

Robot Süpürge Markalarında Dijital İnovasyon Sinyallerinin FASTopic ve Çok Kriterli Karar Verme ile Analizi

Yunus Emre Özenoğlu^{1*}, Okan Alkan² ve Oğuz Emre Balkar³

¹Yönetim Bilişim Sistemleri Anabilim Dalı, Sosyal Bilimler Enstitüsü, Atatürk Üniversitesi, Türkiye

²Bilgisayar Teknolojileri Bölümü, Meslek Yüksekokulu, Erzincan Binali Yıldırım Üniversitesi, Türkiye

³Yönetim ve Organizasyon Bölümü, Meslek Yüksekokulu, Erzincan Binali Yıldırım Üniversitesi, Türkiye

*(yunusemre.ozenoglu12@ogr.atauni.edu.tr)

Özet – Bu çalışma, robot süpürge uygulamalarına ait Google Play kullanıcı yorumlarında görünür hâle gelen kullanıcı değer alanlarını, dijital hizmet sorunlarını ve geri bildirim temelli inceleme önceliklerini belirlemeyi amaçlamaktadır. Karma kullanımlı uygulamalardan kaynaklanabilecek ürün uygunluğu yanlılığını azaltmak için analiz yalnızca robot süpürge yönetimine ayrılmış yedi uygulama paketiyle sınırlandırılmıştır. Bu uygulamalardan 1 Ocak 2023–5 Temmuz 2026 döneminde elde edilen 31.594 benzersiz yorumun 17.912’si ön işleme ölçütlerini karşılamıştır. Yorumlar FASTopic ile 30 konuya modellenmiş; her yorumun en yüksek olasılıklı üç konuya kesirli katkı sağlamasına dayanan çoklu konu yaklaşımı kullanılmıştır. Konuların yıldız gruplarındaki ağırlıkları Mann-Whitney U testi ve rank-biserial etki büyüklüğüyle karşılaştırılmıştır. Eyleme dönük 19 konu; yaygınlık, düşük puan yoğunlaşması, uygulamalar arası yayılım ve zamansal süreklilik ölçütleriyle CRITIC-TOPSIS kullanılarak sıralanmıştır. Düşük puanlı yorumlarda bağlantı kopmaları, uygulama çökmesi, güncelleme sonrası işlev bozulması ve uyumluluk sorunları; yüksek puanlı yorumlarda ise otonomi, zaman tasarrufu, temizlik performansı ve genel memnuniyet öne çıkmıştır. İlk inceleme öncelikleri tekrarlayan bağlantı kopmaları, güncelleme sonrası işlev bozulması, uygulama uyumluluğu, harita kaybı ve uygulama çökmesi olarak belirlenmiştir. Çoklu konu ataması, düşük puan eşiği ve kriter ağırlıklarına ilişkin duyarlılık analizleri ilk sıralardaki konuların büyük ölçüde korunduğunu göstermiştir. Sonuçlar, firma inovasyon performansı veya ürün yol haritası göstergesi değil, incelenen platform ve dönem kapsamında kullanıcı geri bildirimlerinden türetilen dijital inovasyon sinyalleridir.

Anahtar Kelimeler – Robot süpürge, kullanıcı yorumları, dijital inovasyon, FASTopic, çok kriterli karar verme, CRITIC, TOPSIS

I. GİRİŞ

Robot süpürgeler, fiziksel temizlik donanımı ile mobil uygulama, haritalama, bağlantı, kullanıcı hesabı ve yazılım güncellemelerinin birlikte çalıştığı dijital ürün-hizmet sistemlerine dönüşmüştür. Bu bütünleşik yapı içinde kullanıcı deneyimi yalnızca emiş veya paspaslama performansına değil, uygulama kararlılığına, cihazla iletişimin sürekliliğine, harita yönetimine ve satış sonrası hizmetlere de bağlıdır [1], [2]. Dolayısıyla uygulama mağazası yorumları, dijital hizmet bileşenlerinde ortaya çıkan değer ve sorun alanlarını gözlemlemek için önemli bir dış bilgi kaynağıdır.

Kullanıcı tarafından üretilen geri bildirimler, işletmelerin sürekli iyileştirme süreçlerinde ürün dışından gelen bilgi sinyalleri sağlar. Yorum metinleri, yıldız puanlarının tek başına açıklayamadığı özellik taleplerini, kullanım bağlamlarını ve sorun kümelerini görünür hâle getirebilir [3], [4]. Mobil uygulama geri bildirimlerinin güncellemelerle ilişkilendirilmesi ve iteratif geliştirme süreçlerinde kullanılması, kullanıcı bilgisinin teknik karar süreçlerine aktarılabilmesini göstermektedir [5], [6], [7]. Bununla birlikte yorumlar, kendi kendini seçen kullanıcılar tarafından oluşturulduğu için firma performansının veya genel ürün kalitesinin doğrudan ölçümü olarak görülmemelidir.

Geniş yorum koleksiyonlarının manuel olarak incelenmesi güç olduğundan, konu modelleme yöntemleri tekrarlayan anlam yapılarını otomatik biçimde ortaya çıkarmak için kullanılmaktadır. Gömme

temelli FASTopic, bağlamsal belge temsilleri üzerinden yorumlanabilir konu yapıları üretmektedir [8]. Ancak tek bir baskın konuya dayalı sınıflandırma, aynı yorumda birlikte bulunan bağlantı, haritalama ve güncelleme gibi birden fazla sorunu göz ardı edebilir. Bu çalışmada bu bilgi kaybını azaltmak için her yorumun en güçlü üç konuya kesirli katkı sağladığı çoklu konu yaklaşımı benimsenmiştir.

