

Received : 05.12.2018

Editorial Process Begin: 13.02.2018

Published: 15.05.2018

E-ticaret Müşteri Bağlılığı Gri İlişkisel Kümeleme Analizi

Hüseyin FİDAN, Öğr. Gör., Mehmet Akif Ersoy Üniversitesi, Mühendislik Mimarlık Fakültesi, Endüstri Mühendisliği, Burdur-TÜRKİYE, huseyinfidan@gmail.com, Orcid: 0000-0002-7482-8922

ÖZ

İnternet teknolojileriyle hayatımızı değiştiren en büyük gelişmelerden olan e-ticaret tüketicilere ve firmalara önemli avantajlar getirmektedir. Günümüzde e-ticaret bir rekabet aracı olmaktan çok firmaların ayakta kalabilmesi için bir zorunluluk haline gelmiştir. Bu bağlamda yeni müşteri kazanmak, müşterileri elde tutmak, güven oluşturmak ve müşteri bağlılığını sağlamak gibi e-ticaret stratejileri, firmalar açısından önemli konular haline gelmiştir. Özellikle müşteri bağlılığını oluşturmak ve sürdürmek firma karlılığını arttırmak için hayati bir konudur. Bu sebeple bağlılık oluşan müşteri gruplarının belirlenmesi, bu gruplara uygulanacak doğru satış stratejilerinin seçilmesi açısından önem arz etmektedir. Müşteri gruplarının belirlenmesi için kümeleme analizleri gerçekleştirilmekte, bu amaçla K-ortalamlar, K-medoids ve bulanık C-ortalamlar algoritmaları veya bu algoritmaları temel alan metotlar kullanılmaktadır. Ancak merkezi kümeleme algoritmaları olarak bilinen bu algoritmalar belirsiz olan küme sayısı ve küme merkezi gibi değerleri analiz öncesi parametre olarak istemektedir. Bu çalışmada, bir e-ticaret sitesinden temin edilen, toplam satın alma işlem sayısı, toplam işlem tutarı, ortalama işlem tutarı, siteye giriş sayısı, şikâyet sayısı ve ürün geri iade sayısı bilgilerini içeren gerçek işlem verileri temel alınarak müşteri bağlılığı kümeleme analizi gerçekleştirilmiştir. Analiz öncesinde küme sayısı ve küme merkezleri belirsiz olduğu için kümeleme işlemi Gri İlişkisel Analiz ile gerçekleştirilmiştir. Araştırma sonuçlarına göre, analiz öncesi küme sayısı ve küme merkezleri belirlenmeksizin kümelenmenin gerçekleştirilebileceği ortaya konulmuş, Gri İlişkisel Kümeleme analizi ile e-ticaret müşterilerinin bağlılık kümelenmeleri gerçekleştirilmiştir.

Anahtar kelimeler: Gri teori, Gri ilişkisel analiz, Gri ilişkisel kümeleme, Bağlılık, E-ticaret

Gray Relational Clustering Analysis of E-commerce Customers Loyalty

ABSTRACT

E-commerce, which is one of the biggest developments that change our life with internet technologies, brings significant advantages to consumers and firm. Nowadays, e-commerce has become a necessity for firms to survive rather than as a competitive tool. In this context, e-commerce strategies such as acquiring new customers, retaining customers, building trust and providing customer loyalty have become important issues in terms of companies. Especially creating and maintaining customer loyalty are crucial issues to increase the profitability of the firm. For this reason, the identification of loyalty of customer groups is important in terms of selecting the right sales strategies to be applied to these groups. Clustering analyzes are performed to determine customer groups, using K-means, K-medoids and fuzzy C-means algorithms or methods based on these algorithms for this purpose. However, these algorithms, known as central clustering algorithms, require values such as cluster number and cluster center, which are uncertain, as parameters before analysis. In this study, a customer loyalty clustering analysis was conducted based on actual transaction data from an e-commerce site, including the total number of purchases, total transaction amount, average transaction amount, number of entries on site, number of complaints, number of product return. Since the

number of clusters and cluster centers are uncertain before the analysis, clustering was performed by Gray Relational Analysis. According to the results of the research, e-commerce customers' loyalty clusters have been realized with Gray Relational Clustering analysis, which shows that the clusters can be realized without determining the number of clusters and cluster centers before analysis.

Keywords: *Gray theory, Gray relational analysis, Gray relational clustering, Loyalty, E-commerce*

GİRİŞ

İnternet üzerinden yapılan ve tüm ticari süreçleri kapsayan işlemler elektronik ticaret olarak tanımlanmaktadır (Fidan & Albeni, 2014: 288). Web teknolojilerinin 1995 yılından sonra gelişmeye başlamasıyla hız kazanan e-ticaret, firmalar ve tüketiciler açısından geleneksel ticaret anlayışında değişikliklere yol açmıştır. E-ticaretin bilgiye erişim çerçevesinde tam rekabet piyasasına yaklaştıran unsurlar içerdiği düşünülse de fiyat farklılaştırması, yüksek fiyat dağılımları, teknolojik eşitsizlikler, monopol eğilimi gibi piyasa aksaklıklarının da bulunduğu bir yapıya sahiptir (Ellison & Ellison, 2005: 148). Bu bağlamda firmaların fiyat ve kaliteden ibaret olan rekabet araçlarına ek olarak güven ve bağlılık gibi kavramlar da eklenmiştir.

E-ticarete firmaların nihai hedefleri müşterilerde güven tesis ederek bağlılıklarını kazanmaktır. Firmaların uzun dönem büyüme hedefleri açısından önemli bir konu olan bağlılık, karlılık düzeyleri ile yakından ilişkili bir kavramdır (Reichheld, 1995: 10). Bağlılık oluşmuş bir müşteri, normal müşterilere nazaran e-ticaret sitesini daha fazla ziyaret etmekte ve daha fazla harcama gerçekleştirmektedir (Rosen, 2001). Firmalar açısından müşteri bağlılığı bir kar aracı olarak görülmektedir (Srinivasan, Anderson & Ponnayolu, 2002: 41). Bu sebeple firmaların, bağlılığını kazandığı müşterilerini belirlemesi ve ürün satış süreçlerini bu müşteri gruplarına göre tespit etmesi, kar maksimizasyonları için önemlidir.

Veri madenciliğinde, verilerin özelliklerine göre gruplandırılması kümeleme analizi olarak adlandırılmakta ve genellikle makine öğrenmesi algoritmaları kullanılmaktadır. Kümeleme analizlerinde yoğun olarak K-ortalamlar algoritması kullanılmaktadır (Jain, 2010; Han, Kamber & Pei, 2012; Huang & Song, 2014; Kaushik, 2016). Ancak kümelemeyi belirlenen bir merkeze göre gerçekleştirmesi, küme sayısını gösteren k parametresinin analiz öncesinde belirlenmesi, uçdeğerleri tanımlayamaması, farklı yoğunluk ve büyüklükteki kümeler ile çalışmada başarısız olması, veri setinin çok sayıda taranması, bu algoritmanın dezavantajları arasında gösterilmektedir (Kim & Ahn, 2005; Tajunisha, 2010; Bafghi, 2017). Kümeleme işleminin etkinliğini azaltan bu problemler kümelerin oluşmasında başarısızlığa yol açmaktadır. Özellikle oluşacak küme sayısının belirsiz olduğu durumlarda etkin sonuçlara ulaşılamamakta ve tekrarların çoğalmasına bağlı olarak analiz sürecini uzatmaktadır. Gerçek zamanlı çalışacak sistemlerin tasarımında sorun oluşturan bu durum, K-ortalamlar ve benzer yaklaşımdaki algoritmaların yetersizliğinden kaynaklanmaktadır. Gri İlişkisel Analiz (GİA), küçük örneklem ve eksik bilgi içeren problemlerin çözümünde kullanılmak üzere önerilmiş bir yöntemdir (Deng, 1982). Belirsizlik ve eksik bilgi durumunda sonuçlar üretebilen yeni bir yaklaşım olan GİA, küme sayısının belirsiz olduğu veya öngörü yapılamadığı durumlarda kullanılabilecek uygun bir analiz yöntemidir. Ayrıca veri setinin tekrar taranmasına ihtiyaç duymaması, işlem süresinin kısa olmasını sağlamaktadır. Bu kapsamda çalışmanın amacı, küme sayısı belirsizliğinde e-ticaret müşterilerini bağlılıklarına göre Gri ilişkisel yaklaşım ile kümeleme analizinin gerçekleştirilmesidir.

Çalışmada e-ticaret sektöründe faaliyet gösteren bir firmadan temin edilen kullanıcıların işlem verileri ile kullanıcıların bağlılık kümelenmesi gerçekleştirilmiştir. Müşteri bağlılığını değerlendirmek için kullanıcıların toplam satın alma işlem sayısı, toplam işlem tutarı, ortalama işlem tutarı, sisteme giriş sayısı, şikâyet sayısı, ürün geri iade sayısı olmak üzere altı değişken kullanılmıştır. Son bir yıl içerisinde gerçekleştirilen işlem verileri ile veri seti oluşturulmuş ve GİA tabanlı Gri İlişkisel Kümeleme (GİK) analizi gerçekleştirilmiştir. Sonuçlara göre GİK'in analiz öncesi küme sayısı ve küme merkezi belirlenmeksizin, kümelenme işlemlerinde kullanılabilecek etkin bir araç olduğu ortaya konulmuştur.

