

**GRUP ÖNERİ SİSTEMLERİ İÇİN YENİ OTOMATİK GRUP TANIMLAMA
VE TERCİH BİRLEŞTİRME YÖNTEMLERİ**

Emre YALÇIN

Doktora Tezi

Bilgisayar Mühendisliği Anabilim Dalı

Bilgisayar Yazılımı Bilim Dalı

Danışman: Doç. Dr. Alper BİLGE

Eskişehir

Eskişehir Teknik Üniversitesi

Lisansüstü Eğitim Enstitüsü

Ekim 2020

ÖZET

GRUP ÖNERİ SİSTEMLERİ İÇİN YENİ OTOMATİK GRUP TANIMLAMA VE TERCİH BİRLEŞTİRME YÖNTEMLERİ

Emre YALÇIN

Bilgisayar Mühendisliği Anabilim Dalı

Bilgisayar Yazılımı Bilim Dalı

Eskişehir Teknik Üniversitesi, Lisansüstü Eğitim Enstitüsü, Ekim 2020

Danışman: Doç. Dr. Alper BİLGE

Grup öneri sistemleri, grup üyelerinin bireysel tercihlerini bir araya getirerek, bireyler yerine bir kullanıcı grubuna tercih edilebilir ürünler/hizmetler önerebilme yeteneğine sahiplerdir. Kullanıcı gruplarının çoğunlukla önceden tanımlanmaması nedeniyle bu sistemlerde ilk görev genellikle kümeleme yaklaşımları ile benzer kullanıcılardan oluşan grupları otomatik olarak tanımlamaktır. Ancak, kullanıcıları gruplara ayırmak, öneri alanındaki içerik genişledikçe seyreklik ve karmaşıklık sorunlarıyla karşı karşıya kalmaktadır. Ayrıca, grup homojenliği ve büyüklüğü, grup üyelerini organize etmek ve memnuniyetlerini arttırmak için kritik unsurlardır. Bu tezde, gelişmiş grup formasyonu ve grup büyüklüğü kısıtlaması için yeni otomatik kullanıcı gruplama yaklaşımları önerilmektedir. Bununla birlikte, kullanıcıların derecelendirmelerini, onları temsil etmek için küçük ve yoğun bir vektöre dönüştüren iki farklı janr-tabanlı eşleme stratejisi önerilmektedir. Bu stratejiler hem gruplama yapmak için gerekli hesaplama süresini iyileştirir hem de veri seyrekliğinin olumsuz etkilerini ortadan kaldırır. Son olarak, daha homojen bir grup oluşumu sağlamak için janr-tabanlı benzerliklerle birlikte demografik özelliklere dayalı benzerlikleri birlikte kullanan iki farklı strateji geliştirilmiştir. İki farklı popüler veri kümesinde yapılan deneyler, önerilen her yöntemin geleneksel rakibinden daha iyi performans sergilediğini göstermektedir.

Grup önerileri genellikle gruptaki kullanıcıların tercihlerini birleştirerek tüm grubun eğilimlerini analiz eden birleştirme teknikleri kullanılarak üretilmektedir. Literatürde çeşitli birleştirme teknikleri bulunmasına rağmen, bu tekniklerin bazı zayıf yanları mevcuttur. Bu zayıflıkları gidermek ve grup önerisi kalitesini arttırmak amacıyla,

bu tezde ilk olarak PwAvg olarak adlandırılan kişiliğe duyarlı bir birleştirme tekniği geliştirilmiştir. Bu teknik, beş temel kişilik özelliği kullanarak her bir üyenin etki derecesini belirler ve daha sonra bunları birleştirme aşamasında tercihleri ağırlıklandırmak için kullanır. Ayrıca, grup üyelerinin üzerinde fikir birliği sağladığı popüler ürünleri öne çıkarmak için iki temel birleştirme yöntemini birlikte uyum içinde kullanan iki farklı melezleştirilmiş teknik önerilmiştir. Son olarak, melezleştirilmiş tekniklerin üzerine inşa edilen ve AwU olarak adlandırılan gelişmiş bir birleştirme tekniği önerilmektedir. AwU tekniği, bilgi entropisini kullanarak grup üyelerinin derecelendirmelerinin dağılımını güçlü bir şekilde değerlendirir ve böylece maksimum sayıda memnun kişi sayısını sağlayan grup önerileri üretir. Farklı veri kümelerinde yapılan deneyler, önerilen birleştirme tekniklerinin mevcut yöntemlere kıyasla daha kaliteli grup önerileri sağladığını göstermiştir.

Anahtar Sözcükler: Grup öneri sistemleri, Otomatik grup tanımlama, Birleştirme teknikleri

ABSTRACT

NOVEL AUTOMATIC GROUP IDENTIFICATION AND PREFERENCE AGGREGATION METHODS FOR GROUP RECOMMENDER SYSTEMS

Emre YALÇIN

Department of Computer Engineering

Programme in Computer Software

Eskişehir Technical University, Institute of Graduate Programs, October 2020

Supervisor: Assoc. Prof. Dr. Alper BİLGE

Group recommender systems are specialized in suggesting preferable products/services to a group of users rather than an individual by aggregating personal preferences of group members. In these systems, the initial task is to identify groups of similar users via clustering approaches, as user groups are usually not present. However, clustering users into groups commonly suffer from sparsity and complexity problems as the content in the domain proliferate. Moreover, group homogeneity and size are the critical parameters for organizing group members and enhancing their satisfaction. In this dissertation, novel automatic user grouping approaches are proposed for enhanced group formation and group size restriction. Furthermore, two different genre-based mapping strategy that transforms user ratings into a tiny and dense vector to represent users is proposed. These strategies both improve the required computation time for grouping and eliminate adverse effects of data sparsity. Finally, to construct more homogeneous group formation, two distinct strategies that utilize the genre-based similarities and the similarities based on demographic characteristics are developed. Experiments performed on two benchmark datasets demonstrate that each proposed method outperforms its traditional rival significantly.

Group recommendations are commonly produced by utilizing aggregation techniques that analyze the propensities of the whole group by combining the preferences of the users in the group. Although there exist various aggregation techniques in the literature, they have some weaknesses. To overcome these weaknesses and improve the quality of group recommendations, In this dissertation, firstly, a personality-aware aggregation technique named as the PwAvg is developed. This technique determines the influence degree of each member in the group using five fundamental personality traits

and then utilizes them to weight the preferences during the aggregation process. Also, to feature popular items on which group members provided a consensus, two different hybridized techniques that employ two prominent aggregation methods together in harmony is proposed. Finally, an enhanced aggregation technique built on top of the hybridized techniques and named as AwU is proposed. The AwU technique robustly considers the distribution of group members' ratings by utilizing the information entropy and consequently produces group referrals that ensure the maximum number of satisfied individuals. Experiments conducted on different datasets indicate that the proposed aggregation techniques provide more-qualified group recommendations compared to existing methods.

Keywords: Group recommender systems, Automatic group identification, Aggregation techniques

TEŞEKKÜR

Öncelikle, tez çalışması sürecinde, deneyimlerini, bilgilerini ve desteğini en üst düzeyde hissettiğim çok iyi bir yol gösterici olduğuna inandığım, danışman hocam Sayın Doç. Dr. Alper BİLGE'ye katkılarından ötürü en derin teşekkürlerimi sunarım.

Çalışmaya dayanak oluşturan tez izleme süreçleri boyunca tez izleme jürisinde yer alan hocalarım Sayın Doç. Dr. Zehra KAMIŞLI ÖZTÜRK'e ve Sayın Dr. Öğr. Üyesi Fırat İSMAİLOĞLU'na verdikleri fikirler ve öneriler sayesinde yeni bakış açıları kazanmamı sağladıkları için çok teşekkür ederim.

Yaşamım boyunca benden hiçbir zaman sevgisini ve desteğini esirgemeyen, her konuda bana güvenen ve destek veren annem Mesrure YALÇIN'a ve babam Alaattin YALÇIN'a teşekkür ederim. Ayrıca, her zaman desteği ile yanımda olan değerli eşim Betül YALÇIN'a sevgilerimi sunarım.

Emre YALÇIN

ETİK İLKE VE KURALLARA UYGUNLUK BEYANNAMESİ

Bu tezin bana ait, özgün bir çalışma olduğunu; çalışmamın hazırlık, veri toplama, analiz ve bilgilerin sunumu olmak üzere tüm aşamalarında bilimsel etik ve kurallara uygun davrandığımı; bu çalışma kapsamında elde edilen tüm veri ve bilgiler için kaynak gösterdiğimi ve bu kaynaklara kaynakçada yer verdiğimi; bu çalışmanın Eskişehir Teknik Üniversitesi tarafından kullanılan “bilimsel intihal tespit programı”yla tarandığını ve hiçbir şekilde “intihal içermediğini” beyan ederim. Herhangi bir zamanda, çalışmamla ilgili yaptığım bu beyana aykırı bir durumun saptanması durumunda, ortaya çıkacak tüm ahlaki ve hukuki sonuçları kabul ettiğimi bildiririm.

Emre YALÇIN

İÇİNDEKİLER

Sayfa

BAŞLIK SAYFASI	i
JÜRİ VE ENSTİTÜ ONAYI.....	ii
ÖZET	iii
ABSTRACT.....	v
TEŞEKKÜR	vii
ETİK İLKE VE KURALLARA UYGUNLUK BEYANNAMESİ.....	viii
İÇİNDEKİLER.....	ix
TABLolar DİZİNİ	xii
ŞEKİLLER DİZİNİ	xiv
SİMGELER VE KISALTMALAR DİZİNİ.....	xv
1. GİRİŞ	1
1.1. Grup Öneri Sistemleri	2
1.1.1. Otomatik tanımlanan gruplar	3
1.1.2. Birleştirme teknikleri	5
1.2. Tezin Amacı ve Katkıları.....	7
1.3. Tez Organizasyonu	10
2. GENEL BİLGİLER.....	12
2.1. KM Kümeleme Algoritması	12
2.2. CF ile Bireysel Tahmin Üretme Süreci	13
2.3. Grup Önerisi Üretme Yaklaşımları.....	15
2.4. Birleştirme Teknikleri	17
2.4.1. Fikir birliği sağlayan teknikler	17
2.4.2. Derecelendirme sıklığını sayan teknikler	19
2.4.3. En yüksek/düşük derecelendirmeye dayanan teknikler	20
2.4.4. Sıralama önceliğine dayanan teknikler	21
2.4.5. Derecelendirmeleri karşılaştıran teknikler	22
2.4.6. Etkili grup üyelerini dikkate alan teknikler	23

2.4.7. Derecelendirmelerin dağılımını dikkate alan teknikler	24
2.5. Veri kümeleri ve Değerlendirme Ölçütü	26
3. GÖS LİTERATÜR TARAMASI	30
4. YENİ OTOMATİK KULLANICI GRUPLAMA YÖNTEMLERİ	36
4.1. Motivasyon.....	36
4.2. Kullanılan Grup Önerisi Yaklaşımının Genel Şeması.....	38
4.3. Otomatik Grup Tanımlanması için Yeni BKM-tabanlı Yaklaşımlar	39
4.3.1 BKM ile BDT oluşturma.....	39
4.3.2 JTP ile BKM yaklaşımının geliştirilmesi.....	43
4.3.3. Demografik korelasyonların JTP ile kullanımı	46
4.4. DeneySEL Çalışmalar	48
4.4.1. Kıyaslama algoritması: P&C	48
4.4.2 Deney metodolojisi	48
4.4.3. Deney sonuçları.....	49
4.4.3.1. <i>BKM yaklaşımının etkisi</i>	49
4.4.3.2. <i>JTP yaklaşımının etkisi</i>	52
4.4.3.3. <i>DTP'i JTP ile birleştirmenin etkileri</i>	55
4.5. Tartışma	58
4.6. Değerlendirme	59
5. BKM-TABANLI ALGORİTMANIN AYRINTILI ANALİZİ.....	61
5.1. Motivasyon.....	61
5.2. Kullanılan Modern CF Algoritmaları.....	61
5.3. DeneySEL Çalışmalar	63
5.3.1. Deney metodolojisi	63
5.3.2. Deney sonuçları.....	64
5.3.2.1. <i>BKM ile standart KM algoritmalarının karşılaştırılması</i>	64
5.3.2.2. <i>BKM ve KM algoritmalarının P&C ile karşılaştırılması</i>	68
5.4. Değerlendirmeler.....	71
6. KİŞİLİK ÖZELLİK-TABANLI YENİ BİR BİRLEŞTİRME TEKNİĞİ.....	72
6.1. Motivasyon.....	72
6.2. PwAvg Birleştirme Tekniğini İçeren Grup Önerisi Şeması.....	73

6.3. Deneysel Çalışmalar.....	75
6.3.1. Deney metodolojisi	75
6.3.2. Deney sonuçları ve tartışma	76
6.4. Değerlendirmeler.....	79
7. ENTROPİ İLE GÜÇLENDİRİLEN MELEZ BİRLEŞTİRME TEKNİĞİ.....	80
7.1. Motivasyon.....	80
7.2. Geliştirilen Birleştirme Tekniklerini İçeren Grup Önerisi Şeması	82
7.2.1. Melezleştirilmiş birleştirme tekniği	82
7.2.2. AwU birleştirme tekniği.....	85
7.2.3. Açıklayıcı örnek.....	88
7.3. Deneysel Çalışmalar	90
7.3.1. Karşılaştırma metotları	90
7.3.2. Deney metodolojisi	91
7.3.3. Deney sonuçları.....	92
7.3.3.1. Melezleştirilmiş tekniklerin değerlendirilmesi.....	92
7.3.3.2. AwU tekniğinin değerlendirilmesi	97
7.3.3.3. Adalet ve memnuniyet analizi	101
7.3.4. Tartışma	104
7.4. Değerlendirme	105
8. SONUÇLAR	107
KAYNAKÇA.....	109
ÖZGEÇMİŞ	

TABLÖLER DİZİNİ

Sayfa

Tablo 2.1. Örnek kullanıcı $\times$ ürün matrisi.....	12
Tablo 2.2. Bir kullanıcı grubu için örnek kullanıcı $\times$ ürün matrisi	17
Tablo 2.3. Avg, AwM, AU ve Mul teknikleri ile grup derecelendirmelerinin hesaplanması	18
Tablo 2.4. SC ve AV teknikleri ile grup derecelendirmelerinin hesaplanması.....	19
Tablo 2.5. MP ve LM teknikleri ile grup derecelendirmelerinin hesaplanması.....	20
Tablo 2.6. PV tekniği ile oluşturulan sıralı öneri listesi.....	21
Tablo 2.7. BC tekniği ile grup derecelendirmelerinin hesaplanması	22
Tablo 2.8. CR tekniği ile grup derecelendirmelerinin hesaplanması	23
Tablo 2.9. MRP tekniği ile grup derecelendirmelerinin hesaplanması.....	23
Tablo 2.10. UL tekniği ile grup derecelendirmelerinin hesaplanması	25
Tablo 2.11. Veri kümelerinin özellikleri.....	27
Tablo 4.1. Küçük bir kullanıcı $\times$ ürün matrisi.....	44
Tablo 4.2. Kitapların janrları.....	45
Tablo 4.3. Demografi-tabanlı profillerin yapısı	47
Tablo 4.4. MLP veri kümesi için KM ve BKM algoritmalarının nDCG sonuçları	50
Tablo 4.5. MLM veri kümesi için KM ve BKM algoritmalarının nDCG sonuçları	50
Tablo 4.6. MLP veri kümesi için DTP ile bezenmiş RBP yönteminin nDCG sonuçları.....	56
Tablo 4.7. MLM veri kümesi için DTP ile bezenmiş RBP yönteminin nDCG sonuçları	57
Tablo 5.1. MLP veri kümesinde KM ve BKM algoritmalarının modern CF algoritmaları için nDCG sonuçları.....	65
Tablo 5.2. MLM veri kümesinde KM ve BKM algoritmalarının modern CF algoritmaları için nDCG sonuçları.....	66
Tablo 5.3. MLP veri kümesinde P&C stratejisi ile KM ve BKM algoritmalarının modern CF algoritmaları için nDCG sonuçları.....	68
Tablo 5.4. MLM veri kümesinde P&C stratejisi ile KM ve BKM algoritmalarının modern CF algoritmaları için nDCG sonuçları.....	69
Tablo 6.1. Per5 veri kümesi için elde edilen nDCG sonuçları.....	77
Tablo 6.2. Per10 veri kümesi için elde edilen nDCG sonuçları.....	78

Tablo 7.1. Üyelerin ürün tercihlerini gösteren örnek grup.....	89
Tablo 7.2. Ürünlerin entropi ve bilgi kazancı değerleri	89
Tablo 7.3. AwU tekniği ile hesaplanan grup derecelendirmeleri.....	90
Tablo 7.4. MLP veri kümesi için elde edilen nDCG sonuçları	93
Tablo 7.5. MLM veri kümesi için elde edilen nDCG sonuçları.....	94
Tablo 7.6. NF veri kümesi için elde edilen nDCG sonuçları	95
Tablo 7.7. AwU tekniğinin karşılaştırma yöntemleri ile MLP veri kümesinde nDCG kıyaslaması	100
Tablo 7.8. AwU tekniğinin karşılaştırma yöntemleri ile MLM veri kümesinde nDCG kıyaslaması	101
Tablo 7.9. AwU tekniğinin karşılaştırma yöntemleri ile NF veri kümesinde nDCG kıyaslaması.....	101
Tablo 7.10. MLP veri kümesinde en iyi-5 grup önerileri için Fairness sonuçları	102
Tablo 7.11. MLM veri kümesinde en iyi-5 grup önerileri için Fairness sonuçları	103
Tablo 7.12. NF veri kümesinde en iyi-5 grup önerileri için Fairness sonuçları.....	103
Tablo 7.13. En iyi-5 grup önerilerinin GSM sonuçları	104

ŞEKİLLER DİZİNİ

Sayfa

Şekil 2.1. BDB yaklaşımının en iyi-N grup önerisi üretme süreci	15
Şekil 2.2. BTB yaklaşımının en iyi-N grup önerisi üretme süreci	16
Şekil 2.3. BÖLB yaklaşımının en iyi-N grup önerisi üretme süreci.....	17
Şekil 4.1. Kullanılan grup önerisi yaklaşımının genel şeması.....	38
Şekil 4.2. BKM ile oluşturulan örnek bir BDT	42
Şekil 4.3. Ali ve Pınar için janr-tabanlı profiller	45
Şekil 4.4. MLP veri kümesi için JTP'in etkisi.....	53
Şekil 4.5. MLM veri kümesi için JTP'in etkisi	54
Şekil 6.1. PwAvg birleştirme tekniğini içeren grup önerisi şeması.....	73
Şekil 7.1. Melezleştirilmiş ve AwU birleştirme tekniklerini içeren grup önerisi şeması.....	82
Şekil 7.2. MLP veri kümesinde AwU için yapılan deneylerin nDCG sonuçları.....	98
Şekil 7.3. MLM veri kümesinde AwU için yapılan deneylerin nDCG sonuçları	99
Şekil 7.4. NF veri kümesinde AwU için yapılan deneylerin nDCG sonuçları.....	100

SİMGELER VE KISALTMALAR DİZİNİ

ÖS	: Öneri Sistemleri
GÖS	: Grup Öneri Sistemleri
KM	: k -ortalama (k -means)
P&C	: Tahmin&Kümele (Predict&Cluster)
SS	: Standart Sapma
BKM	: İkiye Ayıran k -Ortalama (Bisecting k -means)
BDT	: İkili Karar Ağacı (Binary Decision Tree)
JTP	: Janr-tabanlı Profil
MUL	: Çarpımsal Strateji (Multiplicative Strategy)
AUG	: Artıran Strateji (Augmentative Strategy)
CF	: Ortak Filtreleme (Collaborative Filtering)
kNNBase	: k -En Yakın Komşu Temel (k -Nearest Neighbor Baseline)
SO	: Ağırlıklandırılmış Bir Eğilimli Tahmin Edici (Slope One)
SVD++	: Tekil Değer Ayrıştırması (Singular Value Decomposition)
PwAvg	: Kişilik-Ağırlıklı Ortalama (Personality-weighted Average)
AU	: Çoğalan Faydacı (Additive Utilitarian)
AV	: Onay Oylama (Approval Voting)
AwU	: Belirsizlik Olmadan Uzlaşma (Agreement without Uncertainty)
ACM	: Düzenlenmiş-Kosinüs Metriği (Adjusted-Cosine Metric)
PCC	: Pearson Korelasyonu Katsayısı (Pearson's Correlation Coefficient)
BDB	: Bireysel Derecelendirmeleri Birleştirmek
BTB	: Bireysel Tahminleri Birleştirmek
BÖLB	: Bireysel Öneri Listelerini Birleştirmek
Avg	: Ortalama (Average)
AwM	: Sevindiren Ortalama (Average without Misery)
Mul	: Çarpımsal (Multiplicative)
SC	: Basit Hesaplama (Simple Computation)
MP	: En Çok Memnuniyet (Most Pleasure)
LM	: En Az Mutsuzluk (Least Misery)
PV	: Çoğunluk Oylama (Plurality Voting)

BC	: Borda Sayım (Borda Count)
CR	: Copeland Kuralı (Copeland Rule)
MRP	: En Saygın Kişi (Most Respected Person)
UL	: Yukarıya Doğru Tesviye (Upward Leveling)
MLP	: MovieLens-100K
MLM	: MovieLens-1M
NF	: Netflix
Per	: Personality
n DCG	: Normalleştirilmiş İndirimli Kümülatif Kazanç (Normalized Discounted Cumulative Gain)
DCG	: İndirimli Kümülatif Kazanç (Discounted Cumulative Gain)
GSM	: Grup Memnuniyet Ölçütü (Group Satisfaction Metric)
PBP	: Alım-tabanlı Kullanıcı-janr Profili (Possession-based User-genre Profile)
RBP	: Derecelendirme-tabanlı Kullanıcı-janr Profili (Rating-based User-genre Profile)
DTP	: Demografi-tabanlı Profil

1. GİRİŞ

Son yıllarda, İnternete erişimin kolaylaşması ile bireylerin müzik dinlemek, film izlemek, haber takip etmek, alışveriş yapmak gibi çeşitli günlük aktivitelerini gerçekleştirebileceği birçok faydalı çevrimiçi sistem hizmete sunulmuştur. İnternet gelişimi ile paralel olarak sürekli çeşitlendirilen ve geliştirilen bu sistemler, her ne kadar bireylere çeşitli kolaylıklar sağlasa da bazı zorlukları beraberinde getirmektedir. Bu zorluklardan en önemlisi, İnternet kullanımının hızla artması ile depolanması ve işlenmesi gereken veri miktarının dramatik bir şekilde artmasıdır. Bu artış, aşırı bilgi yükleme sorunu olarak isimlendirilen soruna neden olarak, bireylerin ihtiyacına uygun olan ürünlere/hizmetlere erişiminin zorlaşmasına ve karar verme süreçlerinin daha karmaşık bir hale gelmesine neden olur (Adomavicius & Tuzhilin, 2005). Öneri sistemleri (ÖS), bireylerin İnternet üzerinde erişmeye çalıştıkları değerli ve ilgili bilgileri keşfedip onların karar verme sürecini destekleyerek aşırı bilgi yükleme sorunu ile başa çıkılmasına yardımcı olan yazılım araçlarıdır (Ricci, Rokach, & Shapira, 2011).

İlk olarak, 1990'lı yılların ortalarında bağımsız bir alan olarak ele alınarak, üzerinde çeşitli çalışmalar yapıldığı ÖS, günümüzde birçok çevrimiçi sistemle birlikte çeşitli amaçlar için yaygın bir biçimde kullanılmaktadır. Film (örneğin; Netflix¹), video (örneğin; YouTube²), e-ticaret (örneğin; Amazon³), müzik (örneğin; Spotify⁴), haber (örneğin; Google News⁵), konaklama (örneğin; Booking⁶) gibi alanlar, ÖS'nin etkin olarak kullanıldığı alanlara örnek olarak verilebilir. Geleneksel ÖS'nde temel amaç, kullanıcıların geçmişte çeşitli ürünler/hizmetler için yapmış olduğu derecelendirme değerlendirmelerinden yola çıkarak, onların ilgilenebilecekleri ancak henüz deneyimlemedikleri ürünleri/hizmetleri kapsayan bireysel tavsiyeler üretebilmektir (Herlocker, Konstan, Borchers, & Riedl, 2017). Bunun yanında, bireyler doğası gereği birçok aktivite için tanıdığı kişiler ile etkileşim içinde olarak beraber hareket etme eğilimi gösterirler (örneğin; iş arkadaşları ile öğle yemeği yemek, aile üyeleri ile film izlemek vb.) (Masthoff, 2011). Ayrıca, bir tur organizasyonu ile seyahat etmek (Ardissono, Goy, Petrone, Segnan, & Torasso, 2003), bir spor salonunda egzersiz yapmak (McCarthy &

¹ <https://www.netflix.com/>

² <https://www.youtube.com/>

³ <https://www.amazon.com/>

⁴ <https://www.spotify.com/>

⁵ <https://news.google.com/>

⁶ <https://www.booking.com/>

Anagnost, 1998) gibi bazı durumlar, bireyler istemese bile onların bir grupla hareket etmesini zorunlu kılar. Bu gibi senaryolarda, hedef kitle, bireysel kullanıcı yerine çeşitli nedenlerden dolayı bir araya gelen ve birlikte hareket etmek zorunda olan kullanıcıların oluşturduğu bir topluluktur. Dolayısıyla, topluluktaki bireylerin tamamını aynı anda memnun edebilecek tavsiyelerin üretilmesi için geleneksel ÖS'den daha gelişmiş öneri mekanizmalarını gerekmektedir.

1.1. Grup Öneri Sistemleri

Grup öneri sistemleri (GÖS), aynı ortamı paylaşarak beraber hareket eden bir kullanıcı topluluğunun özelliklerini ve eğilimlerini analiz ederek, onları maksimum düzeyde memnun edecek grup tavsiyelerinin üretilmesini sağlayan mekanizmalardır (Jameson & Smyth, 2007). Son yıllarda, film (O'Connor, Cosley, Konstan, & Riedl, 2001), müzik (Zhiwen, Xingshe, & Daqing, 2005) (McCarthy & Anagnost, 1998), turizm aktivitesi (McCarthy, ve diğerleri, 2006) (Jameson, 2004), restoran (McCarthy J. F., 2002) ve TV programı (Goren-Bar & Glinansky, 2004) gibi birçok alanda, kullanıcı gruplarına öneriler üretilmesine imkân tanıyan birçok GÖS geliştirilmiştir. Ayrıca, GÖS'de, kullanıcı grupları, oluşum şekillerine göre genellikle aşağıda listelenen dört farklı kategoride değerlendirilir (Boratto, Carta, Chessa, Agelli, & Clemente, 2009).

- a) *Rastgele oluşmuş gruplar*: Belirli bir anda bir ortamı tesadüfi olarak paylaşan bireyleri ifade etmektedir (örneğin; bir alışveriş merkezinde gezen insanlar) (Crossen, Budzik, & Hammond, 2002).
- b) *Tesadüfen oluşmuş gruplar*: Bir aktiviteyi kısa vadeli bir birliktelik ile tesadüfi olarak gerçekleştiren bireyleri ifade etmektedir (örneğin; spor salonunda egzersiz yapan insanlar) (O'Connor, Cosley, Konstan, & Riedl, 2001) (Quijano-Sanchez, Recio-Garcia, & Diaz-Agudo, 2011).
- c) *Önceden belirlenmiş gruplar*: Uzun vadeli birlikteliği olan bir grubun parçası olmayı tercih eden bireyleri ifade etmektedir (örneğin; bir ailenin üyeleri) (Kim, Kim, Oh, & Ryu, 2010).
- d) *Otomatik tanımlanan gruplar*: Tavsiye süreci için bir öneri sistemi tarafından getirilen çeşitli sınırlamaların üstesinden gelmek amacıyla, bireylerin tercihleri veya mevcut kaynaklar göz önünde bulundurularak otomatik olarak algılanan kullanıcı gruplarını ifade eder (Baltrunas, Makcinskas, & Ricci, 2010) (Boratto, Carta, & Fenu, 2017).

1.1.1. Otomatik tanımlanan gruplar

Yapılan birçok GÖS çalışmasında, öneri maliyeti ve içeriği gibi nedenlerden dolayı kullanıcılara bireysel öneriler sunmak yerine benzer tercih ve eğilimleri olan kullanıcıları bir grup olarak ele alıp, grup önerisinde kullanılan yaklaşımlar ile tavsiyeler üretilmektedir (Baltrunas, Makcinskias, & Ricci, 2010). Bu çalışmalarda, bir grup önerisi mekanizması oluşturmada önce gerçekleştirilen ilk işlem, kişisel tercihler bağlamında birbirine benzer bireylerden oluşan kullanıcı topluluklarını otomatik olarak belirlemektir. Bu yöntemle oluşturulan gruplar, otomatik tanımlanan gruplar olarak adlandırılmaktadır (Boratto & Carta, 2010). Kullanıcıların gruplara ayrılmasında bu tür otomatikleştirilmiş mekanizmalar kullanılması, iki temel nedenden dolayı gerekli ve faydalıdır; (i) zaman içerisinde sistemde mevcut kullanıcı sayısının artması ile manuel grupta zorlaşmaktadır ve (ii) kullanıcıları grupta, zaman içinde bireylerin tercihlerindeki değişiklikler nedeniyle düzenli güncellemeler gerektiren sürekli bir süreçtir (Khazaei & Alimohammadi, 2018).

Otomatik tanımlanan grupları kullanan GÖS'ün başarısı, rastgele düşünülen kitlelerden ziyade oluşturulan grupların üyelerinin ne kadar uyumlu olduğu ile güçlü bir şekilde ilişkilidir; çünkü benzer zevkleri olan insan topluluklarını tatmin etmek sezgisel olarak daha kolaydır. Uyumlu kullanıcı gruplarını tanımlamak için kullanılan en yaygın yöntem, herhangi bir kısıtlama gözetmeden kullanıcıların bir kümeleme algoritması kullanılarak gruplara ayrılmasıdır (Boratto, Carta, & Fenu, 2016). Kümeleme işlemini gerçekleştirmek için genellikle kullanıcıların ürünler için yapmış oldukları derecelendirme değerlerinden oluşan kullanıcı profilleri göz önünde bulundurulur. Bu kapsamda, *k*-ortalama (*k*-means-KM), verimliliği ve kolay uygulanması nedeniyle grupların otomatik olarak tanımlanmasında en çok kullanılan kümeleme algoritmasıdır (Amatriain vd., 2011; Boratto, Carta, & Fenu, 2017)). Bununla birlikte, son yıllarda geliştirilen Tahmin&Kümele (Predict&Cluster-P&C) yaklaşımı, bu amaçla kullanılan ve en öne çıkan stratejidir (Boratto, Carta, & Satta, 2010). Bu strateji, başlangıçta bireylerin eksik derecelendirmelerini geleneksel bir öneri algoritması ile tahmin eden ve daha sonra kullanıcı gruplarını tanımlamak için hem kullanıcıların sağladığı derecelendirmeleri hem de tahmin edilen derecelendirme değerlerini içeren doldurulmuş kullanıcı $\times$ ürün matrisi üzerinde KM algoritmasını uygulayan bir yaklaşımdır.

Grupların otomatik olarak tanımlanmasında geleneksel kümeleme yaklaşımlarının kullanılması, yaptıkları tercihler bağlamında benzer kullanıcıları belirlemede etkin bir yöntem olmasına rağmen, çıktı olarak üretilen grupların içerdiği kullanıcı sayıları genellikle ciddi farklılıklar gösterebilir. Bununla birlikte, bazı GÖS uygulamaları, verilen hizmetin özelliğine göre bir grupta olabilecek maksimum kullanıcı sayısı konusunda bir kısıtlamaya sahip olabilir (Khazaei & Alimohammadi, 2018). Dolayısıyla, grupların geleneksel kümeleme yöntemleriyle tanımlanması, farklı gruplarda değişen performanslara sahip olan ve her grup önerisi senaryosuna her zaman uygulanamayan GÖS'e yol açar.

Geleneksel ÖS'de olduğu gibi, bireylerin kümeleme algoritmaları ile gruplara ayrılması süreci de kullanıcı profillerinde çok fazla kişi tarafından derecelendirilmemiş ürüne sahip olmanın neden olduğu veri seyrekliği sorunundan muzdariptir (Boratto vd., 2009; Jeong & Kim, 2019). Bu durum, kullanıcıların derecelendirme vektörleriyle zayıf bir şekilde temsil edilmesine ve kümelemedeki en kritik aşamalardan biri olan kullanıcılar arasındaki benzerliklerin doğru bir şekilde belirlenmesinde başarısızlığa yol açar. Bununla birlikte, sistemde mevcut olan içerik sayısının artması, kullanıcılar arasındaki benzerliklerin hesaplanması için gereken hesaplama süresini ciddi oranda arttırarak karmaşıklık problemi olarak da adlandırılabilen soruna neden olur (Bilge & Polat, 2013). Ayrıca, içerik sayısının artması, kişilere ait derecelendirme vektörlerinin uzayda birbirinden uzaklaşmasına neden olur. Bu durum, boyutun laneti olarak da adlandırılan probleme neden olarak vektörler arasındaki uzaklıkların doğru bir biçimde hesaplanamamasına yol açar (Boratto & Carta, 2014).

Kullanıcı gruplarını otomatik tanımlanmasında, grup üyelerinin bireysel derecelendirmelerini kullanmak makul bir yöntem gibi görünse de, kullanıcı derecelendirmeleri, bireylerin profillerini doğru bir şekilde yansıtmak için tek başlarına yeterli değildir. Kullanıcılar, ürünlere yapmış oldukları derecelendirmeler bakımından benzer olsa da, cinsiyet, yaş aralığı ve meslek gibi demografik özellikler açısından oldukça farklı olabilirler (Sridevi & Rao, 2017; Vozalis & Margaritis, 2004). Dolayısıyla, kullanıcıları, sadece derecelendirme vektörlerini göz önünde bulundurarak gruplara ayırmak, farklı demografik özelliklere sahip bireylerin aynı grupta tanımlanmasına ve bu da grup homojenliğine zarar vererek üretilen önerilerden genel memnuniyetin azalmasına neden olabilir.

1.1.2. Birleştirme teknikleri

Grup önerilerinde temel amaç, grubu oluşturan üyelerin tercih ve eğilimlerini bir bütün olarak ele alarak, bütün üyeleri maksimum derecede memnun edebilecek ürünlerin veya hizmetlerin belirlenmesidir. Bu öneriler genellikle grup üyelerinin ürünler/hizmetler için yapmış oldukları derecelendirme değerlerinin, birleştirme teknikleri olarak adlandırılan yöntemler kullanılarak bir araya getirilmesi ile oluşturulan ve grubun ilgili ürünlerdeki/hizmetlerdeki tercihlerini yansıtan grup profilleri göz önünde bulundurularak üretilir (Boratto, Carta, & Fenu, 2016). Bir grup öneri mekanizmasında kullanılan birleştirme tekniği, uygun grup tercihlerinin belirlenmesinde kritik bir role sahiptir ve üretilen grup önerilerinin kalitesini doğrudan etkilerler (Seo vd., 2018).

Son yıllarda, grup önerileri sağlamak için birçok birleştirme tekniği geliştirilmiştir. Bu tekniklerin büyük çoğunluğu genellikle, grup üyeleri tarafından sağlanan derecelendirme değerlerinin ortalaması (Jameson, 2004; McCarthy vd., 2006), sayısı (Lieberman, Van Dyke, & Vivacqua, 1999; McCarthy J. F., 2002), sıralaması (Boratto, Carta, & Fenu, 2016; Marquez & Ziegler, 2016; Masthoff, 2011), en yükseği (Boratto, Carta, & Fenu, 2016) veya en düşüğü (Christensen & Schiaffino, 2011; O'connor vd., 2001; Marquez & Ziegler, 2016) gibi yalnızca tek bir yönü göz önünde bulundurarak birleştirme sürecini gerçekleştirir. Her ne kadar bu teknikler, makul grup önerilerinin üretilmesine yardımcı olsa da, özellikle grubu oluşturan üyelerin tercih bağlamında farklılaşması ve çeşitlenmesi durumunda, grup eğilimlerini doğru bir şekilde belirlemekte tek başlarına yetersiz kalmaktadırlar. Diğer bir deyişle, tercih bağlamında homojen olmayan gruplarda, grup profili oluşturmak için tek bir birleştirme tekniği kullanmak, grup tercihlerinin doğru bir şekilde belirlenmesinde ve sonuç olarak uygun grup önerilerinin üretilmesinde yetersiz kalmaktadır.

Grup önerilerinde temel amaç, grup içerisindeki tek bir birey yerine, bütün grup için uygun olabilecek ürünleri belirleyerek genel memnuniyeti arttırmaktır. Dolayısıyla, öneri sonuçlarından memnun olan üyelerin sayısını en düzeye çıkarmak, bireyler için uygun ürünlerin belirlenmesi kadar önemli bir unsurdur. Bu doğrultuda, birleştirme sürecini gerçekleştirirken, grup üyelerinin üzerinde belirli bir uzlaşma sağlamadığı ürünlerin belirlenerek öneri olarak üretilmemesi gerekir. Üzerinde uzlaşma sağlanamayan ürünlerin belirlenmesi için, birleştirme aşamasında üyeler tarafından sağlanan değerlendirme derecelendirmelerinin dağılımını göz önünde bulundurmak

gerekmektedir. Örneğin, bir ürün için grup üyelerinin sağladığı derecelendirmelerin farklı değerleri benzer frekanslara sahipse (yani, eşit olarak dağılmışsa), tüm grup üyelerinin söz konusu ürün üzerinde aynı fikirde olduğunu iddia etmek zordur

Standard sapma (SS), grup üyelerinin derecelendirmelerinin birleştirilmesi sürecinde derecelendirmelerin dağılımını analiz etmek için kullanılan etkin bir yöntemdir (Seo vd., 2018). Eğer, bir ürün için üyeler tarafından sağlanan derecelendirmelerin sapması yüksek hesaplanırsa, üyelerin ilgili ürün üzerinde bir uzlaşma sağlayamadığı kabul edilir. Her ne kadar bu yaklaşım makul sonuçların elde edilmesine imkan sağlasa da, bazı ekstrem senaryolarda yetersiz kalmaktadır. SS, doğası gereği sürekli tipteki değişkenlerin dağılımını analiz etmede etkin bir yöntemdir. Ancak, derecelendirmeler ile temsil eden kullanıcı tercihleri istisnalar dışında genellikle ayırık tiptedir. Bu durum, derecelendirmelerin $\{0, 1\}$ veya $\{1, 5\}$ gibi dar bir yelpazede değişmesi durumunda daha fazla ortaya çıkmaktadır. Bununla birlikte, on yıldızlı bir derecelendirme sisteminde bile, SS, derecelendirmelerin dağılımının genel davranışını yakalayamayabilir; çünkü derecelendirme yelpazesinin geniş olması durumunda birden fazla pik içeren bir dağılım (yani çok-doruklu dağılım) oluşması muhtemeldir. Bu sorun, büyük ve tercih bağlamında heterojen olan grupların varlığında daha belirgin hale gelir, çünkü benzer tercihlere sahip farklı kullanıcı alt gruplarının oluşma ihtimali yüksektir.

Mevcut birleştirme teknikleri genellikle gruptaki bireylerin psikolojik yönlerini göz önünde bulundurmadan yalnızca sağladıkları derecelendirmeleri dikkate alarak grup profillerini oluştururlar. Ancak, karar verme sürecinin, bireylerin duygusal ve kişisel özellikleri ile güçlü bir şekilde ilişkili olduğu bilinen bir olgudur, çünkü benzer kişisel özelliklere sahip kullanıcıların benzer tercihlere sahip olma olasılığı daha yüksektir (Hu & Pu, 2010). Ayrıca, mevcut tekniklerin çoğu, gruptaki her bir üyenin grubun nihai kararında aynı oranda etkili olduğu varsayımına dayanmaktadır. Ancak, sezgisel olarak, bireyler, grubun karar verme sürecinde farklı derecelerde etkilere sahip olabilir ve bu etki dereceleri onların kişilik özelliklerine bağlı olabilir (Rossi & Cervone, 2016). Örneğin, gruptaki bazı kişiler, çıkarlarını korumak için veya kendi seçimlerinin grubun her üyesi için ideal kararlar olduğuna inandıkları için düşüncelerinde ısrarcı olabilirken, bazı bireyler, diğer kullanıcıların memnuniyetini önemseyerek, grupta bir fikir birliğine ulaşılması için kişisel çıkarları pahasına seçimlerinde değişikliğe gidebilirler.

Dolayısıyla, kullanıcıların bu tür davranışsal ve kişilik özellikleri, birleştirme sürecinde dikkate alınması gereken önemli unsurlardır.

1.2. Tezin Amacı ve Katkıları

Bu tez çalışmasında, önceki bölümlerde belirtilen grupların otomatik olarak tanımlanması ve bireysel tercihlerin birleştirilmesine ilişkin sorunlara çeşitli çözümler önerilmiştir. Tezin literatüre kazandırdığı çözümler ve katkılar, aşağıda listelenen ve iki başlıkta ele alınan araştırma problemlerine cevap olma niteliği taşımaktadır.

Otomatik tanımlanan gruplar bağlamında:

- a) Tanımlanan grupların içerebileceği maksimum üye sayısı ile ilgili kısıtlamaya sahip kümeleme tabanlı bir kullanıcı gruplama yaklaşımı nasıl geliştirilir?
- b) Geliştirilecek gruplama yaklaşımında, kullanıcılar arasındaki benzerliklerin doğru bir şekilde belirlenmesini engelleyen kullanıcı profillerindeki seyreklik sorununa ve artan içerik sayısının neden olduğu yüksek hesaplama süresi ve boyutun laneti problemlerine nasıl bir çözüm getirilebilir?
- c) Geliştirilecek gruplama yaklaşımında, daha homojen gruplar tanımlayabilmek için hem kullanıcıların demografik bilgileri hem de ürünler için sağladıkları derecelendirmeler, kullanıcıların birbirlerine benzerliklerinin belirlenmesi aşamasında nasıl birlikte kullanılabilir?
- d) Geliştirilecek gruplama yaklaşımı ile birlikte modern tahmin algoritmalarının kullanıldığı yüksek kalitede grup önerilerinin üretilmesini sağlayan mekanizmalar nasıl geliştirilir?

Birleştirme teknikleri bağlamında:

- a) Grup üyelerinin sağladığı derecelendirmelere ek olarak, onların kişilik özelliklerini birleştirme aşamasında göz önünde bulunduran bir birleştirme tekniği nasıl geliştirilir?
- b) Birden fazla birleştirme tekniğinin uyum içerisinde birlikte kullanıldığı melezlenmiş birleştirme yaklaşımları nasıl geliştirilebilir?
- c) Birleştirme aşamasında, grup üyelerinin derecelendirmelerinin dağılımını analiz edebilmek için SS'dan farklı olarak yeni bir yaklaşım geliştirilebilir mi? Geliştirilecek yaklaşımı, melezlenmiş teknikler ile birlikte kullanarak, yüksek kalitede grup önerilerinin üretilmesini sağlayan gelişmiş bir birleşme tekniği geliştirilebilir mi?

Yukarıda belirtilen araştırma problemlerine çözüm geliştirmek için hazırlanan bu tez çalışmasının, grupların otomatik olarak tanımlanması bağlamındaki ilk ve en önemli katkısı, grupların maksimum boyutu ile ilgili önceden tanımlanmış bir kısıtlamayı göz önünde bulunduran ikiye ayıran k -ortalama (Bisecting k -means-BKM) temelli yaklaşımıdır. Önerilen bu yaklaşım, kullanıcıların derecelendirme vektörleri üzerinde BKM algoritmasını uygulayarak bir İkili Karar Ağacı (Binary Decision Tree-BDT) oluşturur. Ardından, bu ağacı hem kullanıcı gruplarını tanımlamak hem de sisteme sonradan katılan bir kullanıcının ait olabileceği grubu bulmak üzere kullanır. İki farklı veri kümesinde yapılan çeşitli deneyler, geliştirilen BKM yaklaşımının standart KM algoritmasına kıyasla grup öneri kalitesi açısından daha iyi performans sağladığını göstermektedir.

Tezin, grupların otomatik olarak tanımlanması bağlamındaki ikinci katkısı, BKM yaklaşımındaki benzerlik hesaplama sürecine ilişkin seyreklik ve karmaşıklık sorunlarının etkilerini hafifletmek amacıyla iki farklı Janr-tabanlı Profil (JTP) oluşturma stratejisinin kullanılmasıdır. Bu stratejiler, kümelemenin ayırma becerilerini güçlendirmek amacıyla kullanıcıların derecelendirme vektörlerini, derecelendirilen ürünlerin janrlarını göz önünde bulundurarak JTP'lere dönüştürmeyi amaçlamaktadır. Spesifik olarak, JTP oluşturma işlemi, ürünlerin ait olduğu janr kategorilerini dikkate alarak, kümeleme sürecinde kullanmak üzere tamamen yoğun ve derecelendirme vektörlerine kıyasla çok daha düşük boyuttaki kullanıcı profilleri üretilmesini sağlar. Ayrıca, bu işlem ile üretilen profillerin boyutu sabittir ve sistemdeki ürünlerin mevcut janr sayısına bağlıdır; böylece, benzerliklerin hesaplanması için gereken süre önemli ölçüde azaltılır ve stabil hale gelir. Gerçekleştirilen deneylerin sonuçları bu tür JTP'lerin kullanılmasının, üretilen grup önerilerinden olan memnuniyetin artmasına katkıda bulunduğunu göstermiştir.

Tez çalışmasının grupların otomatik olarak tanımlanması bağlamındaki üçüncü katkısı, geliştirilen BKM yaklaşımında kullanıcılar arasındaki ilişkileri daha doğru belirlemek için demografi-tabanlı korelasyonları benzerlik hesaplama sürecine dâhil eden Çarpımsal Strateji (Multiplicative Strategy-MUL) ve Artıran Strateji (Augmentative Strategy-AUG) olarak isimlendirilen iki farklı stratejidir. Bu stratejiler, JTP benzerlikleri ile kullanıcılar arasındaki demografik benzerlikleri aynı anda dikkate alarak nihai benzerlikleri belirleyen yaklaşımlardır. Spesifik olarak, MUL stratejisi, iki benzerliği eşit

olarak ağırlıklandırarak birleştirirken; AUG stratejisi, demografik benzerliklere göre janr-tabanlı benzerlikleri daha fazla önemseyerek nihai benzerlikleri hesaplar. Deneysel sonuçlar, otomatik grup tanımlama süreci boyunca bu stratejilerden herhangi birinin kullanılmasının sistemin doğruluğunu önemli ölçüde arttırdığını ve MUL stratejisinin büyük gruplar için, AUG stratejisinin ise küçük ve orta ölçekteki gruplar için daha etkili olduğunu göstermiştir. Ayrıca, sonuçlar her iki stratejinin de uygun kullanıcı gruplarını belirlemede, literatürde öne çıkan P&C stratejisine kıyasla daha başarılı olduğunu göstermektedir. Sonuç olarak, kullanıcılarının demografik bilgilerinden faydalanarak, daha homojen grupların tanımlanması sağlanmış ve bireysel üyelerin genel memnuniyeti arttırılmıştır.