Kullanıcı geri bildirimlerinden elde edilen sorunların sıklığı tek başına teknik müdahale kararı için yeterli değildir. Bir konunun düşük puanlarda yoğunlaşması, farklı uygulamalarda görülmesi ve zaman içinde süreklilik göstermesi de inceleme gereksinimini artırabilir. Bu nedenle konu modelleme çıktıları, CRITIC ile veri yapısına dayalı ağırlıklandırma ve TOPSIS ile göreceli sıralama yaklaşımıyla bütünleştirilmiştir [9], [10]. Burada üretilen sıralama “ürün geliştirme yol haritası” değil, teknik ve yönetsel ekiplerin daha ayrıntılı inceleme yapabileceği geri bildirim alanlarının göreceli görünürlüğüdür.

Tablo 1. İlgili çalışmaların yaklaşım ve kapsam bakımından karşılaştırılması

Çalışma	Veri bağlamı	Yöntem	Temel odak	Bu çalışma açısından ayrım
Carames vd. [1]	Robot süpürge perakende yorumları	Çevrimiçi yorum analizi	İnsan-otomasyon etkileşimi	Mobil uygulama odaklı karar sıralaması yok
Pınarbaşı ve Zeyrek Pınarbaşı [11]	Akıllı ev uygulamaları	Konu modelleme	Pazar ve kullanıcı deneyimi temaları	Robot süpürgeye özgü çoklu konu analizi yok
Aydin Gokgoz vd. [7]	Mobil uygulamalar	Geri bildirim-güncelleme analizi	Güncellemelerin puanlarla ilişkisi	Konu düzeyinde inceleme önceliği yok
Wang vd. [12]	Çevrimiçi ürün yorumları	Özellik çıkarımı ve önceliklendirme	Ürün iyileştirme alanları	Uygulama ekosistemi ve yumuşak konu ataması yok
Ng vd. [13]	Sürdürülebilir ürün inovasyonu	Duygu, LLM ve ÇKKV	Çok yöntemli önceliklendirme	Farklı ürün ve karar bağlamı
Bu çalışma	Yedi özel robot süpürge uygulaması	FASTopic, çoklu konu ataması ve CRITIC-TOPSIS	Değer, dijital sorun ve inceleme öncelikleri	Platforma ve döneme özgü geri bildirim sinyalleri

Literatürde robot süpürge deneyimini, uygulama yorumlarını ve ürün iyileştirme önceliklerini inceleyen çalışmalar bulunmakla birlikte, konu olasılıklarını çoklu atamayla kullanarak yıldız grupları ve sürekli ölçütlere dayalı inceleme öncelikleriyle birleştiren uygulamalar sınırlıdır. Bu çalışma bir firma inovasyon performansı ölçümü önermemekte; uygulama mağazası yorumlarında görünür hâle gelen değer ve sorun örüntülerini sistematik inceleme sinyallerine dönüştürmeyi amaçlamaktadır.

Bu amaç doğrultusunda aşağıdaki araştırma sorularına yanıt aranmıştır:

AS1. Robot süpürge uygulama yorumlarında hangi kullanıcı değer alanları ve dijital hizmet sorunları görünür hâle gelmektedir?

AS2. Konu olasılıkları düşük ve yüksek yıldızlı yorum grupları arasında nasıl farklılaşmaktadır?

AS3. Yaygınlık, düşük puan yoğunlaşması, uygulamalar arası yayılım ve zamansal süreklilik birlikte değerlendirildiğinde hangi konular geri bildirim temelli inceleme önceliği taşımaktadır?

II. MATERYAL VE YÖNTEM

A. Araştırma Tasarımı ve Kavramsal Sınırlar

Araştırma nicel, keşifsel ve betimsel bir metin analizi olarak tasarlanmıştır. Çalışmada ölçülen yapı; uygulama mağazası yorumlarında görünür olan kullanıcı değerleri, dijital hizmet sorunları ve bu sorunların göreceli inceleme görünürlüğüdür. Yeni özellik sayısı, güncelleme hızı, kullanıcı tutma oranı, uygulama kullanım yoğunluğu veya firma müdahale süresi ölçülmediğinden “dijital inovasyon performansı” değerlendirmesi yapılmamıştır. Sonuçlar nedensel etki, genel ürün üstünlüğü veya firma sıralaması olarak yorumlanmamıştır.

B. Veri Kaynağı, Uygulama Seçimi ve Ön İşleme

Veriler Google Play'in Amerika Birleşik Devletleri, Birleşik Krallık, Kanada ve Avustralya yerellerinden google-play-scraper aracılığıyla toplanmıştır. Veri dönemi 1 Ocak 2023–5 Temmuz 2026'dır. Önceki analiz hattında yer alan karma kullanımlı Roborock, eufy ve Philips HomeRun paketleri, ürün uygunluğu filtresinin markalar ve uygulamalar arasında eşdeğer olmayan bir seçim süreci oluşturma riskini azaltmak amacıyla revize edilmiş ana örneklemden çıkarılmıştır. Ana analiz yalnızca robot süpürge yönetimine ayrılmış yedi uygulama paketi üzerinden ve paketler birleştirilmeden yürütülmüştür.