E-ticaret ve Müşteri Bağlılığı

İnternetin 1960 yılında keşfedilmesinden sonra hızlı bir gelişim süreci yaşanmış, iletişim amaçlı geliştirilen sistemler ile tüm dünyayı saran bir ağ haline gelmiştir. Başlangıçta bilgi ve dosya paylaşımı amaçlanan internet ile ilk olarak mesaj ve elektronik posta gönderim sistemleri tasarlanmıştır (Gromov, 2012). Tim Barners Lee tarafından geliştirilen HTML sistemi sayesinde temelleri oluşturulan web kavramı, interneti bambaşka bir sürece taşımıştır (Spector, 2001). Bu sayede sadece düz yazı ortamında gerçekleştirilen iletişim daha kolay hale gelmiş, resim, ses ve görüntü bilgilerinin de aktarılmasına imkân vermiştir (Fidan, 2014). Kısıtlı da olsa kullanıcı ile etkileşimin yolunu açan bu gelişmelerle birlikte ilk e-ticaret işlemleri görülmeye başlanmış, ancak bilgi güvenliği endişeleri, kullanıcıların e-ticaretten uzak durmalarına yol açmıştır (Spector, 2001). 1995 yılında iletişim güvenliğini sağlayan sistemlerin geliştirilmesinden sonra hızla yaygınlaşan e-ticaret, günümüzde firmaların yeni bir satış kanalı haline gelmiştir.

E-ticaret için başlangıç yıllarında kusursuz piyasa, tam rekabet piyasası gibi nitelermeler yapılmasına karşın, geleneksel ticarete çok fazla önem verilmeyen bazı sorunların, bu yeni ticaret biçiminde büyük problemlere yol açtığı zamanla görülmüştür. Bu sorunlar arasında bilgi asimetrisi (Fidan & Albeni, 2014), sayısal bölünme (Fidan, 2016), mahremiyet (Ellison & Ellison, 2005), dolandırıcılık gibi unsurlar gösterilebilir. Bu sorunlara çözüm getirebilecek başlıca faktör ise kullanıcılarda güven tesis etmek ve bu sayede müşteri bağlılığını kazanmak olarak görülmektedir (Fidan, 2014). E-ticaret müşteri bağlılığında fiyat mekanizmasının yeterli olmadığı, müşteri bağlılığının öncelikle güven oluşturulması ile ilgili olduğu vurgulanmaktadır (Reicheld & Scheffer 2000:107). Bu çerçevede, müşteri bağlılığının oluşturulması e-ticaret firmalarının başlıca hedefleri arasındadır.

Geleneksel ticaret ve e-ticaret için önemli bir kavram olan müşteri bağlılığı, tüketicinin satıcı hakkında olumlu düşüncelerinin gelecekte tekrarlayan satın almalar ile sonuçlanması şeklinde tanımlanmaktadır (Srinivasan, Anderson & Pannavolu, 2002: 42; Islam, Khadem & Sayem, 2012: 213). Bu çerçevede bağlılık, bir sonraki satış işleminin hedeflendiği bir yaklaşımdır. Uzun ve zorlu bir süreç olan müşteri bağlılığının oluşturulması, taraflar arasında fiziksel etkileşim olmadığı için e-ticarete daha zorlu bir süreçtir. E-ticaretin piyasa yapısı gereği risk algılarının yüksek olması, müşterilerde bağlılık oluşturulmasını zorlaştırmaktadır. Fiziksel olarak inceleyemediği bir ürünü, hiç tanımadığı bir satıcıdan, peşin ödeme ile satın alan kullanıcı, anlatılan biçim ve kalitede ürünün kendisine teslim edilmesini beklemesi, e-ticarete alıcıların risk algılarının artmasına neden olmaktadır. Satış alma işlemlerini doğrudan etkileyen bir faktör olan risk algısını en aza indirmenin yolu ise müşteri bağlılığını oluşturmaktır. E-ticaret sitesine bağlılığı olan kullanıcılarının satın alma eğilimleri yüksek olmasına karşın, alternatif e-ticaret sitelerinden ürün inceleme ve araştırma eğilimleri daha az olmaktadır (Srinivasan, Anderson & Pannavolu, 2002: 45). Reicheld ve

Schefter (2000)'a göre tüketici bağlılığının oluşturulması firma karını maksimizasyonundan daha çok e-ticaret firmasının yaşam mücadelesidir.

Müşteri bağlılığı ile ilgili çalışmalar müşteri ilişkileri yönetimi (CRM) çerçevesinde yapılmakta çalışmaların, tüketicilerin marka bağlılığı ve satıcı bağlılığı konularına yoğunlaştığı belirtilmektedir (Brown & Peterson, 1993; Corstjens & Lal, 2000; Uncles, Dowling & Hammond, 2003). Özellikle internet üzerinden yapılan ticari işlemlerin artmasıyla birlikte müşteri bağlılığı ile ilgili çalışmalar bu alana yönelmiş ve e-ticarete müşteri bağlılığı konusu popüler konulardan biri olmuştur (Toufaily, Ricard & Perrien, 2013). Çalışmaların büyük bir kısmı müşteri bağlılığını etkileyen faktörlerin belirlenmesi üzerine olup bu amaçla birçok farklı ölçek geliştirilmiştir. Kullanılan değişkenler davranışsal (behavioral), tutumsal (attitudinal) ve bilişsel (cognitive) olmak üzere üç grupta toplanmaktadır (Toufaily, Ricard, & Perrien, 2013). Davranışsal değişkenlerle konuya yaklaşımların daha yoğun olduğu analizlerde müşterilerin satın alma işlem sayısı, satın alma işlem tutarı ve satın alma işlem güncelliği gibi değişkenler kullanılmaktadır (Lee & Overby, 2004; Yun & Good, 2007). Sadece tek bir grup değişkenlerin analizde kullanılmasının eksik bir yaklaşım olacağını vurgulayan araştırmalar, konunun geniş bir bakış açısı ile ele alınması gerektiğini belirterek birleşik modeller önermektedir (Verona & Prandelli, 2002; Ilsever, Cyr & Parent, 2006). Bu görüşe göre davranışsal değişkenlerin yanı sıra tutumsal ve bilişsel değişkenlerle analizlerin gerçekleştirilmesi daha uygun bir yaklaşım olacaktır.

E-ticarete müşteri bağlılığı üzerine yapılan çalışmalar incelendiğinde, literatürde genellikle B2C üzerinde durulduğu, bağlılık kavramının satıcı ve web sayfaları çerçevesinde incelendiği ve analizlerde genellikle anket yönteminin kullanıldığı görülmektedir (Toufaily, Ricard & Perrien, 2013: 1438). Oluşturulan anketler ile elde edilen verilere faktör analizi yada doğrudan anket değişkenleri kullanılarak korelasyon analizi ve regresyon analizleri uygulanmaktadır. Bu sayede tüketici bağlılığını oluşturacak etkenler arasındaki anlamlı ilişkiler belirlenmekte, regresyon ifadeleri ile tüketicilerin bağlılık düzeylerinin matematiksel ifadeleri tespit edilmektedir.

E-ticaret müşteri bağlılığının faktörlerini belirlemeye yönelik çalışmaların yanı sıra, bağlılık oluşun ya da oluşabilecek müşteri veya müşteri gruplarının belirlenmesi firmalar açısından önemli diğer konudur. Bu sayede firmalar fiyat farklılaştırma, özendirme, teşvik gibi satış stratejilerini belirlerken hedef kitleyi isabetli biçimde belirleyebilirler. Benzer özellikte verilerin aynı gruplar altında toplanması olarak tanımlanan kümeleme analizleri (Han, Kamber & Pei, 2012), firmaların potansiyel müşterilerini tanıyabilmeleri açısından önemlidir. Bu sayede firmalar müşteri bağlılığı oluşun müşterilerini belirleyebildikleri gibi, bağlılık potansiyeline sahip müşterileri de tespit etme imkânına kavuşurlar. Bu amaçla, müşterilerin davranışsal bağlılık çerçevesinde geçmiş dönem işlemleri temel alınarak kümeleme analizleri kullanılır.

E-ticarete müşteri bağlılığı ile ilgili çalışmalarda, kümelenme analizlerinin fazla sayıda olmadığı görülmektedir. Kümeleme analizleri gerçekleştirilen çalışmalarda genel olarak K-ortalamlar algoritmasının kullanıldığı, bu algoritmanın eksikliklerini gidermek için farklı yaklaşımlarla desteklendiği görülmektedir (Davidson, 2002). E-ticaret sitesinden mobil telefon satın alan müşterilerin kümeleme analizini gerçekleştiren Huang ve Song (2014), K-ortalamların en uygun algoritma olduğunu öne sürmüştür. Özellik seçiminde dalgacık

dönüşümünü tavsiye eden araştırmacılar bu şekilde etkinliğin arttırılacağını savunmuştur. E-ticaret müşteri bağlılığının davranışsal bağlılık değişkenleri ile incelendiği bir başka çalışmada, kümeleme ve birliktelik kuralları analizlerini beraber kullanan Cheng ve Chen (2009), K-ortalamlar ile oluşan kümelerin birliktelik ilişkilerini belirlemiştir. Brown, Pope ve Voges (2003), e-ticaret müşterilerinin K-ortalamlar algoritması kullanarak kümelenmesini gerçekleştirmiş ve oluşan kümelerin cinsiyet, ürün tipi ve geçmiş satın almaları ile ilişkilerini incelemiştir. Küme sayısı ve küme merkezleri seçiminin K-ortalamlar için oldukça önemli olduğunu vurgulayan Niknam ve Amiri (2010), müşteri bağlılığının kümelenmesinde K-ortalamlar temelli bulanık ve karınca kolonisi algoritmalarını kullanmış ve karma bir yöntem geliştirmiştir. Kalaiselvi (2015), internet bankacılığı müşterilerin kümelenmesini K-medoids algoritmasını kullanarak gerçekleştirmiş, Park ve Jun (2009) küme merkezi seçimlerine yeni bir bakış açısı getirerek geliştirdikleri K-medoids algoritmasını Iris ve Soybean veri kümelerine uygulamıştır. Yuliari, Putra ve Rusjayanti (2015), davranışsal değişkenler çerçevesinde bulanık C-ortalamlar algoritmasını kullanarak kullanıcıların kümelenmesini analiz etmişlerdir. K-ortalamlar algoritmasında küme merkezi seçimlerinin analiz öncesinde yapılmasını bir eksiklik olan gören Tajunisha (2010), bu eksikliği gidermek için temel bileşenler analizi ile küme merkezlerinin belirlenmesini önermiş, bu sayede K-ortalamlar algoritmasının daha sağlıklı sonuçlara ulaştığını belirlemiştir. Aynı yaklaşımı hem K-ortalamlar hem de bulanık C-ortalamlar algoritmasında kullanarak uygulayan Afrin, Al-Amin ve Tabassum (2015), temel bileşenler analizi ile merkezlerin belirlendiği bulanık C-ortalamlar algoritmasının daha iyi sonuçlar verdiğini tespit etmiştir. K-ortalamlar algoritmasının etkinliğini arttırmak için genetik algoritma ile birlikte kullanan Kim ve Ahn (2005), sadece K-ortalamlar algoritmasının kullanıldığı sonuçlara göre daha sağlıklı sonuçlar elde etmiştir. Aynı yaklaşımı gerçek işlem verilerine uygulayan Bafghi (2017), genetik algoritmayla desteklenen K-ortalamların veri miktarı yüksek düzeyde olduğunda sağlıklı sonuçlar verdiğini ortaya koymuştur.