Tez çalışmasının grupların otomatik olarak tanımlanması bağlamındaki son katkısı, geliştirilen BKM yaklaşımının performansının üç farklı Ortak Filtreleme (Collaborative Filtering-CF) algoritması ile ayrıntılı biçimde analiz edilmesi ve grup önerisi doğruluğunun arttırılmasıdır. Bu yapılırken, grup önerilerinin üretilmesi amacıyla geleneksel CF algoritması yerine, k -En Yakın Komşu Temel (k -Nearest Neighbor Baseline-kNNBase) (Koren, 2010), Ağırlıklandırılmış Bir Eğilimli Tahmin Edici (Slope One-SO) (Lemire & Maclachlan, 2005) ve Tekil Değer Ayrıştırması (Singular Value Decomposition-SVD++) (Koren, 2008) olmak üzere olmak üzere üç farklı modern CF algoritması kullanılmıştır. Gerçekleştirilen deneylerde, BKM yaklaşımının standart KM algoritmasına göre üstünlüğü hem orijinal kullanıcı $\times$ ürün matrisinin hem de P&C stratejisi ile doldurulmuş kullanıcı $\times$ ürün matrisinin kullanıldığı durumda ortaya konmuştur. Ayrıca, BKM yaklaşımı ile üretilen grup önerilerinin doğruluğu SVD++ algoritması aracılığı ile anlamlı biçimde arttırılmıştır.

Bu tez çalışmasının birleştirme teknikleri bağlamındaki ilk katkısı, grup derecelendirmelerinin hesaplanması sürecinde grup üyelerinin derecelendirmelerine ek olarak onların kişilik özelliklerini de dikkate alan Kişilik-Ağırlıklı Ortalama (Personality-weighted Average-PwAvg) birleştirme tekniğidir. Spesifik olarak bu teknik, *açıklık*, *uyumluluk*, *duygusal istikrar*, *vicdan* ve *dışadönüklük* olmak üzere beş temel kişilik özelliğe göre gruptaki bireylerin grubun nihai kararı üzerindeki etkisini ölçer ve bunları daha sonra birleştirme aşamasında derecelendirmeleri ağırlıklandırmak için kullanır. İki farklı veri kümesinde yapılan deneyler, önerilen tekniğin özellikle büyük kullanıcı

grupları için, kıyaslama tekniği olarak seçilen üç farklı birleştirme tekniğinden daha iyi performans sergilediğini göstermiştir.

Tez çalışmasının birleştirme teknikleri bağlamındaki ikinci katkısı, grup üyeleri tarafından sağlanan derecelendirmelerin farklı yönlerini ele almak amacıyla geliştirilen ve AU_{AV} ve AV_{AU} olarak adlandırılan melez birleştirme teknikleridir. Spesifik olarak, bu teknikler, grup üyelerinin üzerinde fikir birliği sağladığı ve grup içerisinde popüler olan ürünleri Çoğalan Faydacı (Additive Utilitarian-AU) ve Onay Oylama (Approval Voting-AV) birleştirme tekniklerini bir arada kullanarak belirler. Geliştirilen melez tekniklerden AU_{AV} , nihai grup derecelendirmelerini hesaplarken belirleyici faktör olarak AU tekniğini kullanırken, AV_{AU} yönteminde AV birleştirme tekniği daha etkindir. Gerçekleştirilen deneysel çalışmaların sonuçları, özellikle AU_{AV} tekniğinin, grup öneri kalitesi açısından tüm temel birleştirme tekniklerinden daha iyi sonuçlar elde ettiğini göstermiştir.

Tez çalışmasının birleştirme teknikleri bağlamındaki son ve en önemli katkısı, birleştirme sürecinde, sağlanan derecelendirmelerin dağılımını dikkate alan ve geliştirilen melez tipteki birleştirme tekniklerin üzerine inşa edilen Belirsizlik Olmadan Uzlaşma (Agreement without Uncertainty-AwU) birleştirme tekniğidir. Spesifik olarak, AwU, grup üyeleri tarafından sağlanan fikir birliğinin derecesini ölçmek için, derecelendirmelerin dağılımını entropi kullanarak analiz eden bir tekniktir. Entropi hesaplaması ile AwU tekniği, üzerinde uzlaşma sağlanamayan ürünleri göz ardı ederek yüksek kalitede grup önerilerinin üretilmesini sağlamaktadır. Deneysel çalışmalarda elde edilen sonuçlar, grup derecelendirmeleri hesaplanırken, AwU tekniğinin kullanılmasının, geliştirilen melez tekniklerden ve diğer tüm temel birleştirme tekniklerinden daha yüksek doğrulukta grup önerilerinin üretilmesini sağladığını göstermiştir.

1.3. Tez Organizasyonu

Bu tez, organizasyon açısından toplam sekiz bölümden oluşmaktadır. İkinci bölümde grupların otomatik tanımlanmasında yaygın kullanılan KM algoritması, CF yöntemi ile bireysel tahminlerin üretilme süreci, grup önerilerinin üretilmesinde kullanılan temel yaklaşımlar, mevcut birleştirme teknikleri, kullanılan veri setleri ve değerlendirme ölçütleri hakkında genel bilgiler verilmektedir. Sonraki bölümde, GÖS bağlamında literatürde öne çıkan çalışmalar ayrıntılı olarak açıklanmıştır. Dördüncü bölümde grupların otomatik olarak tanımlanması amacıyla geliştirilen BKM yaklaşımı ile bu yaklaşımın JTP tabanlı benzerliklerle ve demografik korelasyonlarla geliştirildiği

yöntemler yer almaktadır. Sonraki bölümde, BKM yaklaşımının grup öneri kalitesi bağlamında literatürde öne çıkan yaklaşımlara olan üstünlüğü modern CF algoritmaları yardımıyla ayrıntılı biçimde analiz edilmiştir. Altıncı bölümde, birleştirme sürecinde grup üyelerinin kişilik özelliklerini dikkate alan PwAvg birleştirme tekniği sunulmaktadır. Yedinci bölümde, bu tez çalışmasının en önemli katkılarından olan melezleştirilmiş ve AwU birleştirme teknikleri açıklanmış ve bu tekniklerin performansları öne çıkan birleştirme teknikleri ile kıyaslanmıştır. Son bölümde ise elde edilen sonuçlarla ilgili genel değerlendirmeler sunulmuştur.



2. GENEL BİLGİLER

Bu bölümde, GÖS’de otomatik tanımlanan gruplar ve grup önerilerinin üretilme sürecine dair genel bilgiler açıklanmaktadır. Öncelikle; bir önceki bölümde ifade edilen, benzer tercihleri olan kullanıcılardan oluşan grupların otomatik tanımlanma sürecinde en çok kullanılan kümeleme algoritmalarından birisi olan KM kümeleme algoritması açıklanmıştır. Ardından, grup önerilerinin üretilmesinde kullanılan üç temel yaklaşım açıklanmıştır. Sonrasında, bireysel tercihlerin veya üretilen bireysel tahminlerin birleştirilerek grup profillerinin belirlenmesi için kullanılan mevcut birleştirme teknikleri ayrıntılı bir biçimde açıklanmıştır. Son olarak, tez kapsamında geliştirilen yaklaşımların performanslarının analizi için kullanılan veri setleri ve değerlendirme ölçütleri açıklanmıştır.

2.1. KM Kümeleme Algoritması

Grupların otomatik tanımlandığı GÖS’de, üretilecek grup önerileri ile grup üyelerini azami ölçüde memnun etmek için tercih bağlamında birbirine benzer kullanıcılardan oluşan grupları uygun bir şekilde tanımlamak çok önemlidir. Bu süreç, bir kümeleme problemi olarak ele alınıp, başarılı bir kümeleme algoritması kullanılarak gerçekleştirilebilir. Bu amaçla, bir önceki bölümde belirtildiği gibi, basitliği ve etkinliği nedeniyle KM algoritması en yaygın kullanılan kümeleme algoritmasıdır (Boratto & Carta, 2010).

Tipik bir öneri sisteminde, kullanıcıların ürünler için yapmış olduğu değerlendirmeler Tablo 2.1’de bir örneğinin sunulduğu bir kullanıcı $\times$ ürün matrisi aracılığı ile saklanmaktadır. Bu tabloda, m adet kullanıcı ve n adet ürün için, kullanıcıların değerlendirme vektörleri $m \times n$ boyutunda bir matris ile gösterilir. Ayrıca, $d_{k\bar{u}}$ ($k = 1, 2, \dots, m$ ve $\bar{u} = 1, 2, \dots, n$), k kullanıcısının $\bar{u}$ ürünü için yapmış olduğu değerlendirmeyi ve (-), değerlendirilmemiş ürünleri temsil etmektedir.

Tablo 2.1. Örnek kullanıcı $\times$ ürün matrisi

	$\bar{u}_1$	$\bar{u}_2$	$\bar{u}_3$	$\dots$	$\bar{u}_n$
k_1	d_{11}	-	d_{13}		-
k_2	-	d_{22}	-		d_{2n}
k_3	-	d_{32}	d_{33}		-
$\vdots$					
k_m	d_{m1}	d_{m2}	-		d_{mn}

KM kümeleme algoritması, kullanıcıların değerlendirme vektörlerini içeren ve Tablo 2.1’de bir örneğinin sunulduğu kullanıcı $\times$ ürün matrisi üzerinde aşağıda listelenen işlem adımlarını gerçekleştirerek, tercih bağlamında benzer kullanıcıları içeren k adet grubun tanımlanmasını gerçekleştirir.

- a) Kullanıcı vektörlerini göz önünde bulundurarak, başlangıçta, rastgele k adet kullanıcı küme merkezleri olarak seçilir.
- b) Her kullanıcı, kendisine en çok benzeyen (en yakın) küme merkezine atanır. Kullanıcı vektörü ile küme merkezleri arasındaki benzerlikler, ÖS literatüründe en yaygın kullanılan ve iki vektör arasındaki benzerliği en doğru şekilde belirleyen ölçütlerden birisi olduğu varsayılan Düzenlenmiş-Kosinüs Metriği (Adjusted-Cosine Metric-ACM) yardımıyla hesaplanır (Schafer vd., 2007).
- c) Kullanıcıların atanması ile birlikte değişen küme merkezleri güncellenir.
- d) Küme merkezleri değişmeye kadar, b ve c adımları tekrarlanır.

ACM metriği ile iki kullanıcı (a ve u) arasındaki benzerlik, Denklem 2.1’de verilen formül yardımıyla hesaplanır.

$$ACM(a, u) = \frac{(\sum_{\tilde{u} \in \tilde{U}} (d_{a,\tilde{u}} - \bar{d}_{\tilde{u}}) \times (d_{u,\tilde{u}} - \bar{d}_{\tilde{u}}))}{\sqrt{\sum_{\tilde{u} \in \tilde{U}} (d_{a,\tilde{u}} - \bar{d}_{\tilde{u}})^2} \sqrt{\sum_{\tilde{u} \in \tilde{U}} (d_{u,\tilde{u}} - \bar{d}_{\tilde{u}})^2}} \quad (2.1)$$

Burada, $\tilde{U}$, a ve u kullanıcıları arasında ortak olarak derecelendirilmiş ürün kümesini temsil ederken, $\tilde{u} \in \tilde{U}$ olmak üzere, a ve u kullanıcılarının $\tilde{u}$ için sağladıkları derecelendirme değerleri sırasıyla $d_{a,\tilde{u}}$ ve $d_{u,\tilde{u}}$ ile gösterilmiştir. Ayrıca, $\bar{d}_{\tilde{u}}$, $\tilde{u}$ ürünü için tüm kullanıcılar tarafından sağlanan derecelendirmelerin ortalamasını ifade etmektedir.

2.2. CF ile Bireysel Tahmin Üretim Süreci

CF teknikleri, öneri üretim sürecinde oldukça başarılı sonuçlar elde edilmesini sağlayan ve birçok çevrimiçi sistem tarafından kullanılan bilgi filtreleme yaklaşımlarıdır (Adomavicius & Tuzhilin, 2005). Bu teknikler, bireylerin önceden deneyimledikleri çeşitli ürünlere/hizmetlere verdikleri derecelendirme değerlerini kullanarak, onların ilgisini çekebilecek ancak henüz deneyimlemedikleri ürünler/hizmetler hakkında tahmin değerlerinin üretilmesini sağlarlar. CF tekniklerinin dayandığı temel varsayım, geçmişte benzer tercihler yapan kullanıcıların gelecekte de benzer tercihler yapma eğiliminde olduğu varsayımdır. CF teknikleri, öneri üretim sürecini gerçekleştirirken kullandığı

yaklaşımlar temelinde hafıza- tabanlı, model-tabanlı ve bunların birlikte kullanıldığı melez yaklaşımlar olarak üç ana başlık altında değerlendirilir (Schafer vd., 2007).

Hafıza-tabanlı yaklaşımlar, CF tekniklerinin ortaya çıktığı zamandan günümüze kadar yaygın bir biçimde kullanılan, yüksek doğrulukta tahminlerin üretilmesini sağlayan yaklaşımlardır. Bu yaklaşımlar, Tablo 2.1’de bir örneğinin sunulduğu kullanıcı $\times$ ürün matrisi üzerinde işlem yaparak, kullanıcıların veya ürünlerin arasındaki korelasyonlar temelinde seçilen komşulukları kullanarak öneri üretme sürecini gerçekleştirirler. Tipik bir komşuluğa dayalı hafıza-tabanlı CF yaklaşımında, bir aktif kullanıcı (a), geçmişte deneyimlediği ürünler için sağladığı derecelendirmeleri servis sağlayıcısı ile paylaştıktan sonra, bir hedef ürün (q) için öneri talep eder. Bu doğrultuda, öneri üretme süreci aşağıda listelenen iki temel adım adımda gerçekleştirilir.

- a) *Benzerliklerin belirlenmesi*: Komşuluk-tabanlı geleneksel CF yaklaşımlarında, sistem ilk olarak, a kullanıcısı ile kendisine en benzer derecelendirme profiline sahip kullanıcılardan oluşan komşu kullanıcıları belirler. Kullanıcı profilleri arasındaki benzerlikler, birçok yöntem ile belirlenebilirken, Pearson Korelasyonu Katsayısı (Pearson’s Correlation Coefficient-PCC), CF çalışmalarında en çok kullanılan ve öneri üretme doğruluğu bakımından literatürde kabul görmüş bir yöntemdir (Choi & Suh, 2013). PCC ile iki kullanıcı (a ve u) arasındaki benzerlik oranı, Denklem 2.2’de verilen formül yardımıyla hesaplanır.

$$PCC(a, u) = \frac{(\sum_{\tilde{u} \in \tilde{U}} (d_{a,\tilde{u}} - \bar{d}_a) \times (d_{u,\tilde{u}} - \bar{d}_u))}{\sqrt{\sum_{\tilde{u} \in \tilde{U}} (d_{a,\tilde{u}} - \bar{d}_a)^2} \sqrt{\sum_{\tilde{u} \in \tilde{U}} (d_{u,\tilde{u}} - \bar{d}_u)^2}} \quad (2.2)$$

Burada, $\tilde{U}$, a ve u kullanıcıları arasında ortak olarak derecelendirilmiş ürün kümesini temsil ederken, $\tilde{u} \in \tilde{U}$ olmak üzere, a ve u kullanıcılarının $\tilde{u}$ için sağladıkları derecelendirme değerleri sırasıyla $d_{a,\tilde{u}}$ ve $d_{u,\tilde{u}}$ ile gösterilmiştir. Ayrıca $\bar{d}_a$ ve $\bar{d}_u$, sırasıyla a ve u kullanıcılarının ortalama derecelendirme değerlerini temsil etmektedir.

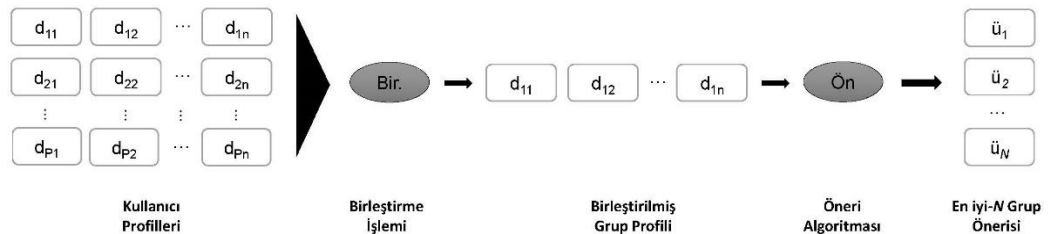
- b) *Tahminlerin hesaplanması*: Denklem 2.2 aracılığı ile a kullanıcısı ile diğer kullanıcılar arasındaki benzerlikler hesaplandıktan sonra, en benzer K tane kullanıcı, a kullanıcısına komşu olarak seçilir. Ardından, öneri istenilen q ürünü için tahmin değeri ($p_{a,q}$), Denklem 2.3’de sunulan formül yardımıyla hesaplanır.

$$p_{a,q} = \bar{d}_a + \frac{\sum_{u \in K} ((d_{u,q} - \bar{d}_u) PCC(a, u))}{\sum_{u \in K} |PCC(a, u)|} \quad (2.3)$$

2.3. Grup Önerisi Üretme Yaklaşımları

GÖS’de, genellikle, bir grup kullanıcı için uygun ürünleri içeren bir öneri listesi, aşağıda ayrıntılı olarak açıklanan üç temel grup öneri stratejisinden birisi kullanılarak üretilmektedir.

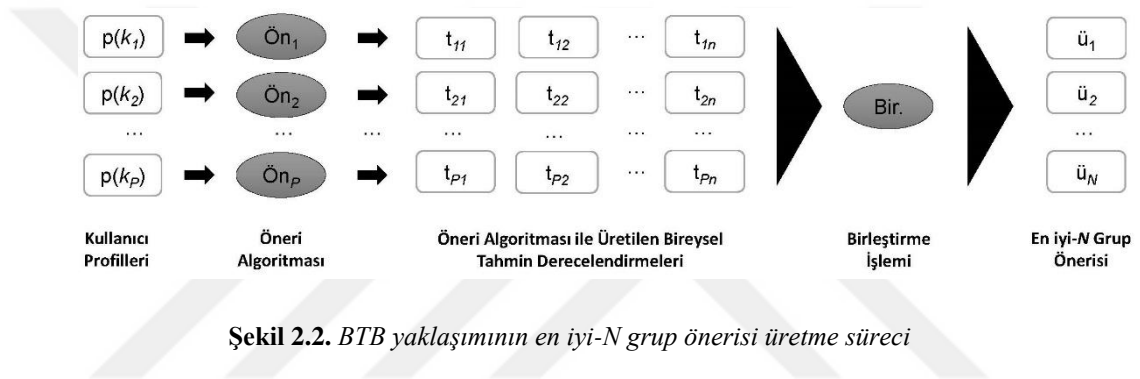
a) *Bireysel Derecelendirmeleri Birleştirmek (BDB)*: Bu yaklaşım, temel olarak, gruptaki bireylerin ürünler için sağladığı derecelendirmelerin bir araya getirilmesine dayanmaktadır. INTRIGUE (Ardissono vd., 2003), Travel Decision Forum (Jameson, 2004) ve Let’s Browse (Lieberman, Van Dyke, & Vivacqua, 1999) gibi birçok GÖS’de, grup önerileri üretmek amacıyla etkin bir biçimde kullanılmıştır. Şekil 2.1’de BDB yaklaşımının çalışma mekanizması ayrıntılı bir biçimde sunulmuştur. Bu yaklaşımda, ilk olarak, bir g grubunda bulunan P adet kullanıcının n ürünleri için sağladığı derecelendirmelerin, $d_{k\bar{u}}$ ($k = 1, 2, \dots, P$ ve $\bar{u} = 1, 2, \dots, n$), uygun bir birleştirme tekniği kullanılarak birleştirilmesi ile grubun ilgili ürünler üzerindeki tercihlerini yansıtan bir grup profili oluşturulur. Bu işlemin ardından her bir grup için oluşturulan bu profiller kullanılarak herhangi bir derecelendirmenin belirlenemediği ürünler için, uygun bir öneri algoritması uygulanarak tahminler üretilir. Son olarak, üretilen tahminler ve gerçek grup derecelendirme değerleri göz önünde bulundurularak, gruplar için en tercih edilir N adet ürünü içeren en iyi- N grup önerileri üretilir.



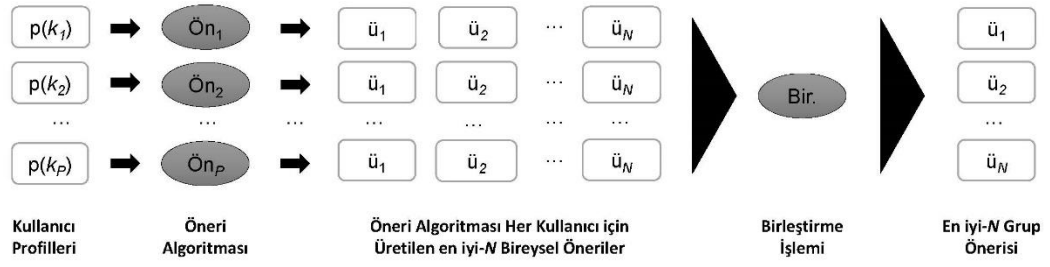
Şekil 2.1. BDB yaklaşımının en iyi- N grup önerisi üretme süreci

b) *Bireysel Tahminleri Birleştirmek (BTB)*: Bu yaklaşım, temel olarak, gruptaki kullanıcılara ürünler için üretilen bireysel tahmin değerlerinin bir araya

getirilmesine dayanmaktadır ve son yıllarda yapılan birçok grup önerisi çalışmasında etkin bir şekilde kullanılmıştır. Şekil 2.2’de BTB yaklaşımının çalışma mekanizması ayrıntılı bir biçimde sunulmaktadır (O’connor vd., 2001) (Amer-Yahia vd., 2009). Buna göre, ilk olarak, uygun bir öneri algoritması yardımıyla bireysel kullanıcıların değerlendirme yapmadıkları ürünler için tahmin derecelendirmeleri üretilir. Ardından, her bir ürün için, üretilen bireysel tahmin derecelendirmeleri ve kullanıcıların sağladığı gerçek derecelendirmeler birleştirilerek, grup için ilgili ürünler için tahmin derecelendirmeleri belirlenir. Son olarak, belirlenen grup derecelendirmeleri kullanılarak, en uygun N adet ürünü içeren en iyi- N grup önerileri üretilir.



- c) *Bireysel Öneri Listelerini Birleştirmek (BÖLB):* Bu yaklaşım, temel olarak, gruptaki kullanıcılara üretilen en iyi- N bireysel öneri listelerinin bir araya getirilmesine dayanmaktadır. Şekil 2.3’de BÖLB yaklaşımının çalışma mekanizması ayrıntılı bir biçimde sunulmuştur. Buna göre, BTB yaklaşımına benzer şekilde, ilk olarak, uygun bir öneri algoritması yardımıyla bireysel kullanıcıların değerlendirme yapmadıkları ürünler için tahmin derecelendirmeleri üretilir. Üretilen tahmin derecelendirmeleri ve kullanıcıların gerçek derecelendirmeleri göz önünde bulundurularak, gruptaki her kullanıcı için bireysel en iyi- N bireysel öneri listeleri üretilir. Ardından, üretilen bireysel öneri listeleri birleştirilerek en iyi- N grup önerileri sağlanır. Bununla birlikte, BÖLB, BDB ve BTB yaklaşımlarına kıyasla GÖS literatüründe daha az kullanılan bir yaklaşımdır (O’connor vd., 2001).



Şekil 2.3. BÖLB yaklaşımının en iyi-N grup önerisi üretme süreci

2.4. Birleştirme Teknikleri

Önceki bölümde belirtildiği üzere, GÖS’de, grup önerileri genellikle grup üyelerinin bireysel tercihleri veya onlara üretilen bireysel tahminler birleştirilerek üretilmektedir. Bu işlem, uygun bir birleştirme tekniği kullanılarak gerçekleştirilir. Literatürde, farklı senaryolardaki ihtiyaçları karşılamak için şimdiye kadar birçok birleştirme tekniği geliştirilmiştir (Seo vd., 2018). Bu bölümde, literatürde öne çıkan en popüler 13 birleştirme tekniğinin çalışma mekanizması ve zayıf yanları ayrıntılı olarak ele alınmıştır. Ayrıca, bu tekniklerin bireysel tercihleri bir araya getirerek grup tercihlerini nasıl belirlediğini gösterebilmek için, Tablo 2.2’de, beş kullanıcıdan oluşan bir kullanıcı grubunun altı ürün için 10-yıldızlı bir derecelendirme yelpazesinde sağladığı derecelendirmeleri içeren örnek bir kullanıcı $\times$ ürün matrisi verilmiştir.

Tablo 2.2. Bir kullanıcı grubu için örnek kullanıcı $\times$ ürün matrisi

	$\bar{u}_1$	$\bar{u}_2$	$\bar{u}_3$	$\bar{u}_4$	$\bar{u}_5$	$\bar{u}_6$
k_1	6	-	7	5	-	7
k_2	-	4	-	4	3	9
k_3	8	5	8	-	-	-
k_4	7	-	-	6	-	10
k_5	10	8	10	8	6	-

Birleştirme teknikleri, grup tercihlerini belirlemede izlediği stratejiye göre sonraki bölümlerde açıklanan yedi farklı kategoride değerlendirilebilir (Seo vd., 2018).

2.4.1. Fikir birliği sağlayan teknikler

Grup önerisi senaryolarındaki yaygın bir yaklaşım, bütün grup üyelerinin grubun nihai kararında eşit oranda etkiye sahip olduğu varsayımına dayanarak, sağladıkları derecelendirmeleri birleştirme aşamasında göz önünde bulundurmaktır. Gruptaki her bireyin tercihinin dikkate alınması, grup üyeleri arasında fikir birliği sağlanan ürünlerin

belirlenmesine olanak tanır. Bu yüzden, bu stratejiyi kullanarak grup tercihlerini belirleyen birleştirme teknikleri fikir birliği sağlayan teknikler olarak adlandırılmaktadır.

Ortalama (Average-Avg) (Ardissono vd., 2003; McCarthy vd., 2006; Jameson, 2004; Zhiwen, Xingshe, & Daqing, 2005), Sevindiren Ortalama (Average without Misery-AwM) (Chao, Balthrop, & Forrest, 2005), Çarpımsal (Multiplicative-Mul) (Christensen & Schiaffino, 2011) ve Çoğalan Faydacı (Additive Utilitarian-AU) (Boratto, Carta, & Fenu, 2016; McCarthy J. F., 2002) bu kategoride değerlendirilen birleştirme teknikleridir. Spesifik olarak, Avg tekniği, bir ürün için grup derecelendirmesini, o ürün için bir değerlendirme yapan grup üyelerinin sağladığı derecelendirmelerin ortalamasını alarak hesaplar. AwM tekniği ise, Avg tekniğinin farklı bir versiyonudur. Somut olarak, AwM, yalnızca üyeler tarafından yapılan derecelendirmelerinin tamamının önceden tanımlanmış bir eşik değerin üzerinde olduğu ürünleri dikkate alır ve bu ürünler için bir ortalama değer hesaplar. Diğer bir deyişle, AwM tekniğinde, en az bir üye tarafından eşik değerin altında bir değer ile değerlendirilmişse bu ürün yok sayılır ve bir grup tercih değeri hesaplanmaz. AU ve Mul tekniklerinde ise, grup derecelendirmeleri, grup üyeleri tarafından sağlanan derecelendirmelerin sırasıyla toplamı ve çarpımı hesaplanarak belirlenir. Tablo 2.2’de verilen derecelendirmeler göz önünde bulundurulduğunda, Avg, AwM, AU ve Mul teknikleri, grup derecelendirmelerini, Tablo 2.3’de gösterildiği gibi hesaplar. Ayrıca, bu örnekte, AwM tekniği için eşik değeri 6 olarak seçilmiştir.

Tablo 2.3. Avg, AwM, AU ve Mul teknikleri ile grup derecelendirmelerinin hesaplanması

		$\bar{u}_1$	$\bar{u}_2$	$\bar{u}_3$	$\bar{u}_4$	$\bar{u}_5$	$\bar{u}_6$
Grup	<i>Avg</i>	7,8	5,7	8,3	5,8	4,5	8,7
	<i>AwM</i>	7,8	-	8,3	-	-	8,7
	<i>AU</i>	31	17	25	23	9	26
	<i>Mul</i>	3.360	160	560	960	18	630

Fikir birliği sağlayan teknikler, temel aritmetik işlemleri kullanarak grup derecelendirmelerini belirlemesi ve bütün üyelerin tercihlerini dikkate alması nedeniyle, uygulaması kolay ve başarılı birleştirme teknikleridir. Ancak bazı ekstrem senaryolarda, bu teknikler, grubun gerçek eğilimlerinin belirlenmesinde yetersiz kalmaktadır. Örneğin, bir ürünün, grubun sadece birkaç üyesi tarafından yüksek derecelendirmeler ile değerlendirilmesi durumunda, Avg tekniği ile hesaplanan grup derecelendirmesi yüksek

olabilir. Ancak, gruptaki üyelerin büyük çoğunluğu, ilgili ürün için bir derecelendirme yapmasa bile bu ürünü tercih etmeyebilir. Benzer bir sorun, AU ve Mul tekniklerinde de ortaya çıkmaktadır. Örneğin, grubun çoğunluğu tarafından düşük derecelendirmelerle değerlendirilen bazı ürünler için, AU ve Mul teknikleri ile hesaplanan grup derecelendirmeleri yüksek olabilir. Bu ürünler, üyeler tarafından düşük derecelendirmeler ile değerlendirilmesine rağmen, gruba öneri olarak üretilebilir. Dahası, çok sayıda grup üyesi tarafından derecelendirilen bir ürün olması durumunda, söz konusu ürün için Mul tekniği ile hesaplanan grup derecelendirmesi sonsuza yaklaşır ve böylece taşma problemi (overflow problem) olarak adlandırılan soruna neden olur.

2.4.2. Derecelendirme sıklığını sayan teknikler

ÖS yaklaşımlarında, kullanıcıların genellikle ilgilendikleri ürünler için derecelendirme sağladığı varsayılır. Bu varsayım temelinde, ürünlere derecelendirme belirten üyelerin sayısı göz önünde bulundurularak, grubun ilgili ürünlere olan ilgisi ölçülebilir. Bu yöntemi kullanarak grup içinde en popüler olan ürünleri belirleyen birleştirme teknikleri, derecelendirme sıklığını sayan teknikler olarak adlandırılabilir. Basit Hesaplama (Simple Computation-SC) (Crossen, Budzik, & Hammond, 2002) ve Onay Oylama (Approval Voting-AV) (Boratto, Carta, & Fenu, 2016; Lieberman, Van Dyke, & Vivacqua, 1999), bu kategoride en öne çıkan birleştirme teknikleridir.

Spesifik olarak, SC tekniği, bir ürünün grup üyeleri tarafından kaç kez derecelendirildiğini sayarak, bu sayma değerini ilgili ürün için grup derecelendirmesi olarak belirler. AV tekniği ise, sayma işlemi sırasında, ilgili ürün için önceden tanımlanmış bir eşik değerin altındaki derecelendirmeleri dikkate almayarak yalnızca üzerindeki derecelendirmelerin sayısını göz önünde bulundurur. Tablo 2.2’de verilen derecelendirmeler göz önünde bulundurulduğunda, SC ve AV teknikleri grup derecelendirmelerini Tablo 2.4’te gösterildiği gibi hesaplar. Ayrıca, bu örnekte AV tekniği için eşik değeri 6 olarak seçilmiştir.

Tablo 2.4. SC ve AV teknikleri ile grup derecelendirmelerinin hesaplanması

		$\bar{u}_1$	$\bar{u}_2$	$\bar{u}_3$	$\bar{u}_4$	$\bar{u}_5$	$\bar{u}_6$
Grup	SC	4	3	3	4	2	3
	AV	4	1	3	2	1	3

SC tekniğinin en büyük dezavantajı, sağlanan derecelendirmelerin büyüklüğünü hiçe saymasıdır. Bu durumda, örneğin grubun büyük çoğunluğu tarafından düşük

derecelendirmeler (yani, negatif derecelendirmeler) ile değerlendirilse bile, bir ürünü tavsiye edilebilir olarak değerlendirmesidir. AV tekniği, yalnızca belirli bir eşik değerin üzerindeki derecelendirmelerin (yani, pozitif derecelendirmeler) sayısını dikkate alarak bu sorunun üstesinden gelir. Ancak, bir ürün için grup derecelendirmesi belirlerken, o ürüne yapılan bütün negatif derecelendirmeleri görmezden gelmek, grubun ilgili ürün üzerindeki eğiliminin doğru bir şekilde analiz edilmesini zorlaştırır.

2.4.3. En yüksek/düşük derecelendirmeye dayanan teknikler

Bu kategoride değerlendirilen birleştirme teknikleri, bir ürün için grup üyeleri tarafından sağlanan ekstrem derecelendirmeleri önemli olarak değerlendirir. Bu ekstrem derecelendirmelerin genellikle en yüksek veya en düşük derecelendirmeler olduğu kabul edilir. En Çok Memnuniyet (Most Pleasure-MP) (Ahmad, Nurjanah , & Rismala, 2017; Boratto, Carta, & Fenu, 2016), En Az Mutsuzluk (Least Misery-LM) (Agarwal, Chakraborty, & Chowdary, 2017; Christensen & Schiaffino, 2011; O’connor vd., 2001) ve Çoğunluk Oylama (Plurality Voting-PV) (Amirali & Boutilier, 2015), bu kategoride öne çıkan birleştirme teknikleridir.

Spesifik olarak, MP tekniği, bir ürün için grup derecelendirmesini belirlerken, grup üyeleri tarafından sağlanan en yüksek derecelendirmeyi dikkate alırken, LM tekniği, en düşük derecelendirmeyi göz önünde bulundurur. Tablo 2.2’de verilen derecelendirmeler göz önünde bulundurulduğunda, MP ve LM teknikleri, grup derecelendirmelerini Tablo 2.5’te gösterildiği gibi hesaplar.

Tablo 2.5. MP ve LM teknikleri ile grup derecelendirmelerinin hesaplanması

		$\ddot{u}_1$	$\ddot{u}_2$	$\ddot{u}_3$	$\ddot{u}_4$	$\ddot{u}_5$	$\ddot{u}_6$
Grup	<i>MP</i>	10	8	10	8	6	10
	<i>LM</i>	6	4	7	4	3	7

MP ve LM teknikleri, küçük kullanıcı gruplarında başarılı grup önerilerinin üretilmesini sağlamasına rağmen, gruplar çok sayıda kullanıcı içerdiğinde, grubun eğilimlerini doğru şekilde analiz etmede yetersiz kalmaktadırlar. Bunun nedeni, grup derecelendirmesini belirlerken, sadece bir üye tarafından sağlanan derecelendirmeyi dikkate alarak diğer üyelerin değerlendirmelerini görmezden gelmeleridir. Bununla birlikte, büyük gruplarda ürünlerin en az bir üye tarafından olabilecek en düşük veya en yüksek derecelendirme ile değerlendirme ihtimalinin çok yüksek olması nedeniyle, her iki teknik de grup içerisinde tercih edilen öğeleri belirlemekte başarılı olamaz.

PV de MP tekniği gibi birleştirme sürecinde en yüksek derecelendirmeleri göz önünde bulunduran bir birleştirme tekniğidir. Ancak, MP tekniği birleştirme aşamasında, bir ürüne yapılan en yüksek derecelendirme değerini dikkate alırken, PV tekniği, bir grup üyesinin en yüksek derecelendirme ile değerlendirdiği ürünü dikkate alır. Ayrıca, yukarıda anlatılan teknikler her ürün için bir grup derecelendirmesi hesaplar, PV tekniği, en yüksek derecelendirmeye sahip ürünleri içeren bir öneri listesi oluşturur.

PV tekniğinde, başlangıçta, her üye için en yüksek derecelendirmeye sahip ürün, ilgili kullanıcının favori ürünü olarak seçilir. Ardından, grup üyeleri arasında en çok seçilen favori ürün, grup için üretilecek öneri listesinin en üst sırasına yerleştirilir. Bu ürün, daha sonra, kullanıcıların favori ürün kümesinden kaldırılır. Sonrasında, geriye kalan ürünler göz önünde bulundurularak, grubun çoğunluğu için favori olan ürün seçilerek öneri listesinin ikinci sırasına eklenir. Bu işlem, öneri listesi tamamlanıncaya kadar tekrarlayarak devam eder. PV tekniğinin çalışma prensibini daha iyi kavrayabilmek için Tablo 2.6’da, Tablo 2.2’deki derecelendirmeleri kullanarak oluşturulan bir sıralı öneri listesi sunulmuştur.

Tablo 2.6. PV tekniği ile oluşturulan sıralı öneri listesi

		Öneri Listesi					
		1	2	3	4	5	6
k_1	$ü_3, ü_6$	$ü_6$	$ü_1$	$ü_4$	-	-	
k_2	$ü_6$	$ü_6$	$ü_2, ü_4$	$ü_2, ü_4$	$ü_2$	$ü_5$	
k_3	$ü_1, ü_3$	$ü_1$	$ü_1$	$ü_2$	$ü_2$	-	
k_4	$ü_6$	$ü_6$	$ü_1$	$ü_4$	-	-	
k_5	$ü_1, ü_3$	$ü_1$	$ü_1$	$ü_2, ü_4$	$ü_2$	$ü_5$	
Grup	PV	$ü_3$	$ü_6$	$ü_1$	$ü_4$	$ü_2$	$ü_5$

2.4.4. Sıralama önceliğine dayanan teknikler

Bu kategoride değerlendirilen birleştirme teknikleri, PV tekniğinde olduğu gibi, ürünlerin sıralamasına dayanmaktadır; ancak, bireylerin derecelendirme sağladıkları her bir ürünü aynı anda dikkate alması bakımından farklılık gösterir. Borda Sayım (Borda Count-BC) (Boratto, Carta, & Fenu, 2016; Marquez & Ziegler, 2016), bu kategoride öne çıkan birleştirme tekniğidir. Spesifik olarak BC, bir kullanıcının değerlendirdiği her bir ürün için, ilgili kullanıcının favori ürün kümesindeki konumuna göre bir sıralama skoru hesaplar. Örneğin, kullanıcının en düşük derecelendirme değeri ile değerlendirdiği ürün (yani favori ürün kümesi içerisinde en son sırada bulunan ürün) için 0 skoru atanır; ikinci en düşük derecelendirmeye sahip ürün için 1 skoru atanır ve bu şekilde kullanıcının

değerlendirdiği her ürün için bir sıralama skoru belirlenir. Bununla birlikte, kullanıcının birden fazla ürünü aynı derecelendirme değeri ile değerlendirmesi durumunda, ilgili ürünlere, sıralama skorlarının toplamı eşit olarak paylaştırılır. Bu işlem, her bir grup üyesi için tamamlandıktan sonra, BC tekniği, bir ürün için grup derecelendirmesini, ilgili ürüne atanmış olan sıralama skorlarının toplamını hesaplayarak belirler. Tablo 2.2'deki değerlendirmeler göz önünde bulundurulduğunda, BC tekniği, grup derecelendirmelerini Tablo 2.7'deki gibi hesaplar.

Tablo 2.7. BC tekniği ile grup derecelendirmelerinin hesaplanması

	$\bar{u}_1$	$\bar{u}_2$	$\bar{u}_3$	$\bar{u}_4$	$\bar{u}_5$	$\bar{u}_6$	
k_1	1	-	2,5	0	-	2,5	
k_2	-	1,5	-	1,5	0	3	
k_3	1,5	0	1,5	-	-	-	
k_4	1	-	-	0	-	2	
k_5	3,5	1,5	3,5	1,5	0	-	
Grup	<i>BC</i>	7	3	7,5	3	0	7,5

BC tekniğinin en büyük dezavantajı, bazı durumlarda, kullanıcıların ürünler için sağladığı derecelendirmeleri sıralayarak birleştirme işlemini gerçekleştirmenin mümkün olmamasıdır. Örneğin, bazı grup üyelerinin ürünler için yapmış olduğu derecelendirmeler eşit dağılım gösterebilir. Bu durumda, aynı derecelendirmeye sahip ürünlerden hangisinin kullanıcı için daha tercih edilebilir olduğu ve sıralamada hangisinin daha önde olması gerektiğini belirlemek zorlaşır. Ayrıca, 5-yıldızlı derecelendirme yelpazesini kullanan tipik bir öneri sisteminde bile, yüzlerce ürünün mevcut olması nedeniyle, kullanıcının aynı değer ile derecelendirdiği birçok ürün olabilir. Bu durumda da, kullanıcının favori ürünlerinin doğru bir şekilde belirlenmesi zorlaşır. Dahası, diğer birleştirme teknikleriyle kıyaslandığında, sıralama stratejisini izleyen tekniklerin, yüksek hesaplama süresi gerektirmesi nedeniyle verimsiz olduğu kabul edilmektedir. Bu problemler, BC tekniğinde olduğu gibi, ürünlere yapılan en yüksek derecelendirmeyi dikkate alarak bir sıralama yöntemini kullanan PV tekniğinde de sıklıkla ortaya çıkmaktadır.

2.4.5. Derecelendirmeleri karşılaştıran teknikler

Grup içerisinde en çok tercih edilen ürünler, ürünlerin birbirlerine göre göreceli önemi dikkate alınarak da belirlenebilir. Bu işlem, grup üyelerinin derecelendirmelerine göre ürünlerin karşılıklı tercih durumu dikkate alınarak gerçekleştirilebilir. Bu kategoride değerlendirilen birleştirme teknikleri derecelendirmeleri karşılaştıran teknikler olarak adlandırılmaktadır. Copeland Kuralı (Copeland Rule-CR) (Masthoff, 2011), bu kategoride öne çıkan birleştirme tekniğidir.

Tablo 2.2’de sunulan derecelendirmeler doğrultusunda, CR tekniği, grup derecelendirmelerini Tablo 2.8’deki gibi hesaplar. Bu tabloda, k ’ncı ve l ’inci sütuna ($k, l \in \{1, 2, \dots, 6\}, k \neq l$) karşılık gelen hücre, grup üyeleri arasında l ’inci sütundaki ürünün k ’ncı satırdaki ürüne göre kaç kez tercih edildiğini ifade etmektedir. Örneğin, birinci ürün ($ü_1$) ile ikinci ürün ($ü_2$) karşılaştırıldığında, dört grup üyesi $ü_1$ ürününü $ü_2$ ürününe tercih ederken, sadece bir grup üyesi $ü_2$ ürününü $ü_1$ ürününe tercih etmektedir. Dolayısıyla, ($ü_2, ü_1$)’e karşılık gelen hücre $4 - 1 = 3$ ve ($ü_1, ü_2$)’e karşılık gelen hücre $1 - 4 = -3$ olarak hesaplanır. Her ürün için son CR puanı, ilgili ürünün bütün ikili karşılaştırmalarda aldığı puanlar toplanarak elde edilir.

Tablo 2.8. CR tekniği ile grup derecelendirmelerinin hesaplanması

	$\ddot{u}_1$	$\ddot{u}_2$	$\ddot{u}_3$	$\ddot{u}_4$	$\ddot{u}_5$	$\ddot{u}_6$	
$\ddot{u}_1$	■	-3	0	-3	-3	+1	
$\ddot{u}_2$	+3	■	+2	+1	-3	+1	
$\ddot{u}_3$	0	-2	■	-1	-2	0	
$\ddot{u}_4$	+3	-1	+1	■	-4	+2	
$\ddot{u}_5$	+3	+3	+2	+4	■	+2	
$\ddot{u}_6$	-1	-1	0	-2	-2	■	
Grup	<i>CR</i>	+8	-4	+5	-1	-14	+6

CR birleştirme tekniğinin temel dezavantajı, bütün ürünler için bir ikili kıyaslama işlemi gerçekleştirmesidir. Bu işlem, özellikle mevcut ürün sayısının fazla olması durumunda, yüksek hesaplama süresine neden olur.

2.4.6. Etkili grup üyelerini dikkate alan teknikler

Grubun bazı üyeleri, tecrübe ve uzmanlık gibi bazı özellikleri göz önünde bulundurulduğunda, diğer bireylerin karar verme sürecini doğrudan etkileyebilirler. Bu kişiler, grubun etkili bireyleri olarak adlandırılabilir ve sağladıkları derecelendirmeler, birleştirme aşamasında diğer kullanıcılardan farklı şekilde ele alınabilir. Bu stratejiyi izleyen en popüler birleştirme tekniği En Saygın Kişi (Most Respected Person-MRP)’dir (Masthoff, 2011). Bu teknik, birleştirme aşamasında, grup üyeleri arasındaki etkili (saygın) kişiyi farklı şekillerde belirleyerek, bu kişinin sağladığı derecelendirmeleri grup derecelendirmesi olarak seçer. Tablo 2.2’de sunulan derecelendirmelere göre, MRP tekniği grup derecelendirmelerini Tablo 2.9’deki gibi belirler. Burada, grup içindeki en saygın kişi 3. kullanıcı olarak seçilmiştir.

Tablo 2.9. MRP tekniği ile grup derecelendirmelerinin hesaplanması

$ü_1$	$ü_2$	$ü_3$	$ü_4$	$ü_5$	$ü_6$
-------	-------	-------	-------	-------	-------

MRP tekniğinin en temel dezavantajı, birleştirme sürecinde grup içerisindeki sadece bir kullanıcının derecelendirmelerini göz önünde bulundurarak, grubun diğer üyelerinin derecelendirmelerini yok saymasıdır. Bu durum, özellikle büyük kullanıcı gruplarında, grubun genel eğiliminin doğru bir şekilde yansıtılamamasına yol açmaktadır. Ayrıca, gruptaki en saygın kişinin belirlenmesi konusunda, grubun ortalama derecelendirme vektörüne en benzer derecelendirme vektörüne sahip kullanıcının seçilmesi gibi bazı yaklaşımlar olsa da (Masthoff, 2011), grubun nihai kararı üzerinde en çok etkisi olan kullanıcının doğru bir şekilde belirlenmesinin mümkün olmadığı düşünülmektedir.

2.4.7. Derecelendirmelerin dağılımını dikkate alan teknikler

Bir grup önerisinde temel amaçlardan biri, gruptaki en çok sayıda kullanıcıyı memnun edebilecek ürünlerin belirlenerek genel memnuniyeti en üst düzeye çıkarmaktır. Ancak gruplar her ne kadar benzer eğilimleri olan kullanıcılardan oluşsa da, mükemmel olmadıkları için, bazı ürünler için kullanıcılar tarafından sağlanan derecelendirmeler ciddi farklılık gösterebilir. Bu ürünlerin birleştirme sürecinde belirlenerek, grup için üretilecek öneri listesinde olmaması sağlanmalıdır. Diğer yandan, grubun çoğunluğu tarafından benzer derecelendirmelerle değerlendirilen ürünleri seçerek, öneri listesine dâhil etmek, sağlanan öneri listesinden memnun kalan birey sayısının artmasına imkân tanır. Bunu gerçekleştirmek için grup üyelerinin ürünler için sağladığı derecelendirmelerin dağılımlarını dikkate almak gerekir. Bu stratejiyi izleyen birleştirme teknikleri derecelendirme dağılımlarını dikkate alan teknikler olarak adlandırılır.