Tablo 2. Uygulama paketi düzeyinde ham ve nihai örneklem

Uygulama	Ham n	Nihai n	Korunma	Ham ort.	Nihai ort.
iRobot Home (Classic)	9.800	5.623	%57,4	3,19	2,70
SharkClean	9.600	5.604	%58,4	3,24	2,71
ECOVACS HOME	6.400	3.462	%54,1	3,03	2,42
eufy Clean (EufyHome)	3.000	1.575	%52,5	3,80	3,19
Roomba Home	1.394	787	%56,5	3,31	2,72
Dreamhome	1.000	638	%63,8	2,62	2,38
Rowenta Robots	400	223	%55,8	1,85	1,68
Toplam / ağırlıklı ortalama	31.594	17.912	%56,7	3,20	2,67

Not. Uygulamalar marka düzeyinde birleştirilmemiştir. Ham ve nihai ortalamalar yıldız puanlarını göstermektedir.

Yereller arasında yinelenen kayıtlar benzersiz yorum kimliği üzerinden kaldırılmıştır. Metinlere Unicode NFKC normalizasyonu uygulanmış; HTML, URL, e-posta, emoji ve kontrol karakterleri temizlenmiştir. En az 20 karakter ve dört kelime içermeyen yorumlar ile İngilizce olarak sınıflandırılmayan veya dil güven puanı 0,70'in altında kalan kayıtlar analiz dışında bırakılmıştır. Yedi özel uygulamada 31.594 benzersiz yorumdan 17.912'si nihai korpusa alınmıştır. Nihai örneklemin ortalama puanının ham örneklemden daha düşük olması, filtrelerin puan dağılımını etkileyebileceğini göstermektedir; bu durum sınırlılıklar bölümünde dikkate alınmıştır.

C. FASTopic, İnsan Kodlaması ve Model Sağlamlığı

Yorumlar FASTopic 1.0.1 ile modellenmiştir [8]. Belge temsilleri *paraphrase-multilingual-MiniLM-L12-v2* modeliyle oluşturulmuştur [14]. Model, 30 konu ve 10.000 sözcüklük kelime dağarcığıyla 200 eğitim dönemi boyunca 0,002 öğrenme oranıyla çalıştırılmış; rassal başlangıç değeri 42 olarak belirlenmiştir. Her konu için 15 anahtar sözcük ve 10 temsili yorum elde edilmiştir.

Konu adları ve kavramsal boyutlar iki araştırmacı tarafından yıldız puanı, uygulama adı, tarih, önceki etiket ve marka ifadeleri gizlenerek bağımsız biçimde değerlendirilmiştir. Üst boyut sınıflandırmasında 26/30 doğrudan uyum (%86,7; Cohen's $\kappa=0,841$), analiz rolünde 27/30 uyum (%90,0; $\kappa=0,793$) elde edilmiştir. Uyuşmazlıklar anahtar sözcükler ve temsili yorumlar yeniden incelenerek uzlaştırılmıştır. Analiz rolü sınıflandırması, yıldız puanı ilişkisini doğrulamak için kullanılmamış; puan gruplarıyla karşılaştırmalar doğrudan konu ağırlıkları üzerinden yürütülmüştür.

Konu sayısı seçiminin tek bir çözüme bağımlılığını değerlendirmek amacıyla, 3.000 yorumdan oluşan uygulama ve puan açısından tabakalandırılmış örneklem üzerinde tamamlayıcı bir TF-IDF/SVD-MiniBatchKMeans analizi yapılmıştır. Analizde konu sayıları 15, 20, 25, 30, 35, 40 ve 50; rassal başlangıç değerleri ise 11, 23, 42, 67 ve 89 olarak belirlenmiştir. Farklı konu sayıları; Silhouette katsayısı, ilk 10 sözcüğe dayalı konu çeşitliliği, rassal başlangıçlar arasındaki düzeltilmiş Rand indeksi ve yeterli büyüklükteki kümelerin kapsama oranı birlikte değerlendirilerek karşılaştırılmıştır.

Tablo 3. Alternatif vektör-kümeleme yaklaşımında farklı konu sayılarının karşılaştırılması

K	Silhouette	Çeşitlilik	ARI kararlılığı	Kapsama	Ort. en küçük küme büyüklüğü (n)
15	0,103	0,603	0,272	%100,0	97,8
20	0,122	0,577	0,350	%100,0	78,0
25	0,133	0,537	0,383	%100,0	52,2
30	0,141	0,513	0,442	%100,0	39,8
35	0,146	0,473	0,484	%100,0	34,8
40	0,151	0,438	0,530	%100,0	31,0
50	0,143	0,373	0,557	%97,2	15,6

Not. Değerler beş rassal başlangıç değerinin ortalamasıdır. Kapsama, örneklemin en az %0,1'ine karşılık gelen ve en az 20 yorum içeren kümelerin toplam kümeler içindeki oranıdır.