E-ticarete müşteri bağlılığının kümelenmesi ile ilgili çalışmalar incelendiğinde genellikle merkezi yaklaşımlı kümeleme algoritmalarının etkinliğinin arttırılması hedeflenmektedir. K-ortalamlar, K-medoids ve bulanık C-ortalamlar temelinde veri madenciliği ve makine öğrenmesi algoritmaları ile desteklenen metotlar geliştirilerek daha etkin kümelenmeler elde edildiği savunulmaktadır. Ancak literatürde genel kabul görmüş bir çözüm yaklaşımı bulunmamaktadır. Bu algoritmalar ile ilgili literatür incelendiğinde, analiz öncesi küme sayısının parametre olarak alınması, geçici küme merkezlerinin analiz öncesi belirlenmesi, her merkez değişiminde işlemlerin tekrarlanması konularının, söz konusu algoritmaların ortak sorunları olduğu görülmektedir.

Makine Öğrenmesi ve Kümeleme

Bir veri seti içerisinde yer alan verileri kullanarak yeni bilgilere ulaşma yöntemi olan Makine Öğrenmesi (Machine Learning - ML) 1950'li yıllarda temelleri oluşturulan bir araştırma konusudur. Özellikle 1990'lı yıllardan sonra başlı başına bir araştırma alanı haline gelen ML, günümüzde bilgisayar bilimleri, iş analitiği, veri bilimi gibi alanlar başta olmak üzere hem teknik bilimler hem de sosyal bilimlerde kullanılmaktadır. Geçmiş verilere dayanarak sonraki uygulanacak işlemin veya oluşacak değer belirlenmesini amaçlayan algoritmalar tahminleme, en iyileme ve kümeleme analizlerinde etkin biçimde kullanılmaktadır. Günümüzde birçok ML algoritması geliştirilmekle beraber bu algoritmalar denetimli,

denetimsiz ve takviyeli öğrenme olmak üzere üç grup altında toplanmaktadır (Han, Kamber & Pei, 2012).

Danışmanlı öğrenme olarak da isimlendirilen denetimli öğrenme, sahip olunan verilere dayanarak sonuçta oluşacak değerlerin belirlenmesi amacı ile oluşturulan algoritmalar. Sonuç bilgisi için danışılacak olan, öğrenme için kullanılacak veri setidir. Denetimli öğrenme çerçevesinde kullanılan algoritmalar sınıflama ve tahminleme amacıyla kullanılan algoritmalar olup karar ağaçları, basit bayes sınıflandırıcı, doğrusal regresyon, lojistik regresyon, rastgele orman ve destek vektör makinalarıdır (Min & Han, 2005; Bishop, 2007).

Denetimsiz öğrenme, çıktı değeri söz konusu olmaksızın, veri seti içerisindeki ilişkileri ve örüntüleri tespit etmek amacıyla gerçekleştirilen algoritmalar. Temel amaç veri seti içerisindeki benzer nitelikte olan vektörlerin gruplandırılmasıdır. Bu sebeple kümeleme analizi olarak isimlendirilir. Bu algoritmalar arasında temel bileşenler analizi, birliktelik kuralları, kümeleme, tekil değer ayrışımı algoritmaları sayılabilir (Bishop, 2007).

Takviyeli öğrenme yapı itibari ile denetimli öğrenmeye benzemekle birlikte, çıktı sonuçlarının kontrolü amacıyla geri beslemeyi içeren bir yaklaşımdır. Deneme yanılma ile doğru sonuca ulaşılma süreci, sonucun kontrolü ile gerçekleştirilir. Bir nevi sebep sonuç ilişkisine dayanan kendi kendine öğrenme sürecidir. Bu öğrenme sürecinde ilk olarak sonuç tahmini yapılır. İşlem sonucu ile karşılaştırılan bu tahmin değeri, sisteme uygun değilse, sonuca göre yeni bir tahmin ile istenen sonuca ulaşana kadar süreç tekrarlanır. Monte Carlo ve TD algoritmaları bu öğrenme modeline örnek olarak verilebilir (Sutton & Barto, 2005).

Kümeleme algoritmaları

Makine öğrenmesinde önemli konular arasında olan kümeleme analizi, makine öğrenmesinde bir denetimsiz öğrenme yaklaşımıdır. Veri seti içerisindeki ilişkiler ve benzerlikler ile gruplandırmalar oluşturmak için kullanılan kümeleme algoritmaları psikoloji, biyoloji, örüntü tanıma, istatistik, makine öğrenmesi, veri madenciliği gibi birçok alanda kullanılmaktadır (Tan, Steinbach & Kumar, 2013: 487). Verilerin gruplandırılması için kullanılan kümeleme algoritmaları 100'ün üzerindedir (Kaushik, 2016). Farklı yaklaşımlara göre birçok gruplandırılmanın yapıldığı algoritmalar literatürde genellikle hiyerarşik (hierarchical) ve bölücü (partitional) algoritmalar olmak üzere iki grup altında toplanmaktadır. Bölücü algoritmalar içerisinde yer alan ve merkezi kümeleme algoritmaları olarak isimlendirilen K-ortalamlar, K-medoid ve bulanık C-ortalamlar algoritmaları literatürde yaygın olarak yer almaktadır. K-ortalamlar algoritması, hızlı olması ve işlem mantığının basit olması sebebiyle uygulamalarda en fazla tercih edilen algoritmadır (Işık & Çamurcu, 2007; Jain, 2010; Han, Kamber & Pei, 2012; Kaushik, 2016).

K-ortalamlar algoritması 1967 yılında MacQueen tarafından yazılan makalede anlatılmıştır (MacQueen, 1967). Daha önce Ball ve Wallace tarafından gerçekleştirilen veri analizi çalışmalarında benzer yaklaşımlar kullanılsa da literatüre K-ortalamlar algoritmasını kazandıran MacQueen kabul edilmektedir (Tan, Steinbach & Kumar, 2013: 555). Veriler hakkında ön bilgi veya etiket olmaksızın veriler arasındaki benzerlikler ile gruplandırmalar yapabilen algoritma, k tane küme oluşturabilmek için her bir küme için seçilen bir merkeze en yakın üyeleri belirlemeyi amaçlamaktadır. Üyeler arası uzaklıkların belirlenmesinde Öklid, Manhattan, kosinüs gibi uzaklık ölçütleri kullanılmakta, kümedeki üye uzaklık

değerlerinin ortalaması küme merkezi kabul edilmekte ve bu nokta kümeyi temsil etmektedir (Han, Kamber & Pei, 2012). Her bir küme, üyelerin merkeze olan uzaklıkları ile oluşturulduğu için veriler sadece bir kümede yer almakta, bir üye başka bir kümeye ait olamamaktadır.

K-ortalamalar algoritması, uygulaması basit ve az miktardaki verilerde işlem hızı yüksek bir algoritma olmasına karşın bazı eksiklikleri de bulunmaktadır. K-ortalamalar algoritmasının temel problemi başlangıç kümesinin uygun seçilemediği durumlarda etkin sonuçlara ulaşılamamasıdır (Jain, Murty & Flynn, 1999: 18). Sonraki süreçler, seçilen başlangıç kümesine göre yürütüleceği ve sonucu direk olarak etkileyeceği için, bu kümenin seçimi oldukça önemlidir. Bir diğer eksiklik, algoritma başlatılmadan önce küme sayısı (k) parametresinin belirlenme zorunluluğudur. K değerinin analiz öncesi belirlenmesi veri setinden oluşabilecek küme sayısında belirsizlik olduğunda sorun teşkil etmektedir. Diğer taraftan küçük k değerlerinde oldukça hızlı çalışan algoritma, k değeri arttıkça yavaşlamaktadır. K-ortalamalar algoritmasının bir diğer eksikliği aykırı değerlerin kümelenmesinde başarısız olmasıdır. Ayrıca, küme yoğunluğu ve küme büyüklüklerinin farklı olması K-ortalamalar algoritmasının etkinliğini azaltmaktadır.

K-ortalamalar algoritmasındaki eksiklikleri gidermek için merkezi kümeleme yaklaşımı kullanan farklı algoritmalar geliştirilmiştir. K-medoids, veri setindeki aykırı değerlerde başarısız olan K-ortalamalar algoritmasının bu eksikliğini ortadan kaldırmak için geliştirilmiştir. Kauffman ve Rousseeuw tarafından geliştirilen algoritma, ilk merkezlerin belirlenmesi bağlamında K-ortalamalar algoritmasından ayrılmaktadır (Kauffman & Rousseeuw, 1990). K-medoids'te veri seti içerisindeki verilerden rastgele merkezler seçilir ve daha sonra merkeze en yakın olan veri ile değiştirilir. Bu işlem merkez değişmeyene kadar sürdürülür. K-medoids algoritmasının önemli sorunu, tekrar sayısının artması, böylece kümeleme işleminin uzun zaman almasıdır. K değeri yine parametre olarak kümeleme öncesi belirlenmektedir. Bu durum, oluşacak küme sayısı belirsizliğine yine çözüm bulunamadığı anlamına gelmektedir (Wentian, Qingju, Sheng & En, 2013).