Yukarıya Doğru Tesviye (Upward Leveling-UL) (Seo vd., 2018), son yıllarda geliştirilen ve ürünlere sağlanan derecelendirmelerin dağılımlarını, SS kullanarak dikkate alan modern bir birleştirme tekniğidir. Bu teknikte, yüksek SS değerine sahip ürünlerin, grup içerisinde bir uzlaşma sağlanamadığı ürünler, diğer yandan, düşük SS değerine sahip ürünlerin kullanıcıların çoğunluğu tarafından üzerinde fikir birliği sağladığı ürünler olduğu düşünülür. Spesifik olarak, UL, bir ürün için nihai grup derecelendirmesi hesaplarken, kullanıcılar tarafından sağlanan derecelendirmelerin SS değeri ile AV ve Avg teknikleri ile hesaplanan grup derecelendirmelerin ağırlıklı ortalamalarını hesaplayarak kombine eden melez bir tekniktir. Ayrıca, UL tekniği, bu üç metot ile grup derecelendirmelerini hesaplamadan önce, bu teknikler ile hesaplanacak grup

derecelendirmelerinin farklı yelpazelerdeki değerler olacağını ve karşılaştırılabilir olmayacağını göz önünde bulundurarak, ilk işlem olarak bireysel derecelendirmeleri [0, 1] ölçeğine normalize eder. Bununla birlikte, AV tekniğinin grup derecelendirmesi hesaplarken değerlendirilen ürünlerin sayısını dikkate alması nedeniyle, bu teknik ile hesaplanan grup derecelendirme için bir min-maks normalizasyon işlemi uygulayarak bu değerleri [0, 1] ölçeğinde indirger. Sonuç olarak, hem yüksek ortalamaya sahip (Avg ile), hem kullanıcılar tarafından sıkça değerlendirilen (AV ile), hem de grup üyelerinin üzerinde uzlaştığı (SS ile) ürünler aynı anda belirlenerek, gruba üretilecek öneri listesine dâhil edilmiş olur. UL birleştirme tekniğinin bir ürün için nasıl grup derecelendirmesi hesapladığı Denklem 2.4'te sunulmuştur.

$$D_{g,\ddot{u}} = \alpha(1 - SS_{g,\ddot{u}}) + \beta(AV_{g,\ddot{u}}) + \gamma(Avg_{g,\ddot{u}}) \quad (2.4)$$

Burada, $SS_{g,\ddot{u}}$, $AV_{g,\ddot{u}}$, ve $Avg_{g,\ddot{u}}$, g grubunda $\ddot{u}$ ürünü için, sırasıyla SS, AV ve Avg ile hesaplanan grup derecelendirmelerini ifade etmektedir. $D_{g,\ddot{u}}$ ise UL tekniği ile hesaplanan grup derecelendirmesini temsil etmektedir. Denklemdaki ağırlık değerleri (yani α , β ve γ), grup önerisi senaryosuna göre, Denklem 2.5'te verilen koşulu sağlayacak şekilde rastgele seçilir.

$$(\alpha + \beta + \gamma = 1) \quad (2.5)$$

Tablo 2.2'de sunulan derecelendirmeler göz önünde bulundurulduğunda, UL tekniği grup derecelendirmelerini Tablo 2.10'daki gibi hesaplamaktadır. Not olarak, AV tekniği için eşik değeri 6 olarak, α , β , ve γ katsayıları ise eşit biçimde ve 1/3 olarak seçilmiştir.

Tablo 2.10. UL tekniği ile grup derecelendirmelerinin hesaplanması

		$\ddot{u}_1$	$\ddot{u}_2$	$\ddot{u}_3$	$\ddot{u}_4$	$\ddot{u}_5$	$\ddot{u}_6$
Grup	Avg	0,78	0,57	0,83	0,58	0,45	0,87
	AV	1	0	0,67	0,33	0	0,67
	1-SD	0,83	0,81	0,85	0,83	0,79	0,85
	UL	0,87	0,46	0,78	0,58	0,41	0,80

Her ne kadar UL tekniği, birleştirme sürecinde ürünlere yapılan derecelendirmelerin dağılımını dikkate alma konusunda türünün ilk örneği olsa da, dağılımları analiz etmek için SS'i kullanması nedeniyle, dağılımların yalnızca tek bir tepe

noktasına (yani, tek yönlü-dağılım) sahip olması durumunda iyi sonuçlar vermektedir. Ancak, kullanıcı gruplarının boyunun büyük olması ve geniş bir derecelendirme yelpazesine sahip olunması durumunda (Ör. 10-yıldızlı derecelendirme yelpazesi), derecelendirmelerin dağılımı birden fazla tepe noktası (yani, çok-yönlü dağılım) içerebilir. Böyle bir durumda, üzerinde farklı seviyelerde ve güçlü bir şekilde bir uzlaşma sağlanmasına rağmen, hesaplanan sapma değeri yüksek olur ve ilgili ürün grup için uygun bir öneri olarak değerlendirilmez. Ayrıca katsayıların rastgele seçilmesi UL tekniğinin bir başka önemli dezavantajıdır.

2.5. Veri Kümeleri ve Değerlendirme Ölçütü

Tez kapsamında yapılan çalışmaların performansını değerlendirmek için, öneri sistemleri üzerine yapılan çalışmalarda yaygın olarak kullanılan ve Minnesota Üniversitesinde bulunan GroupLens⁷ araştırma ekibi tarafından toplanan, kullanıcıların çeşitli filmler için sağladığı derecelendirmeleri içeren MovieLens (Harper & Konstan, 2015) ve Personality (Nguyen vd., 2017) veri kümeleri kullanılmıştır. Bununla birlikte, kullanıcıların filmler için yapmış olduğu derecelendirmeleri içeren bir başka veri kümesi olan popüler Netflix⁸ prize veri kümesi de bu tez çalışmasında kullanılmıştır. Orijinal Netflix veri kümesinin oldukça büyük ve seyrek yapıda olması nedeniyle, deneysel çalışmalarda, 10.000 kullanıcının örneklediği orijinal Netflix veri kümesinin bir alt kümesi (NF) kullanılmıştır. Örneklenen NF veri kümesi, orijinal Netflix veri kümesindeki filmlerin ve kullanıcıların yoğunluk özelliklerini aynı oranda temsil edecek şekilde oluşturulmuştur.

MovieLens veri kümesinin, içerdiği derecelendirme sayısına göre iki versiyonu mevcuttur. Bunlar; MovieLens-100K (MLP) ve MovieLens-1M (MLM) veri kümeleridir. MLP, MLM ve NF veri kümelerinde, kullanıcı derecelendirmeleri, ayrık ve 5-yıldızlı bir derecelendirme yelpazesinde temsil edilmektedir. Bununla birlikte, MLP ve MLM veri kümelerinde, filmlerin, en az birine ait olduğu 18 adet film janrası tanımlanmıştır. Ayrıca, bu iki veri kümesinde, sistemdeki kullanıcıların yaş, cinsiyet ve meslek gibi demografik bilgileri mevcuttur.

Personality (Per) veri kümesinde ise, kullanıcıların, filmler için yapmış olduğu derecelendirmeler yine ayrık ve 10-yıldızlı bir derecelendirme yelpazesindedir. Per veri

⁷ <http://grouplens.org/>

⁸ <http://www.netflixprize.com/>

kümesi ayrıca, her kullanıcı için beş temel kişisel özellik (yani *açıklık*, *uyumluluk*, *duygusal istikrar*, *vicdanlılık* ve *dışadönüklük*) hakkında 7-yıldızlı bir derecelendirme ölçeğinde değerlendirme puanları içermektedir. Ham haliyle Per veri kümesi, mevcut filmlerin yalnızca %1,6'sının derecelendirildiği seyrek bir veri kümesidir. Bu nedenle, yapılan çalışmada, orijinal veri setinin farklı doluluk oranlarına sahip Per5 ve Per10 olarak isimlendirilen iki farklı alt kümesi kullanılmıştır. Veri kümesini adlandırmak için kullanılan sayısal değerler, veri setlerindeki ürünler ve kullanıcılar için minimum miktardaki derecelendirme sayısını ifade etmektedir. Örneğin, Per5 veri kümesinde, her kullanıcı en az 5 farklı ürünü derecelendirmiştir ve her ürün en az 5 farklı kullanıcı tarafından derecelendirilmiştir.

Bu tez çalışmasında kullanılan MLP, MLM, NF, Per5 ve Per10 veri kümelerinin sahip oldukları derecelendirme, kullanıcı ve ürün sayıları ile birlikte seyreklik oranları Tablo 2.11'de ayrıntılı olarak sunulmaktadır.

Tablo 2.11. Veri kümelerinin özellikleri

	<i>MLP</i>	<i>MLM</i>	<i>NF</i>	<i>Per5</i>	<i>Per10</i>
Kullanıcı Sayısı	943	6.040	10.000	1.819	1.819
Ürün Sayısı	1.682	3.952	17.700	14.868	10.375
Derecelendirme Sayısı	100.000	1.000.000	2.337.295	984.499	954.985
Seyreklik Oranı	%93,7	%93,8	%98,68	%96,4	%94,9

Bu çalışmada, grup önerileri olarak üretilen en iyi- N sıralı ürün listelerinin kalitesini ölçmek için literatürde yaygın olarak kullanılan değerlendirme ölçütlerinden birisi olan Normalleştirilmiş İndirimli Kümülatif Kazanç (Normalized Discounted Cumulative Gain- n DCG) kullanılmaktadır. Somut olarak, n DCG, hem ürünlere sağlanan gerçek derecelendirme değerlerini, hem de ürünlerin öneri listesindeki konumlarını göz önünde bulundurarak sıralı bir öneri listesinin kalitesini ölçen bir metriktir (Seo vd., 2018).

Bir g grubu için öneri listesi olarak üretilen ürünler $\{\bar{u}_1, \bar{u}_2, \dots, \bar{u}_N\}$ ile temsil edildiğinde, g içindeki bir k üyesi için İndirimli Kümülatif Kazanç (Discounted Cumulative Gain-DCG) ve n DCG değerleri sırasıyla Denklem 2.6 ve 2.7'deki gibi hesaplanmaktadır.

$$DCG_N^k = d_{k,\bar{u}_1} + \sum_{n=1}^N \frac{d_{k,\bar{u}_n}}{\log_2(n)} \quad (2.6)$$

$$nDCG_N^k = \frac{DCG_N^k}{IDCG_N^k} \quad (2.7)$$

Burada, $d_{k,\bar{u}}$, k grup üyesinin $\bar{u}$ ürün için yapmış olduğu gerçek derecelendirme değerini ifade etmektedir. Ayrıca, $IDCG_N^k$, k üyesi için N adet ürünün yeniden sıralanmasıyla elde edilebilecek maksimum kazancı ifade etmektedir. Her bir grup üyesi için $nDCG$ değerlerinin hesaplanmasından sonra, g grubunun $nDCG$ değeri, üyelerin $nDCG$ değerlerinin ortalaması alınarak hesaplanmaktadır.

GÖS'lerin birden fazla paydaş içermesi göz önünde bulundurulduğunda, bu tür sistemlerin genellikle adalet sorunlarına yol açtığı söylenebilir (Felfernig vd., 2018). Dolayısıyla, tez çalışmasında geliştirilen bazı birleştirme tekniklerinin grup üyelerini eşit şekilde tatmin etmeyi amaçladığı dikkate alındığında, adalet performanslarının incelenmesi gerekmektedir. Bu amaçla, adaleti, önerilen en iyi- N öneri listesindeki ürünlerden en az m tanesi için yüksek bir derecelendirme sağlayan grup üyelerinin, grubun tamamına oranı olarak yorumlayan m -orantılılık tabanlı adalet ölçütü (Fairness) kullanılmıştır (Serbos vd., 2017). Bu bağlamda, Fairness ölçütü, g grubunun tamamı için üretilen N ürünlerinin ne oranda adil olduğunu Denklem 2.8'de verilen formül yardımıyla hesaplar.

$$Fairmess(g) = \frac{|g_N|}{|g|} \quad (2.8)$$

Burada, $|g_N|$, m -orantılılık koşulunu sağlayan grup üyelerinin sayısını göstermektedir. Diğer bir deyişle, $|g_N|$, öneri listesindeki ürünlerden en az m tanesi için yüksek bir derecelendirme sağlayan grup üyelerinin sayısını ifade etmektedir.

Ayrıca, gruba üretilen bir öneri listesinin kalitesini kapsamlı bir şekilde değerlendirmek için, gruptaki her üyenin memnuniyet derecesini ölçen Grup Memnuniyet Ölçütü (Group Satisfaction Metric-GSM) kullanılmıştır (Barzegar Nozari & Koohi, 2020). GSM ölçütü, Denklem 2.9'da sunulan formül ile hesaplanmaktadır.

$$GSM(g) = \frac{\sum_{k \in g} |\bar{U}_k \cap N|}{|g| \times |N|} \quad (2.9)$$

Burada, $\bar{U}_k$, gruptaki k kullanıcısının, grup için üretilen sıralı N ürünleri içerisinde tatmin olduğu ürünlerin kümesini ifade etmektedir. Not olarak, hem Fairness hem de GSM ölçütlerini hesaplarken, bir grup üyesinin önerilen üründen memnun olup olmadığını belirlemek için eşik değeri olarak 3,5 seçilmiştir. Bunun nedeni, beş yıldızlı bir değerlendirme yelpazesinde, olumlu derecelendirmelerin, sezgisel olarak 4 ve 5'e karşılık gelmesidir (Bobadilla, Serradilla, & Bernal, 2010).



3. GÖS LİTERATÜR TARAMASI

Son yıllarda, müzik (Chao, Balthrop, & Forrest, 2005; Christensen & Schiaffino, 2011; Crossen, Budzik, & Hammond, 2002; McCarthy & Anagnost, 1998; Zhiwen, Xingshe, & Daqing, 2005), film (Christensen & Schiaffino, 2011; O'connor vd., 2001; Quijano-Sanchez, Recio-Garcia, & Diaz-Agudo, 2011), restoran (McCarthy J. F., 2002), turizm (Ardissono vd., 2003; Jameson, 2004; Marquez & Ziegler, 2016; McCarthy vd., 2006) ve çevrimiçi tarama (Lieberman, Van Dyke, & Vivacqua, 1999) gibi çeşitli alanlarda farklı senaryolar için birçok GÖS uygulaması geliştirilmiştir. Tek bir birleştirme tekniği tüm uygulamalarda yüksek performans sağlayamadığından, bu uygulamalarda grup önerileri üretmek amacıyla kullanılan birleştirme tekniği genellikle farklılık göstermektedir.

MusicFX (McCarthy & Anagnost, 1998), bir spor salonunda egzersiz yapan bir grup kullanıcı için, onların ilgisini çeken müzik türlerini tahmin ederek, uygun radyo istasyonu seçen akıllı bir ortamdır. Bu sistemde, birleştirme tekniği olarak AwM kullanılarak yalnızca kullanıcı derecelendirmelerini belirli bir eşik değerin üstünde olan müzik türleri göz önünde bulundurulur. FlyTrap (Crossen, Budzik, & Hammond, 2002), müzik alanı için geliştirilen GÖS için diğer bir popüler uygulamadır. Bu uygulama, bir odada bulunan bir grup birey için çalma listesi üreten sanal bir DJ oluşturur. Üretilen sanal DJ, bireylerin geçmişteki müzik tercihlerini, müzik türleri, müzikler arasındaki geçiş ve sanatçılar arasındaki etkiler gibi bilgileri harmanlayarak öneriler üretir. Ayrıca, bu uygulamada, birleştirme tekniği olarak SC tekniğinin farklı bir versiyonu kullanılır. Müzik alanı için diğer popüler örnekler Adaptive Radio (Chao, Balthrop, & Forrest, 2005) ve Adaptive In-Vehicle Multimedia System'dir (Zhiwen, Xingshe, & Daqing, 2005). Spesifik olarak, Adaptive Radio, bir grup kullanıcı için, onların ne tür müzik dinlemek istemediklerine dair bilgilere dayanarak çalınacak müziği seçen bir sistemdir. Bu sistemde, grup üyelerinin müzik türleri hakkındaki bu tür olumsuz tercihler, AwM tekniği yardımıyla birleştirilir. Diğer yandan, Adaptive In-Vehicle Multimedia System, adından da anlaşılabilir gibi, bir araçta birlikte seyahat eden bir grup yolcu için bir dizi şarkıyı içeren öneri listesi üreten bir sistemdir. Bu sistemde, Avg tekniği ile yolcuların seçilen özellikleri birleştirilerek grup önerisi üretmek amacıyla bir grup profili oluşturulur.

PolyLens (O’connor vd., 2001), bireysel kullanıcılar için film önerileri sağlayan popüler MovieLens (Harper & Konstan, 2015) sisteminin gelişmiş bir versiyonudur. Bu sistemde hedef kitle, bireysel kullanıcılar yerine kullanıcı gruplarıdır. Kullanıcıların filmler için sağladığı bireysel derecelendirmeleri birleştirmek amacıyla LM tekniği kullanılmaktadır. Ayrıca, PolyLens’de kullanıcıların sistemden memnuniyeti, tasarlanan bir anket aracılığı ile araştırılmaktadır. jMusicRecommender ve jMovierecommender (Christensen & Schiaffino, 2011), sırasıyla müzik ve film alanları için grup önerilerinin üretilmesini sağlayan bir eğlence GÖS’üdür. Bu sistemler, birleştirme teknikleri olarak Mul, Avg ve LM’nin bir uyum içinde kullanıldığı çeşitli grup önerisi yaklaşımları kullanılmaktadır. Film alanı için diğer bir örnek, birleştirme tekniği olarak LM ve MRP tekniklerinin kullanıldığı gRecs (Ntoutsis vd., 2012) sistemidir. gRecs, ilk olarak, benzer ilgi alanlarına sahip kullanıcıları gruplarını belirleyen ve ardından bir CF stratejisi yardımıyla kullanıcılara en iyi- N film önerileri üreten bir sistemdir. Son olarak, bir Facebook uygulaması olarak geliştirilen HappyMovie (Quijano-Sanchez, Recio-Garcia, & Diaz-Agudo, 2011), grup üyelerinin ilgi alanları ve kendi aralarındaki güvene dayalı olarak film önerileri üretir. Bu sistemde, bireysel derecelendirmeleri bir araya getirmek için Avg ve LM teknikleri kullanılmaktadır.

Pocket Restaurant Finder (McCarthy J. F., 2002), bir grup kullanıcı için önerilen restoranların listesini üreten bir sistemdir. Bu sistemde, grup önerileri üretebilmek amacıyla, kişilerin yemek (tercih, tat, fiyat kategorisi, restoran olanakları, mutfak türü vb.) ve konum konusundaki tercihleri AU tekniği yardımıyla birleştirilir. INTRIGUE (Ardissone vd., 2003), Travel Decision Forum (Jameson, 2004) ve CATS (McCarthy vd., 2006) turizm alanındaki popüler GÖS’lerdir. Spesifik olarak, INTRIGUE, rehberli tur organizasyonları için, katılımcıların özelliklerini göz önünde bulundurarak bir dizi turistik grup önerileri sağlayan bir sistemdir. Travel Decision Forum ise, birlikte tatil yapmayı planlayan bir grup insan için nereye gitmeleri gerektiği konusunda tek bir grup önerisi üreterek karar verme süreçlerine yardımcı olan bir sistemdir. Son olarak, CATS, bir grubun bir bütün olarak ihtiyaçlarına en uygun kayak paketlerini öneren bir sistemdir. Bu sistemdeki grup profilleri, bireylerin eleştirilerinin Avg tekniği yardımı ile birleştirilmesiyle oluşturulur. İnternet sayfalarının fiziksel mekânlar olarak düşünülmesi durumunda, bir grup kullanıcının hangi sayfaları ziyaret etmesi konusunda öneriler sağlayan Let’s Browse (Lieberman, Van Dyke, & Vivacqua, 1999), kullanıcı grupları ile internet sayfalarını karşılaştırılması neticesinde elde edilen eşleştirme puanlarını

kullanarak, çevrimiçi taramada kullanıcı gruplarının karar verme sürecine yardımcı olur. Bu sistem, AV tekniğinde olduğu gibi eşleştirme puanı sadece belirli bir eşik değerin üzerinde olan internet sayfalarını dikkate alır.

Kassak ve ark. (Kassak, Kompan, & Bielikova, 2016) içerik-tabanlı ve CF yaklaşımlarını birlikte kullanan karma tipte bir grup öneri mekanizması geliştirmişlerdir. Bu mekanizma, nihai grup önerilerini, bu iki algoritma ile üretilen sıralı öneri listelerini birleştirerek belirler. CoGrec (Liu vd., 2016), kullanıcı profillerini negatif olmayan matris çarpanlarına ayırma yoluyla oluşturan ve daha sonra bu profilleri örtüşen topluluklardan oluşan grupları belirlemek için kullanan bir sistemdir. Kullanıcı tercihlerinin birleştirilmesi söz konusunda olduğunda ise, sistem, örtüşen kullanıcıların üyeliklerinin ortalamasını dikkate alır. Mahyar ve ark. (Mahyar vd., 2017) bir gruptaki üyelerin grubun nihai kararı üzerinde ne kadar etkili olduğunu tahmin etmek amacıyla çizge teorisindeki merkeziyet kavramının, birleştirme aşamasında kullanılması sağlayan bir grup önerisi yaklaşım geliştirilmiştir. Bu yaklaşım, kullanıcı merkeziyet ölçümlerinin, birleştirme aşamasında, bir grubun üyelerinin derecelendirmelerinin ağırlıklandırılmasını sağlayarak başarılı grup önerilerinin üretilmesini sağlar. Seo ve ark. (Seo vd., 2018) kullanıcı derecelendirmelerindeki sapmayı grup önerisi için kritik bir unsur olarak kabul eden UL birleştirme tekniğini geliştirmişlerdir. Bu teknik, AV ve Avg ile elde edilen grup derecelendirmeleri ile derecelendirmelerdeki sapmanın ağırlıklı ortalamasını kullanarak maksimum sayıda grup üyesini memnun edebilecek yüksek kalitede grup önerilerinin üretilmesini sağlar. Ayrıca, birleştirme tekniklerinin başarısını kıyaslamak için, literatürde göze çarpan iki deneysel çalışma mevcuttur. Boratto ve ark. (Boratto, Carta, & Fenu, 2016) AU, AV, BC, LM ve MRP dâhil olmak üzere bazı temel birleştirme tekniklerinin çeşitli grup büyüklüklerindeki performansını analiz etmektedirler. Diğer yandan, Yalçın ve ark., (Yalçın, Bilge, & Yüksek, 2019) daha fazla sayıda birleştirme tekniği ile ilgilenerek performanslarını araştırır. Bu çalışmalar, GÖS’nde, grupların ve üretilen öneri listelerinin boyutunun, kullanılan birleştirme tekniğinin performansı üzerinde doğrudan bir etkiye sahip olduğunu göstermektedir. Ayrıca, deneysel olarak, tüm uygulama senaryolarında grup derecelendirmelerini en doğru şekilde belirleyen bir birleştirme tekniği bulunmadığı gösterilmiştir.

Literatürde, uygun grup önerileri sağlamak için kullanıcı kişiliğini dikkate alan yalnızca birkaç çalışma mevcuttur. Örneğin, Recio-Garcia ve ark. (Recio-Garcia vd.,

2009) grup üyeleri arasındaki ihtilafların öneri sürecini nasıl etkilediğini araştırarak kişiliğe duyarlı grup önerileri sunan bir yaklaşım geliştirmiştir. Bu yaklaşımda, bireylerin ihtilaf durumlarındaki davranışlarını ölçmek için, ihtilaf modu ağırlığı adında bir ölçüm kullanılarak gruptaki her üye için bir profil oluşturulur. Bu profiller, kullanıcıların derecelendirmelerini ağırlıklandırmak için birleştirme sürecine entegre edilir. Ayrıca, Rossi ve ark. (Rossi & Cervone, 2016) *uygunluk* kişilik özelliğine göre bir kullanıcının özgecil davranış düzeyini modelleyen, kişilik temelli bir fayda işlevini kullanarak kullanıcı grupları için uygun öğeleri öneren bir grup önerisi yaklaşımı önermişlerdir. Son olarak, Sanchez ve ark. (Quijano-Sanchez, Recio-Garcia, & Diaz-Agudo, 2010) üretilen grup önerilerinin kalitesini arttırmak amacıyla, hem grup kişilik kompozisyonu, hem de grup üyeleri arasındaki sosyal ilişkileri dikkate alan bir yöntem geliştirmişlerdir.

Mevcut bölümde anlatıldığı üzere, şimdiye kadar farklı amaçlar için çeşitli GÖS yaklaşımları geliştirilmiştir. Bununla birlikte, bu yaklaşımların büyük çoğunluğu genellikle tek bir birleştirme tekniği kullanarak grup üyeleri tarafından sağlanan derecelendirmeleri birleştirmektedir. Bu nedenle, grup profillerinin farklı yönlerini yakalamak için birden fazla birleştirme tekniğinin bir uyum içinde kullanılmasına ihtiyaç vardır. Ayrıca, grup üyeleri tarafından sağlanan derecelendirmelerin dağılımı, Giriş Bölümü'nde açıklandığı üzere, grup üyelerinin eğiliminin doğru bir biçimde analiz edilmesinde önemli bir rol oynamaktadır. Derecelendirme dağılımları, UL gibi modern birleştirme tekniklerinde SS ile göz önünde bulundurulmasına rağmen bu yalnızca tek-yönlü dağılımlarda başarılı olmaktadır. Ancak, özellikle büyük kullanıcı gruplarının ve büyük ölçekli derecelendirme sistemlerinin varlığında ortaya çıkabilecek çok-yönlü bir dağılımın analiz edilmesi ve grup üyeleri arasında tercih edilebilir ürünlerin belirlenebilmesi için yeni birleştirme mekanizmalarının geliştirilmesine ihtiyaç vardır. Ayrıca, grup üyelerinin davranışsal özellikleri, grubun eğilimlerinin analizinde önemli bir unsurdur. Her ne kadar son yıllarda yapılan birkaç çalışma bu konuyu ele alsa da, birleştirme sürecinde kişilik özelliklerinin rollerini kapsamlı bir şekilde araştıran ve bunları göz önünde bulunduran yeni birleştirme mekanizmalarının geliştirilmesine ihtiyaç olduğu söylenebilir.

GÖS'de, grupların otomatik olarak tanımlanması bağlamında son yıllarda öne çıkan bazı çalışmalar mevcuttur. Baltrunas ve ark. (Baltrunas, Makcinskas, & Ricci, 2010) kullanıcılar arasındaki bütün ikili korelasyonu hesaplayarak önceden tanımlanmış bir

grup için benzerlik eşik değerini aşan grupların kullanılmasını önermişlerdir. Ardından, ürünleri matris çarpanlara ayırmaya dayalı bir CF algoritması yoluyla üretilen tahminlere göre sıralayarak, tanımlanan gruplara en iyi- N önerileri üretmişlerdir. Bununla birlikte, grupların tanımlanması için tüm kullanıcılar arasındaki benzerliklerin hesaplanmasındaki temel problem, kullanıcıların ortak olarak değerlendirdiği ürünleri bulmanın zor olması ve özellikle büyük bir veri kümesi kümesine uygulandığında yüksek hesaplama süresi gerektirmesidir (Bilge & Polat, 2013). Ayrıca, benzer tercihlere sahip kullanıcı topluluklarını belirlemek için çizge kümeleme yaklaşımları (Wang & Fleury, 2011) (Fortunato & Castellano, 2012), hiyerarşik kümeleme yaklaşımları (Lancichinetti, Fortunato, & Kertesz, 2009) ve modülerliğe dayalı yöntemler (Blondel vd., 2008) (Newman & Girvan, 2004) kullanılmıştır. Fatemi ve Tokarchuk (Fatemi & Tokarchuk, 2012) bu yöntemlerin öne çıkanlarını (Blondel vd., 2008; Lancichinetti, Fortunato, & Kertesz, 2009; Wang & Fleury, 2011) karşılaştırmalı olarak analiz ederek en başarılı sonuçların modülerlik tabanlı bir topluluk algılama algoritması olan Louvain tarafından elde edildiğini göstermişlerdir. Bu algoritma, kullanıcılar arasındaki benzerliklere dayalı olarak oluşturulan bir ağı işleterek, kullanıcıların topluluklarda artan ayrıntı düzeyine sahip hiyerarşik bölümleri içeren bir ağaç oluşturur. Ayrıca, GÖS literatüründe, kullanıcı gruplarını otomatik olarak tanımlamak için Louvain algoritmasını kullanan birçok başka çalışma mevcuttur (Boratto vd., 2009; Boratto & Carta, 2010; Fatemi & Tokarchuk, 2013).

Bununla birlikte, Boratto ve ark. (Boratto, Carta, & Satta, 2010), Louvain ve k -ortalama (k -means-KM) kümeleme algoritmasını karşılaştırmalı olarak analiz ederek uygun kullanıcı gruplarının tanımlanmasında KM algoritmasının daha başarılı olduğunu göstermişlerdir. Bu çalışmada, başlangıçta, bireyler için eksik derecelendirmeleri geleneksel bir CF algoritması yoluyla tahmin edilmiştir ve daha sonra, benzer ilgi alanlarına sahip kullanıcılardan oluşan grupları tanımlamak için hem kullanıcıların sağladığı derecelendirmeleri hem de tahmin edilen derecelendirme değerlerini içeren doldurulmuş kullanıcı $\times$ ürün matrisi üzerinde KM algoritması uygulanmıştır. Bu yaklaşım, yapılan diğer bir çalışmada P&C olarak adlandırılarak, orijinal kullanıcı $\times$ ürün matrisi üzerinde KM algoritmasının uygulanmasına göre üstünlüğü ortaya konmuştur (Boratto & Carta, 2014; Boratto & Carta, 2015a). Ayrıca, P&C yaklaşımı, GÖS literatüründe öne çıkan birçok çalışmada kullanıcı gruplarını otomatik olarak

tanımlanmasında etkin bir biçimde kullanılmıştır (Boratto, Carta, & Fenu, 2016; Boratto, Carta, & Fenu, 2017).

Bazı çalışmalarda, kullanıcı gruplarının KM algoritması aracılığı ile otomatik olarak tanımlanmasında, kullanıcıların temsili için sağladığı derecelendirmeleri içeren vektörler yerine, matris çarpanlarına ayırma yöntemi ile çıkarılan gizli faktörler kullanılmıştır (Liu vd., 2016; Shi, Wu, & Lin, 2015). Alternatif olarak, Cantador ve Castells (Cantador & Castells, 2011) yaptıkları çalışmada, kullanıcı gruplarını, kullanıcıların ontoloji-tabanlı profilleri üzerinde hiyerarşik kümeleme yaklaşımlarını gerçekleştirerek tanımlamışlardır. Khazaei ve Alimohammadi (Khazaei & Alimohammadi, 2018), konum-tabanlı sosyal ağlarda kullanıcı tercihlerinin/bilgilerinin onların profillerinden dolayı olarak çıkarılmasıyla motive edilmiş bir *k*-medoids kümeleme tekniğine dayanan otomatik bir kullanıcı gruplama modeli önermişlerdir. Bu model, kullanıcıları belirli bir boyuta sahip gruplara ayırmada etkili bir yöntem olmasına rağmen, kullanıcılar hakkında yalnızca ayrıntılı mekânsal bilgilerin (örneğin; kullanıcıların mekânsal yakınlıkları, kullanıcıların tatil günleri ve kullanıcılar arasındaki sosyal ilişkiler) varlığında uygulanabilir.

Grup büyüklüğü, grup yapısını yönetmede ve kullanıcıların üretilen önerilerden olan memnuniyetini arttırmada önemli bir unsurdur. Ancak, mevcut yaklaşımların büyük çoğunluğu grup büyüklüğü üzerinde herhangi bir kısıtlamaya sahip değildir. Ayrıca, kullanıcıların derecelendirme vektörlerinde ortaya sıklıkla çıkan seyreklik sorununu ele alan bazı çalışmalar olmasına rağmen, yeni çözümlerin geliştirilmesine ihtiyaç olduğu söylenebilir. Son olarak, gruplandırma yaklaşımlarında genellikle kullanıcıları tam bir biçimde temsil etmekten yoksun olan derecelendirme vektörleri dikkate alınır. Bu nedenle, gruplama aşamasında kullanıcıların temsil süreci, onlar hakkındaki ek bilgilerle güçlendirilerek daha uygun gruplar elde tanımlanabilir.

4. YENİ OTOMATİK KULLANICI GRUPLAMA YÖNTEMLERİ

Bu bölümde, grup önerisi yaklaşımlarında kullanılmak üzere kullanıcı gruplarının otomatik olarak tanımlanması için yeni yaklaşımlar önerilmektedir.

4.1. Motivasyon

Daha önce Giriş Bölümü'nde açıklandığı üzere, grupların otomatik olarak tanımlanmasında ortaya çıkabilecek bazı zorluklar mevcuttur. Bu zorluklardan ilki, otomatik grup tanımlama bağlamında yapılan çalışmalarda genellikle tanımlanacak grupların boyutu ile ilgili herhangi bir belirlemenin söz konusu olmamasıdır. Ayrıca, grupların tanımlanması için en yaygın yöntemlerden birisi olan kümeleme algoritmalarının uygulanmasında kritik bir basamak olan kullanıcılar arasındaki benzerlik hesabı, kullanıcı derecelendirme vektörlerinin boşluklu (yani seyrek) yapıda olması sorunundan muzdariptir (Bilge & Polat, 2013). Bu sorun, benzerliklerin doğru bir biçimde hesaplanmasına engel olarak, uygun kullanıcı gruplarının tanımlanma sürecini zorlaştırmaktadır. Son olarak, benzerlik hesaplama aşamasında, kullanıcıların yalnızca ürünler için sağladığı derecelendirmeleri içeren vektörlerin dikkate alınması, benzer kullanıcıların belirlenmesinde tek başına yeterli olmayabilir. Bu bağlamda, kullanıcıların yaş, cinsiyet ve meslek gibi demografik bilgilerinin kümeleme aşamasında dikkate alınması, daha homojen grupların üretilmesine ve böylece üretilcek grup önerilerinden olan genel memnuniyetin artmasına katkı sağlayacaktır. Bu doğrultuda, aşağıda, grup önerilerinden elde edilen memnuniyetin, kişilerin demografik özelliklerine bağlı olabileceği ve tanımlanacak grupların boyutu ile ilgili bir kısıtlamanın tanımlanabileceği üç somut uygulama senaryosu listelenmiştir.

- a. Steam⁹ gibi çevrimiçi oyun platformlarında, oyuncular için diğer bireyler ile birlikte oyun oynayabilecekleri odalar mevcuttur. Oyunların çoğu için, sınırlı sayıda oyuncu ve önerilen maksimum oyuncu sayısı tanımlanabilir. Bu nedenle, platform oyuncuları başka bireylerle oynamak için farklı odalara ayrılır. Bununla birlikte, bu oyuncuları oyun zevklerine, yaş gruplarına, cinsiyetlerine ve mesleklerine göre gruplandırmak, bu oyuncuların birlikte bir oyun oynamaktan daha fazla zevk almasına olanak sağlar. Bu da, çevrimiçi oyunların sosyal bir etkinlik olarak değerlendirebileceği nedeni ile oyuncuların memnuniyetini artırır. Ayrıca, platform, tanımadığı ancak benzer ilgileri olan kullanıcılarla yeni oyun deneyimleri arayan

⁹ <https://store.steampowered.com/>

tecrübesiz oyunculara çok oyunculu oyunlar önerebilir. Böyle bir senaryoda, tüm grup üyeleri oyunları beraber oynamaktadır.

- b. Uydu sistemleri ve mobil IPTV servis sağlayıcıları, bazı video akışları diğerlerinden daha yüksek bit hızları gerektirdiğinden, değişken bant genişliği gereksinimlerine sahip birden çok video hizmeti için kullanıcıların talebini karşılamaya çalışır. Bu sistemlerde, bant genişliği kısıtlaması nedeniyle, hem hizmet hem de deneyim kalitesini sağlamak için belirli bir kanalı görüntüleyen kullanıcı sayısının kısıtlanması gerekebilir. Ayrıca, servis sağlayıcının kanal üyeleri için grup olarak kişiselleştirilmiş TV programı önerileri sunduğu varsayıldığında, üyelerin kişisel tercihlerinin yanında aile durumu ve çalışma saatleri gibi demografik özelliklerin dikkate alınması gerekir. Böyle bir senaryoda, kanal üyeleri üretilen önerilerden bireysel olarak faydalanır.
- c. Çok sayıda turiste hizmet veren bir seyahat acentesi, tapınaklar, panoramik mekânlar ve müzeler gibi farklı özelliklere sahip çeşitli turistik noktalara ziyaret amacıyla günlük şehir gezi turları düzenlemektedir. Bu turlar, sınırlı sayıda koltuk içeren otobüslerle yapılması nedeniyle, müşterilerin turlara katılması için belirli bir maksimum boyutta küçük alt gruplara ayrılmasını gerektirmektedir. Ayrıca, düzenlenen turlardan olan genel memnuniyeti en üst seviyeye çıkarmak için, gruplar bireylerin kişisel tercihlerinin yanında yaş grupları ve eğitim düzeyi gibi demografik özellikler göz önünde bulundurularak üretilir. Böyle bir senaryoda, grup üyeleri turları bir grup olarak birlikte deneyimlemektedir.

Bu çalışmada, otomatik kullanıcı gruplandırma, bu bölümün başında özetlenen problemlerin çözümü ve yukarıda örneklendirilen bazı uygulama senaryolarında, grup önerisi ihtiyaçlarını karşılayabilmek için birbiri üzerine inşa edilmiş çeşitli çözümler üretilmiştir. Bu kapsamda, literatüre kazandırılan katkılar aşağıda listelenmiştir.

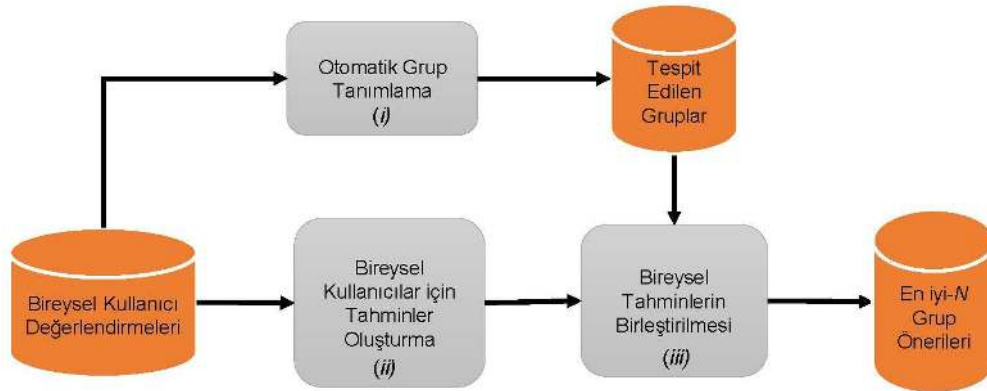
- a. Önceden tanımlanan bir eşik değerini aşmayan, yani belirlenen maksimum boyuttaki grupların otomatik olarak tanımlanmasını sağlayan BKM algoritması temelli yeni bir otomatik kullanıcı gruplandırma yaklaşımı geliştirilmiştir. Bu yaklaşım, aynı zamanda, sisteme sonradan dâhil olan yeni bir kullanıcı için en uygun grubun belirlenmesine yardımcı olarak özellikle büyük veri kümelerinde karşılaşılabilecek potansiyel ölçeklenebilirlik sorununa çözüm olma potansiyeline sahiptir.

- b. Büyük ve çoğunlukla seyrek yapıda olan kullanıcı derecelendirme vektörlerinin, kümelemedeki benzerliklerin belirlenmesi sürecinde neden olabileceği hem seyreklik hem de karmaşıklık sorununun üstesinden gelmek için, bu kullanıcıların tamamen yoğun ve daha küçük boyutta vektörlerle temsilini sağlayan iki farklı ürün JTP oluşturma yaklaşımı kullanılmıştır.
- c. Grup homojenliğini arttırmak ve sonuç olarak üretilecek grup önerilerinden olan genel memnuniyeti arttırmak için kullanıcıların demografik özelliklerini benzerlik hesaplama sürecine uygun bir şekilde dâhil eden iki farklı strateji geliştirilmiştir.

4.2. Kullanılan Grup Önerisi Yaklaşımının Genel Şeması

Literatürdeki bazı deneysel çalışmalar (Amer-Yahia vd., 2009; Boratto & Carta, 2015b), Bölüm 2.3'te ayrıntılı bir biçimde açıklanan grup önerisi üretme stratejilerini karşılaştırmalı olarak analiz ederek, BTB yaklaşımının diğerlerine kıyasla daha başarılı grup önerileri ürettiğini göstermesi nedeniyle, bu çalışmada BTB yaklaşımı benimsenmiştir. Bununla birlikte, GÖS senaryolarının büyük çoğunluğunda kullanıcı gruplarının önceden bilinmemesi nedeniyle, benzer tercihlere sahip kullanıcılardan oluşan grupların otomatik olarak tanımlanması gerekmektedir.

Bu çalışmada kullanılan grup önerisi yaklaşımının genel şeması Şekil 4.1'de sunulmaktadır. Buna göre, ilk olarak, (i) kullanıcı gruplarının KM algoritması yardımıyla tanımlanır, ardından (ii) bireysel kullanıcılar için Bölüm 2.2'de açıklanan CF algoritması ile tahminlerin üretilir ve son olarak (iii) en iyi- N grup önerileri sağlamak amacıyla üretilen bireysel tahminlerin birleştirilir. Burada, birleştirme tekniği olarak Avg, LM ve CR teknikleri seçilmiştir.



Şekil 4.1. Kullanılan grup önerisi yaklaşımının genel şeması

4.3. Otomatik Grup Tanımlanması için Yeni BKM-tabanlı Yaklaşımlar

Bu bölümde, uygun kullanıcı gruplarını tespit etmek için yeni yaklaşımlar sunulmaktadır. İlk olarak, BKM kümeleme algoritması uygulanarak, kullanıcıları gruplara ayırmak için kullanılacak bir BDT'nin nasıl oluşturulacağı açıklanmıştır. Ardından, JTP'ler kullanarak önerilen yaklaşımın nasıl iyileştirileceği açıklanmıştır. Bu profiller, kullanıcı $\times$ ürün matrisinin seyrekliğinin neden olduğu olumsuz etkilerden kurtulmak ve kullanılan kümeleme algoritmasının ayırma becerilerini geliştirmek için, kullanıcıların derecelendirmeye dayalı vektörlerinin janrlara göre eşleştirilmesi ile üretilmektedir. Son olarak, kullanıcıların arasındaki ilişkileri daha doğru şekilde tespit etmek ve daha homojen gruplar oluşturmak için JTP temelli benzerliklerin demografik korelasyonlarla nasıl entegre edileceği açıklanmıştır.

4.3.1 BKM ile BDT oluşturma

Kümeleme algoritmaları, genellikle, n adet nesneyi, yüksek küme içi ve düşük kümeler arası benzerliklere sahip k adet kümeye ayırma işlemini gerçekleştirmektedir. GÖS bağlamında, kullanıcı gruplarının otomatik olarak tanımlanması amacıyla bu algoritmalar çevrimdışı bir işlem olarak yaygın olarak kullanılmaktadır. Bu işlem, ayrıca, sisteme sonradan dâhil olan bir kullanıcının ait olabileceği grubun belirlenmesi de sağlar. Bu bağlamda en çok kullanılan algoritmalarından birisi, kullanıcıları önceden tanımlanmış k adet kümeler ayırmayı sağlayan KM algoritmasıdır (Boratto, Carta, & Fenu, 2017). Bu algoritma, başlangıçta k rastgele kullanıcıları küme merkezleri olarak seçer ve ardından tüm kullanıcılar ve her küme merkezi arasındaki benzerlikleri hesaplayarak kullanıcıları kendisine en yakın kümeye atar. Ardından, küme merkezlerini tekrar hesaplar. Küme merkezleri belirgin bir şekilde değişmeyene kadar bu işlem tekrarlanır. Bununla birlikte, KM algoritması ile oluşturulan gruplar genellikle farklı boyutlarda olur ve küme sayısının artması ile bu kümeleri oluşturmak için gereken hesaplama süresinin dramatik bir şekilde artar (Bilge & Polat, 2013).

Diğer yandan, BKM kümeleme algoritması, genellikle öneri sistemleri (Bilge & Polat, 2013), görüntü işleme (Zhao, Li, & Cang, 2015), doküman kümeleme (Zhao, Deng, & Ngo, 2018) gibi çok miktarda verinin incelendiği alanlarda kullanılan hiyerarşik bir kümeleme algoritmasıdır. Başlangıçta, bu algoritma, tüm kullanıcıları tek bir küme olarak değerlendirir ve ikiye ayırma adımı olarak da bilinen KM algoritmasını kullanarak onları iki alt kümeye bölme işlemini gerçekleştirir. Bu işlem, istenen sayıda kümeye ulaşana

veya belirli bir kriter sađlanana kadar tekrarlanır. BKM algoritması, özellikle küme sayısı çok olduđuunda KM algoritmasına kıyasla daha etkilidir. Bununla birlikte, KM algoritması ile elde edilen grupların boyutlar büyük ölçüde farklılık gösterirken BKM ile elde edilen gruplar benzer boyutlara sahip olmaktadır.

Kümeleme algoritmaları, benzer kullanıcılardan oluşan grupları belirlemek için yararlı bir yaklaşımdır (Sarwar vd., 2002). Kullanıcıları, bir kriteri göz önünde bulundurarak gruplara ayırmada başarılı olmalarının yanında, en iyi performanslarını bir grup kullanıcıyı iki alt gruba ayırırken sergilerler. Bunun nedeni, yalnızca iki grup bulunması durumunda, kullanıcıların gruplara ait olan üyelik değerlerinin yakın olmasının pek olası olmaması sayesinde algoritmanın hassasiyetinin artmasıdır. Bu nedenle, benzer kullanıcı gruplarını saptamak için gruplandırma sonuçlarına göre kullanıcıları tekrarlayan bir biçimde iki alt gruba ayırarak bir BDT oluşturan, ikiye ayırma k -ortalama kümeleme tabanlı bir yaklaşım geliştirilmiştir. Bu yaklaşım kısaca BKM-tabanlı olarak adlandırılmaktadır.

