]	[0,91 0,99]	[0 0]
Çok hızlı artan trendlerin çoğunun değişkenliği azdır	[0,95 0,99]	[0,15 0,2]	[0,91 0,99]	[0,24 0,3]
Çok az değişkenliğe sahip trendlerin çoğu sabittir	[0,95 0,97]	[0,3 0,38]	[0,93 0,94]	[0,35 0,43]
Çok fazla değişkenliğe sahip trendlerin yarısı hızla artandır	[0 0]	[0,76 0,84]	[0 0]	[0,74 0,78]
Çok kısa trendlerin yarısının değişkenliği çok azdır	[0,09 0,13]	[0,91 0,91]	[0,15 0,17]	[0,76 0,87]
Çok kısa trendlerin çoğunun değişkenliği çok azdır	[0,93 0,95]	[0,52 0,52]	[0,91 0,91]	[0,52 0,56]
Kısa trendlerin çoğunun değişkenliği çok azdır	[0,87 0,92]	[0,28 0,28]	[0,8 0,83]	[0,3 0,37]
Orta uzunluktaki trendlerin çoğunun değişkenliği çok azdır.	[0,9 0,95]	[0,2 0,2]	[0,88 0,89]	[0,25 0,32]
Uzun trendlerin çoğunun değişkenliği çok azdır	[0,83 0,94]	[0,15 0,16]	[0,75 0,87]	[0,24 0,31]
Çok uzun trendlerin çoğunun değişkenliği çok azdır	[0,91 0,98]	[0,06 0,07]	[0,8 0,89]	[0,11 0,18]
Değişkenliği çok az olan trendlerin yarısı çok kısadır	[0,82 0,87]	[0,8 0,81]	[0,81 0,83]	[0,75 0,79]

Tablo 6. THYAO hissesine ait dilsel özetlerin doğruluk dereceleri (Truth degrees of linguistic summaries of THYAO stock)

THYAO	Eşit Bölümleme	AT2BCO Bölümleme	Eşit Bölümleme	AT2BCO Bölümleme
	Olasılık Temelli Yaklaşım		Olabilirlik Temelli Yaklaşım	
Çok kısa trendlerin çoğu sabittir	[0,77 0,8]	[0,25 0,34]	[0,74 0,76]	[0,32 0,39]
Kısa trendlerin çoğu sabittir	[0,73 0,82]	[0,29 0,4]	[0,76 0,79]	[0,39 0,5]
Orta uzunluktaki trendlerin çoğu sabittir	[0,65 0,77]	[0,27 0,36]	[0,76 0,76]	[0,36 0,47]
Hızla azalan trendlerin yarısı çok kısadır	[0,21 0,42]	[0,89 0,95]	[0,26 0,53]	[0,81 0,82]
Hızla azalan trendlerin çoğu çok kısadır	[0,79 0,9]	[0,45 0,51]	[0,6 0,74]	[0,46 0,48]
Sabit trendlerin yarısı çok kısadır	[0,83 0,91]	[0,54 0,69]	[0,68 0,89]	[0,6 0,63]
Çok hızlı artan trendlerin yarısı çok kısadır	[0,42 0,74]	[0,79 0,87]	[0,47 0,74]	[0,79 0,81]
Azalan trendlerin çoğunun değişkenliği çok azdır	[0,76 0,81]	[0,49 0,57]	[0,73 0,77]	[0,49 0,56]
Sabit trendlerin çoğunun değişkenliği çok azdır	[0,95 0,97]	[0,63 0,72]	[0,94 0,94]	[0,59 0,67]
Artan trendlerin yarısının değişkenliği çok azdır	[0,54 0,77]	[0,79 0,9]	[0,55 0,66]	[0,75 0,81]
Hızla artan trendlerin yarısının değişkenliği ortadır	[0,9 1]	[0,52 0,65]	[0,8 1]	[0,49 0,55]
Değişkenliği çok az olan trendlerin çoğu hızla azalandır	[0,86 0,89]	[0,44 0,55]	[0,83 0,84]	[0,45 0,55]
Değişkenliği fazla olan trendlerin yarısı çok hızla artandır	[0 0]	[0,74 0,82]	[0 0]	[0,67 0,76]
Değişkenliği çok fazla olan trendlerin yarısı çok hızla artandır	[0 0]	[0,89 0,94]	[0 0]	[0,8 0,86]
Çok kısa trendlerin çoğunun değişkenliği çok azdır	[0,9 0,93]	[0,62 0,67]	[0,89 0,9]	[0,61 0,64]
Kısa trendlerin çoğunun değişkenliği çok azdır	[0,78 0,87]	[0,33 0,42]	[0,86 0,86]	[0,41 0,48]
Orta uzunluktaki trendlerin çoğunun değişkenliği çok azdır	[0,66 0,79]	[0,28 0,36]	[0,76 0,82]	[0,38 0,47]
Değişkenliği çok az olan trendlerin yarısı çok kısadır	[0,81 0,92]	[0,84 0,92]	[0,68 0,85]	[0,79 0,82]