Merkezi kümeleme yaklaşımında geliştirilen bir başka algoritma olan C-ortalamalar, bulanık mantık tabanlı bir algoritmadır. Bu sebeple bulanık C-ortalamalar olarak isimlendirilen algoritma fikir olarak Dunn tarafından 1973 yılında ortaya atılmış, 1981 yılında Bezdek tarafından geliştirilmiş ve literatüre kazandırılmıştır (Höppner, Klawonn, Kruse, & Runkler, 2000). Bir verinin sadece bir kümeye ait olamayacağı, aynı anda bir başka kümeye üyeliğinin mümkün olabileceği düşüncesiyle geliştirilen bir algoritmadır. Verinin merkezlere olan uzaklıkları $[0,1]$ arasındaki reel değerlerle belirlenir, verinin merkezlere olan uzaklık değerleri toplamı 1'e eşittir. Algoritma işleyişi K-ortalamalar algoritması ile benzer olmasına karşın verimliliğinin daha yüksek olduğu belirtilmektedir (Suganya & Shanthi, 2012). Ancak veri miktarındaki artışa daha fazla duyarlı olması ve analiz sürecinin uzun olması uygulamalarda tercih edilme oranını azaltmaktadır. Diğer taraftan, k değerinin yine kümeleme işlemi başlamadan önce belirlenmesi, oluşacak küme sayısı belirsizliğinin bu algoritma da söz konusu olduğu anlamına gelmektedir. Ayrıca K-ortalamalar algoritmasına benzer şekilde aykırı değerlere duyarlı olması, kümelenme etkinliğini azaltan bir diğer unsurdur.

Gri ilişkisel analiz

Gri Sistem Teorisi ilk olarak Prof. Ju Long Deng tarafından 1982 yılında "Control Problems of Grey System" adlı çalışmada, küçük örneklem ve eksik bilgi içeren problemlerin

çözümünde kullanılmak üzere önerilmiş bir yöntemdir (Deng, 1982). Söz konusu çalışma belirsizlik ve eksik bilgi durumunda sonuçlar üretebilen yeni bir yaklaşım olarak kabul görmüştür (Liu & Lin, 2006). Gri sistem teorisinde kesin bilinmeyen bilgi siyah, kesin bilinen bilgi beyaz ile temsil edilirken, bu iki uç nokta arasındaki kısmi bilgi gri ile temsil edilir (Liu, Forrest & Yang, 2012).

Gri İlişkisel Analiz (GİA) gri sistem teorisi kullanılarak geliştirilmiş, ilişkisel derecelendirme, sınıflandırma ve karar verme yöntemidir (Liu & Lin, 2006). Bir gri sistemdeki her bir faktörle kıyas yapılan referans faktör arasındaki ilişki derecesini belirlemeyi amaçlayan metotta, faktörler arası ilişki düzeyi gri ilişkisel derece olarak isimlendirilir (Yıldırım, 2015). GİA yöntemi, seçenekler arasından çeşitli kriterlere göre en uygun seçimin yapılabilmesi için kullanılan bir yöntemdir. Bu yönüyle çok kriterli karar verme aracı olan GİA, en iyi alternatiflerin seçilmesini amaçlayan bir metottur (Hinduja & Pandey, 2017).

Gri ilişkisel kümeleme

Gri İlişkisel Kümeleme (GİK), GİA gibi Gri Sistem Teorisinin konularından biridir. Karar kriterlerine göre alternatif gözlemler arasından benzer olanlarını tespit etme amacını taşıyan GİK, 1987 yılında Deng tarafından tasarlanan bir kümeleme yaklaşımıdır. GİK yaklaşımında kümeler belirli bir kurala göre gruplanmış nesnelerden oluştuğu için kümeler homojenliğe sahiptir (Wu, Lin, Peng & Huang, 2012). Basit algoritması ve esnek yapısı ile etkin bir kümeleme metodu olan GİK, yeniden hesaplama yapmaksızın nesnelerin kendi içerisinde ayrımını yapabilen bir yaklaşımdır. Bu nedenle işlem süresi açısından avantaj getirmektedir. Ayrıca küme sayısı analiz öncesinde değil, kümeleme gerçekleştirildikten sonra belirlenebileceği için yaygın kullanılan diğer kümeleme algoritmalarına göre daha gerçekçi bir yaklaşımdır.

GİK algoritmasının işleyiş yapısı GİA ile benzerlik göstermektedir. Alternatifler arasından kriterlere göre birbirine en yakın olanları belirlemek için gerçekleştirilecek adımlar şu şekildedir;

1. Karar matrisinin oluşturulması
2. Normalizasyon matrisinin oluşturulması
3. Mutlak farklar tablosunun oluşturulması
4. Gri ilişkisel katsayı matrisinin oluşturulması
5. Gri ilişkisel derecelerin hesaplanması
6. Gri benzerlik matrisinin oluşturulması
7. Küme elemanlarının belirlenmesi

Adım 1: GRA yönteminde ilk yapılacak işlem alternatif (a) ve karar kriterlerini (c) içeren karar matrisinin oluşturulmasıdır. (1) nolu denklemde gösterilen X karar matrisinde Xa her bir alternatifi gösterirken, Xac a alternatfinin c kriterinin değerini temsil etmektedir. Karar matrisindeki alternatif sayısı m ile ve kriter sayısı n ile gösterilmektedir.

$$\begin{aligned}
 a &= 1, 2, 3, \dots, m \\
 c &= 1, 2, 3, \dots, n \\
 X &= \begin{bmatrix} X_{11} & X_{12} & \dots & X_{1n} \\ X_{21} & X_{22} & \dots & X_{2n} \\ \vdots & \vdots & \dots & \vdots \\ X_{m1} & X_{m2} & \dots & X_{mn} \end{bmatrix} \quad (1)
 \end{aligned}$$

Adım 2: Bir serideki verilerin küçük aralıklara indirgenmesi işlemine normalizasyon adı verilmektedir. Normalizasyon işlemi sonrası serilerin birbirleriyle karşılaştırılması mümkün hale gelmektedir. Normalizasyon yöntemi, serideki kriter değerlerinin maksimum veya minimum olmasına göre değişmektedir. Karar için seri değerinin maksimum olması (utility based) olumlu katkı sağlıyorsa (2), minimum olması (cost based) olumlu katkı sağlıyorsa (3) nolu denklem kullanılır. Seri değerinin istenen bir optimum değerinin (optimal based) karara olumlu katkısı varsa (4) nolu denklem kullanılarak normalizasyon gerçekleştirilir ve (5) nolu denklemdeki normalizasyon matrisi elde edilir.

$$X_a^* = \frac{X_{ac} - \min X_{ac}}{\max X_{ac} - \min X_{ac}} \quad (2)$$

$$X_a^* = \frac{\max X_{ac} - X_{ac}}{\max X_{ac} - \min X_{ac}} \quad (3)$$

$$X_a^* = \frac{|X_{ac} - X_{0c}|}{\max X_{ac} - X_{0c}} \quad (4)$$

$$X^* = \begin{bmatrix} X_{11}^* & X_{12}^* & \dots & X_{1n}^* \\ X_{21}^* & X_{22}^* & \dots & X_{2n}^* \\ \vdots & \vdots & \dots & \vdots \\ X_{m1}^* & X_{m2}^* & \dots & X_{mn}^* \end{bmatrix} \quad (5)$$

Adım 3: Referans serisi ile her bir alternatif serinin mutlak farkının alındığı (6) nolu denklem yardımıyla (7) nolu matristeki mutlak farklar tablosu oluşturulur. Denklem 6'da yer alan $i, j \in a = (1, 2, \dots, m)$ ve X_{ic}^* referans serisini, X_{jc}^* farkı alınacak seriyi c ise kriteri temsil etmektedir.

$$\Delta_{jc} = |X_{ic}^* - X_{jc}^*| \quad (6)$$

$$\Delta_j = \begin{bmatrix} \Delta_{11} & \Delta_{12} & \dots & \Delta_{1n} \\ \Delta_{21} & \Delta_{22} & \dots & \Delta_{2n} \\ \vdots & \vdots & \dots & \vdots \\ \Delta_{m1} & \Delta_{m2} & \dots & \Delta_{mn} \end{bmatrix} \quad (7)$$

Adım 4: Karar matrisinin, karşılaştırma matrisine olan yakınlık derecesini gösteren Gri ilişkisel katsayıları hesaplayabilmek için mutlak değer tablosu değerleri kullanılır. Gri ilişkisel katsayı (10) nolu denklem yardımıyla hesaplanır.

$$\Delta_{max} = \max_a \max_c \Delta_{ac} \quad (8)$$

$$\Delta_{min} = \min_a \min_c \Delta_{ac} \quad (9)$$

$$\gamma_{ac} = \frac{\Delta_{min} + \rho \Delta_{max}}{\Delta_{ac} + \rho \Delta_{max}} \quad (10)$$

ρ değeri ayırım katsayısı olarak isimlendirilir ve $\rho \in [0,1]$. ρ değeri katsayı aralığının genişletme veya daraltılması için kullanılır (Hinduja & Pandey, 2017). Literatürdeki çalışmalarda ρ değerinin genellikle 0,5 alındığı ifade edilmektedir (Ertuğrul, Öztaş, Özçil & Öztaş, 2016). Ancak analiz edilen veriler arasındaki farkların fazla olması durumunda ρ değeri 0'a yakın olacak biçimde belirlenmelidir (Yıldırım, 2015).