Silhouette katsayısı 40 konuya kadar yükselmiş ve 50 konulu çözümde azalmıştır. Rassal başlangıçlar arasındaki kararlılık konu sayısı arttıkça yükselirken, konu çeşitliliği azalmış ve 50 konulu çözümde küçük kümeler belirginleşmiştir. Otuz konulu çözüm, ayrışma, konu çeşitliliği ve yorumlanabilirlik arasında makul bir ara seçenek olarak değerlendirilmiştir. Bu bulgu, 30 konunun tek ve kesin optimum çözüm olduğunu göstermemekte; yalnızca seçilen konu sayısının alternatif konu sayıları arasında savunulabilir bir ara çözüm olduğunu desteklemektedir. FASTopic konuları ile alternatif vektör kümeleri arasındaki sözcüksel örtüşmenin kısmi düzeyde kalması nedeniyle model bağımlılığı bir sınırlılık olarak korunmuştur.

D. Çoklu Konu Ataması ve Yıldız Gruplarıyla Karşılaştırma

Her yorumun yalnızca baskın konuya zorlanmasından kaynaklanan bilgi kaybını azaltmak için, en yüksek olasılığa sahip üç konu korunmuş ve bu üç olasılık yorum içinde toplamı 1 olacak biçimde yeniden normalize edilmiştir. Böylece bir yorum aynı anda birden fazla konuya kesirli katkı sağlamıştır. Ana analiz üç konulu atamaya dayanmış; ilk iki ve ilk beş konuya dayalı sonuçlar duyarlılık analizi olarak karşılaştırılmıştır.

Yıldız puanları, konu adlandırmasından bağımsız dış değerlendirme değişkeni olarak kullanılmıştır. Bir-iki yıldızlı yorumlar düşük, dört-beş yıldızlı yorumlar yüksek puan grubu olarak tanımlanmıştır. Her konu için iki gruptaki ortalama kesirli ağırlıklar Mann-Whitney U testiyle karşılaştırılmış; etki büyüklüğü rank-biserial korelasyonla, çoklu testler Benjamini-Hochberg FDR düzeltmesiyle raporlanmıştır. Pozitif etki değeri konunun düşük puan grubunda, negatif değer yüksek puan grubunda daha görünür olduğunu göstermektedir.

E. Geri Bildirim Temelli İnceleme Öncelikleri

Teknik düzeltilebilirlik, maliyet, güvenlik riski ve stratejik uyum gibi firma içi ölçütler bulunmadığından, sonuçlar “geliştirme önceliği” yerine “geri bildirim temelli inceleme önceliği” olarak tanımlanmıştır. Genel duygu ve gürültü konuları ile yalnızca değer bildiren konular çıkarılmış; eyleme dönük 19 sorun konusu dört sürekli ölçütle değerlendirilmiştir: (i) kesirli konu ağırlıklarının ortalamasına dayalı yaygınlık, (ii) konu ağırlığının bir-iki yıldızlı yorumlarda yoğunlaşma oranı, (iii) konu kütesinin yedi uygulama arasındaki normalize Shannon entropisine dayalı uygulama yayılımı ve (iv) aylık konu yaygınlığının değişim katsayısından türetilen $1/(1+CV)$ zamansal süreklilik değeri.

Bütün ölçütler fayda yönlü olarak düzenlenmiştir. CRITIC yönteminde min-maks standartlaştırması sonrasında standart sapma ve kriterler arası korelasyon kullanılmış; elde edilen ağırlıklarla TOPSIS yakınlık skorları hesaplanmıştır. Bu sıralama uygulama veya marka performansı değildir. Duyarlılık analizlerinde iki ve beş konulu atamalar, yalnızca bir yıldız ve bir-üç yıldız aralığını esas alan alternatif düşük puan tanımları, eşit ağırlıklar ve ağırlıkların $\pm\%20$ rastgele değiştirilmesine dayalı 5.000 senaryo kullanılmıştır.

F. Etik ve Gizlilik

Araştırmada yalnızca kamuya açık Google Play yorumları kullanılmış; kullanıcılarla doğrudan iletişim kurulmamış ve özel nitelikli verilere erişilmemiştir. Bu nedenle etik kurul onayı gerektirmemiştir. Kullanıcı adları tek yönlü karma değerlerle takma adlandırılmış, profil görselleri kullanılmamış; e-posta, telefon ve URL gibi tanımlayıcılar maskelenmiştir. Bulgular yalnızca toplulaştırılmış düzeyde raporlanmıştır.

III. BULGULAR

A. Örneklem ve Konu Yapısı

Nihai örneklem 17.912 yorumdan oluşmaktadır. Ortalama yıldız puanı 2,67'dir; yorumların %52,50'si bir-iki yıldız, %35,96'sı dört-beş yıldızdır. Çoklu konu atamasında en yüksek kesirli yaygınlık paspaslama ve zemin temizleme performansında (%6,31) görülmüştür. Bunu uygulama uyumluluğu ve arayüz hataları (%5,63), tavsiye ve memnuniyet (%4,93), güncelleme sonrası işlev bozulması (%4,84), tekrarlayan bağlantı kopmaları (%4,84), otonomi ve zaman tasarrufu (%4,81) ile harita kaybı ve çok katlı harita yönetimi (%4,69) izlemiştir. Konuların tamamı Ek Tablo A1'de sunulmuştur.