Tablo 7. Sonuçların Karşılaştırılması (Comparison of Results)

	Olasılık Temelli Yaklaşım		Olabilirlik Temelli Yaklaşım	
	Eşit Bölümleme	AT2BCO Bölümleme	Eşit Bölümleme	AT2BCO Bölümleme
THYAO	0,0607	0,06039	0,0598	0,054341
GARAN	0,0725	0,05134	0,07013	0,05292
PETKM	0,07058	0,056547	0,0704	0,0533

Boran ve Akay [47] ve Özdoğan vd. [50] tarafından önerilen eşit bölümleme yöntemi ile elde edilen sonuçlarda THYAO için olasılık temelli yaklaşım ile 10 adet, olabilirlik temelli yaklaşım ile 8 adet dilsel özetin elde edildiği bunların 5 tanesinin aynı dilsel özet olduğu görülmektedir. PETKM için olasılık temelli yaklaşım ile 16 tane, olabilirlik temelli yaklaşım ile 10 tane dilsel özet elde edildiği bunlardan 10 tanesinin aynı dilsel özet olduğu görülmektedir. GARAN için olasılık temelli yaklaşım ile 22 adet, olabilirlik temelli yaklaşım ile 21 adet dilsel özet elde edildiği bunlardan 21 tanesinin aynı dilsel özet olduğu görülmektedir.

Doğruluk derecelerinin kararlılığı, Hu ve Hu (2020) tarafından yapılan çalışmada aralık değerli verilerin varyans hesaplama yöntemi (Eş. 22) kullanılarak incelenmiştir. Varyans değeri küçük olan yöntemlerin daha kararlı sonuç verdiği öngörülmüştür.

$$Var(X) = Var(mid(x)) + Var(rad(x)) + \frac{2}{n} \sum_{i=1}^n |\Delta m_i \Delta r_i| \quad (22)$$

Oluşturulan bütün dilsel özetlerin doğruluk dereceleri dikkate alınarak aralıklı değerlere ait varyans hesabı Tablo 7’de verilmiştir. Tablo 7’den de kolayca anlaşılabileceği gibi olasılık ve olabilirlik temelli yaklaşımda AT2BCO bölümleme ile elde edilen doğruluk derecelerinin varyansının daha düşük dolayısıyla önerilen yöntemin daha kararlı olduğu ve dört farklı yöntem karşılaştırıldığında da en

tutarlı sonucu “AT2BCO Bölümleme ile Olabilirlik Temelli Yaklaşım”ının verdiği görülmektedir. AT2BCO bölümleme ile oluşturulan bulanık kümelerin ardından olasılık ve olabilirlik temelli yaklaşımla elde edilen bazı dilsel özetleri dikkate alındığında alım satım stratejileri aşağıdaki şekilde belirlenebilir.