Adım 5: Gri ilişkisel derecelerin hesaplanmasında iki seçenek söz konusudur. Eğer karar için kullanılacak kriterler eşit öneme sahipse (11), diğer durumda (12) nolu denklem kullanılır (Yıldırım, 2015). (13) nolu denklemde yer alan w_{ac} değeri karar kriterinin ağırlığını gösterir.

$$\delta_a = \frac{1}{n} \sum_{c=1}^n \gamma_{ac} \quad (11)$$

$$\delta_a = \sum_{c=1}^n w_{ac} \gamma_{ac} \quad (12)$$

Adım 6: Alternatifler arasındaki benzerlikleri belirlemek için kullanılacak olan gri benzerlik matrisi G , gri ilişkisel derece matrisi temel alınarak (13) nolu denkleme göre oluşturulur. Denklem incelendiğinde ilişkisel derecelerin çapraz değerlerinin ortalaması olduğu görülecektir. Eğer verilere normalizasyon işlemi uygulanmışsa, çapraz değerler birbirine eşit olacağından dolayı ortalama almaya gerek duyulmaz.

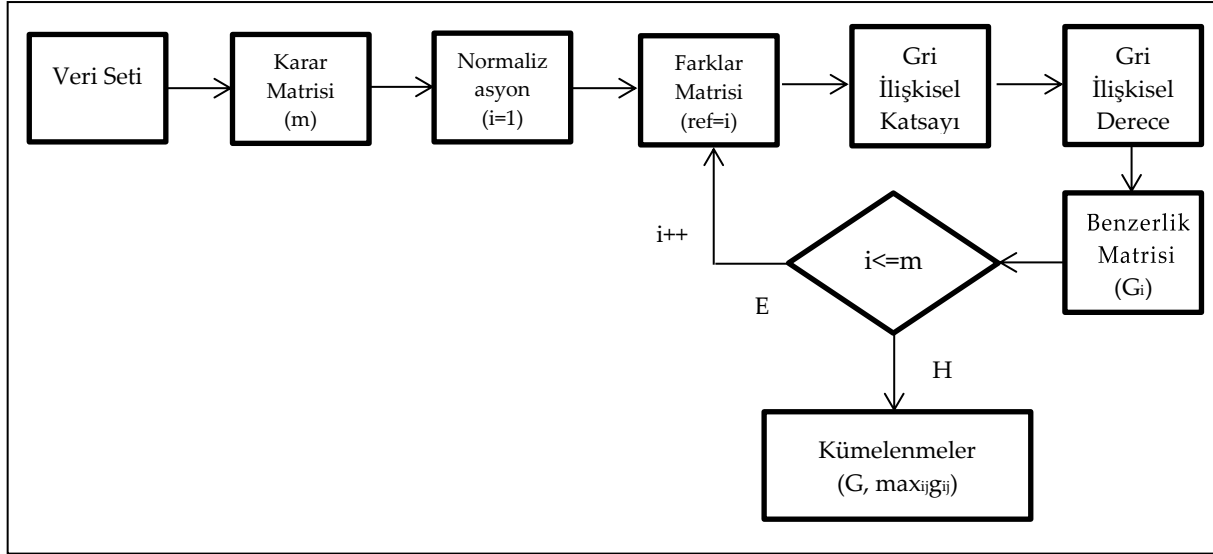
$$G = [g_{ij}] = (\delta_{ij} + \delta_{ji})/2 \quad (13)$$

Adım 7: G benzerlik matrisi içerisindeki en yüksek değere sahip olanlar ile küme belirlenir. $\max_{ij}(g_{ij})$ ile ifade edilebilecek bu işlem ile kriterler arasından birbirine en fazla benzeyen alternatifler belirlenerek küme oluşturulur. Oluşan küme tüm veri setini içeriyorsa kümeleme işlemi tamamlanmıştır. Aksi halde, bir sonraki alternatif referans olarak belirlenir ve işlemler 2.adımdan itibaren tekrarlanarak, kümeleme işlemine devam edilir.

Araştırma

Araştırma amacı ve analiz yöntemi

Çalışmanın amacı müşteri bağlılığı oluşabilecek müşteri gruplarının tespit edilmesidir. Benzer özelliklere sahip olan verilerin gruplanması kümeleme analizleri ile gerçekleştirilebilir. Literatürdeki kümeleme analizlerinde, oluşacak küme sayısının önceden belirlenmesine karşın, bu çalışmada analiz öncesi küme sayısının belirlenmeden kümelenmelerin gerçekleştirilebileceği amaçlanmaktadır. Bu bağlamda çalışmada müşterilerin satın alma sayıları, toplam satın alma tutarı, ortalama satın alma tutarı, sitede giriş sayısı, şikayet sayısı ve ürün geri iade sayısı olmak üzere altı değişkene göre kümeleme analizinin gerçekleştirilmektedir. Araştırmada analiz yöntemi olarak GRA yaklaşımı kullanılmış ve gri ilişkisel tabanlı GİK uygulanmıştır. Elde edilen veri setinin kümeleme süreci Şekil 1'de gösterilmiştir.



Şekil 1: GİK analiz süreci (Kaynak: Wu, Lin, Peng & Huang, 2012'den türetilmiştir)

Araştırma veri seti

E-ticaret sektöründe faaliyet gösteren bir firmadan satışlar ve müşteriler ile ilgili son bir seneye ait veriler temin edilmiştir. Gizlilik ilkelerine uygun olarak firmadan alınan 3 tablo içerisinde müşterilerin özel bilgileri bulunmamaktadır. Müşteriler ile yaptıkları işlemler müşteri id'leri ile gerçekleştirilmiştir. Bir yıl içerisinde gerçekleştirilen toplam satın alma işlem sayısı 2845 olup bu süre içerisinde işlem yapan müşteri sayısı 986 dır. İki veya daha az satın alma işlemi gerçekleştiren müşteri sayısı 730'dur.

GİK sürecini basit biçimde anlatabilmek amacıyla rastgele 12 müşteri belirlenmiş ve bu müşterilerin bir yıl içerisindeki satın alma sayıları, toplam satın alma tutarları, ortalama satın alma tutarları, siteye giriş sayıları, şikâyet sayıları ve ürün geri iade sayıları ile veri seti oluşturulmuştur. Analiz için kullanılacak veri seti Tablo 1'de verilmektedir.

Tablo 1: Veri seti

Müşteri	İşlem sayısı	Toplam tutar	Ortalama tutar	Site login	Kullanıcı şikâyet	Geri iade
m1	7	1.221,57	174,51	45	0	0
m2	3	587,94	195,98	56	1	0
m3	22	2.971,40	135,06	85	0	1
m4	14	8.829,63	630,69	123	1	0
m5	21	15.566,57	741,27	148	2	0
m6	18	1.681,23	93,40	72	0	2
m7	4	168,95	42,24	65	0	1
m8	6	2.035,76	339,29	120	0	0
m9	4	1.230,97	307,74	38	1	0
m10	8	516,23	64,53	93	0	3
m11	5	1.455,70	291,14	24	0	0
m12	4	202,50	50,63	32	1	0

Tablo 1’de en fazla alış veriş yapan müşterinin m3 olduğu görülmektedir. E-ticaret sitesini en fazla kullanan müşteri ise 148 defa sisteme giriş yapan m5’dir. Satın aldığı ürünlerden en fazla geri iade gerçekleştiren m10 kullanıcısı 3 kez, m6 ise 2 kez ürün iadesi yapmıştır.

Kümeleme bulguları

Tablo 1 kullanılarak gerçekleştirilen GİK analizinde ilk olarak veri setinin normalizasyonu gerçekleştirilmiştir. Veriler arasında değer farklılıklarının yüksek olması halinde analiz sonuçları olumsuz etkilenecektir. Ayrıca veri setinde, maksimum veya minimum olması istenen, kriterler söz konusudur. Örneğin müşteri bağlılığı çerçevesinde işlem sayısı, toplam tutar, ortalama tutar ve siteye giriş sayısı değerlerinin maksimum olması istenen bir durum iken şikâyet sayısı ve geri iade sayısının minimum olması istenir. Bu bağlamda oluşturulan normalizasyon matrisi Tablo 2’de sunulmuştur.

Tablo 2: Normalizasyon matrisi

Müşteri	İşlem sayısı	Toplam tutar	Ortalama tutar	Site login	Kullanıcı şikâyet	Geri iade
m1	0,2105	0,0684	0,1892	0,1694	1	1
m2	0	0,0272	0,2199	0,2581	0,5000	1
m3	1	0,1820	0,1328	0,4919	1	0,6667
m4	0,5789	0,5625	0,8418	0,7984	0,5000	1
m5	0,9474	1	1	1	0	1
m6	0,7895	0,0982	0,0732	0,3871	1	0,3333
m7	0,0526	0	0	0,3306	1	0,6667
m8	0,1579	0,1212	0,4250	0,7742	1	1
m9	0,0526	0,0690	0,3798	0,1129	0,5000	1
m10	0,2632	0,0226	0,0319	0,5565	1	0
m11	0,1053	0,0836	0,3561	0	1	1
m12	0,0526	0,0022	0,0120	0,0645	0,5000	1

Normalizasyon matrisindeki değerler kullanılarak mutlak farklar matrisi elde edilir. Farkları hesaplamak için bir referans seriye ihtiyaç vardır. Bu seri kullanıcı tarafından istenen değerlere göre oluşturulabildiği gibi, matris içerisindeki bir seri de seçilebilir. Analizde matris içerisindeki tüm seriler sırayla referans serisi olarak belirlenecektir. m1 serisi referans olarak seçilerek hazırlanan mutlak farklar matrisi Tablo 3’te verilmiştir.