B. Konu Ağırlıklarının Yıldız Gruplarına Göre Farklılaşması

Tablo 4. Düşük ve yüksek puan gruplarında en belirgin konu farklılıkları

Yön	Konu	Düşük grup (%)	Yüksek grup (%)	Rank-biserial r	FDR p
Düşük puanda yüksek	Tekrarlayan bağlantı kopmaları	7,53	1,18	0,188	<0,001
Düşük puanda yüksek	Algılanan yazılım tasarımı ve güvenilirlik yetersizliği	6,81	0,97	0,173	<0,001
Düşük puanda yüksek	Uygulama çökmesi, açılmama ve yeniden kurulum	6,36	0,94	0,161	<0,001
Düşük puanda yüksek	Güncelleme ve firmware sonrası işlev bozulması	7,12	1,80	0,159	<0,001
Düşük puanda yüksek	Wi-Fi eşleştirme, oturum açma ve cihaz bağlantısı	5,91	1,09	0,146	<0,001
Yüksek puanda yüksek	Otonomi, kolaylık ve zaman tasarrufu	0,48	11,69	-0,335	<0,001
Yüksek puanda yüksek	Tavsiye, memnuniyet ve beklenti aşımı	1,15	11,69	-0,307	<0,001
Yüksek puanda yüksek	Genel olumlu temizlik ve sahipelik deneyimi	0,14	8,88	-0,284	<0,001
Yüksek puanda yüksek	Paspaslama ve zemin temizleme performansı	2,42	12,38	-0,281	<0,001
Yüksek puanda yüksek	Beklentilerin ve vaat edilen işlevlerin karşılanması	2,39	7,36	-0,149	<0,001

Not. Yüzdelere, yorum içindeki ilk üç konu ağırlığının grup ortalamasıdır. Pozitif r düşük puan grubunda, negatif r yüksek puan grubunda daha yüksek ağırlığı gösterir.

Düşük puanlı yorumlarda tekrarlayan bağlantı kopmaları en belirgin ayrışmayı göstermiştir ($r=0,188$). Algılanan yazılım tasarımı ve güvenilirlik yetersizliği ($r=0,173$), uygulama çökmesi ($r=0,161$), güncelleme sonrası işlev bozulması ($r=0,159$) ve Wi-Fi/oturum bağlantısı ($r=0,146$) düşük puan grubunda daha yüksek ağırlığa sahiptir. Etki büyüklükleri küçük ile küçük-orta arasında olup, büyük örneklem nedeniyle yalnızca p değerlerine dayalı güçlü yorum yapılmamıştır.

Yüksek puanlı yorumlarda otonomi, kolaylık ve zaman tasarrufu ($r=-0,335$), tavsiye ve beklenti aşımı ($r=-0,307$), genel olumlu temizlik deneyimi ($r=-0,284$), paspaslama ve zemin temizleme performansı ($r=-0,281$) ile beklentilerin karşılanması ($r=-0,149$) daha görünürdür. Bu sonuçlar konu ve değerlendirme yönünün ayrı düzlemlerde ele alınmasının, aynı korpusta hem değer hem sorun sinyallerini göstermeye imkân verdiğini ortaya koymaktadır.

C. Geri Bildirim Temelli İnceleme Öncelikleri

CRITIC ağırlıkları Yaygınlık %22,66, Düşük puan yoğunlaşması %31,68, Uygulama yayılımı %17,91, Zamansal süreklilik %27,76 olarak hesaplanmıştır. Ağırlıklar yönetsel önem veya teknik kritiklik değil, 19 konu arasındaki değişkenlik ve farklılaşma yapısını göstermektedir.

Tablo 5. Kullanıcı geri bildirimlerine dayalı ilk 10 inceleme önceliği

Sıra	Konu	Yaygınlık	Düşük puan	Uygulama yayılımı	Zamansal süreklilik	TOPSIS
1	Tekrarlayan bağlantı kopmaları	%4,84	%81,71	0,812	0,826	0,842
2	Güncelleme ve firmware sonrası işlev bozulması	%4,84	%77,25	0,818	0,791	0,814
3	Uygulama uyumluluğu, yükleme ve arayüz hataları	%5,63	%67,74	0,822	0,839	0,809
4	Harita kaybı, bozulması ve çok katlı harita yönetimi	%4,69	%66,67	0,752	0,824	0,704
5	Uygulama çökmesi, açılmama ve yeniden kurulum	%4,05	%82,46	0,738	0,696	0,673
6	Wi-Fi eşleştirme, oturum açma ve cihaz bağlantısı	%3,85	%80,59	0,796	0,846	0,671
7	Gizlilik, hesap güvenliği ve parola politikaları	%3,22	%77,11	0,859	0,843	0,563
8	Reklam, ücretlendirme ve iş modeli eleştirileri	%2,66	%82,26	0,828	0,769	0,509
9	Engel algılama ve otonom navigasyon davranışı	%3,36	%64,65	0,783	0,857	0,495
10	Harita oluşturma, kaybı ve yeniden haritalama	%3,15	%68,91	0,749	0,820	0,487

Not. Sıralama, ürün yol haritası veya firma performansı olarak değil, daha ayrıntılı teknik ve yönetsel inceleme gerektiren kullanıcı geri bildirim alanları olarak yorumlanmalıdır.

İlk üç sırada tekrarlayan bağlantı kopmaları, güncelleme ve firmware sonrası işlev bozulması ile uygulama uyumluluğu, yükleme ve arayüz hataları yer almıştır. Harita kaybı ve çok katlı harita yönetimi dördüncü, uygulama çökmesi ve yeniden kurulum beşinci sıradadır. Reklam ve ücretlendirme eleştirilerinin yüksek düşük-puan yoğunluğuna rağmen daha düşük yaygınlık ve süreklilik nedeniyle sekizinci sırada kalması, tek bir şikâyet oranının sıralamayı belirlemediğini göstermektedir.