Kullanıcı $\times$ ürün derecelendirme matrisinde, kullanıcıların orijinal derecelendirme tabanlı vektörlerinin $U_{m \times n}$ ile ve maksimum grup boyutunun P ile temsil edildiđi durumda, BKM algoritmasının temel işlem adımları Prosedür 4.1’de sunulmuştur. Buna göre, her levelde, ikiye ayırma adımı olarak adlandırılan U üzerinde KM algoritması çalıştırılarak kullanıcılar iki alt gruba ayrılır. Aynı zamanda, bu grupların merkezleri de sisteme sonradan dâhil olacak bir kullanıcının grubunun belirlenmesinde kullanılmak üzere kaydedilir. Oluşturulan alt grupların herhangi birinde kullanıcı sayısı, P değerinden yüksek olursa, söz konusu grubu ikiye bölmek amacıyla BKM algoritması rekürsif olarak tekrar çalıştırılır. Bu işlem, tüm gruplarda en fazla P adet kullanıcı olana kadar tekrarlı bir biçimde devam eder. Bu nedenle, P değeri BKM algoritması için bir durma kriteri olarak kabul edilebilir. Böylece, grup merkezlerinin dal düğümlerde, benzer kullanıcı gruplarının yaprak düğümlerde kaydedildiđi bir BDT oluşturulmuş olur.

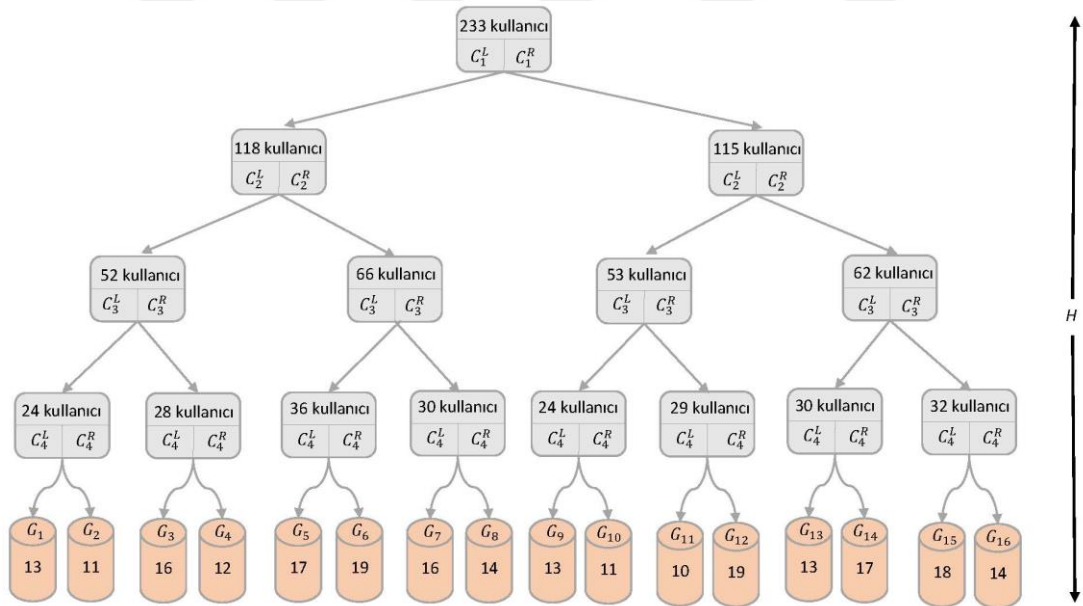
Prosedür 4.1. BKM-tabanlı yaklaşım

```
1:  function BKM( $P, U_{m \times n}$ ) ► BKM-tabanlı gruplama
    Başlatma:
2:     $idx(x) \leftarrow null$  ► Yönlendirme (sağ veya sol)
3:     $BDT.merkezler \leftarrow null$  ► Grup merkezleri
4:     $BDT.sol \leftarrow null$  ► Sol alt-ağaç kullanıcıları
5:     $BDT.sağ \leftarrow null$  ► Sağ alt-ağaç kullanıcıları
6:     $BDT.RST \leftarrow \text{yeni bir BDT}$  ► Sol alt-ağacı
7:     $BDT.LST \leftarrow \text{yeni bir BDT}$  ► Sağ alt-ağacı
    Gruplama:
8:     $[idx, BDT.merkezler] = KO(2, U)$  ► İkiye ayırma adımı
9:    for all kullanıcı  $k_i$  in  $U (i \leftarrow 1 \text{ to } m)$  do
10:     if  $idx(k_i) = 'sol'$  then
11:        $k_i$ 'yi  $BDT.LST$ 'e ekle
12:     else
13:        $k_i$ 'yi  $BDT.RST$ 'e ekle
14:     end if
15:   end for
16:   if  $\text{boyut}(BDT.sağ) > P$  then
17:      $BDT.RST = BKM(P, BDT.sağ)$ 
18:   end if
19:   if  $\text{boyut}(BDT.sol) > P$  then
20:      $BDT.LST = BKM(P, BDT.sol)$ 
21:   end if
22:   return BDT
23: end function
```

Standart KM algoritmasında, tüm kullanıcılar ve tüm küme merkezleri her yinelemede, benzerlik hesaplamasına dâhil olurken, BKM'deki her ikiye ayırma adımında, sadece bir grup kullanıcı ve iki grup merkezi dikkate alınmaktadır. Bu durum hesaplama süresini önemli ölçüde azaltmaktadır. Ayrıca, BKM aracılığı ile bir BDT oluşturulduktan sonra, sisteme sonradan dâhil olan yeni bir kullanıcının ait olduğu grup,

oluşturulan BDT’de basitçe yukarıdan aşağıya doğru gezinerek belirlenir. Bu yapılırken, her levelde, iki grup merkezi için benzerlik değeri hesaplanır. Daha yüksek benzerlik değerine karşılık gelen düğüm, Prosedür 4.1’de idx olarak tanımlanan bir değişkende saklanarak, gezinmedeki bir sonraki düğümü temsil eder. Bu nedenle, yeni kullanıcının ait olduğu nihai grup en fazla $2 \times (H - 1)$ benzerlik hesabı ile belirlenir. Burada, H değeri, oluşturulan BDT’nin yüksekliğini ifade eder. H değeri toplam kullanıcı sayısını ifade eden (m) değerine bağlı olsa da, bu değerden genellikle oldukça daha küçüktür. Ayrıca, sezgisel olarak, H değeri, ölçeklenebilirlik sorunundan mustarip sistemlerde, belirlenecek toplam grup sayısından (k) çok daha küçüktür. Dolayısıyla, yeni gelen bir kullanıcının grubu m (Baltrunas, Makcinskas, & Ricci, 2010) veya k yerine en fazla $2 \times (H - 1)$ benzerlik hesaplaması ile tanımlanabilir ve bu da toplam hesaplama süresini oldukça azaltır.

BKM algoritması tarafından oluşturulan bir BDT örneği Şekil 4.2’de gösterilmektedir. Bu örnekte, başlangıçtaki kullanıcı sayısı 233’tür ve P değeri 20 olarak seçilmiştir.



Şekil 4.2. BKM ile oluşturulan örnek bir BDT

İlk levelde, kullanıcılar, 118 ve 115 üye içeren iki gruba ayrılmıştır. Grup merkezleri, kök düğümün alt tarafında C_1^L ve C_1^R olarak kaydedilir. Burada, alt indisler, her bir sonraki aşamada bir arttırılan mevcut BDT’nin yüksekliğini ifade etmektedir. Ayrıca, üst indisler, grubun hangi alt ağaca (sol veya sağ) doğru ilerleneceğini ifade

etmektedir. Benzer şekilde, kök düğümün sağ alt ağacındaki 115 kullanıcı C_2^L ve C_2^R merkezleri olmak üzere 53 ve 62 kullanıcıdan oluşan iki gruba ayrılmıştır. Bu şekilde, tüm dal düğümleri en fazla P kullanıcısı içerene kadar ikiye ayrılma işlemi gerçekleştirilir. Bu işlemin sonunda, yaprak düğümlerde, maksimum P büyüklüğündeki gruplar elde edilir. Örneğin, 16 indeksine sahip en sağdaki yaprak düğüm 14 kullanıcı içerir ve G_{16} olarak gösterilmiştir. Nihayetinde, ilerleme prosedürünü gerçekleştirmek için her bir dal düğümünde iki grup merkezinin ve her bir yaprak düğümünde benzer boyutlardaki uygun kullanıcı gruplarına sahip bir BDT inşa edilmiş olur. Ayrıca, bu örnekte, yeni bir kullanıcının ait olduğu grubun belirlenmesi için 16 yerine en fazla $2 \times (4 - 1) = 6$ benzerlik hesabı ile belirlenmiş olur. Bu da, BKM yaklaşımının, hesaplama sayısını yaklaşık 3 kat azalttığını göstererek, benzer boyutlarda grupları tanımlayabilmesinin yanında, özellikle büyük veri kümelerinde karşılaşılabilecek ölçeklenebilirlik sorununa bir çözüm getirdiği söylenebilir.

4.3.2 JTP ile BKM yaklaşımının geliştirilmesi

Kullanıcı profilleri arasındaki benzerliklerin doğru bir şekilde hesaplanması, benzer kullanıcı gruplarının tanımlanmasının önemli bir unsurudur. Bu işlem, genellikle, Denklem 2.1'de gösterildiği gibi ortak olarak derecelendirilen ürünler göz önünde bulundurularak gerçekleştirilir. Ancak, sistemdeki mevcut ürün sayısının artması, kullanıcı profillerinin seyrekleşmesine neden olup ortak olarak derecelendirilen ürünleri bulmayı zorlaştırır. Bu durumda, ortak olarak derecelendirilen ürünler bulunsa bile, bu ürünlerin sayıları kullanıcı benzerliklerinin doğru şekilde hesaplanması için yeterli olmaz. Ayrıca, ürün sayısının artması, benzerliklerin belirlenmesi için gereken hesaplama süresini dramatik bir şekilde arttırarak ölçeklenebilirlik sorununa neden olur.

Veri seyrekliği ve zaman karmaşıklığı sorunlarının üstesinden gelebilmek için, kullanıcı profilleri, derecelendirme sağladıkları ürünlerin janrları göz önünde bulundurularak oluşturulabilir (Bilge & Polat, 2011). Burada temel amaç, seyrek olan büyük boyuttaki kullanıcıların U içerisinde bulunan derecelendirme vektörlerini, tam olarak yoğun ve çok daha küçük JTP'lere eşlemektir. Bu profiller, kullanıcılar arasında ortak olarak derecelendirilmemiş ürünler olmasa bile, ilgili kullanıcılar arasındaki benzerliklerin hesaplanabilmesine olanak tanır. Kullanıcıları, benzerlik hesaplama aşamasında JTP ile temsil edebilmek için, ürünleri janrlarına göre kategorilere ayırmak gerekir. Örneğin, bu janrlar, filmler için *korku*, *drama* veya *komedi*, müzikler için *pop*,

klasik veya *türkü*, kitaplar için *otobiyografi*, *macera* veya *romantizm* olabilir. Ayrıca, sistemdeki bütün ürünlerin ait olabileceği toplam janr sayısının sabit bir değer olması ve toplam ürün sayısından çok daha az olması nedeniyle, benzerliklerin hesaplanması için gereken süre bütün kullanıcılar için aynı olmasına ve önemli ölçüde azalmasına imkan tanır. Bir önceki bölümde açıklanan BKM yaklaşımının performansını artırmak amacıyla, kullanıcıların orijinal derecelendirme vektörlerini kullanmak yerine, aşağıda ayrıntılı olarak açıklanan ve kullanıcılar için JTP üretmeyi sağlayan iki farklı kullanıcı profili oluşturma yaklaşımı benimsenmiştir.

1. Alım-tabanlı Kullanıcı-janr Profili (Possession-based User-genre Profile-

PBP): Bu kullanıcı profilleri, ilgili kullanıcının mevcut ürünleri derecelendirip derecelendirmediğini göz önünde bulundurularak üretilir. Eğer bir ürün derecelendirilmişse, ilgili ürünün ait olduğu janrlar bir arttırılır. Bununla birlikte, PBP yaklaşımı, ürünün sadece değerlendirip değerlendirmedeği ile ilgilenir ve kullanıcı tarafından sağlanan derecelendirmenin büyüklüğünü göz ardı eder.

2. Derecelendirme-tabanlı Kullanıcı-janr Profili (Rating-based User-genre

Profile-RBP): Bu kullanıcı profilleri, PBP yaklaşımına benzer şekilde ilk olarak kullanıcının mevcut ürünler için bir derecelendirme sağlayıp sağlamadığını kontrol eder. Eğer, kullanıcı bir ürün için bir derecelendirme sağlamışsa, ilgili ürünün ait olduğu janrlar derecelendirme değeri kadar arttırılır. Diğer bir deyişle, kullanıcı bir ürünü derecelendirerek bir görüş bildirmesi durumunda, bu ürünün her janrası derecelendirme değeri kadar artar. Dolayısıyla, RBP yaklaşımı, hem ürün için bir derecelendirme sağlanıp sağlanmadığını kontrol eder hem de ilgili ürünün ne oranda beğenildiği ile ilgilenir.

PBP ve RBP yaklaşımlarının nasıl çalıştığını daha iyi anlamak için, Tablo 4.1’de beş kitap için iki kullanıcının 5-yıldızlı bir yelpazede olan derecelendirmelerini içeren küçük bir kullanıcı-ürün matrisi verilmiştir. Ayrıca, her kitap Tablo 4.2’de gösterildiği gibi, {Klasik, Efsane, Fantezi, Gizem, Romantik} kümesinden en az bir janra aittir. Bu tabloda, bir kitap janra ait ise, karşılık gelen hücre 1, ait değil ise 0 ile ifade edilmiştir.

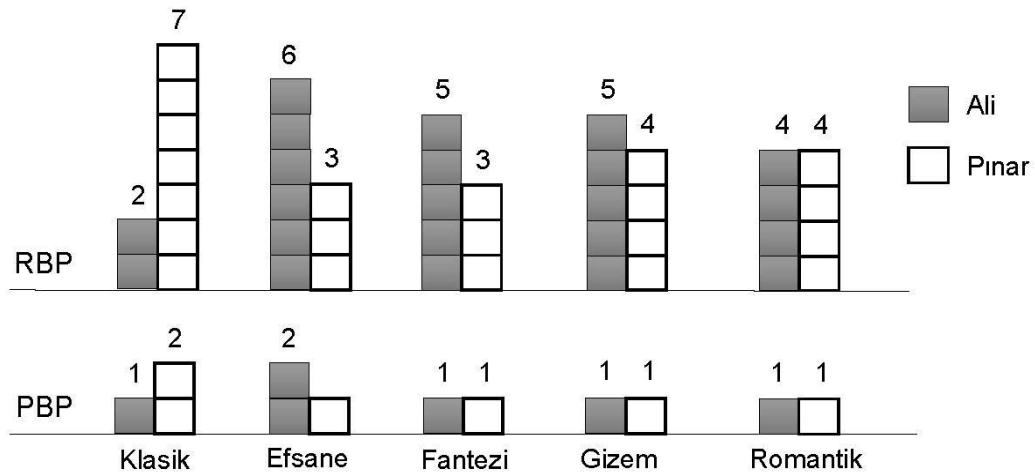
Tablo 4.1. Küçük bir kullanıcı $\times$ ürün matrisi

	<i>kitap₁</i>	<i>kitap₂</i>	<i>kitap₃</i>	<i>kitap₄</i>	<i>kitap₅</i>
<i>Ali</i>	-	2	5	4	-
<i>Pınar</i>	4	-	-	-	3

Tablo 4.2. Kitapların janrları

	<i>Klasik</i>	<i>Efsane</i>	<i>Fantezi</i>	<i>Gizem</i>	<i>Romantik</i>
<i>kitap₁</i>	1	0	0	1	1
<i>kitap₂</i>	1	1	0	0	0
<i>kitap₃</i>	0	0	1	1	0
<i>kitap₄</i>	0	1	0	0	1
<i>kitap₅</i>	1	1	1	0	0

Tablo 4.1 ve 4.2’de verilen örneğe göre, kullanıcılar için oluşturulan PBP’ler ve RPB’ler Şekil 4.3’de gösterilmektedir. Buna göre, Ali için üretilen PBP sırasıyla *Klasik*, *Efsane*, *Fantezi*, *Gizem* ve *Romantik* janrları için [1, 2, 1, 1, 1] olmuştur. Çünkü, Ali, *Klasik* janrasına ait olan üç kitap olan *kitap₁*, *kitap₂*, ve *kitap₅*’den yalnızca birisi için bir derecelendirme sağlamıştır. Benzer şekilde, janrası *Efsane* olan üç kitap olan *kitap₂*, *kitap₄*, ve *kitap₅*’den iki tanesi için bir derecelendirme sağladığı için, üretilen PBP’de *Efsane* janrası için olan hücre 2 değerini almıştır. Benzer şekilde, Pınar’ın PBP’si [2, 1, 1, 1, 1] olarak üretilmiştir. Diğer yandan, Ali ve Pınar için üretilen RBP’ler sırasıyla [2, 6, 5, 5, 4] ve [7, 3, 3, 4, 4] olmuştur. Ali tarafından derecelendirilen kitaplar arasında sadece *kitap₂*, *Klasik* janrasına ait olduğu için, Ali’nin bu kitap için sağladığı derecelendirme değeri onun RBP’sindeki *Klasik* değerini oluşturmuştur. Ayrıca, Ali, *Efsane* janrasına ait kitaplardan *kitap₂* ve *kitap₄* için sırasıyla 2 ve 4 derecelendirmelerini sağladığı için, RBP’sindeki *Efsane* janrası için olan hücresi $2 + 4 = 6$ olarak hesaplanmıştır.

**Şekil 4.3.** Ali ve Pınar için janr-tabanlı profiller

Derecelendirilen ürünlerin sayısının kullanıcıdan kullanıcıya değişme ihtimali ve her ürünün farklı sayıda janraya ait olma ihtimali, üretilen JTP’ler için bir normalizasyon işleminin uygulanmasını gerektirmektedir. Örneğin, Tablo 4.1’den de görülebileceği gibi,

Ali üç kitap için derecelendirme sağlarken, Pınar iki kitap için görüş bildirmiştir. Ayrıca, Tablo 4.2'den görülebileceği gibi, *kitap₁* üç janraya ait iken, *kitap₂* iki janraya aittir. Bu sorunların üstesinden gelmek için, üretilen JTP'lerdeki her değeri, ilgili profildeki toplam değere bölünmesi ile JTP'ler normalleştirilir. Bu doğrultuda Ali için normalleştirilmiş PBP ve RBP'ler sırasıyla $[\frac{1}{6}, \frac{2}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6}]$ ve $[\frac{2}{22}, \frac{6}{22}, \frac{5}{22}, \frac{5}{22}, \frac{4}{22}]$ olmaktadır. Benzer şekilde, Pınar için üretilen PBP ve RBP'ler sırasıyla $[\frac{2}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6}, \frac{1}{6}]$ ve $[\frac{7}{21}, \frac{3}{21}, \frac{3}{21}, \frac{4}{21}, \frac{4}{21}]$ şeklinde normalleştirilir.

Tüm kullanıcıların JTP'leri oluşturulduktan sonra, bu küçük boyuttaki profiller üzerinde benzerlikler hesaplanır. Bu örnekte, Ali ve Pınar arasındaki benzerlik, ortak olarak derecelendirme yaptıkları herhangi bir ürün olmadığı halde, üretilen PBP veya RBP'ler aracılığıyla hesaplanabileceği açıkça görülmektedir.

4.3.3. Demografik korelasyonların JTP ile kullanımı

Önceki bölümde, BKM yaklaşımında kullanmak üzere, kullanıcıların orijinal derecelendirmelerden oluşan vektörleri yerine, benzerlik hesaplama sürecinde çeşitli avantajlara sahip JTP'nin nasıl üretileceği açıklanmıştır. Bununla birlikte, JTP, sadece kullanıcıların ürünler için sağladığı derecelendirmeler göz önünde bulundurularak üretildiği için kullanıcılar hakkında başka bir bilgiyi yansıtmazlar. Bu durum, kullanıcıların, benzerlik belirleme sürecinde hala zayıf bir şekilde temsil edilmesine ve sonuç olarak üretilen grup önerilerden eşit derecede memnun olmayacak heterojen grupların oluşmasına neden olur. Bu sorunu çözmek için, kullanıcıların demografik bilgiye dayalı korelasyonlarının, JTP-tabanlı benzerliklerine uygun bir biçimde entegre edilmesi gerekmektedir. Böylece, kullanıcılar arasındaki demografik korelasyon, kullanıcılar arasındaki benzerlik hesaplamasında önemli bir rol oynayarak daha homojen grupların tanımlanmasına yardımcı olur.

Bu amaçla, iki farklı demografik korelasyon temelli yaklaşım geliştirilmiştir. Her iki yaklaşımda da ilk adım, kullanıcıların demografik kategorik vektörlerinin, mevcut demografik bilgileri göz önünde bulundurularak çıkartılmasıdır. Bu çalışmada kullanılan, MovieLens veri kümesi göz önünde bulundurulduğunda, bu tür demografik bilgiler, yaş, cinsiyet ve veri kümesindeki her kullanıcı için 21 olası meslekten oluşan bir seçimden oluşmaktadır. Dolayısıyla, demografik özelliklere dayalı benzerlikler hesaplanmadan

önce, Tablo 4.3’de açıklandığı gibi, kullanıcıların 27 özelliğe sahip Demografi-tabanlı Profillerinin (DTP) oluşturulması gerekmektedir.

Tablo 4.3. Demografi-tabanlı profillerin yapısı

#Özellik	İsim	İçerik	Açıklama
1	Yaş	$Yaş \leq 18$	Her kullanıcı, yalnızca bir yaş kategorisine aittir. İlgili hücre 1, geri kalanlar 0 olarak atanır.
2		$18 < Yaş \leq 29$	
3		$29 < Yaş \leq 49$	
4		$49 < Yaş$	
5	Cinsiyet	Erkek	Kullanıcı cinsiyetini temsil eden hücre. İlgili hücre 1, diğeri 0 olarak atanır.
6		Kadın	
7	Meslek	Yönetici	Kullanıcının mesleği, tek bir hücre ile temsil edilmektedir. İlgili hücre 1, geri kalanlar 0 olarak atanır.
8		Sanatçı	
9		Doktor	
:		:	
27		Yazar	

Tüm kullanıcılar için JTP (PBP veya RBP) ile birlikte DTP oluşturduktan sonra, her bir kullanıcı çifti için benzerliklerin hesaplanması için aşağıda açıklanan iki farklı strateji geliştirilmiştir. Bu stratejilerden birisi, DTP tabanlı benzerlikleri, kullanıcılar arasındaki nihai benzerlik için daha önemli olarak değerlendirirken, diğeri JTP tabanlı benzerliklerin değerli olduğunu varsaymaktadır.

- MUL Stratejisi:** Bu stratejide, iki kullanıcı (a ve u) arasındaki nihai benzerlik değeri ($w(a,u)$), Denklem 4.1’de gösterildiği gibi kullanıcıların DTP tabanlı benzerlik değeri ($w(DTP)_{au}$) ile JTP tabanlı benzerlik değerinin ($w(JTP)_{au}$) çarpılması ile hesaplanmaktadır. Böylece, MUL stratejisinde, demografik korelasyonların ve derecelendirme tabanlı benzerliklerin, kullanıcılar arasındaki nihai benzerlikte eşit bir öneme sahip olduğu varsayılmaktadır.

$$w(a,u) = w(DTP)_{au} \times w(JTP)_{au} \quad (4.1)$$

- AUG Stratejisi:** Bu strateji, Denklem 4.2’de gösterildiği gibi kullanıcılar arasındaki nihai benzerlikte, JTP tabanlı benzerliğin itici güç olduğunu varsayarak, DTP tabanlı benzerlik değerini nihai benzerlik üzerinde ek bir faktör olarak değerlendirir. Dolayısıyla, AUG stratejisi, derecelendirme temelli benzerlikleri belirleyici unsur olarak kabul ederken, demografik korelasyonları nihai benzerlik üzerinde daha etkisiz görür.

$$w(a,u) = w(JTP)_{au} + [w(DTP)_{au} \times w(DTP)_{au}] \quad (4.2)$$

4.4. Deneysel Çalışmalar

Bu bölümde, geliştirilen otomatik kullanıcı gruplandırma yaklaşımlarının etkinliğini incelemek için yapılan deneysel çalışmalar sunulmaktadır. İlk olarak, grup önerisi çalışmalarında kullanıcı gruplandırma bağlamında öne çıkan ve bu çalışmada da kıyaslama algoritması olarak seçilen P&C algoritması açıklanmaktadır. Ardından, gerçekleştirilen deneylerde izlenen metodoloji ayrıntılı biçimde açıklanmaktadır. Son olarak, elde edilen deneysel sonuçlar sunularak değerlendirilmektedir.

4.4.1. Kıyaslama algoritması: P&C

Bu çalışmada geliştirilen yöntemlerin verimliliğini analiz etmek amacıyla, kullanıcıların otomatik olarak gruplandırmasında en popüler yaklaşım olan ve son yıllarda yapılan birçok çalışmada etkin biçimde kullanılan P&C algoritması, kıyaslama algoritması olarak seçilmiştir (Boratto & Carta, 2015; Boratto, Carta, & Satta, 2010). Bu algoritma, kullanıcı gruplarını tanımlamak için gerçekleştirilen kümeleme sürecini iyileştirmek için tasarlanmış bir algoritmadır.

P&C, prensip olarak, yalnızca kullanıcı derecelendirmelerini içeren orijinal kullanıcı $\times$ ürün matrisine değil, kullanıcıların sağladığı derecelendirmelere ek olarak, derecelendirilmemiş hücrelerin Bölüm 2.2’de açıklanan CF algoritması aracılığı ile tahmin edildiği tam matris üzerinde, KM algoritmasını uygulayarak benzer kullanıcı gruplarını belirleyen bir algoritmadır. Bu yaklaşım, kümelemedeki seyreklik sorununun olumsuz etkilerini ortadan kaldırarak, daha uygun kümeleri, yani benzer ilgi alanlarına sahip kullanıcılardan oluşan kullanıcı gruplarının belirlenmesine yardımcı olur.

4.4.2 Deney metodolojisi

Bu çalışmada önerilen grupta yaklaşımının performansını değerlendirmek için yapılan çeşitli deneylerde, Bölüm 2.4’de ayrıntılı olarak açıklanan MLP ve MLM veri kümeleri kullanılmıştır. Ayrıca, deneyler, 10-kat çapraz doğrulama deney metodolojisi kullanılarak gerçekleştirilmiştir. Bu amaçla, ürünler kümesi rastgele biçimde on alt kümeye bölünerek, oluşturulan her alt kümenin, ürünlerin %10’unu içermesi sağlanmıştır. Her yinelemede, alt gruptan birisi test kümesi olarak, kalan dokuz alt grubun kombinasyonu da eğitim seti olarak kullanılmıştır. Dolayısıyla, nihai n DCG değerleri, 10-kat deneysel doğruluk sonuçlarının ortalaması alınarak hesaplanmıştır. Ayrıca, kullanıcı tarafından derecelendirilmeyen ürünleri öneren yaklaşımı cezalandırmak için, her kullanıcı için n DCG hesaplarırken, kullanıcının

derecelendirmedeği ürünler için “0” değeri atanmıştır. Böylece, kullanıcıların ilgisini çeken ürünleri öneren bir yaklaşımın, daha yüksek n DCG değeri elde etmesi sağlanmıştır (Gorla, Lathia, Robertson, & Wang, 2013).

Test ve eğitim kümeleri oluşturulduktan sonra, gruplar, eğitim veri kümesine göre standart KM, P&C ve BKM varyantı olan algoritmalar aracılığı ile tanımlanmıştır. BKM varyantı algoritmaların performansını hem KM hem de P&C ile doğru biçimde karşılaştırmak için, ilk olarak KM ve P&C ile oluşturulacak grupların sayısının (k) belirlenmesi gerekmektedir. Bu amaçla, k değeri, toplam kullanıcı sayısının, BKM yaklaşımındaki grupların maksimum büyüklüğünü ifade eden (P) değerine olan oranı kullanılarak hesaplanmıştır. Örneğin, BKM yaklaşımındaki P değeri 50 olarak ayarlandığı ve veri kümesindeki toplam kullanıcı sayısının 600 olduğu durumda, k değeri $600/50 = 12$ olarak ayarlanmıştır. Ayrıca, grup büyüklüğünün önerilen yaklaşımların performansı üzerindeki etkisini değerlendirmek için, deneylerde, 5 ile 200 arasında değişen çeşitli P değerleri kullanılmıştır. Son olarak, deneylerdeki gruplar, içerdiği kullanıcı sayısına göre küçük, orta ve büyük olarak sınıflandırılmıştır. Burada, 10’dan az kullanıcı içerenler küçük grup, 10 ile 100 arasında kullanıcı içerenler orta grup ve son olarak 100’den fazla kullanıcı içerenler büyük grup olarak kabul edilmiştir.

Grupların tanımlanmasının ardından, Bölüm 2.2’de açıklanan CF algoritması kullanarak her ürün için bireysel tahmin değerleri hesaplanmıştır. Ardından, Avg, LM ve CR birleştirme teknikleri ile her grup için, üyelere üretilen tahminler birleştirilerek grup derecelendirmeleri belirlenmiştir. Hesaplanan grup derecelendirmeleri göz önünde bulundurularak, her grup için en yüksek grup derecelendirmelerine sahip en iyi- N sıralı ürün listeleri grup önerileri olarak üretilmiştir. Burada, ayrıca, öneri listesinin boyutunun, geliştirilen yaklaşımların üzerindeki etkisini incelemek amacıyla 5, 10 ve 20 olarak değişen üç farklı N değeri için deneyler gerçekleştirilmiştir. Son olarak, üretilen en iyi- N grup önerilerinin doğruluğu n DCG ile gözlemlenmiştir.

4.4.3. Deney sonuçları

4.4.3.1. BKM yaklaşımının etkisi

Önerilen BKM algoritmasının benzer kullanıcı gruplarını algılamadaki doğruluk performansını incelemek için grup büyüklüğü, kullanılan birleştirme tekniği ve öneri listesinin boyutu (N) dâhil olmak üzere farklı parametreler ile birçok deney gerçekleştirilmiştir. Bununla birlikte, MLP ve MLM veri kümeleri için BKM yaklaşımı

ile elde edilen ampirik sonuçlar, sırasıyla Tablo 4.4 ve 4.5’de gösterildiği gibi standart KM algoritması ile karşılaştırılmıştır.

Tablo 4.4. *MLP veri kümesi için KM ve BKM algoritmalarının nDCG sonuçları*

Birleştirme Tekniği	Grup Boyutu (P)		Küçük		Orta		Büyük	
	en iyi-N	Model	5	10	25	50	150	200
Avg	5	KM	0,159	0,161	0,166	0,167	0,176	0,194
		BKM	0,171 [†]	0,181 [†]	0,190 [†]	0,194 [†]	0,197	0,218*
	10	KM	0,173	0,184	0,194	0,195	0,196	0,188
		BKM	0,185 [†]	0,194*	0,202*	0,213*	0,216	0,222 [†]
	20	KM	0,208	0,218	0,223	0,229	0,230	0,234
		BKM	0,225 [†]	0,223 [†]	0,239 [†]	0,247*	0,250*	0,258 [†]
LM	5	KM	0,146	0,141	0,126	0,115	0,112	0,124
		BKM	0,164 [†]	0,158*	0,162 [†]	0,141 [†]	0,122	0,124
	10	KM	0,158	0,143	0,141	0,125	0,121	0,112
		BKM	0,176 [†]	0,167 [†]	0,156 [†]	0,155 [†]	0,148 [†]	0,127
	20	KM	0,189	0,169	0,150	0,144	0,124	0,127
		BKM	0,199 [†]	0,189 [†]	0,179 [†]	0,170*	0,146*	0,130
CR	5	KM	0,163	0,162	0,170	0,178	0,198	0,192
		BKM	0,173*	0,177*	0,190 [†]	0,202 [†]	0,224*	0,212*
	10	KM	0,176	0,188	0,188	0,198	0,211	0,236
		BKM	0,189 [†]	0,201 [†]	0,207*	0,227 [†]	0,236*	0,232
	20	KM	0,217	0,220	0,236	0,234	0,240	0,231
		BKM	0,221	0,227	0,244*	0,243	0,257	0,255*

* %95 güven aralığında istatistiksel olarak anlamlı

† %99 güven aralığında istatistiksel olarak anlamlı

Tablo 4.5. *MLM veri kümesi için KM ve BKM algoritmalarının nDCG sonuçları*

Birleştirme Tekniği	Grup Boyutu (P)		Küçük		Orta		Büyük	
	en iyi-N	Model	5	10	25	50	150	200
Avg	5	KM	0,139	0,150	0,166	0,172	0,184	0,188
		BKM	0,152 [†]	0,172 [†]	0,197 [†]	0,214 [†]	0,225 [†]	0,221 [†]
	10	KM	0,147	0,152	0,172	0,170	0,187	0,195
		BKM	0,155 [†]	0,171 [†]	0,199 [†]	0,201 [†]	0,223 [†]	0,222 [†]
	20	KM	0,164	0,175	0,190	0,197	0,210	0,222
		BKM	0,171 [†]	0,186 [†]	0,210 [†]	0,220 [†]	0,236 [†]	0,241 [†]
LM	5	KM	0,119	0,113	0,107	0,090	0,082	0,079
		BKM	0,142 [†]	0,144 [†]	0,142 [†]	0,130 [†]	0,117 [†]	0,111 [†]
	10	KM	0,124	0,122	0,110	0,100	0,090	0,086
		BKM	0,142 [†]	0,145 [†]	0,141 [†]	0,129 [†]	0,113 [†]	0,107 [†]
	20	KM	0,136	0,127	0,122	0,111	0,097	0,099
		BKM	0,152 [†]	0,152 [†]	0,147 [†]	0,137 [†]	0,118 [†]	0,110*
CR	5	KM	0,076	0,072	0,059	0,077	0,067	0,051
		BKM	0,085 [†]	0,087 [†]	0,075 [†]	0,101 [†]	0,085 [†]	0,068 [†]
	10	KM	0,067	0,066	0,059	0,071	0,063	0,065
		BKM	0,076 [†]	0,079 [†]	0,073 [†]	0,085 [†]	0,080 [†]	0,079 [†]
	20	KM	0,093	0,087	0,076	0,091	0,087	0,075
		BKM	0,101 [†]	0,099 [†]	0,089 [†]	0,108 [†]	0,101 [†]	0,089 [†]

* %95 güven aralığında istatistiksel olarak anlamlı

† %99 güven aralığında istatistiksel olarak anlamlı

Tablo 4.4'te görülebileceği gibi, MLP için yapılan deneylerin $nDCG$ sonuçları, BKM algoritmasının her bir farklı ayar için KM algoritmasına kıyasla daha iyi olduğunu göstermektedir. Ayrıca, BKM ve KM algoritması ile elde edilen sonuçları karşılaştırmak için istatistiksel anlamlılık testleri yapılmış ve bu testlerin sonuçları tabloların dipnotlarında verilmiştir. Yapılan tek kuyruklu t -testlerinin sonuçları, BKM yaklaşımının, KM algoritması üzerinde sağladığı iyileştirmelerin, özellikle küçük ve orta boyutlardaki gruplar için istatistiksel olarak anlamlı olduğunu göstermektedir. Bununla birlikte, BKM, büyük kullanıcı grupları için KM algoritmasından daha iyi performans göstermesine rağmen, sağladığı iyileştirmeler anlamlı gözükmemektedir.

Benzer şekilde, Tablo 4.5'ten görülebileceği üzere, MLM veri kümesi için yapılan deneylerde de BKM yaklaşımı KM algoritmasından daha iyi performans göstermiştir. MLM veri kümesi MLP'ye kıyasla daha büyük ve daha seyrek olmasına rağmen, kullanılan birleştirme tekniğinin LM olduğu, N ve P değerlerinin 20 olarak ayarlandığı durumda gözlemlenen iyileştirmeler %95, bu ayar dışındaki diğer tüm iyileştirmeler %99 güven düzeyinde istatistiksel olarak anlamlı gözükmemektedir. Dolayısıyla, özellikle büyük veri kümelerinde, önerilen BKM algoritmasının, daha uygun kullanıcı gruplarını tespit ederek yüksek kaliteli grup önerilerinin üretilmesini sağladığı sonucuna varılabilir. Bu sonuç aynı zamanda, BKM algoritmasının ölçeklenebilirlik açısından da uygunluğunu göstermektedir.

Tablo 4.4'te gösterildiği gibi, MLP için yapılan deneylerde, Avg ve CR teknikleri benzer performanslar göstermesine rağmen, en yüksek doğruluk sonuçları CR ile elde edilmiştir. Diğer bir deyişle, CR tekniği, bireysel tahminlerin birleştirilerek grup önerilerinin üretilmesinden diğer iki teknikten daha iyi performans göstermiştir. Ayrıca, ampirik sonuçlar, birleştirme tekniği olarak LM'nin kullanılmasının, en kötü doğruluk sonuçlarının elde edilmesine neden olduğunu göstermiştir. Bununla birlikte, Avg ve CR tekniklerinin doğruluk performansları artan grup büyüklüğü ile iyileşirken, LM tekniğinin performansı kötüleşmiştir. Bu sonucun temel nedeni, LM tekniğinin diğer üyelerin görüşlerini dikkate almayarak yalnızca grup içinde sağlanan en düşük derecelendirmeyi kullanmasıdır. Bu yaklaşım, özellikle büyük gruplar için grubun genel eğiliminin doğru analiz edilememesine ve sonuç olarak uygun olmayan grup önerilerinin üretilmesine neden olur.

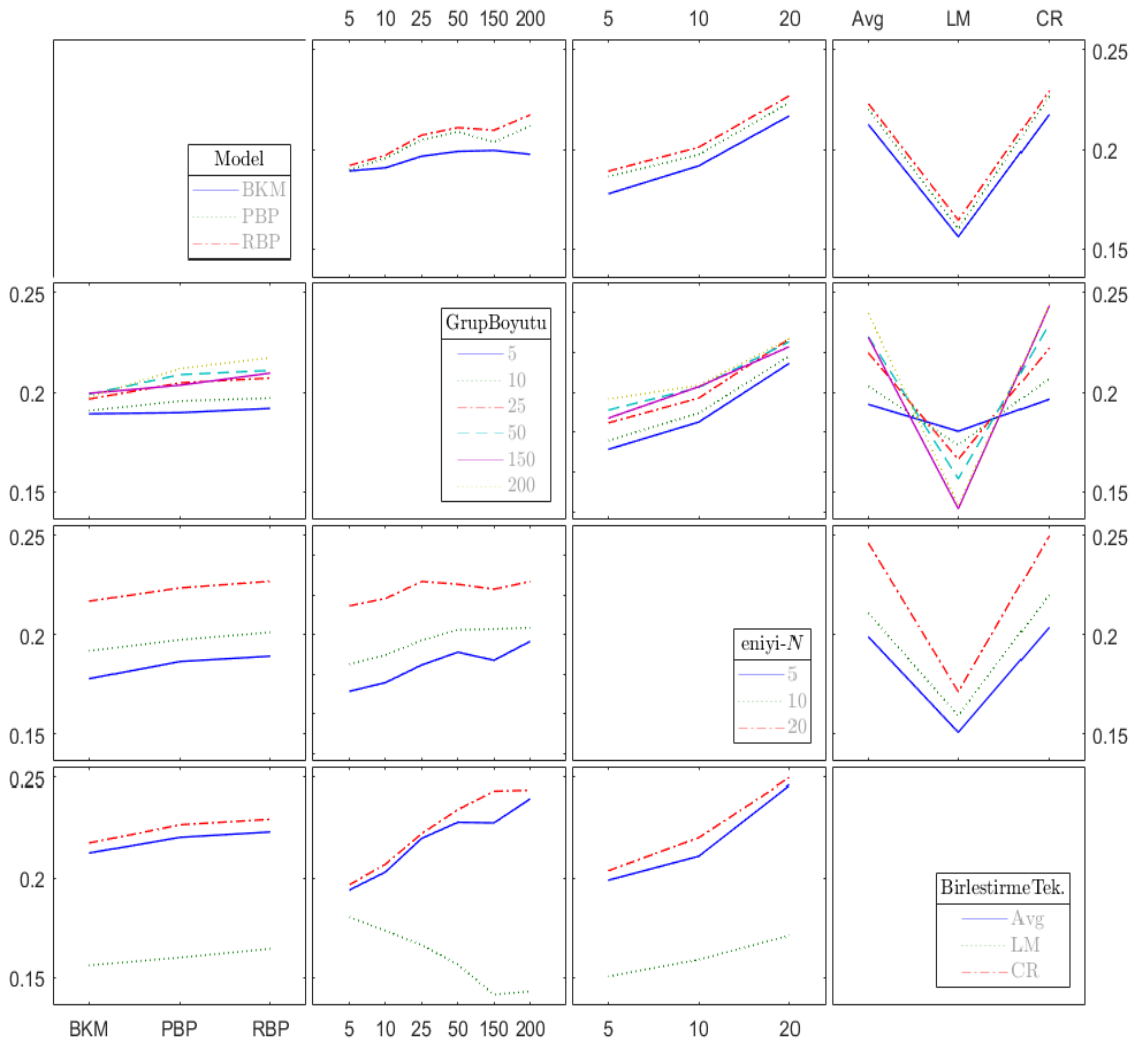
Her ne kadar CR tekniđi, MLP için yapılan deneylerde diđer birleřtirme tekniklerinden daha iyi performans gösterse de, Tablo 4.5’de görüldüğü gibi, MLM için yapılan deneylerde bu tekniđin doğruluk performansı seyreklik sorunları nedeni ile önemli ölçüde düşmüřtür. Bunun nedeni, grup içinde en çok tercih edilen ürünlerin CR tekniđi ile belirleme işleminin başarısı, temel olarak, derecelendirilmemiş ürünler için üretilen tahminlerin güvenilirliğine bađlıdır. Ancak, tahmin üretmek için kullanılan CF algoritması, seyreklik sorunundan önemli ölçüde etkilenecek güvenilir olmayan tahminlerin üretilmesine yol açmaktadır. Ayrıca, MLM veri kümesinde CR tekniđi uygulanırken, ürünlerin ikili olarak kıyaslanması için gereken süre, mevcut ürünlerin sayısı nedeniyle önemli ölçüde artmaktadır. Dolayısıyla, her iki veri kümesi için yapılan deneylerin $nDCG$ sonuçları göz önünde bulundurulduğunda, bireysel tahminlerin birleřtirilmesinde Avg tekniđinin kullanılmasının genellikle diđer tekniklerden daha iyi olduđu sonucuna varılabilir. Çünkü Avg tekniđinde, bir ürüne iliřkin tüm grup üyelerinin derecelendirmeleri, ilgili ürün için hesaplanan grup derecelendirmesinde eřit olarak yansıtılması ile bir fikir birliđi sađlanmaktadır. Son olarak, öneri listesinin boyutunun, elde edilen $nDCG$ sonuçları üzerinde olumlu bir etkisinin olduđu açıkça görülmektedir. Bunun nedeni, daha büyük bir öneri listesinde, grup üyelerinin tercih edeceđi ürünlerin bulunma şansının daha fazla olmasıdır.

4.4.3.2. JTP yaklaşımının etkisi

Bu bölümde, BKM algoritması uygulanırken kullanıcıların orijinal derecelendirmeye dayalı vektörleri (U) yerine JTP vektörlerinin kullanılmasının etkisini arařtırmak için deđiřen parametreler ile gerçekteřtirilen deneylerin sonuçları sunulmaktadır. Her iki veri kümesi için yapılan deneylerden elde edilen $nDCG$ sonuçları dikkate alınarak, U veya JTP (hem PBP hem RBP) kullanmanın BKM yaklaşımının doğruluk performansı üzerindeki etkisi karřılařtırılmıřtır. MLP ve MLM için elde edilen deneysel sonuçlar, etkileřim grafikleri aracılıđı ile sırasıyla řekil 4.4 ve 4.5’de sunulmuřtur.

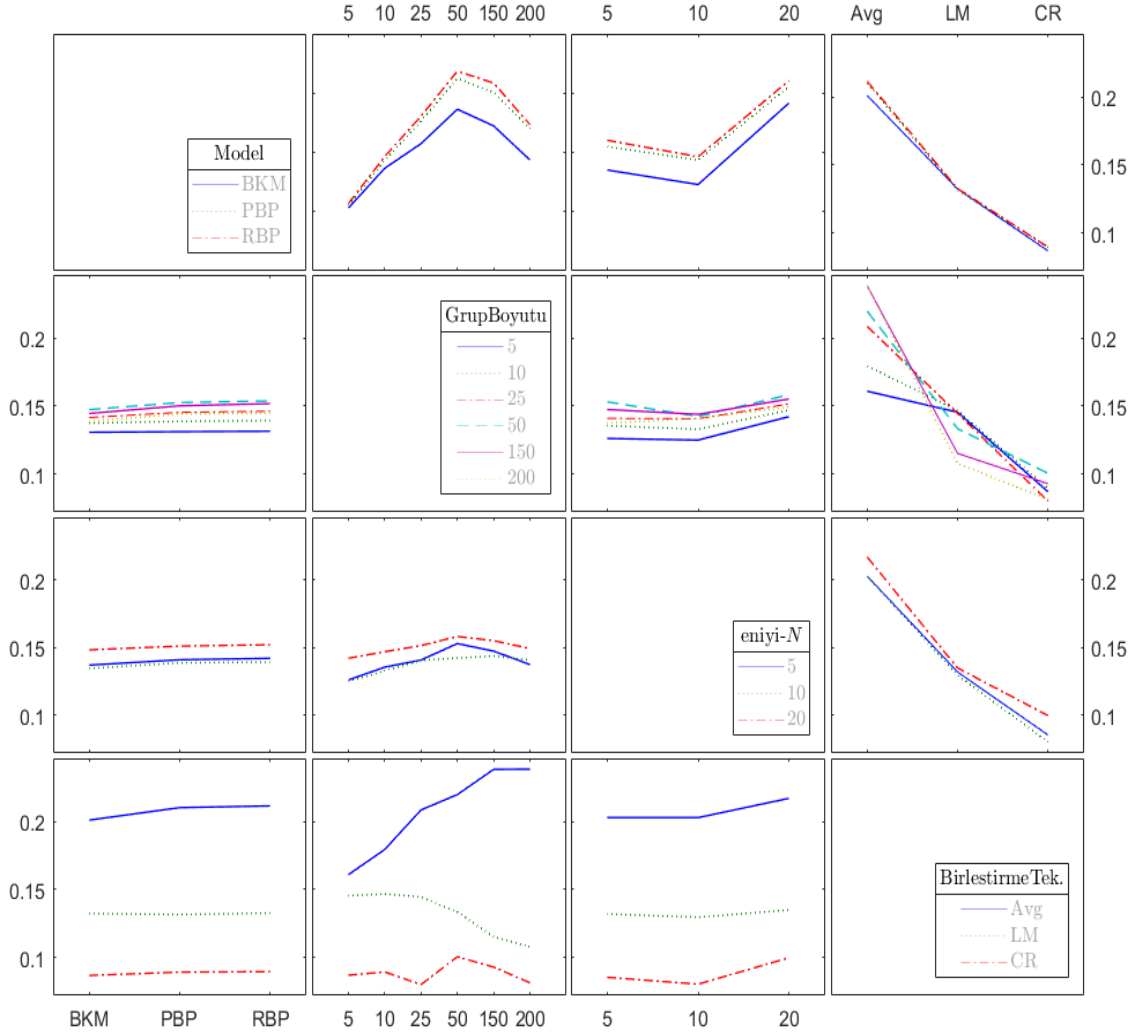
řekil 4.4’ten görüldüğü üzere, MLP veri kümesi için yapılan deneylerin sonuçları dikkate alındığında, BKM yaklaşımında orijinal derecelendirme matrisi yerine JTP’in kullanılması (PBP veya RBP) tüm ayarlamalarda daha iyi sonuçların elde edilmesini sađlamıřtır. Bununla birlikte, RBP kullanımının özellikle büyük kullanıcı gruplarında PBP’ye kıyasla daha uygun kullanıcı gruplarının belirlenmesini sađladıđı açıkça

gözükmektedir. Bunun nedeni, RBP yaklaşımının yalnızca bir ürünün derecelendirip derecelendirilmediğini değil, aynı zamanda kullanıcının ilgili ürünün ait olduğu janrlara karşı ne kadar ilgisi olduğunu veya olmadığını dikkate alarak kullanıcı profillerini oluşturmasıdır. Ayrıca, grup büyüklüğünün 150 olduğu durumlar hariç, hem PBP hem de RBP yaklaşımlarının benzer kullanıcı gruplarını belirlemedeki performansı önemli ölçüde iyileşmiştir. Benzer şekilde, hem BKM hem de JTP'in performansının, öneri listesinin büyüklüğünün artmasıyla biraz daha iyi hale geldiği gözlemlenmiştir.