Tablo 3: Mutlak farklar matrisi (m1 referans seri iken)

Müşteri	İşlem sayısı	Toplam tutar	Ortalama tutar	Site login	Kullanıcı şikâyet	Geri iade
m1	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000
m2	0,2105	0,0412	0,0307	0,0887	0,5000	0,0000
m3	0,7895	0,1136	0,0564	0,3226	0,0000	0,3333
m4	0,3684	0,4941	0,6526	0,6290	0,5000	0,0000
m5	0,7368	0,9316	0,8108	0,8306	1,0000	0,0000
m6	0,5789	0,0299	0,1160	0,2177	0,0000	0,6667
m7	0,1579	0,0684	0,1892	0,1613	0,0000	0,3333
m8	0,0526	0,0529	0,2357	0,6048	0,0000	0,0000

m9	0,1579	0,0006	0,1906	0,0565	0,5000	0,0000
m10	0,0526	0,0458	0,1573	0,3871	0,0000	1,0000
m11	0,1053	0,0152	0,1668	0,1694	0,0000	0,0000
m12	0,1579	0,0662	0,1772	0,1048	0,5000	0,0000

m1 referans seri iken mutlak farkların hesaplandığı tabloda m1 referans serisi, m2 mutlak farkın alınacağı seri ve işlem sayısı kriteri farkın alınacağı değer ($c=1$) iken denklem (6) yardımıyla $\Delta_{21} = |X_{11}^* - X_{21}^*| = |0,2105 - 0| = 0,2105$ şeklinde hesaplanır.

Mutlak farklar matrisi yardımıyla Gri ilişkisel katsayıları belirlemek için öncelikle farklar matrisindeki en büyük ve en küçük değerlerin belirlenmesi gerekir. Tablo 3'te yer alan değerlerden en yüksek değer $\Delta_{max}=1$ ve en küçük değerin $\Delta_{min}=0$ olduğu görülmektedir. Tablo 1'de görüldüğü üzere, veriler arasındaki farklılıklar yüksektir. Bu nedenle ayırım katsayısı $\rho = 0,1$ olarak alınmıştır. Denklem 10 kullanılarak hesaplanan gri ilişkisel katsayılar Tablo 4'te görülmektedir.

Tablo 4: Gri ilişkisel katsayılar (m1 referans serisi iken)

Müşteri	İşlem sayısı	Toplam tutar	Ortalama tutar	Site login	Kullanıcı şikâyet	Geri iade
m1	1,0000	1,0000	1,0000	1,0000	1,0000	1,0000
m2	0,3220	0,7085	0,7650	0,5299	0,1667	1,0000
m3	0,1124	0,4681	0,6393	0,2366	1,0000	0,2308
m4	0,2135	0,1683	0,1329	0,1372	0,1667	1,0000
m5	0,1195	0,0969	0,1098	0,1075	0,0909	1,0000
m6	0,1473	0,7701	0,4629	0,3147	1,0000	0,1304
m7	0,3878	0,5940	0,3458	0,3827	1,0000	0,2308
m8	0,6552	0,6541	0,2979	0,1419	1,0000	1,0000
m9	0,3878	0,9939	0,3441	0,6392	0,1667	1,0000
m10	0,6552	0,6858	0,3886	0,2053	1,0000	0,0909
m11	0,4872	0,8680	0,3747	0,3713	1,0000	1,0000
m12	0,3878	0,6017	0,3607	0,4882	0,1667	1,0000

Her bir alternatif seri için hesaplanacak Gri ilişkisel dereceler, her bir alternatifteki kriter değerlerinin aritmetik ortalamasıdır. m1 referans seri iken hesaplanan gri ilişkisel dereceler Tablo 5'te verilmiştir. Bu değerler tüm alternatiflerin referans serisine olan ilişki derecesini göstermektedir. m1 de yer alan değerin 1 olması bu sebeptir. Bu değerlere göre m1 ile en yüksek ilişkili olan alternatifin 0,6835 ile m11 olduğu görülmektedir.

Tablo 5: Gri ilişkisel dereceler (m1 referans seri iken)

Müşteri	Gri ilişkisel dereceler
m1	1,0000
m2	0,5820
m3	0,4479
m4	0,3031
m5	0,2541
m6	0,4709
m7	0,4902
m8	0,6248

m9	0,5886
m10	0,5043
m11	0,6835
m12	0,5008

Tablo 5’te verilen gri ilişkisel derecelerin hesaplama yöntemi tüm alternatiflerin sırasıyla referans seri olarak belirlenmesi ile tekrarlanır. Her referans seri ile mutlak farklar, katsayılar ve derece değerleri hesaplanarak Gri ilişkisel dereceler matrisi oluşturulur. Tablo 6’da verilen bu matris tüm alternatiflerin birbirleriyle olan gri ilişkisel derece değerleri ve Tablo 7’de gri ilişkisel benzerlik değerleri göstermektedir.

Tablo 6: Gri ilişkisel dereceler (tüm alternatifler)

Müşteri	m1	m2	m3	m4	m5	m6	m7	m8	m9	m10	m11	m12
m1	1,0000	0,5820	0,4479	0,3031	0,2541	0,4709	0,4902	0,6248	0,5886	0,5043	0,6835	0,5008
m2	0,5820	1,0000	0,2858	0,4332	0,2646	0,3060	0,4551	0,4267	0,6922	0,3478	0,4994	0,6867
m3	0,4479	0,2858	1,0000	0,1945	0,2256	0,5353	0,5437	0,4126	0,2432	0,4568	0,3856	0,2488
m4	0,3031	0,4332	0,1945	1,0000	0,3808	0,1845	0,1651	0,4237	0,4389	0,1761	0,2993	0,4231
m5	0,2541	0,2646	0,2256	0,3808	1,0000	0,1578	0,1223	0,2934	0,2674	0,1133	0,2535	0,2578
m6	0,4709	0,3060	0,5353	0,1845	0,1578	1,0000	0,5119	0,4178	0,2839	0,5064	0,4328	0,2973
m7	0,4902	0,4551	0,5437	0,1651	0,1223	0,5119	1,0000	0,4241	0,4187	0,5556	0,4804	0,5904
m8	0,6248	0,4267	0,4126	0,4237	0,2934	0,4178	0,4241	1,0000	0,5218	0,4332	0,6813	0,4048
m9	0,5886	0,6922	0,2432	0,4389	0,2674	0,2839	0,4187	0,5218	1,0000	0,2783	0,6620	0,7479
m10	0,5043	0,3478	0,4568	0,1761	0,1133	0,5064	0,5556	0,4332	0,2783	1,0000	0,4146	0,4022
m11	0,6835	0,4994	0,3856	0,2993	0,2535	0,4328	0,4804	0,6813	0,6620	0,4146	1,0000	0,5344
m12	0,5008	0,6867	0,2488	0,4231	0,2578	0,2973	0,5904	0,4048	0,7479	0,4022	0,5344	1,0000

Tablo 7: Gri ilişkisel benzerlik matrisi

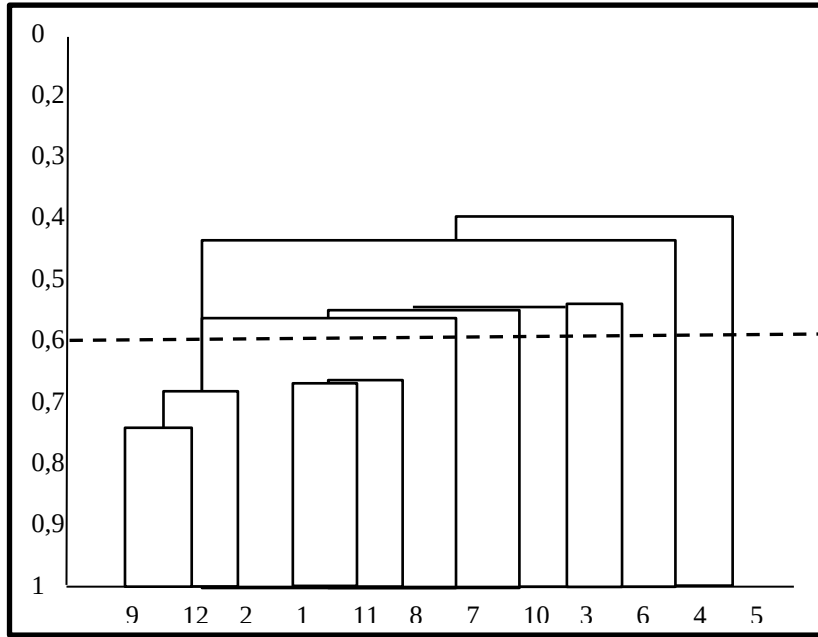
Müşteri	m1	m2	m3	m4	m5	m6	m7	m8	m9	m10	m11	m12
m1	1,0000											
m2	0,5820	1,0000										
m3	0,4479	0,2858	1,0000									
m4	0,3031	0,4332	0,1945	1,0000								
m5	0,2541	0,2646	0,2256	0,3808	1,0000							
m6	0,4709	0,3060	0,5353	0,1845	0,1578	1,0000						
m7	0,4902	0,4551	0,5437	0,1651	0,1223	0,5119	1,0000					
m8	0,6248	0,4267	0,4126	0,4237	0,2934	0,4178	0,4241	1,0000				
m9	0,5886	0,6922	0,2432	0,4389	0,2674	0,2839	0,4187	0,5218	1,0000			
m10	0,5043	0,3478	0,4568	0,1761	0,1133	0,5064	0,5556	0,4332	0,2783	1,0000		
m11	0,6835	0,4994	0,3856	0,2993	0,2535	0,4328	0,4804	0,6813	0,6620	0,4146	1,0000	
m12	0,5008	0,6867	0,2488	0,4231	0,2578	0,2973	0,5904	0,4048	0,7479	0,4022	0,5344	1,0000

Analizdeki son adım benzerlik değerlerine göre kümelerin tespit edilmesidir. Yüksek değerler benzerliğin yüksek olduğunu, düşük değerler ise benzerliğin az olduğunu ifade eder. Buna göre en yüksek benzerlik m9 ile m12 arasında, en yüksek farklılık ise m5 ile m7 arasındadır. Kümelenmeler en yüksek ilişki derecesine sahip olan alternatif eşleştirmeleri ile gerçekleştirilir. Buna göre Tablo 7’deki m9 ile m12 benzerlik değeri 0,7479 ile en yüksek benzerlik değerine sahiptir. Bunun anlamı belirlenen kriterlere göre birbirine benzerliği en fazla olan müşteriler m9 ve m12 dir. Birinci küme tespit edildikten sonra ikinci en büyük

benzerlik değerine göre ikinci küme tespit edilir. Bu işlem eşleme yapılmayan alternatif (kullanıcı) kalmayana kadar devam eder. Böylece Tablo 1’de ki verilere göre kullanıcıların GİK analizi ile kümelenmeyi belirleyen en büyük gri ilişkisel benzerlik değerleri Tablo 8’de ve kümelenmenin ağaç gösterimi Şekil 2’de verilmiştir.