Duyarlılık sonuçları sıralamanın ana örüntüsünün korunduğunu göstermiştir. İlk iki, üç ve beş konuya dayalı atamalar arasındaki Spearman korelasyonları 0,984-0,995; bir-iki yıldız tanımı ile yalnızca bir yıldız ve bir-üç yıldız tanımları arasındaki korelasyonlar 0,958 ve 0,960'tır. CRITIC ve eşit ağırlıklı sıralamalar arasındaki korelasyon 0,982'dir. Ağırlıkların ± 20 değiştirildiği 5.000 senaryoda ilk üç konu bütün tekrarların %100'ünde ilk beşte kalmıştır; dördüncü ve beşinci sıralardaki konuların ilk beşte kalma olasılıkları sırasıyla %86,2 ve %65,7'dir. Bu bulgu ilk üç önceliğin daha kararlı, beşinci-altıncı sıra sınırının ise ağırlık seçimine duyarlı olduğunu göstermektedir.

IV. TARTIŞMA

A. Kullanıcı Değeri ve Dijital Hizmet Sorunları

Temizlik, otonomi ve zaman tasarrufunun yüksek puanlı yorumlarda öne çıkması, robot süpürgelerde kullanıcı değerinin yalnızca teknik yenilikten değil, gündelik iş yükünün azaltılmasından ve görevin öngörülebilir biçimde tamamlanmasından doğduğunu göstermektedir. Bu sonuç, robot süpürgelerde işlevsellik ve akıllılık boyutlarının kullanıcı deneyimindeki önemini vurgulayan çalışmalarla uyumludur [1], [2]. Paspaslama performansının en yaygın konu olması, fiziksel temizlik çıktısının hâlen temel değer çekirdeği olduğunu ortaya koymaktadır.

Buna karşılık bağlantı kopmaları, uygulama çökmesi, güncelleme sonrası bozulma ve uyumluluk sorunlarının düşük puanlarda yoğunlaşması, fiziksel cihazın değerinin dijital hizmet sürekliliğine bağlı olduğunu göstermektedir. Bağlantı ve harita yönetimi gibi işlevler arızalandığında kullanıcı yalnızca bir yazılım hatası değil, cihazın özerkliğini ve gündelik kullanım değerini kesintiye uğratan bütünsel bir hizmet sorunu yaşamaktadır. Akıllı ev uygulamalarında yazılım, bağlantı ve destek süreçlerinin kullanıcı memnuniyetini belirlediğine ilişkin bulgular bu yorumu desteklemektedir [11], [2].

B. Yöntemsel ve Yönetimsel Çıkarımlar

Çalışmanın yöntemsel katkısı, yorumları tek bir baskın konuya zorlamak yerine çoklu konu ağırlıklarıyla ele almasıdır. Bağlantı kopması, harita kaybı ve güncelleme sorunu aynı yorumda birlikte bulunabildiğinden, ilk üç konuya kesirli atama konu birlikteliğini kısmen korumuştur. İlk iki ve ilk beş konuya dayalı sıralamaların yüksek düzeyde benzer olması, temel sonuçların seçilen çoklu atama genişliğine aşırı bağımlı olmadığını göstermektedir. Bununla birlikte FASTopic olasılıklarının mutlak kalibrasyonu doğrulanmadığından bu ağırlıklar kesin olay oranları olarak yorumlanmamalıdır.

CRITIC-TOPSIS sonucu, firma içi ürün yol haritası yerine inceleme gündemi sunmaktadır. Teknik ekipler ilk sıralardaki bağlantı, güncelleme, uyumluluk, harita ve çökme konularını hata kayıtları, sürüm günlükleri, cihaz telemetrisi ve müşteri hizmetleri verileriyle doğrulayabilir. Gerçek geliştirme önceliğinin belirlenmesi için teknik maliyet, güvenlik riski, etkilenen aktif kullanıcı sayısı, stratejik uyum ve uygulama süresi gibi firma içi ölçütlerin ayrıca değerlendirilmesi gerekir. Dolayısıyla kullanıcı yorumları kararın kendisi değil, daha ayrıntılı incelemenin başlangıç sinyalidir.

C. Alternatif Açıklamalar

Bulgular yorum yazan kullanıcıların bileşiminden etkilenebilir. Sorun yaşayan kullanıcıların yorum bırakma olasılığı daha yüksek olabilir; başarılı güncellemeler ise sessiz kalabilir. Uygulamaların puan isteme mekanizmaları, kullanıcı tabanlarının büyüklüğü, cihaz modeli karmaşıklığı, ülke dağılımı ve alternatif destek kanalları uygulamalar arasında farklılaşabilir. Ayrıca uygulama mağazası yorumları fiziksel cihaz deneyimiyle dijital uygulama deneyimini her zaman ayırmaz. Bu nedenle düşük puan yoğunlaşması belirli bir teknik hatanın yaygınlığını veya gerçek kullanıcı nüfusundaki oranını doğrudan göstermez.