Şekil 4.4. MLP veri kümesi için JTP'in etkisi

Önceki bölümdeki sonuçlara benzer şekilde, CR ve Avg teknikleri LM'ye göre nispeten daha iyi sonuçlar elde etmiş ve performansları artan grup büyüklüğüyle de iyileşmiştir. Öte yandan, LM tekniğinin doğruluğu grup büyüklüğü arttıkça kötüleşerek, özellikle büyük gruplar için Avg ve CR tekniklerinin çok gerisinde kalmıştır.



Şekil 4.5. MLM veri kümesi için JTP'in etkisi

Şekil 4.5'te gösterilen MLM için yapılan deneylerin $nDCG$ sonuçları da JTP'in uygun kullanıcı gruplarını belirlemede genellikle yararlı bir yaklaşım olduğunu göstermektedir. Bununla birlikte, MLM'deki mevcut ürün sayısının MLP'ye kıyasla oldukça fazla olması nedeniyle, PBP ve RBP yaklaşımları birbirine yakın performans göstermiştir. Bu sonuç, artan kullanıcı ve ürün sayısı ile birlikte her iki profil oluşturma mekanizmasının da belirli bir iyileştirme seviyesine yaklaştığını göstermektedir. Buna ek olarak, küçük ve orta boyuttaki gruplarda artan grup büyüklüğü, üç modelin de uygun kullanıcı gruplarını belirlemedeki doğruluğunu artırmıştır; ancak, büyük gruplarda, grupların büyüklüğünün artması performansın kötüleşmesine neden olmuştur. Ayrıca, öneri listesinin boyutunun artırılması, üç modelin hepsinin de performansının artmasını sağlamıştır; ancak, en iyi-10'luk önerilerde hafif bir düşüş olduğu gözlemlenmiştir. Bununla birlikte, Şekil 4.5'ten görüldüğü üzere, Avg birleştirme tekniği her üç model için

de diğer tekniklerden daha yüksek doğruluk sonuçlarının elde edilmesini sağlamıştır. Öte yandan, CR tekniği, diğerlerinden nispeten daha kötü bir performans göstermiştir. Son olarak, Avg tekniği, daha büyük gruplar için grup modelleme işleminin kalitesini artırırken, önceki bölümde tartışıldığı gibi grup COP ve LM tekniklerinin başarısı grup büyüklüğü arttıkça azalmıştır.

4.4.3.3. DTP’i JTP ile birleştirmenin etkileri

Bu bölümdeki son deneyler, kullanıcılar arasındaki demografik korelasyonları kullanmanın, uygun kullanıcı gruplarını belirleme işlemi üzerindeki etkisini incelemek için gerçekleştirilmiştir. Bir önceki bölümde gösterildiği gibi, RBP yaklaşımının PBP’e kıyasla daha iyi performans göstermesi nedeniyle, bu deneylerde RBP tabanlı benzerlikler üzerine inşa edilmiş MUL ve AUG stratejileri kullanılmıştır. Bu stratejilerin MLP ve MLM veri kümeleri için elde ettiği doğruluk sonuçları hem RBP hem de kıyaslama algoritması olarak seçilen P&C yaklaşımı ile karşılaştırılarak sırasıyla Tablo 4.6 ve 4.7’de gösterilmiştir. Ayrıca, gözlemlenen iyileşmelerin istatistiksel açıdan anlamlı olup olmadığı, yapılan *t*-testler aracılığı ile analiz edilmiş ve sonuçlar tabloların dipnotlarında gösterilmiştir.

Tablo 4.6’da sunulan MLP veri kümesi için elde edilen *n*DCG sonuçlarına göre, küçük kullanıcı grupları için, MUL ve AUG stratejileri RBP yönteminin doğruluk performansını biraz arttırmış olsa bile, sağlanan iyileştirmeler tüm ayarlamalarda istatistiksel açıdan anlamlı görünmemektedir. Diğer yandan, orta ve büyük boyuttaki gruplar için elde edilen sonuçlar, MUL ve AUG stratejilerinin oldukça etkili olduğunu ve doğruluğu önemli ölçüde arttırdığını göstermektedir. Dolayısıyla, grup büyüklüğü arttıkça, kullanıcılar arasındaki benzerliklerin belirlenmesi sürecinde kullanıcı derecelendirmelerine ek olarak onların demografik bilgilerini de göz önünde bulundurmanın, daha homojen kullanıcı gruplarının elde edilmesini sağladığı söylenebilir.

Her iki stratejinin (MUL ve AUG) hem MLP hem de MLM için RBP yönteminin doğruluğunu arttırmasının yanında, Tablo 4.7’den de görüleceği üzere, MLM veri kümesi için çok daha iyi performans göstermişlerdir. Burada, seçilen bütün ayarlamalar ile gözlenen hemen hemen tüm iyileştirmeler istatistiksel açıdan anlamlı bulunmuştur. Dolayısıyla, benzerlik belirleme sürecinde, demografik korelasyonların kullanılması

kullanılan birleştirme tekniğinden ve grup büyüklüğünden bağımsız olarak kullanıcıların daha iyi gruplandırılmasına katkıda bulunabilir sonucu çıkarılabilir.

Tablo 4.6. MLP veri kümesi için DTP ile bezenmiş RBP yönteminin $nDCG$ sonuçları

Birleştirme Tekniği	Grup Boyutu (P)		Küçük		Orta		Büyük	
	en iyi-N	Model	5	10	25	50	150	200
Avg	5	<i>P&C</i>	0,161	0,168	0,189	0,185	0,190	0,186
		<i>RBP</i>	0,176	0,186	0,204	0,220	0,205	0,230
		<i>RBP_{MUL}</i>	0,172 [†]	0,191 [†]	0,219 ^{†*}	0,233 ^{†*}	0,221 ^{†*}	0,245 ^{†*}
		<i>RBP_{AUG}</i>	0,177 [†]	0,194 ^{†*}	0,217 ^{†*}	0,223 [†]	0,230 ^{†*}	0,245 ^{†*}
	10	<i>P&C</i>	0,184	0,191	0,202	0,199	0,199	0,209
		<i>RBP</i>	0,187	0,199	0,217	0,225	0,229	0,231
		<i>RBP_{MUL}</i>	0,186	0,199 [†]	0,229 ^{†*}	0,234 ^{†*}	0,249 ^{†*}	0,245 ^{†*}
		<i>RBP_{AUG}</i>	0,189	0,199 [†]	0,229 ^{†*}	0,234 ^{†*}	0,249 ^{†*}	0,247 ^{†*}
	20	<i>P&C</i>	0,221	0,222	0,220	0,236	0,248	0,231
		<i>RBP</i>	0,221	0,231	0,253	0,260	0,265	0,275
		<i>RBP_{MUL}</i>	0,223	0,232	0,267 ^{†*}	0,271 ^{†*}	0,280 ^{†*}	0,283 [†]
		<i>RBP_{AUG}</i>	0,224	0,234	0,260 ^{†*}	0,268 ^{†*}	0,282 ^{†*}	0,285 ^{†*}
LM	5	<i>P&C</i>	0,164	0,164	0,168	0,159	0,149	0,147
		<i>RBP</i>	0,169	0,165	0,158	0,152	0,137	0,154
		<i>RBP_{MUL}</i>	0,165	0,170 [*]	0,175 [*]	0,163	0,155 [*]	0,159
		<i>RBP_{AUG}</i>	0,165	0,164	0,164	0,167 [*]	0,150	0,163
	10	<i>P&C</i>	0,182	0,185	0,187	0,175	0,188	0,166
		<i>RBP</i>	0,180	0,170	0,168	0,159	0,153	0,156
		<i>RBP_{MUL}</i>	0,183	0,173	0,173 [*]	0,176 [*]	0,166	0,167
		<i>RBP_{AUG}</i>	0,180	0,173	0,176 [*]	0,171	0,163	0,166
	20	<i>P&C</i>	0,210	0,209	0,204	0,183	0,165	0,168
		<i>RBP</i>	0,197	0,195	0,181	0,164	0,150	0,156
		<i>RBP_{MUL}</i>	0,202	0,199	0,195 [*]	0,187 [*]	0,165 [*]	0,162
		<i>RBP_{AUG}</i>	0,205 [*]	0,202 [*]	0,190	0,189 [*]	0,169 [*]	0,165
CR	5	<i>P&C</i>	0,169	0,174	0,179	0,192	0,209	0,176
		<i>RBP</i>	0,180	0,185	0,201	0,222	0,230	0,229
		<i>RBP_{MUL}</i>	0,179	0,193 ^{†*}	0,217 ^{†*}	0,230 [†]	0,243 ^{†*}	0,246 ^{†*}
		<i>RBP_{AUG}</i>	0,180 [†]	0,195 ^{†*}	0,208 ^{†*}	0,223 [†]	0,245 ^{†*}	0,243 ^{†*}
	10	<i>P&C</i>	0,186	0,194	0,201	0,215	0,205	0,226
		<i>RBP</i>	0,191	0,206	0,225	0,234	0,243	0,249
		<i>RBP_{MUL}</i>	0,192	0,208	0,229 [†]	0,246 ^{†*}	0,259 ^{†*}	0,256 ^{†*}
		<i>RBP_{AUG}</i>	0,190	0,212 [*]	0,227 [†]	0,245 ^{†*}	0,262 ^{†*}	0,258 ^{†*}
	20	<i>P&C</i>	0,225	0,227	0,234	0,240	0,239	0,232
		<i>RBP</i>	0,226	0,237	0,258	0,262	0,275	0,277
		<i>RBP_{MUL}</i>	0,228	0,245 ^{†*}	0,269 ^{†*}	0,271 ^{†*}	0,295 ^{†*}	0,289 ^{†*}
		<i>RBP_{AUG}</i>	0,229	0,248 ^{†*}	0,269 ^{†*}	0,264 [†]	0,291 ^{†*}	0,287 ^{†*}

* *RBP*'ye göre %95 güven aralığında istatistiksel olarak anlamlı

† *P&C*'ye göre %99 güven aralığında istatistiksel olarak anlamlı

Tablo 4.6'dan görüldüğü üzere, LM birleştirme tekniği kullanıldığında, önerilen MUL ve AUG stratejileri kıyaslama algoritması olarak seçilen *P&C* ile benzer sonuçlar elde etmiştir. Öte yandan, birleştirme tekniği olarak Avg veya CR kullanıldığında, özellikle orta ve büyük gruplar için her iki strateji de *P&C* yaklaşımından daha iyi performans göstermiştir. MLM için yapılan deneylerin sonuçları gözlemlendiğinde ise,

birleştirme tekniği olarak LM, grup boyutunun büyük olduğu durum hariç, diğer tüm ayarlamalarda MUL ve AUG stratejileri *P&C* algoritmasından önemli ölçüde daha iyi performans göstermiştir. Bu nedenle, bu çalışmada önerilen nihai grup tanımlama yöntemlerinin (yani MUL ve AUG), uygun bir birleştirme tekniği ile birlikte kullanılması durumunda, bu anlamda literatürdeki temel yaklaşım olan *P&C* algoritmasından daha iyi sonuçların elde edilmesine imkân tanıdığı söylenebilir.

Tablo 4.7. *MLM veri kümesi için DTP ile bezenmiş RBP yönteminin nDCG sonuçları*

Birleştirme Tekniği	Grup Boyutu (P)		Küçük		Orta		Büyük	
	en iyi-N	Model	5	10	25	50	150	200
Avg	5	<i>P&C</i>	0,108	0,113	0,120	0,122	0,153	0,155
		<i>RBP</i>	0,154	0,177	0,205	0,229	0,238	0,237
		<i>RBP</i> _{MUL}	0,156 ^{†*}	0,182 ^{†*}	0,214 ^{†*}	0,235 ^{†*}	0,244 ^{†*}	0,245 ^{†*}
		<i>RBP</i> _{AUG}	0,155 [†]	0,181 ^{†*}	0,213 ^{†*}	0,230 [†]	0,242 ^{†*}	0,240 [†]
	10	<i>P&C</i>	0,123	0,129	0,136	0,143	0,146	0,160
		<i>RBP</i>	0,159	0,176	0,210	0,215	0,243	0,240
		<i>RBP</i> _{MUL}	0,159 [†]	0,180 ^{†*}	0,219 ^{†*}	0,220 ^{†*}	0,247 [†]	0,247 ^{†*}
		<i>RBP</i> _{AUG}	0,160 [†]	0,181 ^{†*}	0,217 ^{†*}	0,217 [†]	0,245 [†]	0,243 [†]
	20	<i>P&C</i>	0,149	0,154	0,158	0,162	0,173	0,173
		<i>RBP</i>	0,173	0,191	0,220	0,231	0,253	0,255
		<i>RBP</i> _{MUL}	0,175 ^{†*}	0,195 ^{†*}	0,228 ^{†*}	0,239 ^{†*}	0,258 ^{†*}	0,261 ^{†*}
		<i>RBP</i> _{AUG}	0,176 ^{†*}	0,195 ^{†*}	0,229 ^{†*}	0,239 ^{†*}	0,253 ^{†*}	0,258 ^{†*}
LM	5	<i>P&C</i>	0,118	0,117	0,116	0,126	0,123	0,123
		<i>RBP</i>	0,141	0,143	0,145	0,137	0,120	0,114
		<i>RBP</i> _{MUL}	0,143 ^{†*}	0,148 ^{†*}	0,154 ^{†*}	0,148 ^{†*}	0,131 [*]	0,128 [*]
		<i>RBP</i> _{AUG}	0,144 ^{†*}	0,149 ^{†*}	0,155 ^{†*}	0,143 ^{†*}	0,132 [*]	0,119
	10	<i>P&C</i>	0,123	0,127	0,130	0,124	0,124	0,117
		<i>RBP</i>	0,143	0,145	0,142	0,129	0,112	0,106
		<i>RBP</i> _{MUL}	0,145 ^{†*}	0,148 ^{†*}	0,150 ^{†*}	0,139 ^{†*}	0,124 [*]	0,114 [*]
		<i>RBP</i> _{AUG}	0,146 ^{†*}	0,151 ^{†*}	0,150 ^{†*}	0,140 ^{†*}	0,117	0,114 [*]
	20	<i>P&C</i>	0,141	0,141	0,144	0,140	0,128	0,151
		<i>RBP</i>	0,153	0,153	0,150	0,139	0,112	0,103
		<i>RBP</i> _{MUL}	0,154 [†]	0,157 ^{†*}	0,155 ^{†*}	0,146 ^{†*}	0,116	0,109 [*]
		<i>RBP</i> _{AUG}	0,157 ^{†*}	0,157 ^{†*}	0,153 [†]	0,143 [*]	0,115	0,108 [*]
CR	5	<i>P&C</i>	0,082	0,080	0,069	0,099	0,083	0,064
		<i>RBP</i>	0,084	0,089	0,078	0,104	0,095	0,071
		<i>RBP</i> _{MUL}	0,086 [*]	0,091 ^{†*}	0,079 ^{†*}	0,107 ^{†*}	0,095 [†]	0,075 ^{†*}
		<i>RBP</i> _{AUG}	0,086 [*]	0,092 ^{†*}	0,079 ^{†*}	0,107 ^{†*}	0,094 [†]	0,076 ^{†*}
	10	<i>P&C</i>	0,075	0,076	0,072	0,078	0,080	0,077
		<i>RBP</i>	0,075	0,080	0,075	0,090	0,084	0,083
		<i>RBP</i> _{MUL}	0,076	0,082 ^{†*}	0,078 ^{†*}	0,096 ^{†*}	0,088 ^{†*}	0,087 ^{†*}
		<i>RBP</i> _{AUG}	0,076	0,082 ^{†*}	0,078 ^{†*}	0,094 ^{†*}	0,086 ^{†*}	0,086 ^{†*}
	20	<i>P&C</i>	0,101	0,098	0,087	0,105	0,102	0,083
		<i>RBP</i>	0,100	0,101	0,091	0,111	0,108	0,093
		<i>RBP</i> _{MUL}	0,102 [*]	0,104 ^{†*}	0,093 ^{†*}	0,114 ^{†*}	0,110 ^{†*}	0,098 ^{†*}
		<i>RBP</i> _{AUG}	0,102 [*]	0,104 ^{†*}	0,093 ^{†*}	0,113 ^{†*}	0,108 [†]	0,098 ^{†*}

* *RBP*'ye göre %95 güven aralığında istatistiksel olarak anlamlı

† *P&C*'ye göre %99 güven aralığında istatistiksel olarak anlamlı

4.5. Tartışma

Bu bölümdeki temel amaç, GÖS'lerin kullanıcı gruplarını tatmin edecek önerileri üretme yeteneğini geliştirmek için otomatik gruplandırma yaklaşımları geliştirmektir. Bu doğrultuda, geleneksel KM algoritmasının kullanılması yerine, önceden belirlenmiş maksimum üye sayısına sahip kullanıcı gruplarının oluşturulmasını sağlayan BKM algoritmasının kullanılması önerilmiştir. Elde edilen deneysel sonuçlar, BKM gruplama yaklaşımının, kullanıcıları, özellikle büyük veri kümelerinde küçük ve orta boy gruplara uygun şekilde ayırmakta faydalı olduğunu göstermiştir. Ayrıca, önerilen yaklaşım, yeni kullanıcıların ait olabileceği uygun gruplara atanması ve grup yapısının yeniden düzenlenmesini sağlamak için gerekli hesaplama süresinin KM algoritmasına kıyasla önemli ölçüde azalmasına olanak sağlamıştır. Bu da, kümeleme yaklaşımlarında sıklıkla ortaya çıkan boyutsallığın laneti problemi ile başa çıkmak için büyük bir avantajdır.

JTP yönteminin performansını değerlendirmek için yapılan deneysel çalışmalar, grupların belirlenmesi aşamasında PBP veya RBP kullanımının, grup oluşturma kalitesi, seyrek veri kümelerini işleme ve hesaplama süresini azaltma açısından, kullanıcıların orjinal derecelendirme vektörlerinden daha etkili olduğu göstermektedir. Ayrıca, ampirik sonuçlar, RBP profil oluşturma mekanizmasının PBP'den daha iyi performans sergilediğini göstermiştir. Bununla birlikte, değerlendirilen ürünlerin sayısı ne kadar fazla olursa, oluşturulan profil üzerindeki etkisinin o kadar az olması nedeniyle, bu iki yaklaşımın performansı, beklendiği gibi, kullanıcı başına düşen derecelendirilmiş ürün sayısının artması ile azalarak birbirine yakınsamaktadır.

Ayrıca, her iki veri kümesi için yapılan deneylerin sonuçları, kullanıcıların bireysel tercihlerinin demografik özelliklerinden çok daha değerli olduğunu varsayan AUG stratejisinin özellikle küçük gruplar için daha etkili olduğu göstermiştir. Diğer yandan, sonuçlar, MUL stratejisinin daha homojen yapıdaki orta ve büyük grupların oluşturulmasını sağladığını göstermiştir. Bu da, grupların demografik yapısının, gruplar kalabalıklaştıkça daha kritik hale geldiğini ve dikkate alınması gereken bir unsur haline geldiğini göstermiştir. Ayrıca, sonuçlar, otomatik gruplama işlemi süresince bu iki stratejiden birinin kullanılmasının, karşılaştırma algoritması olarak seçilen P&C stratejisinden daha yüksek kalitede grup önerilerinin üretilmesine imkân tanıdığını göstermiştir.

Literatürde GÖS bağlamında yapılan çoğu çalışmada, üretilen grup önerilerinin doğruluğu genellikle artan grup büyüklüğü ile azalır (Boratto & Carta, 2010). Çünkü bu çalışmalar, mevcut tüm kullanıcı derecelendirmelerini tek tek tahmin etmeye odaklanır ve önerilerin yaklaşımları değerlendirmek için gerçek derecelendirme değerleri ile karşılaştırır. Dolayısıyla, grup büyüklüğündeki bir artış, grup öneri yaklaşımları ile mevcut tüm derecelendirmeleri doğru bir şekilde tahmin etmeyi zorlaştırır ve bu da sistemin doğruluk açısından performansının düşmesine yol açar. Ancak, bu çalışmada, grup üyelerinin sağladığı her bir derecelendirme için tahminler üretmek yerine, bir grup için N adet tercih edilen ürünü içeren bir öneri listesinin üretilmesine odaklanılmıştır. Bu ürünler, kullanılan birleştirme tekniği ile karakterize edilen üyeler arasında bir tür mutabakat mekanizması ile belirlenmiştir. Bu nedenle, bir grupta ne kadar çok insan varsa, daha geniş bir topluluğun görüşlerinden elde edildiklerinden, en iyi- N önerileri o kadar güvenilir olur. Dolayısıyla, üretilen en iyi- N önerilerinin kalitesi, özellikle kullanılan birleştirme tekniği Avg (Baltrunas, Makcinskas, & Ricci, 2010; Kassak, Kompan, & Bielikova, 2016) veya CR (Seo vd., 2018) olduğunda, artan grup boyutu ile daha iyi hale gelmektedir. Sonuç olarak, gruplardaki üye sayısına bakılmaksızın, önerilen BKM varyantı algoritmalar, uygun birleştirme tekniğinin kullanılması durumunda yüksek kalitede en iyi- N grup önerileri sağlayabilir.

Son olarak, deney sonuçlarını bir bütün olarak değerlendirildiğinde, önerilen tüm gruplama yaklaşımlarının, kullanılan birleştirme tekniğinden, öneri listesinin boyutundan ve değişen grup boyutlarından bağımsız olarak faydalı olduğu sonucuna varılabilir.

4.6. Değerlendirme

Bu bölümde, üyelerin genel memnuniyeti arttırmak için grup önerilerinde kullanılmak üzere, uygun kullanıcı gruplarını otomatik olarak belirleyen BKM-tabanlı bir gruplama yaklaşımı geliştirilmiştir. Önerilen kümeleme yaklaşımı, maksimum grup boyutunu sağlamakla birlikte, doğruluğun arttırılması ve ölçeklenebilirlik sorunlarının hafifletilmesine yardımcı olur. Spesifik olarak, bu yaklaşım, ikili bir karar ağacı oluşturmak için kullanıcıların derecelendirme vektörleri üzerinde BKM kümeleme algoritması uygular. Ardından, oluşturulan karar ağacını, hem kullanıcı gruplarını tanımlamak hem de sisteme sonradan dâhil olan bir kullanıcının ait olabileceği uygun grubu bulmak üzere kullanır. Önerilen yaklaşımın etkinliğini ölçmek için iki farklı veri kümesinde yapılan çeşitli deneyler, istatistiksel anlamlılık testleriyle desteklendiği gibi,

önerilen BKM yaklaşımının geleneksel KM algoritmasına kıyasla grup öneri kalitesi açısından daha iyi performans sağladığını göstermektedir.

Ayrıca, BKM yaklaşımındaki benzerlik hesaplama sürecine ilişkin seyreklik ve ölçeklenebilirlik sorunlarının etkilerini hafifletmek amacıyla, iki farklı janr-tabanlı profil oluşturma yönteminin kullanılması önerilmektedir. Bu yöntemler, kümelemenin ayırma becerilerini güçlendirmek amacıyla, kullanıcıların derecelendirme vektörlerini, derecelendirilen ürünlerin janrlarını göz önünde bulundurarak, JTP'lere dönüştürmeyi amaçlamaktadır. Spesifik olarak, JTP oluşturma işlemi, ürünlerin standart özelliklerini dikkate alarak, kümeleme sürecinde kullanmak üzere tamamen yoğun ve derecelendirme vektörlerine kıyasla çok daha düşük boyuttaki kullanıcı profilleri üretir. Ayrıca, bu işlem ile üretilen profillerin boyutu sabittir ve sistemdeki ürünlerin mevcut janr sayısına bağlıdır; böylece, benzerliklerin hesaplanması için gereken süre önemli ölçüde azaltılır ve stabil hale gelir. Ek olarak, deneysel sonuçlar bu tür JTP'in kullanılmasının, üretilen grup önerilerinden olan genel memnuniyetin artmasına katkıda bulunduğunu göstermiştir.

Önerilen BKM yaklaşımında, kullanıcılar arasındaki ilişkileri daha doğru bir biçimde belirlemek için, demografi-tabanlı korelasyonları benzerlik hesaplama sürecine dâhil eden, MUL ve AUG olarak isimlendirilen iki farklı strateji önerilmektedir. Bu stratejiler, JTP benzerlikleri ile kullanıcılar arasındaki demografik benzerlikleri aynı anda dikkate alarak nihai benzerlikleri belirleyen yaklaşımlardır. Spesifik olarak, MUL stratejisi, iki benzerliği eşit olarak ağırlıklandırarak birleştirirken, AUG stratejisi, demografik benzerliklere göre janr-tabanlı benzerlikleri daha fazla önemseyerek nihai benzerlikleri hesaplar. Deneysel sonuçlar, otomatik grup tanımlama süreci boyunca bu stratejilerden herhangi birinin kullanılmasının sistemin doğruluğunu önemli ölçüde arttırdığını ve MUL stratejisinin büyük gruplar için, AUG stratejinin ise küçük ve orta ölçekteki gruplar için daha etkili olduğunu göstermektedir. Ayrıca, ampirik sonuçlar her iki stratejinin de uygun kullanıcı gruplarını belirlemede, literatürde öne çıkan P&C algoritmasına kıyasla daha başarılı olduğunu göstermektedir. Sonuç olarak, sistem kullanıcılarının demografik bilgilerinden faydalananarak, daha homojen grupların tanımlanması sağlanmış ve bireysel üyelerin genel memnuniyetini arttırılmıştır.

5. BKM-TABANLI ALGORİTMANIN AYRINTILI ANALİZİ

Bu bölümde, bir önceki bölümde kullanıcı gruplarının otomatik olarak tanımlanması amacıyla geliştirilen BKM yaklaşımının performansı, modern CF algoritmaları yardımıyla daha kapsamlı bir biçimde analiz edilmiştir.

5.1. Motivasyon

Bölüm 2.1’de açıklandığı üzere grup önerilerinde kullanıcı gruplarının otomatik olarak belirlenmesinde en yaygın kullanılan kümeleme yaklaşımı KM algoritmasıdır. Bir önceki bölümde, geliştirilen BKM yaklaşımının bu bağlamda KM algoritmasına göre üstünlüğü, gerçekleştirilen kapsamlı deneysel çalışmalar ile gösterilmiştir. Bununla birlikte, BKM yaklaşımı, JTP’ler ve demografik korelasyonlar yardımıyla geliştirilerek Bölüm 4.4.1’de ayrıntılı olarak açıklanan KM-tabanlı P&C stratejisinden daha başarılı grup önerilerinin üretilmesi sağlanmıştır. Bu kapsamda yapılan deneysel çalışmalarda, kullanıcı $\times$ ürün matrisindeki derecelendirilmemiş hücrelerin tahmininde ve daha sonra grup önerilerinin üretilmesinde Bölüm 2.2’de açıklanan en temel ve geleneksel CF algoritmalarından birisi kullanılmıştır.

Her ne kadar söz konusu CF algoritması başarılı öneriler üretilmesini sağlasa da son yıllarda öneri doğruluğu açısından daha etkin ve daha modern CF algoritmaları geliştirilmiştir. Bu algoritmalar daha kaliteli bireysel tahminler üretilmesini sağlayarak kullanıcıların genel memnuniyetinin artmasına katkıda bulunur. GÖS bağlamında geliştirilen bir sistemin temel hedeflerinden birisi üretilen grup önerilerinden memnun olan toplam kullanıcı sayısını arttırmak olduğu göz önünde bulundurulduğunda, bir önceki bölümde geliştirilen BKM yaklaşımının bu tür modern algoritmalar ile kullanılması akıllıca olacaktır. Bu bölümde, BKM yaklaşımının grup önerisi performansının modern CF algoritmalar yardımıyla arttırılması hedeflenmiştir. Bununla birlikte, BKM yaklaşımının hem orijinal kullanıcı $\times$ ürün matrisi üzerindeki KM algoritmasına, hem de KM-tabanlı P&C stratejisine olan üstünlüğü, bu tür modern CF algoritmalar kullanılarak daha kapsamlı bir biçimde analiz edilmiştir.

5.2. Kullanılan Modern CF Algoritmaları

BKM yaklaşımının performansının daha kapsamlı analizi için bu bölümde yapılan çalışmada, geleneksel öneri sistemlerinde yaygın olarak kullanılan ve yüksek doğrulukta tahminler üretilmesini sağlayan üç farklı modern CF algoritması kullanılmıştır. Bu algoritmalar aşağıda ayrıntılı olarak açıklanmıştır. Ayrıca, söz konusu algoritmalar

Surprise¹⁰ olarak isimlendirilen ve öneri sistemlerinde kullanılan popüler CF algoritmalarını içeren bir Python kütüphanesi ile uygulanmıştır.

- **kNNBase algoritması:** Bu algoritma, CF öneri sistemlerinde kullanılan popüler hafıza tabanlı algoritmalarından birisidir ve komşuluk hesabı üzerine çalışmaktadır (Koren, 2010). CF algoritmalarında kullanılan veriler, kullanıcıların ürünleri değerlendirirken gösterdikleri eğilimleri yansıtır. Yani, bazı kullanıcılar sistematik olarak diğerlerinden daha yüksek derecelendirme verme eğilimlerinde olurlar ve bazı ürünler diğerlerinden daha yüksek derecelendirmeler alırlar. kNNBase algoritması, kullanıcıların ve ürünlerin bu tür eğilimlerini, ürettikleri bireysel tahmin değerleri içerisinde kapsülleyerek öneri üretme sürecini gerçekleştiren bir algoritmadır. Ayrıca, bu algoritma uygulanırken, kullanıcılar arasındaki benzerlikler, Bölüm 2.2’de ayrıntılı olarak açıklanan PCC metriği ile hesaplanmış ve komşuluk sayısı 40 olarak ayarlanmıştır (Koren, 2010).
- **SO algoritması:** Bu algoritma, karmaşık ve yüksek hesaplama gerektiren CF algoritmaları kadar yüksek doğrulukta tahminler üretebilen, sade ve uygulaması kolay model-tabanlı bir CF öneri algoritmasıdır (Lemire & Maclachlan, 2005). SO algoritması, ürün derecelendirme tabanlı CF algoritma ailesinden gelmekte ve aşırı yüklenme sorununu büyük ölçüde çözme yeteneğine sahiptir. Prensip olarak, SO, ürün vektörleri arasındaki ilişkileri, doğrusal regresyonlar ($f(x) = ax + b$) yerine tek bir serbest parametre ile daha basit bir regresyon biçiminde ($f(x) = x + b$) tanımlar. Bu durumda serbest parametre (b), iki ürünün derecelendirmeleri arasındaki ortalama fark olarak değerlendirilir. Bununla birlikte, bazı durumlarda doğrusal regresyondan çok daha doğru sonuçlar elde etmektedir ve depolama maliyeti açısından daha verimlidir (Lemire & Maclachlan, 2005).
- **SVD++ algoritması:** Çarpanlara ayırma temelli başarılı bir model tabanlı CF algoritması olan SVD++, kullanıcıları ve ürünleri daha düşük boyuta sahip bir gizli boyutta temsil etmektedir (Koren, 2008). Bu algoritma, çarpanlara ayırma temelli diğer bir popüler SVD algoritmasının gelişmiş bir versiyonudur (Mnih & Salakhutdinov, 2008). SVD algoritması çok iyi bir öneri kalitesi sunsa da, kullanıcı/ürün etkileşimlerini yalnızca açıkça sağlanan derecelendirmeleri kullanarak belirler. Diğer yandan, günümüz öneri sistemleri, kullanıcılar/ürünler arasındaki hem

¹⁰ <http://surpriselib.com/>

açık (örneğin; sayısal derecelendirmeler) hem de örtülü (örneğin; değerlendirmeler, beğeniler, satın almalar ve yer imleri) etkileşimlerden yararlanmalıdır. SVD algoritmasının bir türev modeli olarak geliştirilen SVD ++ algoritması, bir kullanıcının değerlendirdiği filmler gibi örtülü geri bildirim bilgileri, öneri sürecine dâhil ederek daha yüksek doğruluğa sahip tahminlerin üretilmesini sağlar.

5.3. DeneySEL Çalışmalar

5.3.1. Deney metodolojisi

Bu bölümde yapılan deneySEL çalışmalarda, Bölüm 2.4’de ayrıntılı olarak açıklanan MLP ve MLM veri kümeleri kullanılmıştır. Ayrıca, deneyler, 5-kat çapraz doğrulama deney metodolojisi kullanılarak gerçekleştirilmiştir. Bu amaçla, ürünler kümesi rastgele biçimde beş alt kümeye bölünerek oluşturulan her alt kümenin ürünlerin %20’sini içermesi sağlanmıştır. Her yinelemede, alt gruplardan birisi test kümesi olarak, kalan dört alt grubun kombinasyonu da eğitim seti olarak kullanılmıştır. Dolayısıyla, nihai $nDCG$ değerleri, 5-kat deneySEL doğruluk sonuçlarının ortalaması alınarak hesaplanmıştır. Ayrıca, kullanıcı tarafından derecelendirilmeyen ürünleri öneren yaklaşımı cezalandırmak amacıyla, her kullanıcı için $nDCG$ hesaplarırken, kullanıcının derecelendirmediği ürünler için “0” değeri atanmıştır. Böylece, kullanıcıların ilgisini çeken ürünleri öneren bir yaklaşımın daha yüksek $nDCG$ değeri elde etmesi sağlanmıştır (Gorla vd., 2013).

Geliştirilen BKM yaklaşımının performansını modern CF algoritmaları kullanarak daha ayrıntılı analiz etmek ve grup önerisi bağlamında etkinliğini arttırmak amacıyla iki farklı deneySEL çalışma gerçekleştirilmiştir. İlk deney setinde, standart KM algoritması ile BKM yaklaşımını karşılaştırmak amacıyla, kullanıcı grupları eğitim veri kümesine göre standart KM algoritmasının ve önerilen BKM yaklaşımının orijinal kullanıcı $\times$ ürün matrisi üzerinde uygulanması ile belirlenmiştir. Ardından, grup önerileri Bölüm 2.3’de ayrıntılı olarak açıklanan BTB yaklaşımı kullanılarak üretilmiştir. Burada, bireysel tahminler üretmek için önceki bölümde açıklanan kNNBase, SO ve SVD++ algoritmaları kullanılmıştır. Gerçekleştirilen ikinci deney setinde ise BKM algoritmasının KM-tabanlı P&C yöntemine göre üstünlüğü modern CF algoritmaları aracılığı ile analiz edilmiştir. Bu deneylerde, P&C yönteminin ilk adımı olan kullanıcı $\times$ ürün matrisindeki değerlendirilmemiş hücrelerin tahmin edilerek doldurulması, önceki bölümde açıklanan modern CF algoritmalarından (yani, kNNBase, SO ve SVD++) birisi kullanılarak

gerçekleştirilmiştir. Üretilen tahminlerle doldurulmuş kullanıcı $\times$ ürün matrisi göz önünde bulundurularak, kullanıcı grupları, eğitim veri kümesine göre KM ve BKM algoritmalar aracılığı ile tanımlanmıştır. Grupların belirlenmesinden sonra, grup önerileri yine BTB yaklaşımı ile üretilmiştir. Burada, daha önce kullanıcı $\times$ ürün matrisini doldurmak için modern CF algoritmaları ile üretilen bireysel tahminler birleştirilerek grup önerilerinin üretilmesi sağlanmıştır. Böylece, geliştirilen BKM yaklaşımının performansı KM-tabanlı P&C yöntemi ile daha kapsamlı bir biçimde kıyaslanmıştır.

BKM algoritmasının performansının daha doğru bir şekilde karşılaştırılabilmesi için, standart KM algoritması ve KM temelli P&C yöntemi ile oluşturulacak grupların sayısının belirlenmesi gerekmektedir. Bu amaçla, k değeri, toplam kullanıcı sayısının, BKM yaklaşımındaki grupların maksimum boyutunu ifade eden (P) değerine olan oranı kullanılarak hesaplanmıştır. Örneğin, BKM yaklaşımındaki P değeri 50 olarak ayarlandığı ve veri kümesindeki toplam kullanıcı sayısının 600 olduğu durumda, k değeri $600/50=12$ olarak hesaplanmıştır. Ayrıca, grup büyüklüğünün önerilen yaklaşımların performansı üzerindeki etkisini değerlendirmek için, deneylerde, 5 ile 200 arasında değişen çeşitli P değerleri kullanılmıştır. Son olarak, deneylerdeki gruplar, içerdiği kullanıcı sayısına göre küçük, orta ve büyük olarak sınıflandırılmıştır. Burada, 10'dan az kullanıcı içerenler küçük gruplar, 10 ile 100 arasında kullanıcı içerenler orta gruplar, ve son olarak 100'den fazla kullanıcı içerenler büyük gruplar olarak değerlendirilmiştir.

Yapılan deneylerde, üretilen bireysel tahminlerin birleştirilerek grup derecelendirmelerinin hesaplanması amacıyla önceki bölümdeki deneylere benzer şekilde Avg, LM ve CR birleştirme teknikleri kullanılmıştır. Hesaplanan grup derecelendirmeleri göz önünde bulundurularak, her grup için en yüksek grup derecelendirmelerine sahip en iyi- N sıralı ürün listeleri grup önerileri olarak üretilmiştir. Burada, ayrıca, öneri listesinin boyutunun, geliştirilen yaklaşımların üzerindeki etkisini incelemek amacıyla 5, 10 ve 20 olarak değişen üç farklı N değeri için deneyler gerçekleştirilmiştir. Son olarak, üretilen en iyi- N grup önerilerinin doğruluğu $nDCG$ skoru ile gözlemlenmiştir.

5.3.2. Deney sonuçları

5.3.2.1. BKM ile standart KM algoritmalarının karşılaştırılması

Önerilen BKM algoritmasının benzer kullanıcı gruplarını algılamadaki doğruluk performansını incelemek için grup büyüklüğü, kullanılan toplama tekniği, öneri listesinin boyutu (N) ve bireysel tahminlerin üretilmesi amacıyla kullanılan CF algoritması dâhil

olmak üzere farklı parametreler ile birçok deney gerçekleştirilmiştir. Bununla birlikte, MLP ve MLM veri kümeleri için BKM algoritması ile elde edilen deneysel sonuçlar, sırasıyla Tablo 5.1. ve 5.2’de gösterildiği gibi standart KM algoritması ile karşılaştırılmıştır. Bu bölümde gerçekleştirilen deneylerde, kullanıcı grupları orijinal kullanıcı $\times$ ürün matrisi kullanılarak belirlenmiştir ve modern CF algoritmaları ile üretilen bireysel tahminler yalnızca grup önerileri üretilmesi aşamasında kullanılmıştır.

Tablo 5.1. MLP veri kümesinde KM ve BKM algoritmalarının modern CF algoritmaları için nDCG sonuçları

Birleştirme Tekniği	en iyi-N	Grup Boyutu (P)		Küçük		Orta		Büyük	
		CF Algoritması	Model	5	10	25	50	150	200
Avg	5	kNNBase	KM	0,24	0,13	0,07	0,04	0,02	0,02
			BKM	0,26*	0,16*	0,08*	0,05*	0,02	0,02
		SVD++	KM	0,43	0,34	0,26	0,21	0,19	0,18
			BKM	0,47*	0,38*	0,29*	0,27*	0,22*	0,22*
		SO	KM	0,22	0,14	0,07	0,05	0,03	0,04
			BKM	0,25*	0,17*	0,09*	0,06*	0,04*	0,04
	10	kNNBase	KM	0,24	0,17	0,09	0,07	0,05	0,05
			BKM	0,27*	0,19*	0,13*	0,09*	0,07*	0,07*
		SVD++	KM	0,40	0,32	0,24	0,21	0,18	0,16
			BKM	0,43*	0,36*	0,28*	0,25*	0,22*	0,22*
		SO	KM	0,23	0,17	0,07	0,08	0,06	0,06
			BKM	0,25*	0,18*	0,11*	0,09*	0,07*	0,07*
	20	kNNBase	KM	0,26	0,20	0,13	0,10	0,08	0,09
			BKM	0,28*	0,22*	0,15*	0,13*	0,11*	0,11*
		SVD++	KM	0,40	0,31	0,26	0,21	0,18	0,18
			BKM	0,42*	0,35*	0,28*	0,26*	0,21*	0,20*
		SO	KM	0,24	0,17	0,12	0,11	0,08	0,10
			BKM	0,27*	0,19*	0,14*	0,12*	0,11*	0,11*
LM	5	kNNBase	KM	0,04	0,05	0,04	0,02	0,02	0,02
			BKM	0,05*	0,07*	0,05*	0,03*	0,02*	0,02*
		SVD++	KM	0,17	0,15	0,12	0,08	0,07	0,06
			BKM	0,21*	0,20*	0,16*	0,13*	0,09*	0,07*
		SO	KM	0,05	0,05	0,04	0,03	0,02	0,02
			BKM	0,06*	0,06*	0,06*	0,05*	0,03*	0,02*
	10	kNNBase	KM	0,06	0,06	0,04	0,04	0,02	0,02
			BKM	0,09*	0,08*	0,07*	0,05*	0,03*	0,03*
		SVD++	KM	0,16	0,15	0,12	0,09	0,07	0,05
			BKM	0,21*	0,19*	0,15*	0,12*	0,09*	0,09*
		SO	KM	0,06	0,06	0,05	0,04	0,03	0,02
			BKM	0,09*	0,09*	0,07*	0,06*	0,04*	0,04*
	20	kNNBase	KM	0,08	0,08	0,06	0,05	0,03	0,02
			BKM	0,11*	0,11*	0,08*	0,06*	0,04*	0,04*
		SVD++	KM	0,18	0,16	0,13	0,10	0,07	0,07
			BKM	0,21*	0,20*	0,15*	0,12*	0,09*	0,09*
		SO	KM	0,09	0,09	0,07	0,06	0,04	0,03
			BKM	0,12*	0,11*	0,09*	0,07*	0,06*	0,05*
CR	5	kNNBase	KM	0,06	0,04	0,03	0,02	0,02	0,02
			BKM	0,07*	0,05*	0,04*	0,03*	0,02	0,02
		SVD++	KM	0,11	0,09	0,08	0,07	0,06	0,06
			BKM	0,13*	0,11*	0,09*	0,08*	0,08*	0,08*
		SO	KM	0,06	0,04	0,03	0,03	0,03	0,02
			BKM	0,06	0,04	0,03	0,03	0,03	0,02

		Grup Boyutu (P)		Küçük		Orta		Büyük	
Birleştirme Tekniği	en iyi-N	CF Algoritması	Model	5	10	25	50	150	200
	10	kNNBase	BKM	0,07*	0,05*	0,04*	0,04*	0,03*	0,03
			KM	0,06	0,05	0,04	0,04	0,03	0,03
		SVD++	BKM	0,07*	0,06*	0,05*	0,05*	0,04*	0,04*
			KM	0,11	0,09	0,08	0,08	0,08	0,08
		SO	BKM	0,13*	0,11*	0,10*	0,09*	0,09*	0,09*
			KM	0,07	0,06	0,05	0,04	0,04	0,04
	20	kNNBase	BKM	0,08*	0,07*	0,06*	0,05*	0,05*	0,05*
			KM	0,07	0,06	0,06	0,05	0,05	0,05
		SVD++	BKM	0,08*	0,07*	0,07*	0,06*	0,06*	0,06*
			KM	0,12	0,10	0,09	0,08	0,08	0,08
		SO	BKM	0,13*	0,11*	0,10*	0,09*	0,09*	0,09*
			KM	0,08	0,07	0,05	0,05	0,05	0,05
			BKM	0,09*	0,08*	0,06*	0,06*	0,06*	0,06*

* %95 güven aralığında istatistiksel olarak anlamlı

Tablo 5.2. MLM veri kümesinde KM ve BKM algoritmalarının modern CF algoritmaları için nDCG sonuçları

		Grup Boyutu (P)		Küçük		Orta		Büyük	
Birleştirme Tekniği	en iyi-N	CF Algoritması	Model	5	10	25	50	150	200
Avg	5	kNNBase	KM	0,13	0,09	0,06	0,05	0,04	0,03
			BKM	0,14*	0,10*	0,07*	0,06*	0,05*	0,05*
		SVD++	KM	0,22	0,20	0,17	0,16	0,16	0,16
			BKM	0,25*	0,23*	0,22*	0,21*	0,20*	0,20*
		SO	KM	0,05	0,04	0,04	0,04	0,04	0,04
			BKM	0,06*	0,06*	0,06*	0,06*	0,06*	0,06*
	10	kNNBase	KM	0,15	0,11	0,08	0,07	0,06	0,05
			BKM	0,17*	0,13*	0,10*	0,09*	0,08*	0,07*
		SVD++	KM	0,21	0,19	0,17	0,16	0,15	0,15
			BKM	0,24*	0,22*	0,21*	0,20*	0,19*	0,19*
		SO	KM	0,07	0,07	0,06	0,06	0,06	0,06
			BKM	0,08*	0,08*	0,08*	0,08*	0,08*	0,08*
	20	kNNBase	KM	0,16	0,13	0,10	0,08	0,07	0,07
			BKM	0,18*	0,15*	0,11*	0,10*	0,09*	0,09*
		SVD++	KM	0,21	0,19	0,18	0,16	0,15	0,15
			BKM	0,23*	0,21*	0,20*	0,20*	0,19*	0,19*
		SO	KM	0,09	0,09	0,09	0,08	0,09	0,09
			BKM	0,10*	0,10*	0,10*	0,10*	0,10*	0,10*
LM	5	kNNBase	KM	0,12	0,08	0,04	0,03	0,02	0,02
			BKM	0,14*	0,09*	0,06*	0,04*	0,03*	0,03*
		SVD++	KM	0,19	0,17	0,13	0,12	0,09	0,08
			BKM	0,22*	0,20*	0,18*	0,16*	0,14*	0,13*
		SO	KM	0,07	0,06	0,05	0,04	0,02	0,01
			BKM	0,08*	0,07*	0,07*	0,07*	0,04*	0,04*
	10	kNNBase	KM	0,13	0,08	0,05	0,03	0,03	0,03
			BKM	0,15*	0,10*	0,07*	0,05*	0,04*	0,04*
		SVD++	KM	0,19	0,16	0,13	0,12	0,09	0,08
			BKM	0,21*	0,19*	0,17*	0,15*	0,13*	0,12*
		SO	KM	0,08	0,07	0,06	0,05	0,02	0,02
			BKM	0,09*	0,09*	0,08*	0,07*	0,05*	0,05*
	20	kNNBase	KM	0,13	0,10	0,06	0,04	0,04	0,04
			BKM	0,15*	0,12*	0,08*	0,07*	0,05*	0,05*
		SVD++	KM	0,19	0,16	0,14	0,12	0,10	0,10

		Grup Boyutu (P)		Küçük		Orta		Büyük	
Birleştirme Tekniği	en iyi-N	CF Algoritması	Model	5	10	25	50	150	200
CR	5	SO	BKM	0,21*	0,19*	0,17*	0,15*	0,13*	0,13*
			KM	0,10	0,09	0,07	0,05	0,03	0,01
			BKM	0,11*	0,11*	0,10*	0,08*	0,05*	0,04*
		kNNBase	KM	0,04	0,03	0,03	0,03	0,03	0,03
			BKM	0,05*	0,04*	0,04*	0,04*	0,04*	0,04*
		SVD++	KM	0,05	0,05	0,04	0,04	0,04	0,03
			BKM	0,06*	0,06*	0,05*	0,05*	0,05*	0,05*
		SO	KM	0,04	0,02	0,02	0,02	0,02	0,02
			BKM	0,05*	0,03*	0,03*	0,02	0,02	0,02
	10	kNNBase	KM	0,05	0,04	0,03	0,03	0,03	0,03
			BKM	0,06*	0,05*	0,04*	0,04*	0,04*	0,03
		SVD++	KM	0,05	0,06	0,04	0,04	0,04	0,04
			BKM	0,06*	0,07*	0,05*	0,05*	0,05*	0,05*
		SO	KM	0,04	0,03	0,03	0,03	0,02	0,02
			BKM	0,05*	0,04*	0,04*	0,04*	0,03*	0,03*
	20	kNNBase	KM	0,05	0,05	0,04	0,04	0,04	0,04
			BKM	0,06*	0,06*	0,05*	0,05*	0,05*	0,05*
		SVD++	KM	0,05	0,06	0,05	0,05	0,05	0,05
			BKM	0,07*	0,07*	0,06*	0,06*	0,06*	0,06*
		SO	KM	0,04	0,04	0,04	0,03	0,03	0,03
			BKM	0,05*	0,05*	0,05*	0,04*	0,04*	0,04*

* %95 güven aralığında istatistiksel olarak anlamlı

Tablo 5.1 ve 5.2'den görülebileceği gibi, MLP ve MLM veri kümeleri ile yapılan deneylerin $nDCG$ sonuçları, BKM algoritmasının her bir farklı ayar için KM algoritmasına kıyasla genellikle daha iyi olduğunu göstermektedir. Ayrıca, BKM ve KM algoritması ile elde edilen sonuçları karşılaştırmak için istatistiksel anlamlılık testleri yapılmış ve bu testlerin sonuçları tabloların dipnotlarında verilmiştir. Yapılan tek kuyruklu t -testlerinin sonuçları, BKM yaklaşımının KM algoritması üzerinde sağladığı iyileştirmelerin, hemen hemen tüm ayarlamalarda %95 güven aralığında istatistiksel olarak anlamlı olduğunu göstermektedir.