Olasılık temelli yaklaşımda THYAO için “Hızla azalan trendlerin yarısı çok kısadır [0,89 0,95]” ifadesinin doğruluk derecesi yüksek hesaplanmıştır. Son 10 yılın hisse senedi fiyatlarında hızla azalan trendlerin (x eksenine açısı -7 ve daha küçük olan trendler) yarısı çok kısa sürer (3-5 gün). Bu tarz trendler oluştuğunda hisse senedi alınırsa dip fiyattan alınmış olur. Olabilirlik temelli yaklaşım da çıkan sonuçlara bakıldığında da aynı strateji izlenebilir. Olasılık temelli yaklaşımda PETKM için “Orta uzunluktaki trendlerin yarısının değişkenliği çok azdır [0,83 0,84]” ifadesinin doğruluk derecesi yüksek çıkmıştır. 5-10 gün süren trendlerin yarısının trendden sapma göstermediği, aşağı yukarı hareket etmediği, düz bir çizgi izlediğini göstermektedir. Olabilirlik temelli yaklaşımda “Değişkenliği çok az olan trendlerin yarısı çok kısadır [0,81 0,87]” dilsel özeti, trendi takip eden verilerin yarısının süresinin çok kısa sürdüğünü (5-6 gün) göstermektedir. Olasılık temelli yaklaşımda GARAN için “Hızla artan trendlerin yarısı çok kısadır [0,87 0,9]” ifadesi x eksenine açısı 10 ve daha büyük olan trendlerin yarısının 4-5 gün sürdüğünü göstermektedir. Olabilirlik temelli yaklaşımda “Çok kısa trendlerin yarısının değişkenliği çok azdır [0,76 0,87]” dilsel özeti, 4-5 gün süren trendlerin yarısının trend çizgisini izlediğini göstermektedir.

5. Sonuçlar ve Tartışmalar (Results and Conclusions)

Dilsel özetleme konusunda literatürde kullanılan yöntemlerin birçoğu tip-1 bulanık kümeler kullanılarak oluşturulmuştur. Belirsizlikleri daha iyi yansıtsa da dilsel özetleme konusunda aralıklı tip-2 bulanık kümeler literatürde çok az kullanılmıştır. Bunun nedeni, skaler kardinalite tabanlı değerlendirmelere sahip aralıklı tip-2 bulanık kümelerin tutarlı olmayan sonuçlar üretmesidir. Boran ve Akay [47] yaptığı çalışmada aralıklı tip-2 bulanık kümelere ait alfa kesmelerin kombinasyonu kullanarak ilgili kümelerin temsil edilebileceğini kanıtladı, böylelikle bulanık kardinalite ile dilsel özetleme yapılabileceği gösterilmiş oldu. Boran ve Akay [47] ve Özdoğan vd. [50] yaptıkları çalışmada eşit bölümlere ile bulanık kümeleri oluşturarak olasılık ve olabilirlik temelli yaklaşımlar önerdi. Bu çalışmada ise, aralıklı tip-2 bulanık kümeler oluşturulurken eşit bölümlere yerine AT2BCO bölümlere önerilmiştir. Yapılan uygulamada olasılık ve olabilirlik temelli yöntemler ile dilsel özetlerin doğruluk dereceleri hesaplanmıştır.

Deneyisel çalışmalarda, GARAN, THYAO ve PETKM hisse senetlerinin son 10 yıla ait günlük kapanış fiyatlarının aralıklı tip-2 bulanık küme ile dilsel özetlemesi yapılmıştır. İlk olarak, yerel trendler, MACD-histogram değerine göre çıkarılmıştır, daha sonra 3 niceleyici (tümü, yarısı ve çoğu), süre, MAE ve açığı parametreleri kullanılarak her hisse senedi için 450 dilsel özet oluşturulmuştur. Oluşturulan her dilsel özet için doğruluk derecesi hesaplanmış ve alım satım stratejileri için [0,75 0,75]'ten büyük değerler göz önüne alınmıştır. AT2BCO bölümlere ile elde edilen doğruluk derecelerinin varyansının daha düşük dolayısıyla önerilen yöntemin daha kararlı olduğu görülmektedir, aynı durum kümeleme metriklerinde (Tablo 3) de görülmüştür. AT2BCO dağılım kümeleme metriklerinde daha iyi sonuçlar vermiştir. Belirsizliğin Kapladığı Alan (BKA) bulanık mantık sistemlerinde belirsizliğin kapsamını ve derecesini ifade eder. Zaman serilerinde BKA, bulanık kümelerin sınırlarını belirleyerek belirsizliği modellemekte kullanılır ve tahminlerin doğruluğunu artırır. Bu yöntem, verilerin belirsizliklerini daha iyi yönetir ve zaman serilerinin öngörülebilirliğini artırarak daha sağlam sonuçlar elde edilmesini sağlar. Gelecekteki çalışmalarda yarı-bulanık niceleyici cümleler kullanarak daha karmaşık dilsel özetler oluşturulabilir. Bu tür dilsel özetlere “Yaklaşık %30’u hariç bütün insanlar öğrencidir”, “Koşan erkeklere oranla iki veya üç kat fazla bisiklet kullanan kadın var.” cümleleri örnek verilebilir.