V. SONUÇLAR

Bu çalışma, yedi özel robot süpürge uygulamasına ait 17.912 Google Play yorumunu FASTopic, çoklu konu ataması ve CRITIC-TOPSIS ile incelemiştir. Kullanıcı değeri otonomi, zaman tasarrufu, temizlik performansı ve beklentilerin karşılanması etrafında şekillenirken; bağlantı kopmaları, güncelleme sonrası bozulma, uygulama uyumluluğu, harita kaybı ve uygulama çökmesi düşük puanlarla daha güçlü biçimde birlikte görülmüştür. Geri bildirim temelli sıralama, özellikle bağlantı ve yazılım sürekliliği konularının ayrıntılı teknik inceleme için öncelikli sinyaller oluşturduğunu göstermiştir.

Çalışma firma inovasyon performansını, genel ürün kalitesini veya ürün yol haritasını ölçmemektedir. Veriler Android/Google Play, İngilizce dil, dört yerel ve belirli bir takvim dönemiyle sınırlıdır. Kısa yorumların ve düşük dil güvenli kayıtların çıkarılması, özellikle özlü ve uç değerlendirmelerin temsilini azaltmış olabilir; nihai örneklemedeki ortalama puanın ham örneklemden düşük olması seçim etkisine işaret etmektedir. Dil filtresi, dışlanan kısa yorumlar ve 0,70 eşiği için insan doğrulaması bulunmamaktadır. FASTopic farklı konu sayıları ve rassal başlangıç değerleriyle yeniden eğitilmemiştir; alternatif vektör-kümeleme analizi yalnızca tamamlayıcı bir duyarlılık değerlendirmesi sağlamıştır. İlk üç konuya kesirli atama çoklu içerik kaybını azaltmakla birlikte tüm anlamsal birliktelikleri korumaz. Yıldız puanları duygu ve deneyimin kusursuz ölçümü değildir. İnceleme önceliklerinde teknik maliyet, güvenlik, gelir etkisi ve düzeltilebilirlik bulunmadığından sonuçlar yalnızca geri bildirim görünürlüğü olarak değerlendirilmelidir.

Etik Kurul Beyanı

Çalışma yalnızca kamuya açık Google Play kullanıcı yorumlarına dayandığından, doğrudan insan katılımcı etkileşimi ve özel nitelikli veri kullanımı içermediğinden etik kurul onayı gerektirmemiştir.

EKLER

Ek Tablo A1. FASTopic konu etiketleri ve ilk beş anahtar sözcük

Konu	Nihai etiket	Boyut	İlk beş anahtar sözcük
0	Harita oluşturma, kaybı ve yeniden haritalama	Haritalama ve navigasyon	map; maps; mapping; house; new
1	Süpürme ve emiş performansı	Fonksiyonel performans	cleaner; vacuums; vacuuming; suction; power
2	Genel kullanılabilirlik ve puanlama gürültüsü	Dışlanan/gürültü	translation; rating; potential; score; tuning
3	Tekrarlayan bağlantı kopmaları	Bağlantı güvenilirliği	disconnecting; dropping; repeatedly; shooting; rooba
4	Uygulama uyumluluğu, yükleme ve arayüz hataları	Yazılım kalitesi	tablets; application; fold; root; glitching
5	Algılanan yazılım tasarımı ve güvenilirlik yetersizliği	Yazılım kalitesi	developers; coding; retarded; algorithms; programmers
6	Ekran, düğme ve bildirim arayüzü sorunları	Yazılım kalitesi	screen; button; notifications; google; page
7	Oda bazlı temizlik, rutin ve programlama	Yazılım kalitesi	room; rooms; schedule; clean; set
8	Harita kaybı, bozulması ve çok katlı harita yönetimi	Haritalama ve navigasyon	disappearing; multifloor; losing; corrupted; multimap
9	Genel olumlu temizlik ve sahiplik deneyimi	Kullanıcı değer alanları	love; great; house; floors; vacuum
10	Gizlilik, hesap güvenliği ve parola politikaları	Güven ve sorumlu inovasyon	password; security; language; english; privacy
11	Engel algılama ve otonom navigasyon davranışı	Haritalama ve navigasyon	robots; bot; obstacle; mower; goat
12	Kullanım kolaylığı ve öğrenilebilirlik	Kullanıcı değer alanları	good; easy; use; better; great
13	Temizlik görevi sürekliliği ve tamamlama başarısı	Fonksiyonel performans	cleaning; clean; vacuum; gets; job
14	Evcil hayvan tüyü temizleme değeri	Fonksiyonel performans	hair; dog; cats; cat; pet
15	Müşteri desteği, garanti ve satış sonrası hizmet	Hizmet inovasyonu	product; support; service; customer; irobot
16	Paspaslama ve zemin temizleme performansı	Fonksiyonel performans	mopping; sweeping; wet; leaves; stains
17	Wi-Fi eşleştirme, oturma açma ve cihaz bağlantısı	Bağlantı güvenilirliği	connect; wifi; device; connection; phone
18	Genel olumsuz değerlendirme ve yoğun duygu dili	Dışlanan/gürültü	slow; terrible; constantly; worst; garbage
19	Güncelleme ve firmware sonrası işlev bozulması	Yazılım kalitesi	latest; available; broken; updates; broke
20	Otonomi, kolaylık ve zaman tasarrufu	Kullanıcı değer alanları	convenient; breeze; enjoying; convenience; ease
21	Beklentilerin ve vaat edilen işlevlerin karşılanması	Kullanıcı değer alanları	exactly; expected; supposed; advertised; follow
22	Batarya, şarj, istasyona dönüş ve göreve devam	Fonksiyonel performans	battery; station; charging; charge; docking
23	Uygulama çökmesi, açılmama ve yeniden kurulum	Yazılım kalitesi	crashes; reinstalled; crashing; reinstall; uninstalled
24	Tavsiye, memnuniyet ve beklenti aşımı	Kullanıcı değer alanları	recommended; highly; expectations; recommend; satisfied
25	Aşırı genel uygulama-robot-güncelleme konusu	Dışlanan/gürültü	app; robot; vacuum; update; work