Tablo 5.1'de gösterildiği gibi, MLP için yapılan deneylerde, en yüksek $nDCG$ değerleri birleştirme tekniğinin Avg olarak seçildiği durumda elde edilmiştir. Ayrıca, LM ve CR teknikleri benzer doğruluk performansları göstermişlerdir. Diğer yandan, Tablo 5.2'de sonuçları gösterilen MLM veri kümesinde yapılan deneylerde, Avg ve LM benzer performanslar göstermesine rağmen, en yüksek $nDCG$ sonuçları Avg ile elde edilmiştir. Ayrıca, her iki veri kümesi için yapılan deneylerde, SVD++ algoritması, kNNBase ve SO algoritmalarına kıyasla önemli ölçüde daha başarılı olmuştur. Bununla birlikte, kullanılan üç farklı CF algoritmasının doğruluk performansının, grupların boyutu büyüdükçe düştüğü gözlemlenmiştir. Son olarak, öneri listesinin boyutunun, elde edilen $nDCG$ sonuçları üzerinde olumlu bir etkisinin olduğu açıkça görülmektedir. Bunun nedeni, daha

büyük bir öneri listesinde, grup üyelerinin tercih edeceği ürünlerin bulunma şansının daha fazla olmasıdır.

5.3.2.2. BKM ve KM algoritmalarının P&C ile karşılaştırılması

Önceki bölümde gerçekleştirilen deneylerde kullanıcı grupları, orijinal kullanıcı $\times$ ürün matrisi üzerinde BKM ve KM algoritması uygulanarak tanımlanmıştır. Bu bölümde gerçekleştirilen deneylerde ise geliştirilen BKM yaklaşımının KM algoritmasına göre üstünlüğü P&C stratejisi bağlamında değerlendirilmiştir. Böylece, kullanıcı grupları, modern CF algoritmaları ile üretilen tahminler ile doldurulmuş kullanıcı $\times$ ürün matrisi üzerinde BKM ve KM algoritmaları uygulanarak belirlenmiştir. Diğer bir deyişle, modern CF algoritmaları ile üretilen bireysel tahminler hem kullanıcı gruplarının belirlenmesi aşamasında hem de grup önerilerinin üretilmesi aşamasında kullanılmıştır. P&C stratejisi kullanılarak MLP ve MLM veri kümeleri için elde edilen deneysel sonuçlar, sırasıyla Tablo 5.3 ve 5.4'te gösterildiği gibi KM algoritması ile karşılaştırılmıştır.

Tablo 5.3. MLP veri kümesinde P&C stratejisi ile KM ve BKM algoritmalarının modern CF algoritmaları için nDCG sonuçları

Birleştirme Tekniği	en iyi-N	Grup Boyutu (P)		Küçük		Orta		Büyük	
		CF Algoritması	Model	5	10	25	50	150	200
Avg	5	kNNBase	KM	0,48	0,43	0,37	0,37	0,30	0,35
			BKM	0,50	0,45	0,42*	0,41*	0,35*	0,41*
		SVD++	KM	0,60	0,58	0,51	0,50	0,43	0,41
			BKM	0,61	0,60	0,54*	0,53*	0,51*	0,49*
		SO	KM	0,40	0,35	0,32	0,28	0,20	0,18
			BKM	0,42	0,37	0,36*	0,31*	0,23*	0,19*
	10	kNNBase	KM	0,44	0,43	0,36	0,35	0,35	0,31
			BKM	0,46	0,45	0,38*	0,40*	0,41*	0,37*
		SVD++	KM	0,57	0,51	0,46	0,43	0,39	0,36
			BKM	0,58	0,53*	0,48*	0,45*	0,47*	0,43*
		SO	KM	0,38	0,34	0,30	0,28	0,19	0,16
			BKM	0,40*	0,36*	0,33*	0,31*	0,22*	0,19*
	20	kNNBase	KM	0,41	0,39	0,38	0,32	0,23	0,26
			BKM	0,43	0,41	0,42*	0,36*	0,27*	0,30*
		SVD++	KM	0,50	0,49	0,46	0,43	0,35	0,32
			BKM	0,52	0,49	0,47	0,44	0,38*	0,40*
		SO	KM	0,36	0,34	0,31	0,22	0,20	0,19
			BKM	0,38*	0,36*	0,34*	0,24*	0,22*	0,23*
LM	5	kNNBase	KM	0,04	0,05	0,07	0,04	0,05	0,05
			BKM	0,05*	0,06*	0,08	0,05*	0,06*	0,06*
		SVD++	KM	0,18	0,21	0,22	0,23	0,17	0,18
			BKM	0,20*	0,24*	0,25*	0,26*	0,20*	0,22*
		SO	KM	0,05	0,05	0,06	0,04	0,02	0,05
			BKM	0,06	0,05	0,07*	0,05*	0,03*	0,06*
	10	kNNBase	KM	0,05	0,07	0,07	0,10	0,06	0,03
			BKM	0,06*	0,07	0,08*	0,11*	0,07*	0,04*
		SVD++	KM	0,18	0,22	0,21	0,18	0,17	0,19
			BKM	0,20	0,24	0,24*	0,20*	0,20*	0,22*

Birleştirme Tekniği	en iyi-N	Grup Boyutu (P)		Küçük		Orta		Büyük	
		CF Algoritması	Model	5	10	25	50	150	200
CR	20	SO	KM	0,06	0,08	0,06	0,06	0,07	0,07
			BKM	0,06	0,09	0,07*	0,07*	0,08*	0,09*
		kNNBase	KM	0,08	0,10	0,10	0,11	0,09	0,09
			BKM	0,09*	0,11*	0,12*	0,12*	0,10*	0,10*
		SVD++	KM	0,22	0,22	0,21	0,20	0,19	0,20
			BKM	0,25	0,25	0,24*	0,23*	0,23*	0,24*
	5	SO	KM	0,10	0,09	0,10	0,09	0,08	0,09
			BKM	0,10	0,10	0,11*	0,10*	0,09*	0,10
		kNNBase	KM	0,14	0,14	0,13	0,13	0,13	0,13
			BKM	0,14	0,15	0,14*	0,15*	0,15*	0,15*
		SVD++	KM	0,19	0,18	0,18	0,18	0,17	0,14
			BKM	0,19	0,19*	0,19*	0,19*	0,19*	0,15*
	10	SO	KM	0,11	0,11	0,09	0,10	0,08	0,09
			BKM	0,11	0,12	0,10*	0,11*	0,09*	0,11*
		kNNBase	KM	0,13	0,13	0,15	0,12	0,13	0,15
			BKM	0,13	0,13	0,16*	0,14*	0,16*	0,17*
		SVD++	KM	0,17	0,17	0,17	0,17	0,14	0,15
			BKM	0,17	0,17	0,18*	0,18*	0,16*	0,17*
	20	SO	KM	0,11	0,11	0,10	0,10	0,07	0,08
			BKM	0,11	0,12*	0,12*	0,12*	0,08*	0,09*
		kNNBase	KM	0,13	0,12	0,13	0,12	0,14	0,11
			BKM	0,13	0,13*	0,14*	0,14*	0,16*	0,12*
		SVD++	KM	0,17	0,16	0,16	0,16	0,14	0,15
			BKM	0,17	0,16	0,17*	0,17*	0,15*	0,17*
		SO	KM	0,12	0,11	0,10	0,10	0,09	0,07
			BKM	0,12	0,11	0,11*	0,11*	0,10*	0,09*

* %95 güven aralığında istatistiksel olarak anlamlı

Tablo 5.4. MLM veri kümesinde P&C stratejisi ile KM ve BKM algoritmalarının modern CF algoritmaları için nDCG sonuçları

Birleştirme Tekniği	en iyi-N	Grup Boyutu (P)		Küçük		Orta		Büyük	
		CF Algoritması	Model	5	10	25	50	150	200
Avg	5	kNNBase	KM	0,16	0,13	0,08	0,08	0,06	0,06
			BKM	0,17	0,13	0,09*	0,08	0,07*	0,07*
		SVD++	KM	0,27	0,27	0,27	0,26	0,25	0,24
			BKM	0,28	0,28*	0,28*	0,28*	0,27*	0,26*
		SO	KM	0,04	0,04	0,04	0,04	0,04	0,04
			BKM	0,05	0,05*	0,05*	0,05*	0,05*	0,06*
	10	kNNBase	KM	0,17	0,13	0,11	0,09	0,07	0,07
			BKM	0,18	0,14*	0,12*	0,10*	0,08*	0,08*
		SVD++	KM	0,26	0,25	0,25	0,25	0,25	0,24
			BKM	0,26	0,25	0,26*	0,26*	0,27*	0,26*
		SO	KM	0,06	0,06	0,06	0,06	0,06	0,07
			BKM	0,06	0,06	0,07*	0,07*	0,07*	0,08*
	20	kNNBase	KM	0,18	0,15	0,12	0,11	0,09	0,09
			BKM	0,19	0,16*	0,13*	0,12*	0,11*	0,10*
		SVD++	KM	0,25	0,24	0,24	0,24	0,22	0,23
			BKM	0,25	0,25	0,25*	0,25*	0,24*	0,24
		SO	KM	0,08	0,08	0,08	0,09	0,09	0,08
			BKM	0,09	0,09	0,09	0,10*	0,11*	0,10*
LM	5	kNNBase	KM	0,15	0,11	0,06	0,05	0,03	0,03
			BKM	0,16	0,11	0,07*	0,06*	0,04*	0,04*

Birleştirme Tekniği	en iyi-N	Grup Boyutu (P)		Küçük		Orta		Büyük	
		CF Algoritması	Model	5	10	25	50	150	200
CR	10	SVD++	KM	0,25	0,24	0,23	0,22	0,19	0,17
			BKM	0,26	0,25	0,24*	0,23*	0,20*	0,19*
		SO	KM	0,08	0,08	0,08	0,08	0,09	0,09
			BKM	0,09	0,10*	0,10*	0,10*	0,11*	0,11*
		kNNBase	KM	0,14	0,11	0,08	0,06	0,03	0,04
			BKM	0,15	0,12	0,09*	0,07*	0,03	0,05*
		SVD++	KM	0,24	0,23	0,21	0,20	0,18	0,18
			BKM	0,25	0,24	0,22*	0,22*	0,20*	0,20*
		SO	KM	0,09	0,08	0,09	0,09	0,09	0,10
			BKM	0,10	0,10*	0,11*	0,11*	0,11*	0,13*
	20	kNNBase	KM	0,15	0,12	0,08	0,06	0,06	0,04
			BKM	0,15	0,12	0,09*	0,07*	0,07*	0,05*
		SVD++	KM	0,23	0,22	0,21	0,20	0,17	0,17
			BKM	0,23	0,24*	0,23*	0,21*	0,18*	0,18*
		SO	KM	0,10	0,10	0,10	0,11	0,12	0,11
			BKM	0,12	0,12*	0,13*	0,14*	0,14*	0,14*
	5	kNNBase	KM	0,05	0,03	0,03	0,03	0,03	0,03
			BKM	0,05	0,03	0,04*	0,04*	0,04*	0,04*
		SVD++	KM	0,08	0,07	0,07	0,07	0,06	0,06
			BKM	0,08	0,08	0,07*	0,07*	0,07*	0,06
		SO	KM	0,03	0,02	0,02	0,02	0,02	0,02
			BKM	0,04*	0,02	0,02	0,02*	0,02*	0,02*
	10	kNNBase	KM	0,05	0,04	0,04	0,03	0,03	0,03
			BKM	0,05	0,04	0,05*	0,04*	0,04*	0,04*
		SVD++	KM	0,08	0,08	0,07	0,07	0,07	0,06
			BKM	0,08	0,08	0,08*	0,08*	0,09*	0,08*
		SO	KM	0,04	0,03	0,02	0,02	0,02	0,02
			BKM	0,04	0,03	0,03*	0,03*	0,03*	0,03*
	20	kNNBase	KM	0,05	0,05	0,04	0,04	0,04	0,04
			BKM	0,05	0,05	0,05*	0,05*	0,05*	0,05*
		SVD++	KM	0,07	0,08	0,07	0,08	0,07	0,07
			BKM	0,07	0,08	0,08*	0,09*	0,09*	0,09*
		SO	KM	0,03	0,04	0,04	0,03	0,04	0,04
			BKM	0,03	0,04	0,04*	0,04*	0,05*	0,05*

* %95 güven aralığında istatistiksel olarak anlamlı

Tablo 5.3 ve 5.4'ten görüleceği üzere, P&C stratejisi uygulandığında da, bir önceki bölümdeki sonuçlara benzer şekilde, BKM yaklaşımının farklı ayarlamalarda KM algoritmasına kıyasla daha uygun kullanıcı gruplarını belirlediği ve sonuç olarak daha kaliteli grup önerilerinin üretilmesini sağladığı açıkça görülmektedir. Ayrıca, P&C stratejisi ile uygulanan BKM ve KM algoritması ile elde edilen sonuçları karşılaştırmak için istatistiksel anlamlılık testleri yapılmış ve bu testlerin sonuçları tabloların dipnotlarında verilmiştir. Yapılan tek kuyruklu *t*-testlerinin sonuçları, BKM yaklaşımının KM algoritması üzerinde sağladığı iyileştirmelerin, genellikle büyük kullanıcı gruplarında %95 güven aralığında istatistiksel olarak anlamlı olduğunu göstermektedir.

Önceki bölümlerde gerçekleştirilen deneylerin sonuçlarına paralel şekilde, her iki veri kümesi için yapılan bu deneylerde de en yüksek nDCG değerleri, birleştirme tekniği olarak Avg tekniğini seçildiği durumda elde edilmiştir. Bununla birlikte, üretilen grup önerilerinin doğruluk değerleri artan grup büyüklüğü ile azalmıştır. Ayrıca, her iki veri kümesi için elde edilen sonuçlar, bir önceki bölümde gerçekleştirilen deneylerde olduğu gibi en başarılı CF algoritmasının SVD++ olduğunu göstermiştir.

5.4. Değerlendirmeler

Bu bölümde, Bölüm 4.3'te geliştirilen ve grup önerilerinde kullanıcı gruplarının otomatik olarak belirlenmesini sağlayan BKM yaklaşımının performansı ayrıntılı biçimde analiz edilmiştir. Bu yapılırken, grup önerilerinin üretilmesi amacıyla geleneksel CF algoritması yerine, kNNBase, SVD++ ve SO olmak üzere üç farklı modern CF algoritması kullanılmıştır. Bununla birlikte, iki farklı veri kümesi üzerinde grup büyüklüğü, öneri listesi boyutu ve kullanılan birleştirme tekniği olmak üzere farklı parametrelerle birçok deneysel çalışma gerçekleştirilmiştir.

Gerçekleştirilen ilk deney setinde, BKM yaklaşımının standart KM algoritmasına göre değişen parametrelerden bağımsız olarak daha uygun kullanıcı gruplarını belirlediğini göstermiştir. Bu deneylerde, söz konusu modern CF algoritmaları ile üretilen bireysel tahminler yalnızca grup önerilerinin üretilmesinde kullanılmış ve kullanıcı grupları orijinal kullanıcı $\times$ ürün matrisi aracılığı ile belirlenmiştir. Gerçekleştirilen ikinci deney setinde ise kullanıcı gruplarının otomatik olarak belirlenmesinde başarılı yaklaşımlardan birisi olan P&C stratejisi kullanılmış ve modern CF algoritmaları ile üretilen bireysel tahminler hem grupların belirlenmesi amacıyla hem de daha sonra grup önerilerinin üretilmesi amacıyla kullanılmıştır. Diğer bir deyişle, kullanıcı grupları, bireysel tahminler ile doldurulmuş olan kullanıcı $\times$ ürün matrisi aracılığı ile tanımlanmıştır. Elde edilen sonuçlar, P&C stratejisinde de BKM yaklaşımının KM algoritmasına kıyasla daha uygun grupları belirlediğini ve doğruluk açısından daha iyi grup önerilerinin üretilmesini sağladığını göstermiştir. Ayrıca, gerçekleştirilen deneylerin sonuçları, SVD++ algoritmasının diğer CF algoritmalarından daha iyi performans sergilediğini göstermiştir.

6. KİŞİLİK ÖZELLİK-TABANLI YENİ BİR BİRLEŞTİRME TEKNİĞİ

Bu bölümde, GÖS çalışmalarında kullanılmak üzere, kişilik özellik-tabanlı yeni bir birleştirme tekniği geliştirilmiştir.

6.1. Motivasyon

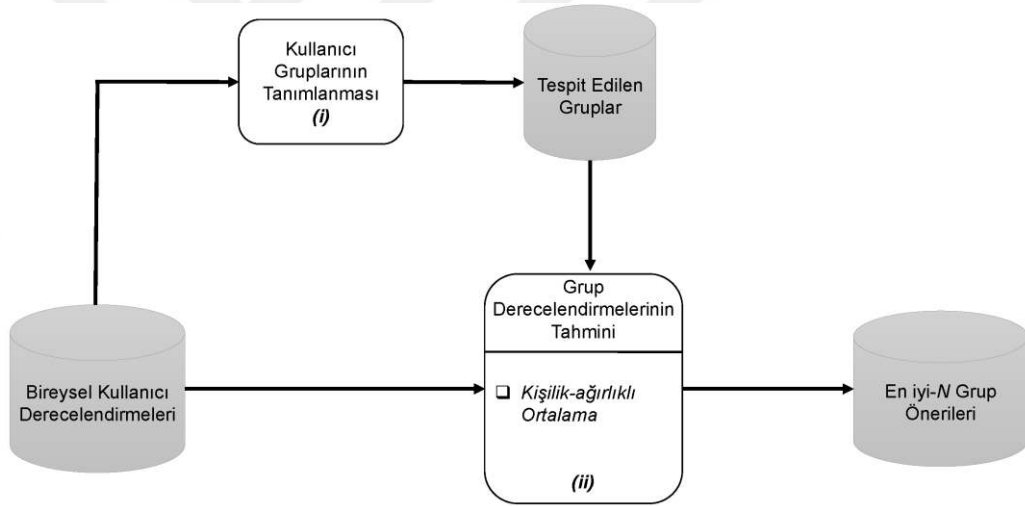
Bölüm 2.3'te ayrıntılı olarak açıklandığı üzere, grup önerileri üretebilmek için birçok birleştirme tekniği geliştirilmiştir. Bu tekniklerin büyük çoğunluğu, grup üyeleri tarafından sağlanan derecelendirmelerin farklı yönlerini dikkate alarak grup profilleri oluşturmaktadır. Diğer bir deyişle, grup profilleri genellikle yüksek ortalama, derecelendirme sayılarının sıklığı, sıralamalar, en yüksek/en düşük derecelendirmeler gibi özellikler göz önünde bulundurularak üretilmektedir. Her ne kadar bu tür birleştirme teknikleri, makul grup önerilerinin üretilmesini sağlasa da, birleştirme sürecinde grup üyelerinin kişilik özelliklerini göz önünde bulundurmazlar. Ancak, benzer kişisel özelliklere sahip kullanıcıların benzer tercihlere/ilgi alanlarına sahip olma olasılığı daha yüksek olduğundan, karar verme sürecinin bireylerin duyguları ve kişilikleri ile doğrudan ilişkili olduğu bilinen bir olgudur (Hu & Pu, 2010).

Bununla birlikte, mevcut birleştirme teknikleri genellikle bir grubun her üyesinin grubun nihai kararı üzerinde eşit derecede etkileri olduğu varsayımına dayanmaktadır. Bununla birlikte, sezgisel olarak, bireylerin, dâhil oldukları topluluğun karar verme sürecinde farklı etkileri olabilir ve gruptaki bir üyenin etki derecesi onun kişilik özelliklerine doğrudan bağlı olabilir (Rossi & Cervone, 2016). Örneğin, grupta, çıkarlarını kaybetmek istemedikleri veya kendi seçimlerinin grubun her üyesi için ideal kararlar olduğuna inandıkları için düşüncelerinde ısrarcı olan bazı üyeler olabilir. Öte yandan, gruptaki diğer insan türleri, diğer tüm üyelerin memnuniyeti konusunda kişisel çıkarları pahasına endişelenebilir ve bu nedenle tüm grup için uygun bir mutabakata varmak amacıyla bazı kişisel çıkarlarından feragat etmeye istekli olabilirler. Grup önerileri üretirken bireylerin bu tür farklı davranışsal özelliklerini göz önünde bulundurmak için, bireylerin kişilik özelliklerini insan bilimleri alanında öne çıkan bazı modeller aracılığı ile incelemek gerekmektedir. Bu anlamda en çok kullanılan modellerden birisi, insan kişiliğini *açıklık*, *uyumluluk*, *duygusal istikrar*, *vicdanlılık* ve *dışadönüklük* olmak üzere beş temel özellik ile tanımlayan beş büyük faktör modelidir (Costa Jr & McCrae, 1992).

Bu bölümde, yüksek kalitede grup önerileri üretmek amacıyla, yukarıda belirtilen tüm kişisel özellikleri, birleştirme aşamasında dikkate alan ve PwAvg olarak adlandırılan yeni bir birleştirme tekniği önerilmektedir. Spesifik olarak, bu teknik, bir gruptaki üyelerin kişiliklerini kullanarak, onların grubun nihai kararı üzerindeki etkisini ölçer ve bu özellikleri, birleştirme aşamasında bireylerin derecelendirmelerini ağırlıklandırmak üzere kullanır.

6.2. PwAvg Birleştirme Tekniğini İçeren Grup Önerisi Şeması

Bu bölümde, Şekil 6.1’de gösterildiği gibi, iki temel adımdan oluşan bir grup önerisi şeması kullanılmaktadır. Buna göre, ilk olarak, (i) benzer ilgi alanlarına sahip kullanıcı grupları, KM algoritması ile tanımlanmaktadır. Ardından, (ii) birleştirme sürecinde grup üyelerinin kişilik özelliklerini kullanan, yeni bir birleştirme tekniği ile hesaplanan grup derecelendirmelerine dayanarak en iyi- N grup önerileri üretilmektedir.



Şekil 6.1. PwAvg birleştirme tekniğini içeren grup önerisi şeması

Grup önerileri genellikle bir grup bireyin ürünler için sağladığı derecelendirmelerin bir araya getirilmesi ile oluşturulan grup profillerine dayanarak üretilmektedir. Birleştirme teknikleri, bu tür grup profillerinin uygun bir biçimde oluşturulması için kullanışlı araçlardır. Kullanıcılar tarafından sağlanan derecelendirmelerin özelliklerini analiz ederek nitelikli grup önerilerinin üretilmesini sağlayan çeşitli birleştirme teknikleri olmasına rağmen, bu teknikler genellikle grup üyelerinin kişisel özelliklerini görmezden gelirler. Ancak, duygu ve kişilik gibi bireylerin psikolojik yönleri, benzer kişisel özelliklere sahip kişilerin benzer tercihlere sahip olma olasılıklarının daha yüksek olması nedeniyle, karar verme süreçlerinde hayati bir rol oynamaktadır (Rossi & Cervone, 2016).

Bu doğrultuda, öneri sistemleri alanında son zamanlarda yapılan birçok çalışmada, kullanıcıların bu tür özellikleri, daha kullanıcı odaklı yaklaşımlar geliştirmek amacıyla kullanılmış ve çeşitli avantajlar elde edilmiştir.

Psikolojide kişilik, farklı durumlarda ifade edilen motivasyon, davranış, duygu, tercih ve düşünce kalıpları olarak tanımlanmaktadır (Rossi & Cervone, 2016). Öneri sistemi bağlamında ise, hem içerikten hem de alandan bağımsız bir kullanıcı profili olarak tanımlanabilir. Yani, konum, zaman veya başka bir alan/içerik ile değişmeyen kullanıcı profilleridir (Quijano-Sanchez, Recio-Garcia, & Diaz-Agudo, 2010). Bu tür kullanıcı profilleri genellikle bir insanı aşağıda sıfatlarının listelendiği beş temel özelliğe göre tanımlayan beş büyük faktör modeli yoluyla oluşturulmaktadır (Costa Jr & McCrae, 1992).

1. *Deneyime açık olmak:* Meraklı, yaratıcı ve anlayışlı olmak.
2. *Kabul edilebilir olmak:* Affedici, cömert, güvenen ve sepatik olmak.
3. *Duygusal yönden istikrarlı olmak:* Kendine güvenen, kararlı ve dengeli olmak.
4. *Vicdanlı olmak:* Güvenilir, organize, verimli ve sorumlu olmak.
5. *Dışadönük olmak:* Enerjik, konuşkan, iddialı ve aktif olmak.

Bu bölümde, grup üyelerinin kişisel özellikleri yansıtan iyi yapılandırılmış grup profilleri oluşturmak ve sonuç olarak öneri kalitesini arttırmak için, üyelerin bu tür kişilik bilgilerini birleştirme aşamasında kullanan ve aşağıda ayrıntılı olarak açıklanan kişilik özelliklerine dayanan yeni bir birleştirme stratejisi geliştirilmiştir.

Mevcut birleştirme teknikleri, genellikle, bir gruptaki her kullanıcının grubun genel zevki üzerinde eşit derecede bir etkiye sahip olduğunu varsayarak bireysel derecelendirmeleri bir araya getirmektedir. Ancak, gruptaki üyelerin grubun karar alma süreci üzerinde farklı etkileri olabilir. Bu sorunun üstesinden gelmek için, bu bölümde, Avg tekniği ile gruptaki üyelerin kişilik özelliklerini bir uyum içerisinde birlikte kullanan PwAvg tekniği geliştirilmiştir. Spesifik olarak, bu teknik, başlangıçta, belirli bir kişilik özellik bağlamında bir grup üyesinin grup üzerindeki göreceli önemini ölçen bir etki skoru hesaplar. Ardından, bu skoru birleştirme sürecinde ilgili kullanıcının tercihini ağırlıklandırmak için kullanır.

Bu çalışmada kullanılan ve özellikleri Bölüm 2.4’de ayrıntılı olarak açıklanan Per veri kümesi göz önünde bulundurulduğunda, bir kullanıcının beş temel kişilik özelliği

(yani, *açıklık* (*aç*), *uyumluluk* (*uy*), *duygusal istikrar* (*di*), *vicdan* (*vi*) ve *dışadönüklük* (*dd*)), 7-yıldızlı bir derecelendirme ölçeğindeki ayırık skorlar ile değerlendirilmektedir. Buna göre, g grubunun üyesi olan k kullanıcısının, $ü$ ürünü için yapmış olduğu derecelendirmenin $d_{k,ü}$, ve $t \in \{aç, uy, di, vi, dd\}$ olmak üzere ilgili kullanıcının t özelliği için verilen skorun $s_{k,t}$ ile temsil edilmesi durumunda, t özelliğine dayanan kişilik ağırlıklı ortalama tekniği olan $PwAvg_t$, Denklem 10'da verilen formülü kullanarak, $ü$ ürünü için bir grup derecelendirmesi $D_{g,ü}$ hesaplar:

$$D_{g,ü} = \frac{\sum_{k \in g} d_{k,ü} \times s_{k,t}}{|g|} \quad (6.1)$$

Burada, grup üyelerinin kişilik özellikleri dikkate alındığında, bu üyeleri temsil eden kişilik skorlarının aralığının her grup için farklı olması olasıdır. Dolayısıyla, grup üyelerinin grup içerisindeki göreceli önemini daha doğru bir şekilde belirlemek için, her grup için bir normleştirme süreci gerekmektedir. Bu amaçla, g grubu için $D_{g,ü}$ değerlerini hesaplamadan önce, ilgili gruptaki $s_{k,t}$ skorları, min-max normalizasyonu yardımıyla $[0, 1]$ ölçeğine dönüştürülmüştür.

6.3. Deneysel Çalışmalar

Bu bölümde, geliştirilen birleştirme tekniğinin etkinliğini incelemek için gerçekleştirilen deneysel çalışmalar sunulmaktadır. İlk olarak, gerçekleştirilen deneylerde izlenen metodoloji ayrıntılı bir biçimde açıklanmaktadır. Ardından, yapılan deneylerin sonuçları sunularak tartışılmaktadır.

6.3.1. Deney metodolojisi

Bu çalışmada önerilen birleştirme tekniğinin performansını değerlendirmek için yapılan çeşitli deneylerde, Bölüm 2.4'de ayrıntılı olarak açıklanan Per5 ve Per10 veri kümeleri kullanılmıştır. Ayrıca, deneyler, 5-kat çapraz doğrulama deney metodolojisi kullanılarak gerçekleştirilmiştir. Bu amaçla, ürünler kümesi rastgele biçimde beş alt kümeye bölünerek, oluşturulan her alt kümenin, ürünlerin %20'ünü içermesi sağlanmıştır. Her yinelemede, alt gruplardan birisi test kümesi olarak, kalan dört alt grubun kombinasyonu da eğitim seti olarak kullanılmıştır. Dolayısıyla, nihai $nDCG$ değerleri, 5-kat deneysel doğruluk sonuçlarının ortalaması alınarak hesaplanmıştır. Ayrıca, önceki bölümlere benzer şekilde $nDCG$ skoru hesaplanırken kullanıcıların derecelendirmeleri Bölüm 2.2'de açıklanan CF algoritması ile tahmin edilmiştir.

Test ve eğitim kümeleri oluşturulduktan sonra, benzer kullanıcılardan oluşan gruplar, eğitim veri kümesi üzerinde standart KM algoritması uygulanarak tanımlanmıştır. Ayrıca, grup büyüklüğünün önerilen yaklaşımın performansı üzerindeki etkisini değerlendirmek için, deneylerde, 2 ile 64 arasında değişen çeşitli sayılardaki (k) gruplar kullanılmıştır. Gruplar tanımlandıktan sonra, her grup sayısı için, önerilen PwAvg ve kıyaslama birleştirme teknikleri kullanarak grup derecelendirmeleri hesaplanmıştır. Burada, kıyaslama birleştirme teknikleri olarak Avg, AwM ve LM teknikleri seçilmiştir. Ayrıca, kullanılan kişilik özellikle göre PwAvg tekniğinin beş farklı varyantı kullanılmıştır. Her bir ürün için grup derecelendirmeleri hesaplandıktan sonra, N değerinin 3, 5, 10 ve 20 olarak ayarlandığı en iyi- N ürünlerini seçerek grup önerileri üretilmiştir. Son olarak, önerilerin en iyi- N listelerine göre bireylerin $nDCG$ değerleri hesaplanmış ve ortalamaları alınarak sistemin genel doğruluk performansı gözlemlenmiştir.

6.3.2. Deney sonuçları ve tartışma

Önerilen kişilik-ağırlıklı ortalama tekniğinin grup derecelendirmelerini hesaplama bağlamındaki doğruluk performansını değerlendirmek için, kullanılan kişilik özelliği, öneri listesinin boyutu (N) ve grup sayısı (k) olmak üzere farklı parametreler ile çeşitli deneyler gerçekleştirilmiştir. Elde edilen deneysel sonuçlar, sırasıyla Tablo 6.1 ve 6.2’de gösterildiği gibi, Per5 ve Per10 veri kümeleri için kıyaslama birleştirme teknikleri olan Avg, AwM ve LM ile karşılaştırılmıştır. Ayrıca, PwAvg ve kıyaslama teknikleri ile elde edilen sonuçları karşılaştırmak için istatistiksel anlamlılık t -testleri gerçekleştirilmiş ve bu testlerin sonuçları tabloların dipnotlarında verilmiştir.

Tablo 6.1 ve 6.2’den görülebileceği gibi, her iki veri kümesi için yapılan deneylerin $nDCG$ sonuçları, tüm kıyaslama tekniklerinin doğruluk performansının grup sayısı arttıkça, yani grupların büyüklüğü azaldıkça artmıştır. Diğer yandan, PwAvg varyant teknikleri, özellikle büyük gruplar söz konusu olduğunda, en yüksek $nDCG$ sonuçlarına ulaşmıştır. Bunun nedeni, artan grup büyüklüğünün farklı özelliklere sahip bireylerden oluşan gruplara sahip olunmasına yol açmasıdır. Dolayısıyla, birleştirme işlemini gerçekleştirirken, sadece grup üyelerinin sağladığı derecelendirmeleri değil, aynı zamanda onların kişiliklerini de dikkate almak, yüksek kaliteli grup önerilerinin üretilmesini sağlamıştır.

Per5 ve Per10 veri kümelerinden elde edilen sonuçlar karşılaştırıldığında, hem PwAvg varyantları hem de tüm kıyaslama teknikleri hemen hemen tüm ayarlamalar için Per10 veri kümesinde daha başarılı olduğu gözlemlenmiştir. Bu bulgunun temel nedeni, Per10 veri kümesindeki bir ürün için sağlanan toplam derecelendirme sayısının Per5’dekinden daha fazla olmasıdır. Bu durum, birleştirme sürecine daha fazla değerinin dâhil olmasına neden olarak, kullanılan birleştirme tekniğinin daha iyi çalışmasını sağlamıştır.

Tablo 6.1. Per5 veri kümesi için elde edilen nDCG sonuçları

en iyi-N	Birleştirme Tekniği	Grup Sayısı (k)					
		Büyük Gruplar		Orta Gruplar		Küçük Gruplar	
		2	4	8	16	32	64
3	Avg	0.647	0.671	0.708	0.732	0.743	0.750
	AwM	0.635	0.692	0.718	0.739	0.745	0.749
	LM	0.652	0.671	0.710	0.736	0.739	0.751
	PwAvg	<i>açıklık</i>	0.629	0.699	0.722	0.730	0.702
		<i>uyumluluk</i>	0.683 [†]	0.708 [†]	0.704	0.723	0.704
		<i>duygusal_istikrar</i>	0.694[†]	0.720[†]	0.722	0.726	0.699
		<i>vicdan</i>	0.681 [†]	0.699	0.692	0.703	0.696
		<i>dışadönüklük</i>	0.668	0.669	0.716	0.721	0.683
							0.678
5	Avg	0.666	0.664	0.689	0.724	0.736	0.751
	AwM	0.653	0.676	0.689	0.728	0.740	0.751
	LM	0.650	0.689	0.707	0.724	0.746	0.742
	PwAvg	<i>açıklık</i>	0.664	0.671	0.691	0.707	0.699
		<i>uyumluluk</i>	0.705 [†]	0.692	0.692	0.722	0.697
		<i>duygusal_istikrar</i>	0.711[†]	0.715[†]	0.705[†]	0.725	0.715
		<i>vicdan</i>	0.689 [†]	0.704 [†]	0.700 [†]	0.715	0.710
		<i>dışadönüklük</i>	0.657	0.694	0.679	0.718	0.706
							0.638
10	Avg	0.640	0.688	0.721	0.721	0.731	0.743
	AwM	0.659	0.690	0.727	0.720	0.730	0.739
	LM	0.642	0.689	0.730	0.727	0.731	0.739
	PwAvg	<i>açıklık</i>	0.641	0.693	0.720	0.708	0.716
		<i>uyumluluk</i>	0.659	0.697	0.738	0.716	0.711
		<i>duygusal_istikrar</i>	0.677[†]	0.723[†]	0.741[†]	0.718	0.717
		<i>vicdan</i>	0.660	0.721 [†]	0.733	0.713	0.713
		<i>dışadönüklük</i>	0.677 [†]	0.707 [†]	0.738	0.705	0.708
							0.681
20	Avg	0.660	0.688	0.719	0.731	0.738	0.735
	AwM	0.659	0.688	0.716	0.734	0.733	0.724
	LM	0.660	0.681	0.714	0.732	0.732	0.723
	PwAvg	<i>açıklık</i>	0.655	0.693	0.717	0.734	0.683
		<i>uyumluluk</i>	0.687 [†]	0.707 [†]	0.715	0.730	0.683
		<i>duygusal_istikrar</i>	0.705[†]	0.719[†]	0.718	0.740	0.680
		<i>vicdan</i>	0.693 [†]	0.718 [†]	0.727	0.737	0.692
		<i>dışadönüklük</i>	0.704 [†]	0.702	0.729	0.735	0.669
							0.658

[†] En iyi performans gösteren kıyaslama tekniğine göre %95 güven aralığında istatistiksel olarak anlamlı

Tablo 6.2. *Per10 veri kümesi için elde edilen nDCG sonuçları*

en iyi-N	Birleştirme Tekniği	Grup Sayısı (k)					
		Büyük Gruplar		Orta Gruplar		Küçük Gruplar	
		2	4	8	16	32	64
3	Avg	0.698	0.736	0.725	0.738	0.742	0.758
	AwM	0.680	0.716	0.722	0.731	0.744	0.760
	LM	0.651	0.708	0.714	0.733	0.739	0.751
	PwAvg	<i>açıklık</i>	0.688	0.718	0.710	0.735	0.689
		<i>uyumluluk</i>	0.707	0.729	0.715	0.726	0.681
		<i>duygusal_istikrar</i>	0.713[†]	0.745[†]	0.752[†]	0.739	0.693
		<i>vicdan</i>	0.697	0.707	0.722	0.727	0.682
		<i>dışadönüklük</i>	0.686	0.718	0.741	0.737	0.678
5	Avg	0.710	0.678	0.728	0.734	0.744	0.757
	AwM	0.659	0.673	0.725	0.733	0.746	0.754
	LM	0.680	0.686	0.705	0.726	0.738	0.746
	PwAvg	<i>açıklık</i>	0.677	0.657	0.722	0.713	0.715
		<i>uyumluluk</i>	0.684	0.689	0.715	0.715	0.722
		<i>duygusal_istikrar</i>	0.727[†]	0.707[†]	0.723	0.714	0.727
		<i>vicdan</i>	0.702	0.693	0.723	0.711	0.732
		<i>dışadönüklük</i>	0.701	0.706 [†]	0.721	0.708	0.722
10	Avg	0.714	0.710	0.710	0.733	0.741	0.755
	AwM	0.662	0.683	0.710	0.735	0.741	0.753
	LM	0.711	0.702	0.722	0.727	0.732	0.747
	PwAvg	<i>açıklık</i>	0.698	0.716	0.716	0.723	0.687
		<i>uyumluluk</i>	0.711	0.716	0.720	0.719	0.681
		<i>duygusal_istikrar</i>	0.726[†]	0.725[†]	0.738[†]	0.720	0.692
		<i>vicdan</i>	0.714	0.719	0.732 [†]	0.729	0.686
		<i>dışadönüklük</i>	0.696	0.724 [†]	0.723	0.708	0.678
20	Avg	0.694	0.702	0.738	0.737	0.738	0.746
	AwM	0.610	0.689	0.719	0.730	0.727	0.731
	LM	0.674	0.700	0.724	0.729	0.730	0.737
	PwAvg	<i>açıklık</i>	0.674	0.700	0.736	0.717	0.698
		<i>uyumluluk</i>	0.715 [†]	0.709	0.737	0.719	0.685
		<i>duygusal_istikrar</i>	0.716[†]	0.716[†]	0.738	0.735	0.693
		<i>vicdan</i>	0.710 [†]	0.709	0.731	0.719	0.701
		<i>dışadönüklük</i>	0.687	0.714	0.735	0.728	0.697

[†] En iyi performans gösteren kıyaslama tekniğine göre %95 güven aralığında istatistiksel olarak anlamlı

Her iki veri kümesi için gerçekleştirilen deneylerin sonuçları, özellikle gruplar büyük olduğunda PwAvg varyant tekniklerinin genellikle kıyaslama tekniklerinden daha iyi doğruluk performansı sergilediğini göstermiştir. Ayrıca, PwAvg tekniğini oluşturmak için tüm kişilik özellikleri arasında *duygusal istikrar*'ın kullanılmasının, diğer kişilik özelliklerine kıyasla daha etkili olduğu gözlemlenmiştir. Bununla birlikte, her iki veri kümesi için, *duygusal istikrar*'a dayalı PwAvg tekniğinin tüm kıyaslama teknikleri üzerinde sağladığı iyileştirmeler, özellikle gruplar büyük olduğunda, %95 güven düzeyinde istatistiksel olarak anlamlı görünmektedir. Dolayısıyla, kararsız kullanıcılar tarafından değerlendirilen ürünlere kıyasla, duygusal yönden istikrarlı olan ve daha

kararlı grup üyeleri tarafından yüksek derecelendirme alan ürünlerin, grup önerileri olarak üretilmesinin, sistemden olan genel memnuniyeti arttırdığı sonucuna varılabilir.

6.4. Değerlendirmeler

GÖS, bireylerden ziyade ortak ilgi alanlarına sahip bir grup kullanıcıya uygun yönlendirmeler üreten kullanışlı araçlardır. Bu tür sistemlerde, grup önerileri genellikle gruptaki üyelerin derecelendirmeleri bir araya getirilerek üretilen grup derecelendirmeleri göz önünde bulundurularak üretilir. Ancak, grupların genellikle farklı kişilik özelliklerine sahip bireylerden oluşması nedeniyle, birleştirme işlemi, karmaşık bir görev olarak değerlendirilebilir. Bu durum, kişilik özelliklerinin dikkate alındığı yeni birleştirme mekanizmalarının geliştirilmesini gerektirir.

Bu çalışmada, sadece grup üyelerinin derecelendirmelerini değil, aynı zamanda bireylerin davranışsal özelliklerini de göz önünde bulundurarak grup derecelendirmelerini hesaplayan ve PwAvg olarak adlandırılan kişiliğe duyarlı bir birleştirme tekniği geliştirilmiştir. Spesifik olarak bu teknik, *açıklık, uyumluluk, duygusal istikrar, vicdan* ve *dışadönüklük* olmak üzere beş temel kişilik özelliğe göre gruptaki bireylerin grubun nihai kararı üzerindeki etkisini ölçer ve bunları daha sonra birleştirme aşamasında derecelendirmeleri ağırlıklandırmak için kullanır.

İki farklı veri kümesinde yapılan deneyler, önerilen tekniğin özellikle büyük kullanıcı grupları için, kıyaslama tekniği olarak seçilen üç farklı birleştirme tekniğinden daha iyi performans sergilediğini göstermiştir. Bu bulgu, istatistiksel anlamlılık testleriyle doğrulandığı üzere, öneri listesinin boyutundan bağımsız bir şekilde elde edilmiştir. Deneysel sonuçlar, ayrıca, *duygusal istikrar* özelliği ile PwAvg tekniğini inşa etmenin, diğerlerine kıyasla öneri doğruluğu açısından daha etkili olduğunu göstermiştir. Sonuç olarak, bireysel derecelendirmeler birleştirilirken grup üyelerinin farklı kişilik özelliklerini dikkate almak, gruplar kalabalıklaştıkça yüksek kalitede grup önerilerinin üretilmesine katkıda bulunduğu çıkarımı yapılmıştır.

7. ENTROPİ İLE GÜÇLENDİRİLEN MELEZ BİRLEŞTİRME TEKNİĞİ

Bu bölümde, GÖS çalışmalarında kullanılmak üzere melezleştirilmiş ve onun entropi ile güçlendirilmiş versiyonu olan iki yeni birleştirme tekniği geliştirilmiştir.