Tablo 8: En yüksek benzerlik değerleri (Benzerlik eşik değeri 0,6)

Kümeler	Müşteriler	En yüksek benzerlikler
küme1	9 - 12	0,7479
	2	0,6922
küme2	1 - 11	0,6835
	8	0,6813
Küme3	7	0,5904
Küme4	10	0,5556
Küme5	3	0,5437
Küme6	6	0,5353
Küme7	4	0,4322
Küme8	5	0,3808



Şekil 2: Kümelenme ağaç gösterimi

GİK ile oluşan kümelenmeler ile analizin tekrarlanmasına gerek kalmadan küme sayısına karar verilebilir. Bunun için karar verici öncelikle müşterilerin ne oranda benzerliklerine göre kümelenme oluşturulacağını belirlemelidir. Eğer müşteriler arasında %60 lık bir benzerlik ile kümelenmenin gerçekleştirileceği kabul edilirse, 1 ile 0,6 arasındaki birleşimler, kümeleri oluşturacaktır. Şekil 2 üzerinde belirlenen 0,6 sınırından yüksek olan benzerlik değerleri ile oluşan kümeler 8 tanedir. Bu kümeler birinci küme 9, 12 ve 2, ikinci küme 1, 11 ve 8, üçüncü küme 7, dördüncü küme 10, beşinci küme 3, altıncı küme 6, yedinci küme 4 ve sekizinci küme 5 ile oluşmuştur.

Analiz bulgularının değerlendirilmesi

GİA tabanlı kümeleme metodunun anlatıldığı çalışmada, öğeler arasındaki ilişki düzeylerini gösteren gri ilişkisel dereceler kullanılarak verilerin kümelenmesi gerçekleştirilmiştir. Literatürde yaygın olarak kullanılan merkezi referans algoritmalarının başlıca özelliği oluşacak küme sayısının önceden belirlenmesidir. GİK ile uygulanan yöntem ile küme sayısına analiz öncesi değil, analiz sonrası karar verilmektedir.

GİK metodunda analiz edilecek veri setinin oluşturulmasından sonra ilk işlem normalizasyon uygulanmasıdır. Veri setinde yer alan değerler arasında farklar büyükse seriler arasında yapılacak karşılaştırmaların verimli sonuçlar vermesi uç değerlerin yüksek olmasından dolayı zordur. Birbirine yakın değerlerde normalizasyon sorun oluşturmamakla birlikte farklar yüksek ise normalizasyon uygulanması sağlıklı sonuçlar verecektir.

Normalizasyon sonrası analiz için hazır hale gelen veri setinde referans küme merkezi seçme işlemi GİK'te yapılmaz. Veriler kendi içerisinde sırasıyla serilerin referans kabul edilmesiyle gerçekleştirilir. Mutlak farkların alınması, maksimum ve minimum değerlerin belirlenmesi, gri ilişkisel katsayıların hesaplanması ve gri ilişkisel derecelerinin hesaplanması bu tekrarlar sürecinde gerçekleştirilecek işlemlerdir. Her bir tekrar yapıldıktan sonra oluşan dereceler bir matris içerisinde birleştirilerek gri benzerlik matrisi elde edilir. Bu süreç analiz öncesi merkez noktanın belirlenmesine ihtiyaç duymaması, veri seti dışında bir referans serinin kullanılmaması ve veri setini tekrar tarama ihtiyacı olmaması noktaları açısından diğer kümeleme yöntemlerinden ayrılmaktadır.

Gri ilişkisel matrisinde yer alan değerler aynı kullanıcılar için iki farklı değer alabilir. Örneğin 2. ve 3. kullanıcıların δ_{23} ve δ_{32} olmak üzere iki farklı gri ilişkisel derece değerine sahiptir. Bu farklılığın tek değere indirilmesi için literatürde iki değer aritmetik ortalaması alınır (Wu, Lin, Peng & Huang, 2012: 7249). Ancak, bu çalışma ortalama almak yerine, mutlak farkların alınmasından önce verilere normalizasyon işlemini önermektedir. Normalizasyon işlemi uygulandığı takdirde söz konusu değer farklılıkları ortadan kalkacaktır. Bu çalışmadaki analiz, normalizasyon yapılmadan tekrarlanmış ve aynı üyelerin iki farklı gri ilişkisel derece değeri olduğu görülmüştür.

Analiz ile elde edilen kümelenme düzeyine, Şekil 2'de verilen ağaç yapısında daha açık görülebileceği üzere, analiz sonrası karar verilebilir. Şekil 2'de eğer bir eşik değer belirlenmezse oluşan küme sayısının 10 tane olduğu, sadece m9-m12 ile m1-m11 kullanıcılarının beraber küme oluşturduğu, diğer kullanıcıların ise tek başına bir kümede yer aldığı görülmektedir. Ancak karar verici, belirli bir benzerlik oranı ile oluşacak kümeler karar verebilir. Bu yönüyle GİK, analizi yenilemeden alternatif kümelenmelere karar verme imkânı sunmaktadır. Bu yönüyle çalışma, özellikle gerçek zamanlı kümelenme uygulamalarına önemli katkı sağlayacaktır.

Analiz sonucunda oluşan kümeler, müşterilerin bağlılık düzeyleri hakkında bilgi vermemektedir. Örneğin en yüksek benzerliğe sahip olan m9 ve m12 kullanıcılarının en yüksek bağlılık düzeyine sahip olduğu söylenemez. Bu çerçevede çalışma küme1'e firmanın nasıl bir politika uygulayacağı konusunda bilgi vermez. Ancak araştırma sonucu, firmanın m9 kullanıcılarına gerçekleştireceği bir uygulamanın m12 kullanıcılarına da uygulanabileceğini ortaya koymaktadır. Başka bir deyişle m9'un bağlılık düzeyi ile m12'nin bağlılık düzeyi

benzer seviyededir. Bu bağlamda çalışma, kümelerin bağlılık düzeylerinin tespit edilmesi yönünde genişletilebilir. Ayrıca analiz sonuçlarının diğer kümelenme analizleri ile karşılaştırılması gelecek çalışmalarda yapılarak, söz konusu yöntemin verimlilik düzeyi belirlenebilir. Bu sayede diğer kümelenme yaklaşımları ile analiz sonuçları kapsamında değerlendirmeler yapılabilir.

SONUÇ

Çalışma Gri İlişkisel Analiz (GİA) yaklaşımının kullanılarak Gri İlişkisel Kümeleme (GİK) olarak isimlendirilen gri analiz tabanlı bir kümeleme yöntemini ortaya koymaktadır. Eksik bilgi ve belirsizlikler anında kullanılan en uygun yöntem olan GİA'nın, müşteri bağlılığı kümelenme uygulamalarında kullanılabileceği bu çalışmayla ortaya konulmaktadır. Klasik kümeleme algoritmalarında oluşacak küme sayısı analiz öncesi belirlenmektedir. Ancak bir veri setinde kaç tane küme oluşacağı analiz sonrası karar verilebilecek bir konudur. Başka bir deyişle analiz öncesi oluşacak küme sayısı, küme üyeleri ve küme merkezi belirsizdir. Bu sebeple GİK, kümeleme işleminde kullanılacak en uygun seçenektir. Çalışmada, GİK ile e-ticaret kullanıcılarının bağlılık kümelenme analizinin küme sayısı varsayımı olmaksızın ve küme merkezleri belirlenmeksizin gerçekleştirilebileceği ortaya konulmuştur.

Müşteri bağlılıklarının kümelenme analizinin yapıldığı çalışmada veriler bir e-ticaret sitesinden temin edilmiştir. Müşteri bağlılığının değerlendirilmesi için literatürde kullanılan kullanıcıların toplam işlem sayısı, toplam işlem tutarı ve ortalama işlem tutarı değişkenlerine ek olarak bu çalışmada kullanıcıların siteye giriş sayıları, şikâyet sayıları ve ürün geri iade sayıları değişkenleri de analize eklenmiştir. Böylece kullanıcıların bağlılık kümelenmeleri altı değişkenle değerlendirilerek gerçekleştirilmiştir. Diğer kümelenme algoritmalarının aksine GİK analizi verilerin tekrar taranmasına ihtiyaç duymaz. Böylece kümelenme analizlerindeki sorunlardan biri olan uzun işlem süresi problemi ortadan kalkmaktadır. Bu sayede gerçek zamanlı kümelenme uygulamaları GİK ile rahatlıkla yapılabilir.