Kaynaklar (References)

- Lecun, Y., Bengio, Y., Hinton, G., Deep learning, *Nature*. 2015.
- Zhao, Z. Q., Zheng, P., Xu, S. T., Wu, X., Object Detection with Deep Learning: A Review, *IEEE Trans Neural Netw Learn Syst*. 2019.
- Wolf, T., Debut, L., Sanh, V., Chaumond, J., Delangue, C., Moi, A., Cistac, P., Rault, T., Louf, R., Funtowicz, M., Davison, J., Shleifer, S., Platen, P. von, Ma, C., Jernite, Y., Plu, J., Xu, C., Scao, T. Le, Gugger, S., Drame, M., Lhoest, Q., Rush, A. M., Transformers: State-of-the-Art Natural Language Processing, 38-45, 2020.
- Qiu, X. P., Sun, T. X., Xu, Y. G., Shao, Y. F., Dai, N., Huang, X. J., Pre-trained models for natural language processing: A survey, *Sci China Technol Sci*, 63 (10), 1872-1897, 2020.
- Liu, P., Yuan, W., Jiang, Z., Hayashi, H., Neubig, G., Fu, J., Yuan, W., Jiang, Z., Hayashi, H., Neubig, G., Fu, J., Pre-train, Prompt, and Predict: A Systematic Survey of Prompting Methods in Natural Language Processing, *ACM Comput Surv*, 55 (9), 1-35, 2023.
- Topol, E. J., High-performance medicine: the convergence of human and artificial intelligence, *Nature Medicine*, 25 (1), 44-56, 2019.
- He, J., Baxter, S. L., Xu, J., Xu, J., Zhou, X., Zhang, K., The practical implementation of artificial intelligence technologies in medicine, *Nat Med*, 25 (1), 30-36, 2019.
- Fernandes, M., Vieira, S. M., Leite, F., Palos, C., Finkelstein, S., Sousa, J. M. C., Clinical Decision Support Systems for Triage in the Emergency Department using Intelligent Systems: A Review, *Artif Intell Med*, 102, 101762, 2020.
- Otter, D. W., Medina, J. R., Kalita, J. K., A Survey of the Usages of Deep Learning for Natural Language Processing, *IEEE Trans Neural Netw Learn Syst*, 2020.
- Xi, L., Zhang, L., Liu, J., Li, Y., Chen, X., Yang, L., Wang, S., A virtual generation ecosystem control strategy for automatic generation control of interconnected microgrids, *IEEE Access*, 8, 94165-94175, 2020.
- Xi, L., Zhou, L., Liu, L., Duan, D., Xu, Y., Yang, L., Wang, S., A deep reinforcement learning algorithm for the order optimization allocation of total power in the interconnected power grids, *CSEE Journal of Power and Energy Systems*, 6, 713-723, 2020.
- Ribeiro, M. T., Wu, T., Guestrin, C., Singh, S., Beyond Accuracy: Behavioral Testing of NLP Models with CheckList, *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, 4902-4912, 2020.
- Zadeh, L. A., Fuzzy sets, *Information and Control*, 8 (3), 338-353, 1965.
- Yager R.R., A new approach to the summarization of data, *Inf Sci (N Y)*, 28 (1), 69-86, 1982.
- Kacprzyk, J., Yager, R. R., Linguistic summaries of data using fuzzy logic, *Int J Gen Syst*, 30 (2), 133-154, 2001.
- Zadeh, L. A., A computational approach to fuzzy quantifiers in natural languages, *Computers and Mathematics with Applications*, 9 (1), 149-184, 1983.
- Kacprzyk, J., Wilbik, A., Zadrozny, S., Linguistic summarization of trends: a fuzzy logic based approach, *11th International Conference Information Processing and Management of Uncertainty in Knowledge-based Systems*, 2166-2172, 2006.