7.1. Motivasyon

GÖS genellikle, öneri maliyeti ve bağlamı nedeniyle kişiselleştirilmiş tavsiyeler sağlamak yerine, grup öneri yöntemleri ile benzer zevklere sahip bir grup kullanıcı için grup önerileri üretmektedir. Ancak, çoğu zaman, benzer ilgileri olan insan grupları önceden bilinmemektedir. Bu nedenle, genellikle, grup önerilerinin üretilmesinden önce kullanıcı toplulukların otomatik olarak algılanması gerekmektedir (Baltrunas, Makcinskias, & Ricci, 2010). Bu tür grupları tespit etmek için, genellikle KM veya hiyerarşik kümeleme gibi iyi bilinen kümeleme yöntemlerinden birisi kullanılır (Boratto vd., 2009). Hiçbir kümeleme algoritmasının mükemmel kümeler üretemeyeceği göz önünde bulundurulduğunda, çünkü kümeleme sübjektif bir konudur ve genel bir doğru yoktur, bu şekilde inşa edilen grupların farklı zevklere sahip kullanıcıları içermesi çok olasıdır. Bu durum, GÖS’de en istenmeyen senaryolardan birisi olan üretilecek grup önerilerinden memnun olan toplam kullanıcı sayısının azalmasına yol açar.

Grupların heterojen bir şekilde oluşturulması durumunda, grup üyelerinin tercihlerini bir araya getiren birleştirme teknikleri iyi çalışmayabilir. Bu durum daha çok tek bir birleştirme tekniği kullanıldığında ortaya çıkar, çünkü her bir teknik grup üyeleri tarafından sağlanan derecelendirmelerin yalnızca belirli bir yönünü dikkate almaktadır. Daha açık bir ifadeyle, Bölüm 2.3’de belirtildiği üzere, birleştirme teknikleri farklı stratejiler izleyerek birleştirme işlemini gerçekleştirmektedir. Örneğin, bazıları derecelendirmelerin sıklıklarını saymaya odaklanırken (Crossen, Budzik, & Hammond, 2002; Lieberman, Van Dyke, & Vivacqua, 1999), bazıları sağlanan derecelendirmelerin ortalamaları (McCarthy vd., 2006; Jameson, 2004) veya sıralamalarını (Marquez & Ziegler, 2016; Masthoff, 2011) göz önünde bulundurur. Birleştirme aşamasında devralınan gruplar uygun şekilde inşa edilmediğinde, tek bir teknik kullanılarak birleştirilmesi mümkün olmayan dengesiz ve farklı özelliklerdeki derecelendirmelerle ilgilenmek durumunda kalınır. Bu durumda, her biri sağlanan derecelendirmelerin farklı bir boyutunu ele alan birden fazla teknik kullanılmalıdır.

GÖS’de homojen olmayan gruplara sahip olmanın bir başka problemi, grup üyelerinin fikir birliğine varamayacağı ürünlerin sistem tarafından grup önerisi olarak

üretileme ihtimalidir. Dolayısıyla, bu tür ürünlerin üretilecek öneri listesine dâhil olma şansını azaltmak önemlidir, çünkü öneri sonuçlarından memnun olan grup üyesi sayısını en üst düzeye çıkarmak, bireyler için mükemmel şekilde tavsiyeler sağlamak kadar önemlidir. Bu amaçla, derecelendirmelerin dağılımlarının birleştirme aşamasından önce analiz edilerek bu tür ürünlerin belirlenmesi gerekmektedir. Örneğin, bir grup içinde bir ürün için grup üyeleri tarafından sağlanan derecelendirmelerin farklı değerleri benzer frekanslara sahipse (yani, eşit olarak dağılmışsa), tüm grup üyelerinin bu ürün üzerinde aynı fikirde olduğunu iddia etmek zordur. SS, bu tür bir analiz için kullanılan temel yöntemdir (Seo vd., 2018). Bir ürün için üyeler tarafından sağlanan derecelendirmelerin sapması yüksek çıkarsa, üyelerin ilgili ürün hakkında farklı görüşleri olduğu çıkarımı yapılmaktadır. Bu, makul bir yaklaşım gibi görünse de, aşağıda ayrıntılı olarak açıklanan bazı durumlarda, SS, birleştirme aşamasında derecelendirmelerin dağılımını doğru bir biçimde analiz etmede yetersiz kalmaktadır.

SS, doğası gereği, sürekli tipteki değişkenlerin dağılımını ölçen bir yaklaşımdır. Ancak, derecelendirmelerle temsil edilen kullanıcı tercihleri bazı istisnalar dışında genellikle ayrık tiptedir. Bu tutarsızlık, özellikle derecelendirmelerin $\{0,1\}$ veya $\{1,5\}$ gibi dar bir aralıkta temsil edildiği sistemlerde daha fazla soruna neden olur. On yıldızlı bir derecelendirme sisteminde bile, SS, derecelendirmelerin dağılımının genel davranışını tespit edemeyebilir. Çünkü derecelendirmelerin daha farklı değerlere sahip olmasına izin verildiğinde, SS'in dağılımı doğru olarak belirlemeyeceği durum olan birden fazla pike sahip bir dağılım oluşması muhtemeldir. Bu sorun, GÖS'de büyük ve heterojen grupların varlığında daha belirgindir, çünkü böyle bir durumda benzer tercihleri olan farklı alt kullanıcı gruplarına sahip olma olasılığı yüksektir.

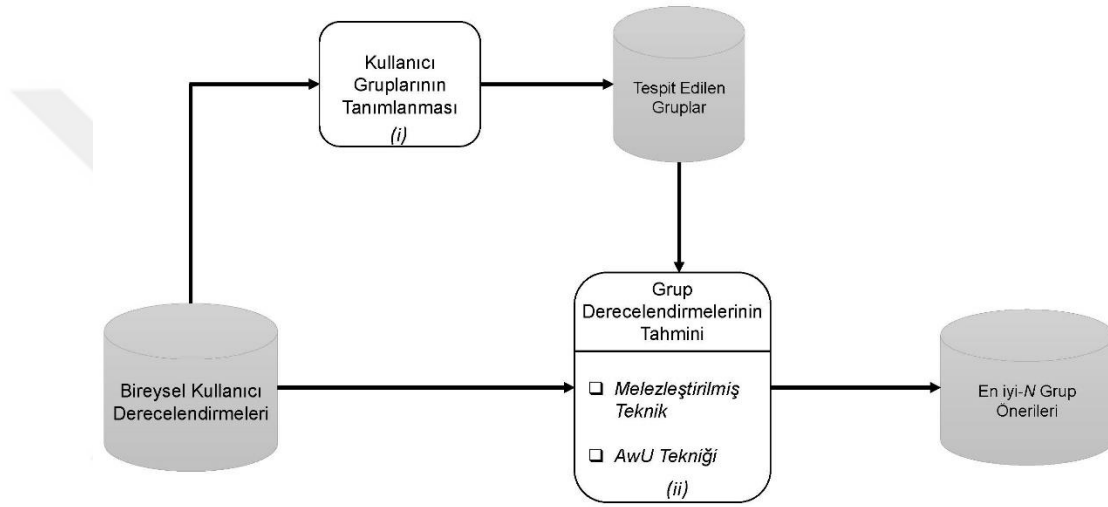
Bu bölümde, yukarıda ifade edilen problemlerin üstesinden gelmek için, yeni birleştirme tekniklerini içeren bir grup öneri şeması önerilmektedir. Aşağıda, bu bölümde literatüre sağlanan temel katkılar özetlenmektedir.

- i. Melezleştirilmiş birleştirme teknikleri olarak adlandırılan, daha önce Bölüm 2.3'de ayrıntılı açıklanan AV ve AU tekniklerini farklı iki şekilde birleştirilen iki yeni birleştirme tekniği geliştirilmiştir.
- ii. Geliştirilen melezleştirilmiş tekniklerin üzerine inşa edilen gelişmiş bir birleştirme tekniği olan AwU tekniği önerilmektedir. Spesifik olarak, AwU, bilgi entropisi kullanarak grup üyelerinin sağladığı derecelendirmelerin dağılımını

güçlü bir şekilde analiz eden ve sonuç olarak maksimum sayıda grup üyesini memnun edebilen grup önerilerinin üretilmesini sağlayan bir tekniktir.

7.2. Geliştirilen Birleştirme Tekniklerini İçeren Grup Önerisi Şeması

Bu bölümde, Şekil 7.1’de gösterildiği gibi, iki temel adımdan oluşan grup önerisi şeması kullanılmaktadır. Buna göre, ilk olarak, (i) benzer ilgi alanlarına sahip kullanıcı grupları, KM algoritması ile tanımlanmaktadır. Ardından, (ii) önerilen melezleştirilmiş tekniğinin iki farklı versiyonu ve geliştirilen nihai teknik olan AwU yardımıyla hesaplanan grup derecelendirmelerine dayanarak en iyi- N grup önerileri üretilmektedir.



Şekil 7.1. Melezleştirilmiş ve AwU birleştirme tekniklerini içeren grup önerisi şeması

7.2.1. Melezleştirilmiş birleştirme tekniği

Birleştirme teknikleri, grup üyelerinin tercihlerinin bir bütün olarak somutlaştırıldığı grup profillerinin oluşturulmasını sağlayan tekniklerdir (Boratto vd., 2016). Grup önerilerinin genellikle bu grup profillerine dayanılarak üretilmesi nedeniyle kullanılacak birleştirme tekniğinin seçimi kritik bir süreçtir. Literatürde farklı türlerde çeşitli birleştirme teknikleri olmasına rağmen, bu tekniklerden her biri grup profili oluşturmak için farklı bir strateji izlemektedir ve her birinin daha önce Bölüm 2.3’de anlatıldığı gibi bazı eksiklikleri mevcuttur. Diğer bir deyişle, her bir teknik, üyelerin derecelendirmelerini farklı bir bakış açısı ile birleştirerek, ürünler için tahmin edilen grup derecelendirmelerinden oluşan bir grup profili oluşturmaktadır. Dolayısıyla, birden fazla birleştirme tekniğini birlikte bir uyum içinde kullanmak, grup içinde tercih edilen ürünlerin belirlenmesi ve sonuç olarak GÖS’den elde edilen genel memnuniyetin artırılmasını sağlayabilir.

Birleştirme teknikleri melezleştirilirken ortaya çıkabilecek ilk soru, melez modelde hangi tekniklerin temel teknikler olarak seçilmesi gerektiğidir. Bu sorunun cevabı, inşa edilecek melez birleştirme tekniğinden olan beklentilerle doğrudan ilişkilidir. Bu çalışmada, grup üyelerinin üzerinde fikir birliği sağladığı ve üyelerin çoğunun gözdesi olan ürünlerin öne çıkarılmasına odaklanılmıştır. İlk beklentiye karşılık için, grubun tüm üyeleri tarafından sağlanan derecelendirmeleri toplayarak bir ürün için grup derecelendirmesi hesaplayan AU tekniğinin kullanılması önerilmiştir. Bu strateji, söz konusu ürün hakkında görüş sahibi olan her üyenin grubun ilgili ürün için olan kararına aynı oranda katkıda bulunmasına izin vererek bir fikir birliğinin oluşmasını sağlamaktadır. Diğer yandan, ikinci beklentiye karşılık için, belirli bir eşik değerinin üzerinde sağlanan derecelendirmelerin sayısını sayan AV tekniğinin kullanılması önerilmiştir. Bu strateji, elde edilecek nihai melez tekniğin grup üyeleri arasında oldukça popüler olan ürünlere karşı hassasiyetini arttırarak ilgili ürünlerin öne çıkarılmasını sağlamaktadır.

Bölüm 2.3’de tartışıldığı üzere, temel birleştirme tekniklerinin bazı eksiklikleri mevcuttur. Spesifik olarak, AU tekniği, bir gruptaki üyelerin büyük çoğunluğundan düşük derecelendirme alan bir ürünü grup önerisi olarak seçme riski taşır. Bunun nedeni, bir ürün için sağlanan derecelendirmelerin toplamı özellikle ilgili ürünü çok sayıda grup üyesi değerlendirmişse yüksek çıkabilmesidir. Bununla birlikte, AV tekniği birleştirme işlemini gerçekleştirirken yalnızca yüksek derecelendirmeler alan ürünleri seçmesi sayesinde bu eksikliği giderir. Diğer yandan, AV tekniği, ürünleri değerlendirirken eşik değerin altında derecelendirme sağlayan üyelerin tercihlerini görmezden gelerek, yalnızca küçük bir üye grubunun derecelendirmelerini göz önünde bulundurur. Dolayısıyla, grup üyeleri arasında bir fikir birliğinin oluşmasını sağlayamaz. Ancak, AU tekniği bütün grup üyelerinin derecelendirmelerini eşit bir biçimde önemseyerek bu sorunun üstesinden gelir. Sonuç olarak, melezleştirilmiş bir teknik inşa ederken, AU ve AV tekniğinin birlikte kullanılması, bir sinerji yaratarak bu tekniklerin bireysel eksikliklerinin karşılıklı olarak giderilmesini sağlar.

Bu bölümde, AU ve AV’nin temel teknikler olduğu, melez tipte iki yeni birleştirme tekniği geliştirilmiştir. Geliştirilen melez teknikler hem (i) gruptaki üyelerin derecelendirmelerin toplamını hem de (ii) yüksek puan alan derecelendirme sayısını aynı anda dikkate alarak grup derecelendirmelerini belirlemektedir. Bunlardan ilki AU ile

ikincisi ise AV tekniğinin kullanılması ile elde edilmektedir. Bu sayede, üzerinde bir şekilde fikir birliği sağlanan popüler ürünlerin belirlenerek grup önerileri olarak üretilmesi sağlanmaktadır.

Geliştirilen melezleştirilmiş birleştirme teknikleri, grup derecelendirmeleri üzerinde hangi temel tekniğin daha belirleyici bir rol oynadığı konusunda farklılık göstermektedir. İtici güç olarak AU tekniğinin kullanıldığı melez teknik AU_{AV} , AV tekniğinin kullanıldığı melez teknik ise AV_{AU} olarak adlandırılmaktadır. Aşağıda, grup derecelendirmelerinin AU_{AV} ve AV_{AU} kullanılarak nasıl hesaplandığı ayrıntılı olarak açıklanmaktadır.

- AU_{AV} tekniği: Bu melez teknik, AU tekniğini itici güç olarak kullanırken Denklem 7.1’de gösterildiği gibi, AV tekniğini grup derecelendirmeleri üzerinde ek bir faktör olarak değerlendirmektedir. Bu formülde, $D_{g,\tilde{u}}^{AU_{AV}}$, g grubunda $\tilde{u}$ ürünü için AU_{AV} melez tekniği ile hesaplanan grup derecelendirmesini, $AU_{g,\tilde{u}}$ ve $AV_{g,\tilde{u}}$ ise $\tilde{u}$ için sırasıyla AU ve AV teknikleri kullanılarak hesaplanan grup derecelendirmelerini ifade etmektedir.

$$D_{g,\tilde{u}}^{AU_{AV}} = AU_{g,\tilde{u}} + (AU_{g,\tilde{u}} \times AV_{g,\tilde{u}}) \quad (7.1)$$

- AV_{AU} tekniği: Bu melez teknik, AV tekniğini, hesaplanacak nihai grup derecelendirmelerinde belirleyici faktör olarak değerlendirirken, Denklem 7.2’de gösterildiği gibi, AU tekniğini yardımcı bir etki unsuru olarak kullanmaktadır. Bu formülde $D_{g,\tilde{u}}^{AV_{AU}}$, g grubunda $\tilde{u}$ ürünü için AV_{AU} melez tekniği ile hesaplanan grup derecelendirmesini ifade etmektedir.

$$D_{g,\tilde{u}}^{AV_{AU}} = AV_{g,\tilde{u}} + (AV_{g,\tilde{u}} \times AU_{g,\tilde{u}}) \quad (7.2)$$

Bir ürün için, AU tekniği ile hesaplanabilecek maksimum grup derecelendirmesi, sistemde kullanılan derecelendirme ölçeğinin maksimum değerinin ve gruptaki mevcut üye sayısının çarpımına karşılık gelmektedir. Bu, tüm üyelerin ilgili ürün için en yüksek derecelendirmeyi verdiği durumda oluşabilir. Diğer yandan, AV tekniği ile hesaplanan grup derecelendirmesinin maksimum değeri gruptaki üye sayısına eşittir. Bu ise, tüm üyelerin belirlenen eşik değerin üzerinde bir derecelendirme sağladığı durumda oluşabilir. Bu nedenle, AU ile hesaplanan grup derecelendirmeleri AV tekniği ile hesaplanan derecelendirmelerden daha geniş bir aralığa sahip olur. Dolayısıyla, nihai

grup derecelendirmelerini hesaplamadan önce grup derecelendirmelerini karşılaştırılabilir hale getirmek için bir normalleştirme işlemi gerekmektedir. Bu amaçla, AU ve AV tarafından hesaplanan grup derecelendirmeleri ilk olarak min-max normalizasyonu yardımıyla $[0,1]$ ölçeğine dönüştürülmüş ve ardından AU_{AV} ve AV_{AU} kullanılarak birleştirilmiştir.

7.2.2. AwU birleştirme tekniği

Bir GÖS'ün başarısı, önerilen ürünlerden memnun olan toplam grup üyesi sayısının en üst düzeye çıkarılması ile güçlü bir şekilde ilişkilidir (Seo vd., 2018). Bu amaca ulaşmak için maksimum sayıda kullanıcı sayısı, önerilen ürünler üzerinde hemfikir olmalıdır. Burada, bir ürün üzerindeki uzlaşma derecesi, ürün için grup üyeleri tarafından sağlanan derecelendirmelerin dağılımı kullanılarak tanımlanabilir. Daha açık bir ifadeyle, bir ürün için yapılan derecelendirmelerin grup içinde eşit olarak dağılması durumunda, ilgili ürün için grup içindeki uzlaşma seviyesinin düşük olduğu kabul edilir. Diğer yandan, eğer dağılım eşit değilse, grup üyelerinin büyük çoğunluğu arasında bir uzlaşmaya varıldığını gösterebilir. Dolayısıyla, üzerinde uzlaşma sağlanan ürünlerin öneri listesine eklenmesi, üretilen önerilerden memnun olan üye sayısını arttırılmasını sağlar.

Bir gruptaki kullanıcıların derecelendirmelerinin nasıl dağıldığını analiz etmek için, son yıllarda öne çıkan bazı çalışmalarda SS kullanılmıştır. Burada, sapmaları yüksek olan ürünler, grup üyeleri arasındaki uzlaşmazlıklar ile ilişkilendirilir ve ilgili ürünler filtrelenerek öneri listelerine eklenmemesi sağlanır (Seo vd., 2018). Derecelendirmeler tek-doruklu dağılıma sahip olduğu sürece, üyelerin arasındaki uzlaşmazlıkları SS kullanarak belirlemek faydalı bir yöntemdir. Ancak, derecelendirmelerin çok-doruklu bir dağılıma sahip olması durumunda, SS, kullanıcılar arasındaki uzlaşmayı doğru bir biçimde analiz etmek için ideal bir araç olmayabilir.

GÖS'deki mevcut gruplar, çok benzer tercihlere sahip alt kullanıcı gruplarından oluştuğunda, çok-doruklu bir dağılım gözlenebilir. Bu durumlarda, bir gruptaki bir alt grup bir üründen büyük ölçüde memnun olurken, o gruptaki başka bir alt grup aynı üründen hoşlanmayabilir. Bu nedenle, ilk alt gruptaki kullanıcılar ilgili ürün için yüksek derecelendirmeler verirken, ikinci alt gruptaki kullanıcılar düşük derecelendirmeler sağlar. Daha da kötüsü, kullanıcı grupları ve/veya kullanılan derecelendirme yelpazesi yeterince büyükse (örneğin; on-yıldızlı derecelendirme yelpazesi), ikiden fazla alt grup

gözlemlenerek çok-doruklu dağılıma neden olur. Örneğin, MovieLens¹¹ veri kümesi göz önünde bulundurulduğunda, derecelendirmeleri çok-doruklu bir dağılım gösteren çeşitli ürünler vardır. Daha spesifik olarak, veri kümesindeki 853. ve 854. ürünler için verilen {1-5} yelpazesindeki derecelendirmeleri göreceli frekansları sırasıyla {0,31, 0, 0, 0,38, 0,31} ve {0,21, 0, 0, 0,29, 0,5}'dir.

Bir ürünün derecelendirmelerinin dağılımı çok-doruklu olduğunda, bu ürün için hesaplanan SS değeri yüksek olur ve bu da ilgili ürünün öneri listesinden çıkarılmasına neden olur; ancak, ürün üzerinde farklı seviyelerde hala bir fikir birliği olabilir. Bu sorunla başa çıkmak için bu çalışmada derecelendirme dağılımlarının analizinde entropinin kullanılması önerilmektedir.

Bilgi teorisinde entropi, belirli bir veri kaynağındaki belirsizliğin miktarını ölçmek için kullanılmaktadır ve son yıllarda yapılan birçok öneri sistemi çalışmasında etkin bir biçimde kullanılmıştır (Kaleli, 2014; Yargıç & Bilge, 2019). Bir örneklemin yüksek entropiye sahip olması, o örneklemini oluşturan olaylar kümesinin görülme sıklığındaki rastlantısallıkla ilişkilidir. Entropi, SS'e kıyasla özellikle çok-doruklu bir dağılım meydana geldiğinde derecelendirmelerin dağılımını daha başarılı şekilde analiz edebilir. Bununla birlikte, SS, daha çok sürekli değişkenler için uygun bir yöntemdir; ancak, derecelendirmeler doğası gereği genellikle ayrık tiptedir. Bu da, dağılımların doğru bir biçimde analizinde entropi kullanmayı gerektirir.

Yukarıda belirtilen nedenlerden dolayı, bu çalışmada, melezleştirilmiş birleştirme tekniklerini entropi ile güçlendirerek nihai bir gelişmiş birleştirme tekniği önerilmektedir. Bu nihai birleştirme tekniği kısaca AwU olarak adlandırılmaktadır. AwU tekniğinin ayrıntılarının açıklanmasından önce, bir derecelendirme sistemi kapsamında entropinin nasıl değerlendirildiği hakkında aşağıda ayrıntılı bilgi verilmektedir.

Bir ürün için grup üyeleri tarafından sağlanan derecelendirmeleri içeren V derecelendirme vektörü için, olası derecelendirmelerin $R = (d_1, d_2, \dots, d_k)$ ile, ve bu derecelendirmelerin görülme olasılıklarının $P = (p_1, p_2, \dots, p_k)$ ile temsil edildiği durumda, bu vektörün entropisi olan $H(V)$ değeri, Denklem 7.3'de verilen formül yardımıyla hesaplanmaktadır.

¹¹ <http://www.grouplens.org/>

$$H(V) = - \sum_{i=1}^k p(d_i) \log_2(p(d_i)) \quad (7.3)$$

Burada, yüksek $H(V)$ değeri, kullanıcıların ilgili üründen emin olmadığı, yani kullanıcılar arasında neredeyse hiçbir uzlaşma olduğu anlamına gelir. Öte yandan, düşük $H(V)$, kullanıcıların ürün hakkında belirli bir uzlaşma sağladığını gösterir.

Bu çalışmada önerilen AwU birleştirme tekniğinin çalışma mekanizması Prosedür 7.1’de ayrıntılı olarak sunulmaktadır. Buna göre, bir grubun derecelendirmeye dayalı kullanıcı $\times$ ürün matrisi ($D_{k \times \bar{U}}$) ve önceden tanımlanmış bir eşik değeri (τ) verildiğinde, AwU tekniği ilk olarak, Denklem 7.3’deki formül yardımıyla tüm ürünler için, grup üyelerinin derecelendirmeleri göz önünde bulundurarak entropi değerlerini hesaplar. Her bir ürün için entropi değerlerinin hesaplanmasından sonra, min-maks normalizasyonu ile her bir entropi değeri $[0-1]$ aralığına indirgenecek şekilde normalleştirilir. Bu işlemin ardından, AwU tekniği normalleştirilmiş entropi değerlerini grup üyeleri arasındaki uzlaşma miktarını ölçmek için 1’den çıkartır, böylece ürünler için bilgi kazancı hesaplanmış olur. Son olarak, bu algoritma, yalnızca ilgili bilgi kazancı, önceden tanımlanmış eşik değerinin (τ) grubun ortalama bilgi kazancı (m) ile çarpımından elde edilen değerin üzerinde olan ürünler için bir önceki bölümde açıklanan melezleştirilmiş teknikleri kullanarak grup derecelendirmelerini hesaplar. Diğer bir deyişle, bu teknik, yüksek entropiye sahip ürünleri göz ardı etmektedir. Böylece, bu ürünler, kullanıcıların üzerinde uzlaşma sağlayamadığı ve önerilmesi uygun olmayan ürünler olarak değerlendirilerek GÖS tarafından önerilmez.

Prosedür 7.1. AwU birleştirme tekniği

Giriş: g grubunun derecelendirme matrisi ($D_{k \times \bar{u}}$), Eşik değeri (τ)

Başlatma:

1: $e_{(1 \times \bar{u})} \leftarrow null$ ► ürünler için entropi vektörü

Her ürün için bilgi kazancı ve entropi hesaplaması:

2: **for** $\bar{u}$ **in** $\{1, 2, \dots, \bar{U}\}$ **do**

3: $e(i) \leftarrow H(D(:, \bar{u}))$ ► Denklem 7.3 yardımıyla

4: **end for**

5: $\bar{e} \leftarrow Normalizasyon(e)$ ► min-maks normalizasyonu

6: **for** $\bar{u}$ **in** $\{1, 2, \dots, \bar{U}\}$ **do**

7: $\bar{e}(i) \leftarrow 1 - \bar{e}(i)$ ► bilgi kazancı

8: **end for**

9: $m \leftarrow ortalama(\bar{e})$ ► ortalama bilgi kazancı

g grubu için grup derecelendirmelerin hesaplanması:

10: **for** $\bar{u}$ **in** $\{1, 2, \dots, \bar{U}\}$ **do**

11: **if** $\bar{e}(i) > (\tau \times m)$ **then**

12: $D_{g, \bar{u}}^{AU} \leftarrow AU_{AV}(\bar{u})$ ► Denklem 7.1 yardımıyla

13: $D_{g, \bar{u}}^{AV} \leftarrow AV_{AU}(\bar{u})$ ► Denklem 7.2 yardımıyla

14: **end if**

15: **end for**

16: **return** $D_{g, \bar{u}}^{AU}$ **or** $D_{g, \bar{u}}^{AV}$

7.2.3. Açıklayıcı örnek

Bu bölümde, geliştirilen AwU tekniğinin birleştirme aşamasında grup derecelendirmelerini nasıl belirlediğini açıklığa kavuşturmak için basit bir örnek sunulmaktadır. Bu amaçla, Tablo 7.1’de gösterildiği gibi, altı ürün için $\{1,5\}$ derecelendirme ölçeğinde değerlendirmeler sağlayan yedi kişilik bir grubun örneği tanımlanmıştır.

AwU tekniği, başlangıçta, her bir ürün için ilgili derecelendirmeleri kullanarak Denklem 7.3’de sunulan formül yardımıyla bir entropi değeri hesaplar ve bu entropi değerlerini $[0-1]$ aralığına min-maks normalizasyonu kullanarak normalize eder.

Ardından, Tablo 7.2’de gösterildiği gibi, grup üyelerinin ürünler üzerinde sağladığı uzlaşmanın derecesini ölçmek için, her bir ürün için hesaplanan entropi değerlerini 1’den çıkartarak bilgi kazancını hesaplar. Bu örnekte hesaplanan kazanç değerleri dikkate alındığında, i_6 , minimum entropi (maksimum bilgi kazancı) değerine sahip olması nedeniyle, grup üyelerinin üzerinde en çok uzlaşma sağladığı ürün olarak kabul edilebilir. Diğer yandan, en yüksek entropi değerine, dolayısıyla en düşük kazanç değerine sahip olan i_5 , grup üyelerinin üzerinde neredeyse hiçbir uzlaşma sağlayamadığı ürün olarak kabul edilebilir.

Tablo 7.1. Üyelerin ürün tercihlerini gösteren örnek grup

	$\ddot{u}_1$	$\ddot{u}_2$	$\ddot{u}_3$	$\ddot{u}_4$	$\ddot{u}_5$	$\ddot{u}_6$
k_1	1	5	-	1	-	3
k_2	2	-	-	2	5	3
k_3	4	-	5	4	4	-
k_4	3	5	5	5	-	3
k_5	4	5	2	4	-	-
k_6	-	4	2	4	1	-
k_7	2	5	5	-	3	-

Tablo 7.2. Ürünlerin entropi ve bilgi kazancı değerleri

	$\ddot{u}_1$	$\ddot{u}_2$	$\ddot{u}_3$	$\ddot{u}_4$	$\ddot{u}_5$	$\ddot{u}_6$
Entropi	1,92	0,72	0,97	1,79	2	0
Entropi(normalize)	0,96	0,36	0,48	0,89	1	0
Bilgi kazancı	0,04	0,64	0,52	0,10	0	1

Tüm ürünler için kazanç değerleri hesaplandıktan sonra, AwU tekniği, eşik değerin belirlenmesinde kullanılan, grubun ortalama kazanç değerini (m) hesaplar. Grubun ortalama kazancı, tüm grup üyeleri tarafından ulaşılan uzlaşmanın derecesinin göstergesi olarak değerlendirilebilir. Somut olarak, yüksek m değeri, grubun benzer zevklere sahip bireylerden oluştuğunu, yani grubun iyi yapılandırılmış olduğunu gösterir. Öte yandan düşük m değeri, grubun farklı zevklere sahip kullanıcılar içerdiği anlamına gelir. Verilen örnek için, grubun ortalama kazanç m değeri 0,338 olarak hesaplanır, bu da kullanıcılar arasında orta düzeyde bir fikir birliği olduğunu gösterir.

Son adımda, AwU, yalnızca kazancı m ile önceden tanımlanmış bir katsayının (τ) çarpılması ile elde edilen eşik değerin üzerinde olan ürünler için, Denklem 7.1 ve 7.2’de sunulan melezleştirilmiş tekniklerden birisini kullanarak grup derecelendirmelerini hesaplar. Burada τ değeri, ürünlerin öneri listesinden kaldırılmasında entropinin etkisini düzenler. Spesifik olarak τ sıfır olarak ayarlandığında, entropinin melez teknikler

üzerinde hiçbir etkisi olmaz, çünkü bu durumda hiçbir ürün elenmez. Bire ayarlandığında ise maksimum etkiye sahip olur ve melez teknikler, az sayıdaki elenmemiş ürün için grup derecelendirmesi hesaplar. Tablo 7.2’de sunulan kazanç değerlerine göre τ değerinin 0,5 olarak ayarlandığı durum için, AwU tekniği ile hesaplanan nihai grup derecelendirmeleri Tablo 7.3’de verilmiştir. Tablo 7.2’den açıkça görülebildiği gibi i_1 , i_4 ve i_5 ürünlerinin kazanç değerleri, $(\tau \times m = 0,385 \times 0,5 = 0,191)$ değerinden daha düşüktür. Bu yüzden, AwU tekniği bu ürünlerin grup için tavsiye edilmeyeceğini kabul eder ve bu ürünleri Tablo 7.3’den de görülebileceği gibi yok sayar. Öte yandan, diğer ürünler için, yani i_2 , i_3 ve i_6 için, sırasıyla Denklem 7.1 ve 7.2’de formüle edilen AU_{AV} ve AV_{AU} tekniklerinden birisini kullanarak bir grup derecelendirmesi hesaplar. Son olarak, hesaplanan grup derecelendirmeleri, GÖS tarafından önerilen ürünlerin bir listesini oluştururken kullanılır.

Tablo 7.3. AwU tekniği ile hesaplanan grup derecelendirmeleri

		$\ddot{u}_1$	$\ddot{u}_2$	$\ddot{u}_3$	$\ddot{u}_4$	$\ddot{u}_5$	$\ddot{u}_6$
Grup	AU_{AV}	-	1,67	1,11	-	-	0
Derecelendirmeleri	AV_{AU}	-	1,33	1,11	-	-	0

7.3. Deneysel Çalışmalar

Bu bölümde, geliştirilen melezleştirilmiş ve AwU birleştirme tekniklerinin etkinliğini incelemek için gerçekleştirilen deneysel çalışmalar sunulmaktadır. İlk olarak, önerilen nihai AwU tekniğinin daha kapsamlı analizi için karşılaştırma metotları olarak seçilen iki modern birleştirme yöntemi açıklanmıştır. Ardından, gerçekleştirilen deneylerde izlenen metodoloji ayrıntılı bir biçimde açıklanmaktadır. Son olarak, yapılan deneylerin sonuçları sunularak tartışılmaktadır.

7.3.1. Karşılaştırma metotları

Deneysel çalışmalarda, AU, Mul, AV, Avg, AwM, BC, CR, LM, SC ve MP yöntemleri, temel birleştirme teknikleri olarak ele alınmıştır. Ayrıca, önerilen nihai AwU tekniğinin performansını daha kapsamlı bir şekilde değerlendirmek için karşılaştırma metotları olarak, IBGR (Barzegar Nozari & Koohi, 2020) ve UL (Seo vd., 2018) olmak üzere iki gelişmiş birleştirme yaklaşımı seçilmiştir.

Spesifik olarak UL tekniği, Bölüm 2.4.7’de ayrıntılı olarak açıklandığı üzere, birleştirme aşamasında grup üyeleri tarafından sağlanan derecelendirmelerin dağılımını dikkate alan gelişmiş bir birleştirme yöntemidir. Bu amaçla, bir ürün için sağlanan

derecelendirmelerin sapmalarını hesaplar ve ardından bunları, ilgili ürün için nihai grup derecelendirmesini belirlemek üzere, Avg ve AV teknikleri ile hesaplanan grup derecelendirmeleriyle harmanlar.

Öte yandan IBGR, sosyal ilişkilerin ve grup üyelerinin etkisinin birleştirme sürecinde önemli unsurlar olarak kabul edildiği bir grup öneri yöntemi olarak geliştirilmiştir. Bu yöntem, ilk olarak, üyeler arasındaki benzerliklerin ve güvenin derecesine bağlı olarak grup üyelerinin birbiri üzerindeki etkisini hesaplar. Ardından bunları, diğer üyelerden daha fazla güvenilen liderleri belirlemek için kullanır. Son olarak, bireysel derecelendirmeleri liderlerin diğer üyeler üzerindeki etkisiyle ağırlıklandırarak nihai grup derecelendirmelerini hesaplar.

7.3.2. Deney metodolojisi

Bu çalışmada önerilen birleştirme tekniklerinin performansını değerlendirmek için yapılan deneylerde, Bölüm 2.5’de ayrıntılı olarak açıklanan MLP, MLM ve NF veri kümeleri kullanılmıştır. Ayrıca, deneyler, 5-kat çapraz doğrulama deney metodolojisi kullanılarak gerçekleştirilmiştir. Bu amaçla, ürünler kümesi rastgele biçimde beş alt kümeye bölünerek, oluşturulan her alt kümenin, ürünlerin %20’ünü içermesi sağlanmıştır. Bununla birlikte, değerlendirme ölçütleri olarak yine Bölüm 2.5’de ayrıntılı olarak açıklanan nDCG, Fairness ve GSM ölçütleri kullanılmıştır. Bu ölçütleri hesaplarken, önceki bölümlere benzer şekilde, Bölüm 2.2’de ayrıntılı olarak açıklanan CF algoritması yardımıyla tahmin edilen bireysel derecelendirmeler kullanılmıştır.

Test ve eğitim kümeleri oluşturulduktan sonra, benzer kullanıcılardan oluşan gruplar, eğitim veri kümesi üzerinde standart KM algoritması uygulanarak tanımlanmıştır. Ayrıca, grup büyüklüğünün önerilen yaklaşımın performansı üzerindeki etkisini değerlendirmek için, deneylerde, 4 ile 128 arasında değişen çeşitli sayılardaki (k) gruplar kullanılmıştır. Gruplar tanımlandıktan sonra, her grup sayısı için grup derecelenmelerini hesaplamak üzere, bu bölümde önerilen birleştirme teknikleri (yani melezleştirilmiş ve AwU), temel teknikler (yani AU, Mul, AV, Avg, AwM, BC, CR, LM, SC ve MP) ve karşılaştırma metotları olarak seçilen IBGR ve UL teknikleri kullanılmıştır.

Deneylerde kullanılan temel birleştirme tekniklerinin parametre ayarları söz konusu olduğunda, AV tekniği ve bu tekniği kullanan UL ile önerilen melezleştirilmiş teknikler (AU_{AV} ve AV_{AU}) için eşik değeri 3 olarak seçilmiştir. Bunun nedeni, 5-yıldızlı bir

değerlendirme yelpazesinde olumlu derecelendirmelerin sezgisel olarak 4 ve 5'e karşılık gelmesidir (Bobadilla, Serradilla, & Bernal, 2010). Benzer şekilde, AwM tekniği için de eşik değeri 3 olarak seçilmiştir. Ayrıca, UL tekniği uygulanırken, bu tekniğin temel bileşenleri olan Avg, AV ve SS'in ağırlıkları, 0,4, 0,2 ve 0,4 olarak seçilmiştir. Bunun nedeni, ilgili ayarlamaların yüksek kalitede grup önerileri sağlamak için UL tekniğinde seçilebilecek en uygun ayarlardan birisi olmasıdır (Seo vd., 2018). Son olarak, farklı eşik değerlerinin AwU tekniği üzerindeki etkilerini görmek için, 0 ile 1 arasında değişen τ değerleri göz önünde bulundurulmuştur.

Grup önerileri üretilirken, yukarıda ifade edilen bir birleştirme tekniği ile hesaplanan grup derecelendirmeleri kullanılmıştır. Hesaplanan grup derecelendirmeleri azalan düzende sıralanarak, N değerinin 1, 3, 5 ve 10 olarak ayarlandığı en iyi- N ürünler grup önerileri olarak seçilmiştir. Son olarak, üretilen en iyi- N listelerinin grup açısından ne kadar kaliteli olduğunu değerlendirmek için, gruptaki her üye için n DCG, Fairness ve GSM değerleri hesaplanarak, ilgili ölçüt için ortalamaları alınmıştır.

7.3.3. Deney sonuçları

7.3.3.1. Melezleştirilmiş tekniklerin değerlendirilmesi

Bu bölümde, önerilen melezleştirilmiş birleştirme tekniklerinin (yani AU_{AV} ve AV_{AU}) grup derecelendirmelerini hesaplamadaki performansını incelemek için, grup sayısı ve öneri listesinin boyutu dâhil olmak üzere farklı parametrelerle geniş bir deney seti gerçekleştirilmiştir. Bununla birlikte, bu deneylerin n DCG sonuçlarını Tablo 7.1, 7.2 ve 7.3'de sunulduğu gibi sırasıyla MLP, MLM ve NF veri kümeleri için on temel birleştirme tekniği ile karşılaştırılmıştır.

Tablo 7.4. *MLP veri kümesi için elde edilen nDCG sonuçları*

en iyi- <i>N</i>	Birleştirme Tekniği	<i>Grup Sayısı (k)</i>					
		<i>Büyük Gruplar</i>		<i>Orta Gruplar</i>		<i>Küçük Gruplar</i>	
		4	8	16	32	64	128
1	AU	0,823	0,816	0,817	0,805	0,813	0,808
	Mul	0,766	0,765	0,773	0,764	0,781	0,779
	AV	0,822	0,814	0,811	0,813	0,814	0,803
	Avg	0,766	0,768	0,773	0,766	0,782	0,780
	AwM	0,766	0,768	0,773	0,766	0,782	0,780
	BC	0,826	0,811	0,808	0,809	0,816	0,814
	CR	0,791	0,792	0,754	0,758	0,745	0,723
	LM	0,766	0,765	0,773	0,765	0,782	0,779
	SC	0,820	0,804	0,799	0,781	0,775	0,763
	MP	0,853	0,757	0,768	0,762	0,780	0,779
	AV _{AU}	0,838 [†]	0,818	0,816	0,821 [†]	0,820 [†]	0,817
3	AU	0,812	0,800	0,799	0,803	0,797	0,795
	Mul	0,761	0,768	0,767	0,776	0,783	0,784
	AV	0,819	0,811	0,806	0,810	0,797	0,795
	Avg	0,763	0,773	0,768	0,778	0,787	0,786
	AwM	0,761	0,771	0,768	0,778	0,786	0,785
	BC	0,818	0,805	0,803	0,805	0,803	0,804
	CR	0,784	0,778	0,742	0,747	0,731	0,709
	LM	0,761	0,769	0,767	0,776	0,784	0,784
	SC	0,805	0,789	0,785	0,782	0,773	0,766
	MP	0,754	0,762	0,762	0,770	0,781	0,784
	AV _{AU}	0,821	0,815	0,811	0,813	0,806	0,803
5	AU	0,803	0,792	0,795	0,792	0,787	0,787
	Mul	0,760	0,768	0,764	0,773	0,777	0,776
	AV	0,809	0,805	0,797	0,795	0,785	0,782
	Avg	0,765	0,776	0,767	0,777	0,780	0,781
	AwM	0,762	0,773	0,766	0,776	0,778	0,778
	BC	0,808	0,800	0,798	0,796	0,792	0,792
	CR	0,773	0,770	0,742	0,735	0,718	0,691
	LM	0,760	0,769	0,765	0,774	0,778	0,777
	SC	0,795	0,782	0,779	0,774	0,764	0,757
	MP	0,752	0,761	0,761	0,768	0,773	0,777
	AV _{AU}	0,814	0,808	0,801	0,802	0,792	0,789
10	AU	0,792	0,787	0,788	0,795	0,782	0,774
	Mul	0,757	0,767	0,768	0,772	0,775	0,767
	AV	0,801	0,797	0,792	0,783	0,776	0,757
	Avg	0,775	0,780	0,777	0,781	0,784	0,774
	AwM	0,768	0,775	0,775	0,775	0,778	0,766
	BC	0,798	0,792	0,792	0,787	0,785	0,778
	CR	0,762	0,764	0,737	0,720	0,704	0,661
	LM	0,759	0,768	0,769	0,773	0,777	0,768
	SC	0,785	0,777	0,775	0,769	0,763	0,751
	MP	0,748	0,760	0,764	0,769	0,775	0,769
	AV _{AU}	0,804	0,799	0,795	0,787	0,782	0,763
	AU _{AV}	0,819 [†]	0,815 [†]	0,811 [†]	0,803 [†]	0,800 [†]	0,787 [†]

[†] En iyi performans gösteren temel tekniğe göre %95 güven aralığında istatistiksel olarak anlamlı

Tablo 7.5. MLM veri kümesi için elde edilen nDCG sonuçları

en iyi-N	Birleştirme Tekniği	Grup Sayısı (k)					
		Büyük Gruplar		Orta Gruplar		Küçük Gruplar	
		4	8	16	32	64	128
1	AU	0,836	0,837	0,824	0,825	0,824	0,821
	Mul	0,738	0,746	0,754	0,759	0,764	0,768
	AV	0,839	0,838	0,828	0,831	0,826	0,824
	Avg	0,745	0,746	0,754	0,759	0,764	0,768
	AwM	0,738	0,746	0,754	0,759	0,764	0,768
	BC	0,831	0,832	0,835	0,828	0,829	0,827
	CR	0,819	0,821	0,814	0,793	0,789	0,793
	LM	0,738	0,746	0,754	0,759	0,764	0,768
	SC	0,834	0,728	0,815	0,813	0,810	0,803
	MP	0,755	0,761	0,761	0,762	0,764	0,768
	AV _{AU}	0,839	0,839	0,831	0,832	0,830	0,828
	AU _{AV}	0,849 [†]	0,844 [†]	0,839 [†]	0,839 [†]	0,836 [†]	0,834 [†]
3	AU	0,835	0,817	0,822	0,814	0,810	0,811
	Mul	0,749	0,753	0,756	0,758	0,764	0,773
	AV	0,839	0,823	0,827	0,819	0,816	0,812
	Avg	0,750	0,757	0,757	0,759	0,764	0,773
	AwM	0,750	0,757	0,757	0,759	0,764	0,773
	BC	0,840	0,821	0,826	0,819	0,814	0,814
	CR	0,769	0,783	0,770	0,776	0,772	0,772
	LM	0,750	0,753	0,756	0,758	0,764	0,773
	SC	0,828	0,808	0,811	0,803	0,798	0,798
	MP	0,761	0,762	0,765	0,767	0,769	0,774
	AV _{AU}	0,840	0,826	0,828	0,822	0,818	0,817
	AU _{AV}	0,850 [†]	0,833 [†]	0,836 [†]	0,830 [†]	0,826 [†]	0,825 [†]
5	AU	0,830	0,813	0,816	0,810	0,805	0,806
	Mul	0,750	0,755	0,757	0,760	0,765	0,775
	AV	0,833	0,819	0,822	0,816	0,810	0,808
	Avg	0,756	0,760	0,759	0,761	0,766	0,775
	AwM	0,751	0,755	0,757	0,760	0,765	0,775
	BC	0,833	0,818	0,820	0,814	0,809	0,810
	CR	0,769	0,775	0,771	0,774	0,767	0,769
	LM	0,751	0,755	0,757	0,760	0,765	0,775
	SC	0,822	0,803	0,806	0,798	0,794	0,794
	MP	0,762	0,760	0,763	0,766	0,770	0,776
	AV _{AU}	0,836	0,822	0,824	0,818	0,813	0,813
	AU _{AV}	0,847 [†]	0,830 [†]	0,832 [†]	0,827 [†]	0,823 [†]	0,821 [†]
10	AU	0,821	0,819	0,812	0,806	0,805	0,803
	Mul	0,756	0,758	0,762	0,768	0,773	0,776
	AV	0,825	0,826	0,819	0,810	0,809	0,806
	Avg	0,770	0,766	0,765	0,770	0,774	0,777
	AwM	0,757	0,759	0,763	0,767	0,773	0,776
	BC	0,823	0,823	0,817	0,809	0,808	0,806
	CR	0,789	0,789	0,789	0,774	0,769	0,767
	LM	0,756	0,759	0,762	0,768	0,773	0,776
	SC	0,813	0,809	0,801	0,795	0,796	0,791
	MP	0,756	0,765	0,767	0,769	0,776	0,779
	AV _{AU}	0,828	0,828	0,821	0,813	0,812	0,810
	AU _{AV}	0,836 [†]	0,839 [†]	0,830 [†]	0,824 [†]	0,823 [†]	0,819 [†]

[†] En iyi performans gösteren temel tekniğe göre %95 güven aralığında istatistiksel olarak anlamlı

Tablo 7.6. *NF* veri kümesi için elde edilen *nDCG* sonuçları

en iyi- <i>N</i>	Birleştirme Tekniği	<i>Grup Sayısı (k)</i>					
		<i>Büyük Gruplar</i>		<i>Orta Gruplar</i>		<i>Küçük Gruplar</i>	
		4	8	16	32	64	128
1	AU	0,770	0,776	0,772	0,773	0,770	0,774
	Mul	0,741	0,734	0,740	0,744	0,753	0,755
	AV	0,784	0,785	0,782	0,782	0,779	0,781
	Avg	0,741	0,734	0,740	0,744	0,753	0,755
	AwM	0,741	0,734	0,740	0,744	0,753	0,755
	BC	0,782	0,787	0,780	0,785	0,779	0,780
	CR	0,764	0,764	0,772	0,767	0,767	0,760
	LM	0,741	0,734	0,740	0,744	0,753	0,755
	SC	0,736	0,737	0,745	0,749	0,751	0,750
	MP	0,736	0,730	0,726	0,736	0,742	0,746
	AV _{AU}	0,784	0,786	0,786	0,786	0,787	0,786
3	AU	0,762	0,761	0,765	0,764	0,766	0,769
	Mul	0,740	0,734	0,744	0,745	0,751	0,753
	AV	0,772	0,771	0,777	0,775	0,775	0,776
	Avg	0,740	0,734	0,744	0,745	0,751	0,753
	AwM	0,740	0,734	0,744	0,745	0,751	0,753
	BC	0,772	0,777	0,771	0,772	0,773	0,775
	CR	0,755	0,756	0,758	0,752	0,749	0,747
	LM	0,740	0,734	0,744	0,745	0,751	0,753
	SC	0,741	0,745	0,747	0,750	0,751	0,752
	MP	0,739	0,730	0,737	0,742	0,747	0,752
	AV _{AU}	0,777	0,777	0,778	0,778	0,779	0,779
5	AU	0,762	0,758	0,763	0,763	0,765	0,766
	Mul	0,741	0,737	0,745	0,746	0,752	0,751
	AV	0,773	0,771	0,773	0,773	0,774	0,774
	Avg	0,741	0,737	0,745	0,746	0,752	0,754
	AwM	0,741	0,737	0,745	0,746	0,752	0,754
	BC	0,772	0,774	0,769	0,771	0,772	0,772
	CR	0,759	0,755	0,755	0,748	0,748	0,745
	LM	0,741	0,737	0,745	0,746	0,752	0,754
	SC	0,745	0,744	0,750	0,749	0,751	0,752
	MP	0,739	0,731	0,741	0,743	0,750	0,754
	AV _{AU}	0,777	0,775	0,776	0,777	0,777	0,777
10	AU	0,761	0,760	0,761	0,762	0,765	0,765
	Mul	0,744	0,740	0,747	0,748	0,754	0,756
	AV	0,770	0,771	0,771	0,772	0,772	0,771
	Avg	0,744	0,740	0,747	0,748	0,754	0,756
	AwM	0,744	0,740	0,747	0,748	0,754	0,756
	BC	0,771	0,773	0,767	0,769	0,770	0,770
	CR	0,757	0,752	0,750	0,746	0,746	0,740
	LM	0,744	0,740	0,747	0,748	0,754	0,756
	SC	0,747	0,746	0,750	0,750	0,752	0,752
	MP	0,742	0,736	0,745	0,747	0,754	0,756
	AV _{AU}	0,772	0,775	0,774	0,775	0,775	0,774
	AU _{AV}	0,788 [†]	0,789 [†]	0,785 [†]	0,784 [†]	0,785 [†]	0,782 [†]

[†] En iyi performans gösteren temel tekniğe göre %95 güven aralığında istatistiksel olarak anlamlı

Tablo 7.1, 7.2 ve 7.3'den görülebileceği gibi, üç veri kümesi için yapılan deneylerden elde edilen $nDCG$ sonuçları, BC, AU ve AV tekniklerinin diğer temel tekniklere kıyasla daha iyi performans sergilediğini göstermiştir. Bu bulgu, AV ve AU tekniklerinin melez bir teknik oluştururken seçilecek temel teknikler için ideal seçenekler olduğunun bir göstergesidir. BC temel teknikler arasında yüksek $nDCG$ değerlerine ulaşan diğer bir teknik olmasına rağmen, bu teknik, içerdiği sıralama süreci nedeniyle hesaplama süresi açısından verimli olmayan bir tekniktir. Dahası, ürünlerin sıralamasını yansıtan grup derecelendirmeleri hesaplaması, bu tekniğin melez bir teknik içerisinde uygun bir şekilde kullanımını zorlaştırır.