Konu	Nihai etiket	Boyut	İlk beş anahtar sözcük
26	Reklam, ücretlendirme ve iş modeli eleştirileri	Hizmet inovasyonu	ads; paid; unacceptable; advertising; marketing
27	Oda bölme, sanal duvar ve sınır düzenleme	Haritalama ve navigasyon	divide; walls; boundaries; dividers; dividing
28	Mobilya etkileşimi ve otonom navigasyon sorunları	Haritalama ve navigasyon	furniture; drive; confused; trouble; motor
29	Planlı temizlik, zamanlama ve rahatsız etmeme otomasyonu	Yazılım kalitesi	scheduled; scheduling; night; schedules; disturb

Not. Konu etiketleri iki araştırmacının bağımsız kodlaması ve uzlaştırması sonucunda belirlenmiştir.

KAYNAKLAR

- [1] K. Carames, K. Mui, A. Azad, and W. C. W. Giang, "Studying robot vacuums using online retailer reviews to understand human-automation interaction," Proceedings of the Human Factors and Ergonomics Society Annual Meeting, vol. 65, no. 1, pp. 1029-1033, 2021, doi: 10.1177/1071181321651106.
- [2] Y. Yu, Q. Fu, D. Zhang, and Q. Gu, "Understanding user experience with smart home products," Journal of Computer Information Systems, 2024, doi: 10.1080/08874417.2024.2408006.
- [3] E. Guzman and W. Maalej, "How do users like this feature? A fine-grained sentiment analysis of app reviews," in 2014 IEEE 22nd International Requirements Engineering Conference (RE), pp. 153-162, 2014, doi: 10.1109/RE.2014.6912257.
- [4] T. Redek and U. Godnov, "From data to decision: Distilling decision intelligence from user-generated content," Kybernetes, vol. 53, no. 13, pp. 1-23, 2024, doi: 10.1108/K-08-2023-1447.
- [5] T. Liu, C. Wang, K. Huang, P. Liang, B. Zhang, M. Daneva, and M. van Sinderen, "RoseMatcher: Identifying the impact of user reviews on app updates," Information and Software Technology, vol. 161, art. 107261, 2023, doi: 10.1016/j.infsof.2023.107261.
- [6] P. Tian and Q. Yang, "The impact of online customer reviews on product iterative innovation," European Journal of Innovation Management, vol. 27, no. 8, pp. 2646-2667, 2024, doi: 10.1108/EJIM-09-2022-0501.
- [7] Z. Aydin Gokgoz, M. B. Ataman, and G. H. van Bruggen, "If it ain't broke, should you still fix it? Effects of incorporating user feedback in product development on mobile application ratings," International Journal of Research in Marketing, vol. 42, pp. 467-486, 2025, doi: 10.1016/j.ijresmar.2024.10.004.
- [8] X. Wu, T. Nguyen, D. C. Zhang, W. Y. Wang, and A. T. Luu, "FASTopic: Pretrained transformer is a fast, adaptive, stable, and transferable topic model," arXiv, 2024, doi: 10.48550/arXiv.2405.17978.
- [9] D. Diakoulaki, G. Mavrotas, and L. Papayannakis, "Determining objective weights in multiple criteria problems: The CRITIC method," Computers & Operations Research, vol. 22, no. 7, pp. 763-770, 1995, doi: 10.1016/0305-0548(94)00059-H.
- [10] C.-L. Hwang and K. Yoon, Multiple Attribute Decision Making: Methods and Applications—A State-of-the-Art Survey. Berlin, Germany: Springer, 1981, doi: 10.1007/978-3-642-48318-9.
- [11] F. Pınarbaşı and F. Zeyrek Pınarbaşı, "Exploring smart home mobile apps market: A topic modeling approach to mobile app reviews," Trends in Business and Economics, vol. 39, no. 4, pp. 430-440, 2025, doi: 10.16951/trendbusecon.1574784.
- [12] K. Wang, A. Lei, Z. Huang, Z. Gao, Q. Ma, C. Zheng, J. Li, B. Eynard, and J. Xiao, "Online review-assisted product improvement attribute extraction and prioritization method for small- and medium-sized enterprises," Systems, vol. 13, art. 149, 2025, doi: 10.3390/systems13030149.
- [13] C. Y. Ng, Y. M. Tang, and G. T. S. Ho, "A blended sentiment analysis, large language model, and multicriteria decision-making approach to sustainable product innovation," Soft Computing, 2026, doi: 10.1007/s00500-026-11402-y.
- [14] N. Reimers and I. Gurevych, "Sentence-BERT: Sentence embeddings using Siamese BERT-networks," in Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing, pp. 3982-3992, 2019, doi: 10.18653/v1/D19-1410.