- Sklansky, J., Gonzalez, V., Fast polygonal approximation of digitized curves, *Pattern Recognit*, 12 (5), 327-331, 1980.
- Kacprzyk, J., Wilbik, A., Zadrozny, S., Capturing the essence of a dynamic behavior of sequences of numerical data using elements of a quasi-natural language, *IEEE International Conference on Systems, Man and Cybernetics*, 4, 3365-3370, 2007.
- Kacprzyk, J., Wilbik, A., Zadrozny, S., A linguistic quantifier based aggregation for a human consistent summarization of time series, *Advances in Soft Computing*, 37, 183-190, 2006.
- Kacprzyk, J., Wilbik, A., Zadrozny, S., On some types of linguistic summaries of time series, *3rd International IEEE Conference Intelligent Systems*, 373-378, 2006.
- Kacprzyk, J., Wilbik, A., Zadrozny, S., Linguistic summarization of time series under different granulation of describing features, *Lecture Notes in Computer Science*, 4585 LNAI, 230-240, 2007.
- Kacprzyk, J., Wilbik, A., Zadrozny, S., Linguistic summarization of time series using a fuzzy quantifier driven aggregation, *Fuzzy Sets Syst*, 159 (12), 1485-1499, 2008.
- Kacprzyk, J., Wilbik, A., Zadrozny, S., Linguistic summaries of time series via a quantifier based aggregation using the Sugeno integral, *IEEE International Conference on Fuzzy Systems*, 713-719, 2006.
- Kacprzyk, J., Wilbik, A., Zadrozny, S., Analysis of Time Series via their Linguistic Summarization: The Use of the Sugeno Integral, *Seventh International Conference on Intelligent Systems Design and Applications (ISDA 2007)*, 262-270, 2007.
- Kacprzyk, J., Wilbik, A., Zadrozny, S., Mining time series data via linguistic summaries of trends by using a modified Sugeno integral based aggregation, *Proceedings of the 2007 IEEE Symposium on Computational Intelligence and Data Mining, CIDM 2007*, 742-749, 2007.
- Kacprzyk, J., Wilbik, A., Zadrozny, S., Linguistic Summarization of Time Series by Using the Choquet Integral, *Lecture Notes in Computer Science (LNCS)*, 4529, 284-294, 2007.
- Kacprzyk, J., Wilbik, A., Zadrozny, S., Linguistic summaries of time series via an OWA operator based aggregation of partial trends, *IEEE International Conference on Fuzzy Systems*, 1-6, 2007.
- Kacprzyk, J., Wilbik, A., Using fuzzy linguistic summaries for the comparison of time series: an application to the analysis of investment fund quotations, *IFSA/EUSFLAT*, 2009, 1321-1326, 2009.
- Castillo-Ortega, R., Marin, N., Martinez-Cruz, C., Sanchez, D., A proposal for the hierarchical segmentation of time series. Application to trend-based linguistic description, *IEEE International Conference on Fuzzy Systems*, 489-496, 2014.
- Castillo-Ortega, R., Marin, N., Martinez-Cruz, C., Sanchez, D., Linguistic comparison of time series using the End-Point Fit algorithm, *IEEE International Conference on Fuzzy Systems (FUZZ-IEEE)*, 1-8, 2015.