KAYNAKÇA

- Afrin, F., Al-Amin M., & Tabassum, M. (2015). Comparative Performance Of Using PCA With KMeans And Fuzzy C Means Clustering For Customer Segmentation. *International Journal of Scientific & Technology Research*, 4(10), 70-74.
- Bafghi, E. P. (2017). Clustering of Customers Based on Shopping Behavior and Employing Genetic Algorithms. *Engineering, Technology & Applied Science Research*, 7(1), 1420-1424.
- Bishop, C. M. (2007). *Pattern Recognition and Machine Learning*. Springer.
- Brown, S. P., & Peterson, R. A. (1993). Antecedents and Consequences of Salesperson Job Satisfaction: Meta-analysis and Assessment of Causal Effects. *Journal of Marketing Research*, 12, 161-173.
- Brown, M., Pope, N., & Voges, K. (2003). Buying or Browsing? An Exploration of Shopping Orientations and Online Purchase Intention. *European Journal of Marketing*, 37(11/12), 1666-1684.
- Cheng, C. H., & Chen Y. S. (2009). Classifying the Segmentation of Customer Value via RFM Model and RS Theory. *Expert Systems with Applications*, 36(3), 4176-4184.
- Corstjens, M., & Lal, R. (2000). Building Store Loyalty Through Store Brands. *Journal of Marketing Research*, 37(3), 281-291.
- Işık, M., ve Çamurcu, A. Y. (2007). K-means, K-medoids ve Bulanık C-means Algoritmalarının Uygulamalı Olarak Performanslarının Tespiti. *İstanbul Ticaret Üniversitesi Fen Bilimleri Dergisi*, 6(11), 31-45.
- Davidson, I. (2002). Understanding K-means Non-hierarchical Clustering. SUNY Albany Technical Report 02-2.

- Deng, J. (1982). Control Problems of Grey Systems. *System and Control Letters*, 1(5), 288-294.
- Ellison, G., & Ellison, S. F. (2005). Lessons About Markets from the Internet. *The Journal of Economic Perspectives*, 19(2): 139-158.
- Ertuğrul, I., Öztaş, T., Özçil, A., & Öztaş, G. Z. (2016). Grey Relational Analysis Approach In Academic Performance Comparison Of University: A Case Study Of Turkish Universities. *European Scientific Journal*, June 2016 SPECIAL edition, 128-139.
- Fidan, H. (2014). *Asimterik Bilginin E-ticaret Üzerindeki Etkileri: Tüketici Güveni Üzerine Bir Uygulama*. Yayınlanmış doktora tezi, Süleyman Demirel Üniversitesi Sosyal Bilimler Enstitüsü, Isparta.
- Fidan, H. (2016). Measurement of the Intersectoral Digital Divide with the Gini Coefficients: Case Study Turkey and Lithuania. *Inzinerine Ekonomika-Engineering Economics*, 27(4), 439–451.
- Fidan, H., & Albeni, M. (2014). Asimetrik Bilginin E-Ticaret Üzerindeki Etkileri: Tüketicilerin Güven Eğilimleri Üzerine Bir Araştırma, *Süleyman Demirel Üniversitesi İktisadi ve İdari Bilimler Fakültesi Dergisi*, 19(2), 287-298.
- Gromov, G. (2012). Internet History with a Human Face. Retrieved November 15, 2017, from http://history-of-internet.com/history_of_internet.pdf.
- Han, J., Kamber, M., & Pei, J. (2012). *Data Mining Concepts and Techniques (Third Edition)*, USA: Morgan Kaufmann Publications.
- Hinduja, A., & Pandey, M. (2017). Multicriteria Recommender System for Life Insurance Plans based on Utility Theory. *Indian Journal of Science and Technology*, 10(14), 1-8, DOI: 10.17485/ijst/2017/v10i14/111376.
- Höppner, F., Klawonn, F., Kruse, R., & Runkler, T., (2000), *Fuzzy Cluster Analysis*, Chichester: John Wiley&Sons.
- Huang, X. ve Song, Z. (2014). Clustering Analysis on E-commerce Transaction Based on K-means Clustering. *Journal of Networks*, 9(2), 443-450.
- Ilsever, J., Cyr, D., & Parent, M. (2006). Extending Models of Flow and E-loyalty. *Journal of Information Science and Technology*, 3(4), 3–22.
- Islam M.A., Khadem, M. & Sayem, A. (2012). Service quality, customer satisfaction and customer loyalty analysis in Bangladesh apparel fashion retail: an empirical study, *International Journal of Fashion Design. Technology and Education*, 5(3), 213-224.
- Jain, A. K. (2010). Data clustering: 50 years beyond K-means. *Pattern Recognition Letters*, 31(8), 651-666.
- Jain, A., Murty, M., & Flynn, P. (1999). Data Clustering: A review, *ACM Comput. Surv*, 31(3), 264–323.
- Kalaiselvi, B. (2015). A Comprehensive Usage of Enhanced K-Medoid Clustering Algorithm in Banking Sector. *International Advanced Research Journal in Science Engineering and Technology*, 2(7), 102-105.
- Kaufman, L., & Rousseeuw, P. J. (1990). *Finding Groups in Data: An Introduction to Cluster Analysis*. Chichester: John Wiley and Sons.
- Kaushik, S. (2016). An Introduction to Clustering and Different Methods of Clustering. Retrieved November 21, 2017, from <https://www.analyticsvidhya.com/blog/2016/11/an-introduction-to-clustering-and-different-methods-of-clustering>.
- Kim K., & Ahn H. (2005). [Using a Clustering Genetic Algorithm to Support Customer Segmentation for Personalized Recommender Systems](#). In: Kim T.G. (eds) *Artificial Intelligence and Simulation* (pp. 409-415), Berlin, Heidelberg: Springer.
- Lee, E. J., & Overby, J. W. (2004). Creating Value for Online Shoppers: Implications for Satisfaction and Loyalty. *Journal of Consumer Satisfaction. Dissatisfaction and Complaining Behavior*, 17, 54–64.
- Liu, S., Forrest J., & Yang, Y. (2012). A Brief Introduction to Grey Systems Theory. *Grey Systems: Theory and Application*, 2(2) 89-104.
- Liu, S., & Lin, Y. (2006). *Grey Information Theory and Practical Applications*. NewYork, USA: Springer Science+Business Media.

- MacQueen, J. (1967). Some Methods for Classification and Analysis of Multivariate Observations. In *Proc. of the 5th Berkeley Symp. on Mathematical Statistics and Probability*, 281-297.
- Min, S. H. & Han I. (2005). Recommender Systems Using Support Vector Machines. In: Lowe D., Gaedke M. (Ed.) *Web Engineering. ICWE 2005. Lecture Notes in Computer Science*, (vol 3579, pp. 387-393). Berlin, Heidelberg: Springer.
- Niknam, T., & Amiri, B. (2010). An Efficient Hybrid Approach Based on PSO, ACO and K-means for Cluster Analysis. *Applied Soft Computing*, 10, 183-197.
- Park, H. S., & Jun, C. H. (2009). A Simple and Fast Algorithm for K-medoids Clustering. *Expert Systems with Applications*, 36, 3336-3341.
- Reichheld, F. (1995). Loyalty and the Renaissance of Marketing. *Marketing Management*, 2(4), 10-21.
- Reichheld, F. & Schefter, P. (2000). E-loyalty: your secret weapon on the Web. *Harvard Business Review*, July-August, 105-113.
- Rosen, S. (2001). Sticky Web site is Key to Success. *Communication World*, 18(3), 36-37.
- Spector, R. (2001). *Amazon.com Ve Yaratıcısı Jeff Bezos*. İstanbul: Scala Yayıncılık.
- Srinivasan, S. S., Anderson, R., & Ponnayolu, K. (2002). Customer Loyalty in E-commerce: An Exploration of its Antecedents and Consequences. *Journal of Retailing*, 78, 41-50.
- Suganya, R., & Shanthi, R. (2012). Fuzzy C- Means Algorithm- A Review. *International Journal of Scientific and Research Publications*, 2(11).
- Sutton, R. S., & Barto, A. G. (2005). *Reinforcement Learning: An Introduction*. London, England: The MIT Press.
- Tajunisha, S. (2010). Performance Analysis of K-means with Different Initialization Methods for High Dimensional Data. *International Journal of Artificial Intelligence & Applications*, 44-52.
- Tan, P. N., Steinbach, M., & Kumar, V. (2013). *Introduction to Data Mining*, USA: Pearson Education Limited.
- Toufaily, E., Ricard, L., ve Perrien, J. (2013). Customer loyalty to a commercial website: Descriptive meta-analysis of the empirical literature and proposal of an integrative model. *Journal of Business Research*, 66, 1436-1447.
- Uncles, M. D., Dowling, G. R., & Hammond, K. (2003). Customer Loyalty and Customer Loyalty Programs. *Journal of Consumer Marketing*, 20(4), 294-316.
- Verona, G., & Prandelli, E. (2002). A Dynamic Model of Customer Loyalty to Sustain Competitive Advantage on the Web. *European Management Journal*, 20(3), 299-309.
- Wentian, J., Qingju, G., Sheng, Z., & En, Z. (2013). Improved K-medoids Clustering Algorithm under Semantic Web. *Proceedings of the 2nd International Conference on Computer Science and Electronics Engineering (ICCSEE 2013)*, 731-733.
- Wu, W. H., Lin, C. T., Peng K. H. & Huang, C. C. (2012). Applying Hierarchical Grey Relation Clustering Analysis to Geographical Information Systems – A Case Study of the Hospitals in Taipei City. *Expert Systems with Applications*, 39, 7247-7254.
- Yıldırım, B. F. (2015). Gri İlişkisel Analiz.. In Yıldırım, B. F., & Önder, E. (Ed.), *Çok Kriterli Karar Verme Yöntemleri* (pp. 229-236). Bursa, Turkey: Dora Basım Yayın.
- Yuliari, N. P. P., Putra, K. G. D., & Rusjayanti, N. K. D. (2015). Customer Segmentation Through Fuzzy C-Means and Fuzzy RFM Method. *Journal of Theoretical and Applied Information Technology*, 78(3), 380-385.
- Yun, Z. S., & Good, L. K. (2007). Developing Customer Loyalty from E-tail Store Image Attributes. *Managing Service Quality*, 17(1), 4-22.