32. Ramos-Soto, A., Bugarin, A., Barro, S., Computing with perceptions for the linguistic description of complex phenomena through the analysis of time series data, 7th ICAART Conference, Doctoral Consortium Session, 7, 2015.
33. Hatipoğlu, H., Boran, F. E., Avci, M., Akay, D., Linguistic summarization of Europe Brent spot price time series along with the interpretations from the perspective of Turkey, *International Journal of Intelligent Systems*, John Wiley and Sons Ltd, 29 (10), 946-970, 2014.
34. Sanchez-Valdes, D., Alvarez-Alvarez, A., Trivino, G., Dynamic linguistic descriptions of time series applied to self-track the physical activity, *Fuzzy Sets Syst*, 285, 162-181, 2016.
35. Kacprzyk, J., Zadrozny, S., Fuzzy logic-based linguistic summaries of time series: A powerful tool for discovering knowledge on time varying processes and systems under imprecision, *Wiley Interdiscip Rev Data Min Knowl Discov*, 6 (1), 37-46, 2016.
36. Kaczmarek, K., Hryniewicz, O., Kruse, R., Human input about linguistic summaries in time series forecasting, ACHI 2015 - 8th International Conference on Advances in Computer-Human Interactions, 1, 9-13, 2015.
37. Moyse, G., Lesot, M. J., Linguistic summaries of locally periodic time series, *Fuzzy Sets Syst*, 285, 94-117, 2016.
38. Marín, N., Sánchez, D., On generating linguistic descriptions of time series, *Fuzzy Sets Syst*, 285, 6-30, 2016.
39. Kaczmarek, K., Hryniewicz, O., Time Series Classification with Linguistic Summaries, 2015 Conference of the International Fuzzy Systems Association and the European Society for Fuzzy Logic and Technology (IFSA-EUSFLAT-15), 2015.
40. Kaczmarek-Majer, K., Hryniewicz, O., Application of linguistic summarization methods in time series forecasting, *Inf Sci (N Y)*, 478, 580-594, 2019.
41. Dündar, B., Akay, D., Boran, F. E., Özdemir, S., Fuzzy Quantification and Opinion Mining on Qualitative Data using Feature Reduction, *International Journal of Intelligent Systems*, 33 (9), 1840-1857, 2018.
42. Jain, A., Popescu, M., Keller, J., Rantz, M., Markway, B., Linguistic summarization of in-home sensor data, *J Biomed Inform*, 96, 103240, 2019.
43. Genç, S., Akay, D., Boran, F. E., Yager, R. R., Linguistic summarization of fuzzy social and economic networks: an application on the international trade network, *Soft comput*, 24 (2), 1511-1527, 2020.
44. Niewiadomski, A., A type-2 fuzzy approach to linguistic summarization of data, *IEEE Transactions on Fuzzy Systems*, 16 (1), 198-212, 2008.
45. Niewiadomski, A., On finity, countability, cardinalities, and cylindric extensions of type-2 fuzzy sets in linguistic summarization of databases, *IEEE Transactions on Fuzzy Systems*, 18 (3), 532-545, 2010.
46. Wu, D., Mendel, J. M., Linguistic summarization using IFTHEN rules and interval Type-2 fuzzy sets, *IEEE Transactions on Fuzzy Systems*, 19 (1), 136-151, 2011.
47. Boran, F. E., Akay, D., A generic method for the evaluation of interval type-2 fuzzy linguistic summaries, *IEEE Trans Cybern*, 44 (9), 1632-1645, 2014.
48. Delgado, M., Sánchez, D., Vila, M. A., Fuzzy cardinality based evaluation of quantified sentences, *International Journal of Approximate Reasoning*, 23 (1), 23-66, 2000.
49. Boran, F. E., Akay, D., Yager, R. R., A probabilistic framework for interval type-2 fuzzy linguistic summarization, *IEEE Transactions on Fuzzy Systems*, 22 (6), 1640-1653, 2014.
50. Özdoğan, İ., Boran, F. E., Akay, D., A possibilistic approach for interval type-2 fuzzy linguistic summarization of time series, *Artificial Intelligence Review* 2021 54:5, 54 (5), 3991-4018, 2021.
51. Klir, G. J., Yuan, Bo., Fuzzy sets and fuzzy logic : Theory and applications. Prentice Hall PTR, 1995.
52. Mendel, J. M., Uncertain Rule-Based Fuzzy Systems. Springer International Publishing, Cham, 2017.
53. Ross, T. J., Fuzzy logic with engineering applications. John Wiley, 2010.
54. Mendel, J. M., John, R. I., Liu, F., Interval type-2 fuzzy logic systems made simple, *IEEE Transactions on Fuzzy Systems*, 14, 808-821, 2006.
55. Hwang, C., Rhee, F. C. H., Uncertain fuzzy clustering: Interval type-2 fuzzy approach to C-means, *IEEE Transactions on Fuzzy Systems*, 15 (1), 107-120, 2007.
56. Wang, J., Kim, J., Predicting stock price trend using MACD optimized by historical volatility, *Math Probl Eng*, 2018, 2018.
57. Xie, X. L., Beni, G., A validity measure for fuzzy clustering, *IEEE Trans Pattern Anal Mach Intell*, 13 (8), 841-847, 1991.
58. Boran, F. E., Akay, D., Yager, R. R., An overview of methods for linguistic summarization with fuzzy sets, *Expert Syst Appl*, 61, 356-377, 2016.
59. Hu, C., Hu, Z. H., On Statistics, Probability, and Entropy of Interval-Valued Datasets, *Communications in Computer and Information Science*, 1239 CCIS, 407-421, 2020.