Her üç veri kümesi için yapılan deneylerin sonuçları ayrıca AU, AV, BC, CR ve SC tekniklerinin grup büyüklüğü arttıkça daha iyi performans sergilediğini göstermektedir. Diğer temel teknikler ise nispeten küçük grupların varlığında en iyi performanslarını göstermişlerdir. Ayrıca, neredeyse bütün birleştirme tekniklerinin başarısı öneri listesinin boyutu (N) arttıkça düşmüştür. Bu bulgunun temel nedeni, $nDCG$ puanlarının hesaplanmasında kullanılacak olan kullanıcıların gerçek derecelendirmelerinin bir önceki bölümde açıklandığı gibi bir CF algoritması ile tahmin edilmesidir. Birleştirme teknikleri ile çok sayıda ürün önerildiğinde, karşılık gelen $nDCG$ puanını hesaplamak için tahmin edilen gerçek derecelendirmelere daha fazla güvenilmesi gerekir. Bu durumda, $nDCG$ değerleri yanıltıcı olabilir.

Üç veri kümesi için yapılan deneylerin sonuçları karşılaştırıldığında, neredeyse tüm birleştirme teknikleri MLM veri kümesinde daha iyi performans göstermiştir. Bunun temel nedeni, MLM veri kümesinde ürün başına düşen derecelendirme sayısının diğerlerinden fazla olmasıdır. Bu durum, kullanılan birleştirme tekniğinin başarısını arttıran bir unsurdur çünkü birleştirme işlemine ne kadar çok derecelendirme girerse elde edilen grup derecelendirmesi o denli doğru olur. Bu, geliştirilen melezleştirilmiş tekniklerin yanı sıra diğer birleştirme teknikleri için de geçerlidir.

Tablo 7.1, 7.2 ve 7.3'de sunulan sonuçlar, bu çalışmada önerilen melezleştirilmiş birleştirme tekniklerinin her iki veri kümesinde de genellikle tüm temel birleştirme tekniklerinden daha iyi performans sergilediğini göstermiştir. Bununla birlikte, elde edilen iyileştirmelerin %95 güven aralığında istatistiksel olarak anlamlı olup olmadığını gözlemlemek için çeşitli tek kuyruklu t -testler gerçekleştirilmiştir. En iyi performans gösteren birleştirme tekniği referans alındığında, her ne kadar AV_{AU} tekniği diğerlerinden

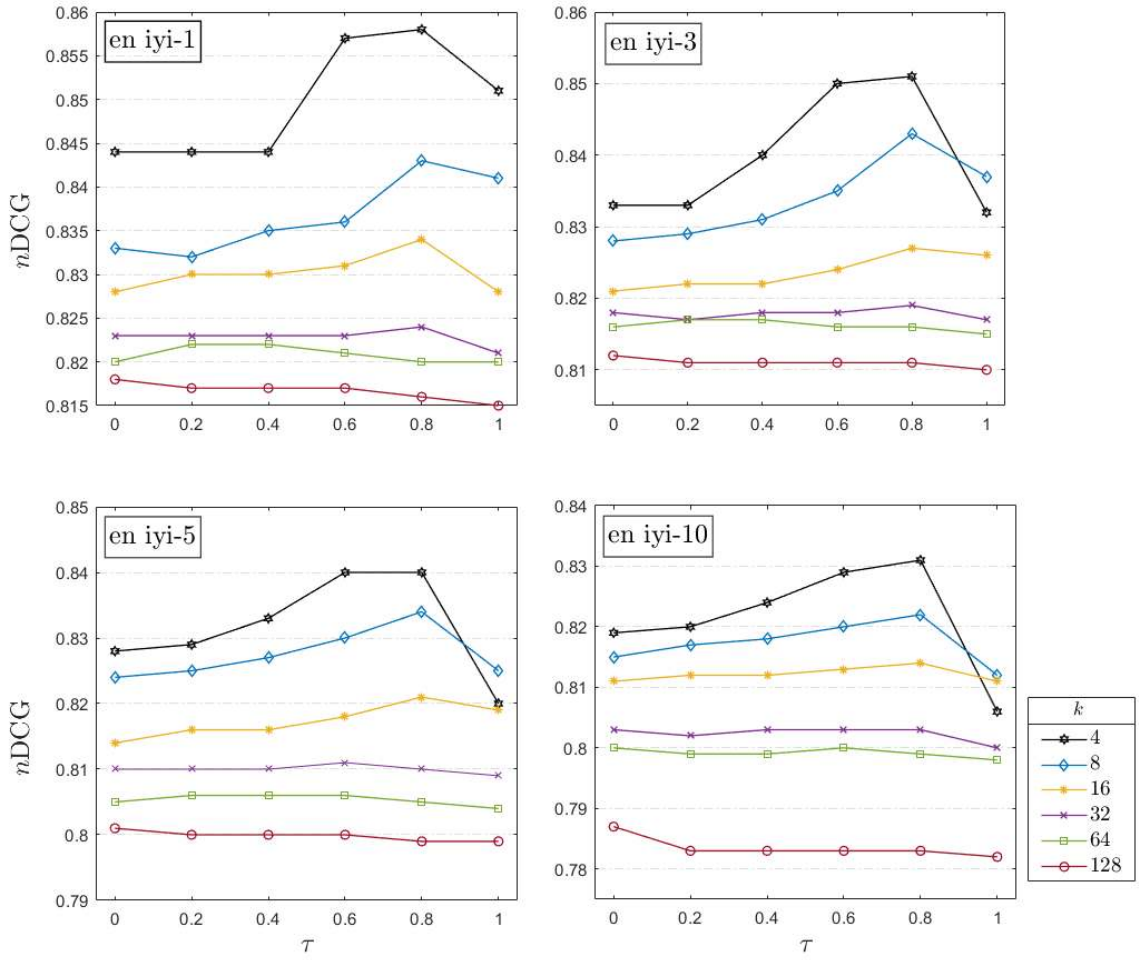
daha üstün olsa da, grup önerilerinin doğruluğunda sağladığı iyileştirmeler istatistiksel açıdan anlamlı bulunmamıştır. Diğer yandan, önerilen diğer melezleştirilmiş teknik olan AU_{AV} , hem AV_{AU} tekniğinden hem de en yüksek doğruluğa sahip temel teknikten önemli ölçüde daha iyi performans göstermiştir. Böylece, melez tipte bir birleştirme tekniğinde belirleyici faktör olarak AU tekniğinin kullanılmasının, AV tekniğine kıyasla yüksek kalitede grup önerileri üretmek açısından daha iyi bir seçim olduğu sonucuna varılmıştır.

7.3.3.2. AwU tekniğinin değerlendirilmesi

Önceki bölümde sonuçları sunulan deneylerde, bu çalışmada önerilen iki melezleştirilmiş birleştirme tekniği arasında AU_{AV} tekniğinin AV_{AU} tekniğinden daha başarılı doğruluk sonuçlarına ulaştığı gözlemlenmiştir. Ayrıca, AU_{AV} tekniğinin diğer temel tekniklerden de daha iyi performans gösterdiği görülmüştür. Bu bölümde, AU_{AV} tekniği kullanmadan önce entropi ölçümü yoluyla belirlenen ürünlerin elimine edilmesinin grup önerilerinin kalitesini artırıp arttırmadığı araştırılmaktadır. Bu amaçla yapılan deneysel çalışmalarda kullanılan AwU tekniği, AU_{AV} tekniğinin entropi ile bezenmiş halidir.

AwU birleştirme tekniğinin etkinliğini araştırmak için grup sayısı, önerilen ürün sayısı ve entropinin etkisini belirleyen eşik değeri (τ) gibi parametrelerin farklı ayarlamaları ile çeşitli deneyler yapılmıştır. Burada, (τ) değeri entropinin etkisi ile pozitif ilişkilidir. Spesifik olarak, (τ) sıfır olarak ayarlandığında entropinin AU_{AV} tekniğinin performansı üzerinde herhangi bir etkisi yoktur; çünkü bu durumda hiçbir ürün elenmez. Diğer yandan, (τ) değeri bir olarak ayarlandığında en az ürün sayısı ile ilgilenen AU_{AV} tekniğine neden olarak en fazla etkiye sahip olur.

Şekil 7.2, 7.3 ve 7.4’de sırasıyla MLP, MLM ve NF veri kümeleri için AwU tekniği ile elde edilen $nDCG$ sonuçları gösterilmektedir. Hemen hemen tüm ayarlamalardaki $nDCG$ değerlerindeki pozitif eğilime dayanarak, derecelendirme dağılımlarının entropisinin dikkate alınmasının AU_{AV} tekniğinin başarısını arttırdığı söylenebilir. Diğer bir deyişle, AwU tekniği genel olarak AU_{AV} tekniğinden daha iyi performans göstermiştir. Bu sonuç özellikle grup sayısının az olduğunda yani grupların büyük olduğu durumda daha belirgin hale gelmiştir.

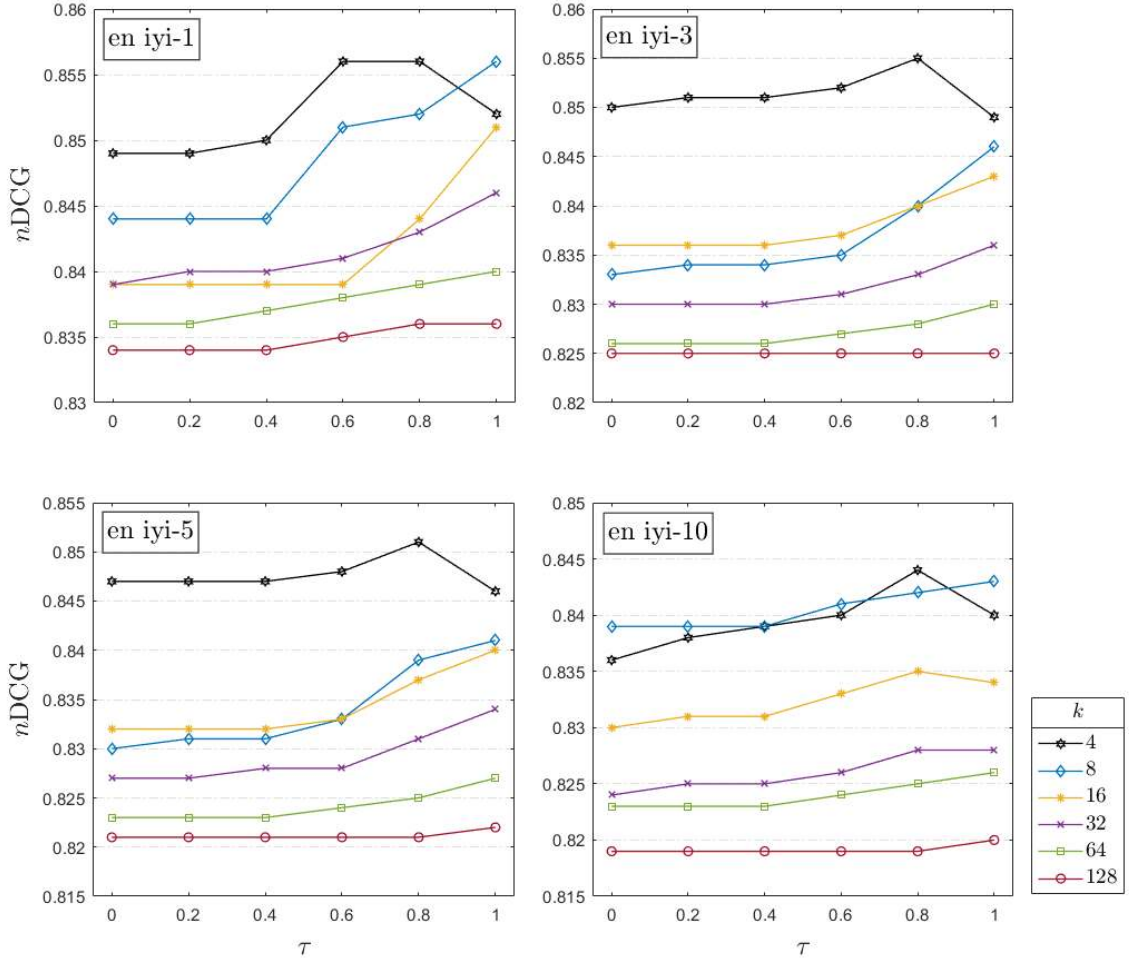


Şekil 7.2. MLP veri kümesinde AwU için yapılan deneylerin nDCG sonuçları

Büyük veri kümeleri için grup sayısının az olduğu, yani grupların nispeten kalabalık olduğu durumlar, bazı ürünler hakkında grup üyeleri arasında farklı görüşlerin oluşmasına neden olur. Bu bölümde elde edilen bulgular, bu türden ürünlerin sahip oldukları yüksek entropi değerleri ile belirlenerek birleştirme adımından önce göz ardı edilmesinin üretilen grup önerilerinin kalitesinin artırılmasına katkıda bulunduğunu göstermiştir.

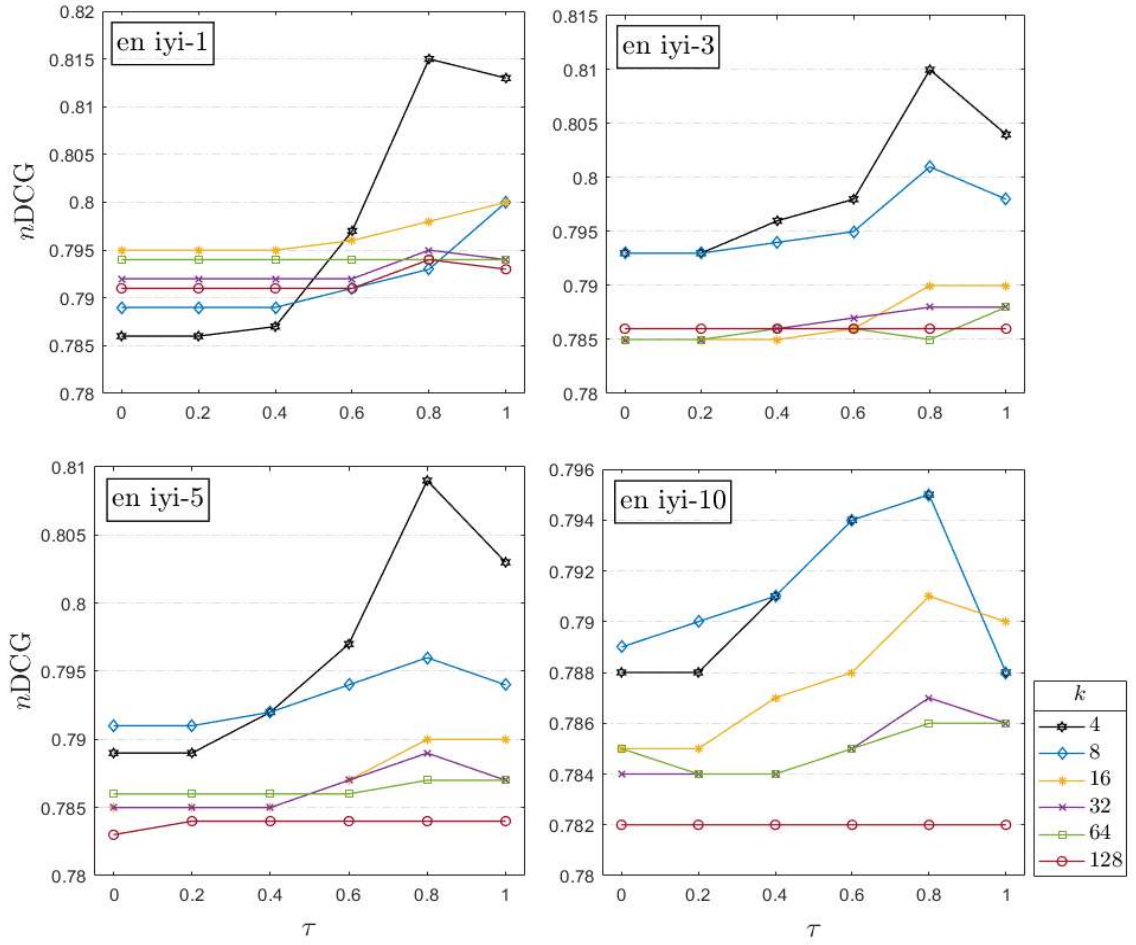
Her bir veri kümesi için yapılan deneylerden elde edilen sonuçlara göre, (τ) parametresinin ayarlanması gereken ideal değerin 0,8 ile 1 arasında olduğu söylenebilir. Bu sonuç da birleştirme aşamasında entropinin kullanılmasına ihtiyaç olduğunun bir kanıtı olarak görülebilir. Özellikle, MLP ve NF veri kümeleri için (τ) değeri 0,8 olarak ayarlandığında, AwU tekniği, grup sayısının 4, 8 ve 16 olduğu durumda AU_{AV} tekniğinden %95 güven aralığında daha başarılı grup önerilerinin üretilmesini sağlamıştır. Bu bulgu aynı zamanda farklı sayıda önerilen ürünleri de kapsamaktadır. Benzer şekilde,

MLM veri kümesi için yapılan deneylerde de AwU tekniği, AU_{AV} tekniğinden daha iyi sonuçlar elde etmiştir ve gözlemlenen tüm iyileştirmeler grup sayısının 128 olduğu ayarlama hariç tüm şemalarda %95 güven aralığında istatistiksel olarak anlamlıdır.



Şekil 7.3. MLM veri kümesinde AwU için yapılan deneylerin nDCG sonuçları

Ayrıca, AwU tekniğinin grup önerilerindeki doğruluk performansı üç veri kümesi için de karşılaştırma metotları olarak seçilen UL ve IBGR yöntemleri ile kıyaslanmıştır. Bu amaçla, MLP, MLM ve NF veri kümeleri için gerçekleştirilen ek deneylerin sonuçları sırasıyla Tablo 7.7, 7.8 ve 7.9’da sunulmaktadır. Bu bölümün başında gerçekleştirilen deneylerin sonuçlarında AwU tekniği için ideal (τ) değerinin 0,8 olması nedeniyle, bu deneylerde AwU tekniği için (τ) değeri 0,8 olarak seçilmiştir. Üç veri kümesinden elde edilen sonuçlara göre nihai AwU tekniğinin her iki karşılaştırma metodunu önemli ölçüde geride bıraktığı sonucuna varılabilir. Bu bulgu, önceki bölümde gerçekleştirilen deneylerin sonuçlarına paralel şekilde orta ve büyük grupların varlığında daha belirgindir.



Şekil 7.4. NF veri kümesinde AwU için yapılan deneylerin nDCG sonuçları

Tablo 7.7. AwU tekniğinin karşılaştırma yöntemleri ile MLP veri kümesinde nDCG kıyaslaması

en iyi-N	Birleştirme Tekniği	Grup Sayısı (k)					
		Büyük Gruplar	Orta Gruplar	Küçük Gruplar			
		4	8	16	32	64	128
1	UL	0,830	0,819	0,812	0,814	0,814	0,805
	IBGR	0,821	0,817	0,810	0,810	0,803	0,801
	AwU ($\tau=0.8$)	0,858 [†]	0,846 [†]	0,834 [†]	0,824 [†]	0,820 [†]	0,816 [†]
3	UL	0,822	0,814	0,808	0,811	0,803	0,799
	IBGR	0,818	0,810	0,808	0,807	0,802	0,795
	AwU ($\tau=0.8$)	0,851 [†]	0,843 [†]	0,827 [†]	0,819	0,815	0,810 [†]
5	UL	0,811	0,807	0,801	0,802	0,793	0,790
	IBGR	0,808	0,806	0,796	0,795	0,788	0,781
	AwU ($\tau=0.8$)	0,840 [†]	0,834 [†]	0,821 [†]	0,810 [†]	0,805 [†]	0,799
10	UL	0,802	0,799	0,797	0,794	0,790	0,778
	IBGR	0,799	0,789	0,786	0,788	0,779	0,771
	AwU ($\tau=0.8$)	0,831 [†]	0,822 [†]	0,814 [†]	0,803 [†]	0,799	0,783

[†] En iyi performans gösteren karşılaştırma yöntemine göre %95 güven aralığında istatistiksel olarak anlamlı

Tablo 7.8. *AwU tekniğinin karşılaştırma yöntemleri ile MLM veri kümesinde nDCG kıyaslaması*

en iyi-N	Birleştirme Tekniği	Grup Sayısı (k)					
		Büyük Gruplar		Orta Gruplar		Küçük Gruplar	
		4	8	16	32	64	128
1	UL	0,841	0,839	0,827	0,830	0,827	0,827
	IBGR	0,827	0,819	0,811	0,810	0,805	0,802
	AwU ($\tau=0.8$)	0,861 [†]	0,854 [†]	0,847 [†]	0,845 [†]	0,839 [†]	0,836
3	UL	0,840	0,824	0,827	0,820	0,816	0,814
	IBGR	0,829	0,826	0,816	0,816	0,805	0,805
	AwU ($\tau=0.8$)	0,855 [†]	0,844 [†]	0,840 [†]	0,833	0,828	0,825 [†]
5	UL	0,835	0,821	0,823	0,817	0,812	0,810
	IBGR	0,821	0,819	0,813	0,808	0,799	0,793
	AwU ($\tau=0.8$)	0,851 [†]	0,840 [†]	0,838 [†]	0,834 [†]	0,827 [†]	0,822
10	UL	0,828	0,829	0,821	0,814	0,812	0,809
	IBGR	0,821	0,819	0,809	0,807	0,796	0,792
	AwU ($\tau=0.8$)	0,844 [†]	0,843 [†]	0,835 [†]	0,829 [†]	0,826	0,820

[†] En iyi performans gösteren karşılaştırma yöntemine göre %95 güven aralığında istatistiksel olarak anlamlı

Tablo 7.9. *AwU tekniğinin karşılaştırma yöntemleri ile NF veri kümesinde nDCG kıyaslaması*

en iyi-N	Birleştirme Tekniği	Grup Sayısı (k)					
		Büyük Gruplar		Orta Gruplar		Küçük Gruplar	
		4	8	16	32	64	128
1	UL	0,784	0,785	0,785	0,782	0,782	0,782
	IBGR	0,781	0,781	0,779	0,775	0,776	0,774
	AwU ($\tau=0.8$)	0,815 [†]	0,798 [†]	0,794 [†]	0,794 [†]	0,793 [†]	0,792
3	UL	0,778	0,778	0,776	0,776	0,773	0,772
	IBGR	0,773	0,772	0,770	0,768	0,764	0,769
	AwU ($\tau=0.8$)	0,811 [†]	0,801 [†]	0,790 [†]	0,789	0,788	0,787
5	UL	0,776	0,775	0,775	0,774	0,774	0,774
	IBGR	0,770	0,771	0,769	0,770	0,768	0,767
	AwU ($\tau=0.8$)	0,809 [†]	0,798 [†]	0,790 [†]	0,788 [†]	0,787	0,784
10	UL	0,775	0,774	0,774	0,773	0,773	0,771
	IBGR	0,773	0,773	0,771	0,770	0,770	0,768
	AwU ($\tau=0.8$)	0,804 [†]	0,795 [†]	0,791 [†]	0,787 [†]	0,786 [†]	0,782

[†] En iyi performans gösteren karşılaştırma yöntemine göre %95 güven aralığında istatistiksel olarak anlamlı

7.3.3.3. Adalet ve memnuniyet analizi

Bu bölümde, adalet ve memnuniyet açısından nihai AwU tekniğinin performansının kapsamlı bir analizini sağlamak için çeşitli ek deneyler gerçekleştirilmiştir. Bu deneylerde, temel birleştirme teknikleri olarak AU, AV ve BC ele alınmıştır; çünkü Bölüm 7.3.3.1’de sunulan sonuçlardan da görülebileceği gibi, temel birleştirme teknikleri arasında en iyi nDCG sonuçlarını bu teknikler elde etmiştir. Ayrıca bu deneylerde AwU tekniğinin adalet ve memnuniyet açısından performansı, karşılaştırma algoritmaları

olarak seçilen UL ve IBGR ile de kıyaslanmıştır. Bu deneylerde ayrıca bütün grup formasyonları, yani büyük, orta ve küçük gruplar, k değerinin sırasıyla 4, 16 ve 64 seçilmesi ile göz önünde bulundurulmuştur. Son olarak m değerinin Fairness ölçütü üzerindeki etkisini gözlemlemek için 1’den 5’e kadar değişen m değerleri kullanılmıştır.

Tablo 7.10, 7.11 ve 7.12, sırasıyla MLP, MLM ve NF veri kümelerinde birleştirme tekniklerinin en iyi-5 grup önerileri için Fairness sonuçlarını sunmaktadır. Elde edilen sonuçlara dayanarak, geliştirilen nihai AwU tekniğinin, hem en iyi n DCG performansı gösteren temel birleştirme tekniklerine (yani AU, AV ve BC) hem de karşılaştırma yöntemlerine (yani UL ve IBGR) nispeten daha adil grup önerileri sağladığı sonucuna varılabilir. Bu bulgu, Bölüm 7.3.3.1’de sunulan n DCG sonuçlarına paralel biçimde grupların büyük ve orta olduğu durumda daha belirgin hale gelmiştir. Ayrıca, elde edilen iyileştirmelerin istatistiksel açıdan anlamlı olup olmadığını analiz etmek için tek kuyruklu t -testleri gerçekleştirilmiştir. Özellikle gruplar büyük ve orta büyüklükte olduğunda AwU tekniğinin diğer yöntemlerden istatistiksel açıdan anlamlı bir şekilde daha iyi olduğu gözlemlenmiştir. Ampirik sonuçlar ayrıca tüm yöntemlerin Fairness skorlarının m değerinin artması ile şaşırtıcı olmayan bir biçimde azaldığını göstermiştir. Bunun nedeni, m değerinin artması ile birlikte, önerilen ürün listesinde bulunan ürünlerden en az m tanesinden memnun olan grup üyesi sayısının genellikle azalmasıdır.

Tablo 7.10. MLP veri kümesinde en iyi-5 grup önerileri için Fairness sonuçları

Grup	Birleştirme Tekniği	m				
		1	2	3	4	5
Büyük ($k=4$)	AU	0,793	0,625	0,448	0,278	0,111
	AV	0,830	0,645	0,473	0,303	0,144
	BC	0,817	0,643	0,462	0,309	0,142
	UL	0,828	0,646	0,473	0,302	0,149
	IBGR	0,821	0,651	0,468	0,315	0,152
	AwU ($\tau=0.8$)	0,849 [†]	0,706 [†]	0,525 [†]	0,352 [†]	0,176 [†]
Orta ($k=16$)	AU	0,798	0,604	0,410	0,248	0,104
	AV	0,823	0,616	0,432	0,277	0,126
	BC	0,810	0,615	0,432	0,262	0,122
	UL	0,822	0,622	0,432	0,278	0,126
	IBGR	0,819	0,619	0,439	0,272	0,130
	AwU ($\tau=0.8$)	0,840 [†]	0,633 [†]	0,459 [†]	0,296 [†]	0,142 [†]
Küçük ($k=64$)	AU	0,789	0,600	0,398	0,242	0,108
	AV	0,802	0,614	0,411	0,255	0,118
	BC	0,795	0,610	0,407	0,249	0,123
	UL	0,804	0,615	0,414	0,254	0,116
	IBGR	0,810	0,613	0,408	0,248	0,113
	AwU ($\tau=0.8$)	0,802	0,611	0,420	0,268 [†]	0,114

[†] En iyi performans gösteren tekniğe göre %95 güven aralığında istatistiksel olarak anlamlı

Tablo 7.11. MLM veri kümesinde en iyi-5 grup önerileri için Fairness sonuçları

Grup	Birleştirme Tekniği	<i>m</i>				
		1	2	3	4	5
Büyük (<i>k</i> =4)	AU	0,876	0,744	0,603	0,430	0,200
	AV	0,880	0,760	0,617	0,439	0,221
	BC	0,880	0,755	0,611	0,450	0,224
	UL	0,881	0,762	0,619	0,445	0,225
	IBGR	0,883	0,764	0,610	0,451	0,230
	AwU ($\tau=0.8$)	0,908 [†]	0,782 [†]	0,637 [†]	0,465 [†]	0,244 [†]
Orta (<i>k</i> =16)	AU	0,863	0,730	0,561	0,389	0,195
	AV	0,869	0,737	0,582	0,398	0,203
	BC	0,868	0,730	0,580	0,391	0,203
	UL	0,869	0,738	0,584	0,403	0,206
	IBGR	0,865	0,743	0,581	0,413	0,210
	AwU ($\tau=0.8$)	0,883 [†]	0,762 [†]	0,600 [†]	0,430 [†]	0,226 [†]
Küçük (<i>k</i> =64)	AU	0,855	0,703	0,541	0,364	0,179
	AV	0,865	0,715	0,547	0,377	0,190
	BC	0,859	0,711	0,544	0,376	0,191
	UL	0,866	0,722	0,548	0,378	0,191
	IBGR	0,871	0,730	0,550	0,371	0,188
	AwU ($\tau=0.8$)	0,874	0,734	0,572 [†]	0,390 [†]	0,204 [†]

[†] En iyi performans gösteren tekniğe göre %95 güven aralığında istatistiksel olarak anlamlı

Tablo 7.12. NF veri kümesinde en iyi-5 grup önerileri için Fairness sonuçları

Grup	Birleştirme Tekniği	<i>m</i>				
		1	2	3	4	5
Büyük (<i>k</i> =4)	AU	0,753	0,488	0,312	0,190	0,083
	AV	0,756	0,498	0,325	0,204	0,092
	BC	0,744	0,477	0,299	0,187	0,092
	UL	0,755	0,499	0,325	0,204	0,094
	IBGR	0,758	0,501	0,331	0,214	0,091
	AwU ($\tau=0.8$)	0,774	0,524	0,344	0,234	0,117
Orta (<i>k</i> =16)	AU	0,745	0,468	0,302	0,192	0,089
	AV	0,751	0,484	0,315	0,196	0,094
	BC	0,744	0,467	0,300	0,196	0,094
	UL	0,750	0,486	0,318	0,201	0,095
	IBGR	0,754	0,484	0,322	0,204	0,094
	AwU ($\tau=0.8$)	0,779	0,507	0,341	0,226	0,096
Küçük (<i>k</i> =64)	AU	0,736	0,470	0,297	0,195	0,090
	AV	0,742	0,477	0,306	0,200	0,097
	BC	0,732	0,466	0,298	0,197	0,094
	UL	0,743	0,479	0,308	0,204	0,098
	IBGR	0,746	0,473	0,312	0,215	0,097
	AwU ($\tau=0.8$)	0,743	0,475	0,303	0,199	0,096

[†] En iyi performans gösteren tekniğe göre %95 güven aralığında istatistiksel olarak anlamlı

Ayrıca AwU tekniğinin GSM sonuçları, Tablo 7.13’de sunulduğu gibi bütün veri kümeleri için hem AU, AV ve BC olmak üzere üç temel teknikle hem de IBGR ve UL olmak üzere iki karşılaştırma yöntemiyle karşılaştırılmıştır. Elde edilen GSM sonuçlarına

göre, AwU tekniğinin yine büyük ve orta büyüklükteki gruplarda, grup üyelerinin üretilen grup önerilerinden memnuniyetini önemli ölçüde arttırdığı sonucuna varılabilir.

Tablo 7.13. *En iyi-5 grup önerilerinin GSM sonuçları*

Veri Kümesi	Birleştirme Tekniği	Grup		
		Büyük ($k=4$)	Orta ($k=16$)	Küçük ($k=64$)
MLP	AU	0,692	0,689	0,674
	AV	0,720	0,714	0,688
	BC	0,717	0,702	0,684
	UL	0,725	0,717	0,695
	IBGR	0,726	0,716	0,690
	AwU ($\tau=0.8$)	0,767 [†]	0,723 [†]	0,698
MLM	AU	0,773	0,788	0,766
	AV	0,790	0,806	0,778
	BC	0,785	0,798	0,773
	UL	0,793	0,810	0,780
	IBGR	0,792	0,804	0,778
	AwU ($\tau=0.8$)	0,826 [†]	0,823 [†]	0,785
NF	AU	0,680	0,677	0,668
	AV	0,708	0,701	0,669
	BC	0,712	0,709	0,674
	UL	0,715	0,708	0,684
	IBGR	0,718	0,711	0,683
	AwU ($\tau=0.8$)	0,754 [†]	0,718 [†]	0,674

[†] En iyi performans gösteren tekniğe göre %95 güven aralığında istatistiksel olarak anlamlı

7.3.4. Tartışma

Bu çalışmada, tüm grubu olabildiğince memnun edebilecek önerilen ürünlerin bir listesini üretmek için gruptaki kullanıcıların eğilimlerini doğru biçimde analiz edebilen birleştirme teknikleri geliştirilmeye odaklanılmıştır. Bu amaçla, öncelikle gruptaki bireylerin tercihlerini, iki farklı temel birleştirme tekniğini farklı şekillerde kullanan iki yeni melez teknik (yani AU_{AV} ve AV_{AU}) geliştirilmiştir. Ardından, grup üyelerinin üzerinde uzlaşma sağlayamadığı ürünlerin entropi yoluyla belirlendiği ve geliştirilen melez tekniklerin üzerine inşa edilen AwU tekniği geliştirilmiştir.

Deneylerde, geliştirilen melez teknikler olan AU_{AV} ve AV_{AU} ile literatürde öne çıkan modern UL tekniği, beklenildiği gibi diğer tüm temel tekniklerden daha yüksek $nDCG$ değerleri elde etmiştir. Bunun nedeni, birden fazla birleştirme tekniğinin birlikte kullanılmasının, temel tekniklerin eksikliklerinin giderilmesine yardımcı olması ve bireylerin derecelendirmelerini farklı açılardan birleştirmesidir. Daha spesifik olarak, AU_{AV} tekniği geliştirilen diğer melez teknik AV_{AU} 'dan ve melez bir teknik olan UL tekniğine kıyasla daha başarılı sonuçlar elde etmiştir. Ayrıca, AU_{AV} tekniğinde baskın

olan taban teknik AU tekniğidir ve dolayısıyla grup üyelerinin üzerinde fikir birliği sağladığı ürünlerin, popüler ürünlere kıyasla daha değerli olduğunu söylenebilir.

Üç farklı veri kümesi için yapılan deneylerin sonuçlarına göre, özellikle gruplar büyük olduğunda geliştirilen AwU tekniğinin en başarılı birleştirme tekniği olduğu gözlemlenmiştir. AwU'nun derecelendirmelerin dağılımını entropi yoluyla analiz eden tek teknik olduğu göz önünde bulundurulduğunda, bu sonuç, derecelendirme dağılımının analiz edilmesinin gruplar kalabalıklaştıkça GÖS'ün birleştirme aşamasının kritik bir unsuru haline geldiğini doğrular. Ayrıca, AwU tekniğinde entropinin etkisi sıfır ile bir aralığında değişen (τ) parametresi ile pozitif korelasyon gösterir. Elde edilen sonuçlar (τ) değerinin 0,8 ile 1 arasında bir değere ayarlandığında AwU tekniğinin en iyi performansına ulaştığını göstermiştir. Bu sonuç entropi hesaplamasının doğru grup derecelendirmelerinin hesaplanmasında çok önemli bir rol oynadığı iddiasını desteklemektedir.

Özet olarak deneylerden elde edilen sonuçlar, önerilen melez teknikler olan AU_{AV} ile AV_{AU} ve entropi ile güçlendirilmiş varyantı olan AwU tekniklerinin, grup üyelerinin büyük çoğunluğunu memnun edebilen yüksek kalitede grup derecelendirmelerin hesaplanmasını sağladığını göstermiştir. Bu tekniklerin başarıları gerek önerilecek ürün listelerinin boyutundan gerekse grupların büyüklüğünden bağımsızdır. Bu da önerilen tekniklerin başarısını ortaya koymaktadır.

7.4. Değerlendirme

GÖS'de, grup önerileri hesaplanan grup derecelendirmelerine dayanarak üretildiği için, bu derecelendirmelerin gruptaki bireylerin tercihlerinin doğru bir şekilde birleştirilerek hesaplanması çok önemlidir. Bu tür birleştirme süreçleri beklenenden çok daha karmaşık olabilir. Bunun nedeni, GÖS'de ilgilenilen grupların genellikle mükemmel olmayışıdır. Diğer bir deyişle, grupların farklı eğilimleri ve zevkleri olan kullanıcılardan oluşmasıdır. Dolayısıyla, bireysel derecelendirmelerin birleştirilmesi için standart yaklaşımlara alternatif daha gelişmiş mekanizmaların geliştirilmesi gerekmektedir.

Bu bölümde, derecelendirmeleri çoklu boyutta birleştirebilmek için birden fazla birleştirme tekniğinin birlikte kullanılması önerilmiştir. Özellikle grup üyelerinin üzerinde fikir birliği sağladığı ve daha fazla derecelendirme sağlanan ürünlerin belirlenerek grup önerisi kalitesini arttırmak için AU ve AV yöntemleri bir arada

kullanılmıştır. Bu amaçla, melez tipte AU_{AV} ve AV_{AU} olarak adlandırılan iki farklı birleştirme tekniği geliştirilmiştir. Bu tekniklerden AU_{AV} nihai grup derecelendirmelerini hesaplarken belirleyici faktör olarak AU tekniğini kullanırken, AV_{AU} tekniğinde AV tekniği daha etkindir.

Farklı boyutlardaki üç farklı veri kümesi üzerinde yapılan deneyler, hem AU_{AV} hem de AV_{AU} tekniklerinin grup öneri kalitesi açısından tüm temel birleştirme tekniklerinden daha iyi sonuçlar elde ettiğini göstermiştir. Bu bulgu, istatistiksel analizlerle kanıtlandığı gibi, grupların büyüklüğü ve öneri listesinin boyutuna bakılmaksızın elde edilmiştir. Önerilen melez teknikler kendi içerisinde karşılaştırıldığında ise, AU_{AV} tekniğinin AV_{AU} tekniğinden daha iyi performans gösterdiği gözlemlenmiştir. Bu sonuç, GÖS'ün başarısı açısından grup üyelerinin üzerinde fikir birliği sağlanan ürünlerin popüler olan ve üyelerin çoğunluğu tarafından derecelendirilen ürünlerden daha değerli olduğunu gösterir.

Ayrıca sağlanan fikir birliğinin derecesini ölçmek için ürünler için sağlanan derecelendirmelerin dağılımlarının kullanımı önerilmektedir. Bu amaçla entropinin kullanılması önerilmiştir. Entropi, derecelendirme dağılımlarını analiz etmek için kullanıldığında SS'ye göre iki temel avantaj sağlamaktadır: (i) bir derecelendirme dağılımında birden fazla doruk olması durumunda daha uygundur ve (ii) ayrık tipteki derecelendirmeler için daha uygundur.

Geliştirilen AU_{AV} tekniği, entropi hesaplamasıyla yüksek entropi değerleri olan ve grup üyeleri arasında üzerinde bir uzlaşma sağlanamayan ürünleri göz ardı edecek şekilde güçlendirilmiş ve bu teknik AwU tekniği olarak adlandırılmıştır. Deneysel çalışmalarda elde edilen $nDCG$ sonuçları, AwU tekniğinin AU_{AV} tekniğinden daha iyi olduğunu göstermiştir. Bu bulgu, temel birleştirme teknikleriyle ve bu çalışmada geliştirilen melez tekniklerle karşılaştırıldığında, AwU tarafından sağlanan grup önerilerinin grup üyeleri açısından daha uygun olduğunu göstermiştir. Ayrıca, deneysel sonuçlar grupların boyutunun büyümesiyle AwU tekniğinin diğer teknikler üzerindeki başarısının daha da arttığını göstermiştir. Sonuç olarak, bireysel derecelendirmelerin birleştirilmesi sürecinde, entropi kullanarak derecelendirme dağılımlarının dikkate alınması, gruplar kalabalıklaştıkça yüksek kalitede grup önerileri sağlamaya katkıda bulunur

8. SONUÇLAR

Bu tez çalışmasında, kullanıcıların otomatik olarak tanımlandığı GÖS’de üretilen grup önerilerinin kalitesini arttırmak için kullanıcı gruplandırma yaklaşımlarının geliştirilmesine odaklanılmıştır. Bu bağlamda, uygulanan kümeleme yaklaşımlarında ortaya çıkan ve kullanıcıların mevcut ürünlerin genelde küçük bir yüzdesine tercih belirtmelerinden dolayı ortaya çıkan seyreklik sorunu için yeni çözümlere odaklanılmıştır. Bunun yanında, artan içerik sayısı ile kümeleme işleminin temel adımlarından birisi olan benzerlik hesaplamasında karşılaşılan karmaşıklık probleminin olumsuz etkilerini hafifletmeye dair çözümler üretilmeye çalışılmıştır. Ayrıca, tanımlanan grupların homojenliği ve kimi uygulama senaryolarında karşılaşılabilecek maksimum grup büyüklüğü ile ilgili kısıtlar dikkate alınarak çözümler geliştirmenin yolları araştırılmıştır.

Yukarıda bahsedilen hedefler doğrultusunda, tez çalışmasının ilk ve en önemli katkılarından birisi önceden tanımlanan maksimum büyüklükteki kullanıcı gruplarının tanımlanmasını sağlayan BKM temelli yaklaşımdır. Spesifik olarak, bu yaklaşım, kullanıcıların derecelendirme vektörleri üzerinde BKM algoritmasını uygulayarak bir ikili karar ağacı oluşturur ve bu ağacı hem kullanıcı gruplarını tanımlamak hem de sisteme sonradan katılan bir kullanıcının ait olabileceği grubun bulmak üzere kullanır. Yapılan deneysel çalışmalar, geliştirilen bu yaklaşımın, literatürde bu amaçla en çok kullanılan standart KM algoritmasına kıyasla daha uygun grupların tanımlanmasını sağladığını göstermiştir. Bununla birlikte, kullanıcı profillerini oluşturmadaki seyreklik ve karmaşıklık sorunlarının olumsuz etkilerini hafifletmek amacıyla, kullanıcı tercihlerinin ürünlerin janrları temelinde temsil edildiği iki farklı JTP oluşturma stratejisinin kullanımı önerilmiştir. Bu sayede, BKM temelli yaklaşımın kullanıcıları gruplandırma bağlamındaki başarısı arttırılmıştır. Ayrıca, kullanıcılar arasındaki ilişkileri daha doğru belirleyen ve benzerlik hesaplama sürecinde demografik bilgileri de dikkate alan iki farklı strateji geliştirilmiştir. Bu stratejiler, JTP-tabanlı benzerlikler ile demografik korelasyonları aynı anda dikkate alarak grup homojenliğinin sağlanmasına katkıda bulunmuş ve böylece en uygun grupların tanımlanmasını sağlamıştır. Son olarak, farklı parametreler ile gerçekleştirilen kapsamlı analizler ve deneysel çalışmalar, önerilen yaklaşımların grupların otomatik olarak tanımlanmasında etkin bir biçimde kullanılabileceğini göstermiştir.

Bu tez çalışmasında, ayrıca, grup önerilerinin üretilmesi amacıyla grup üyelerinin bireysel tercihlerini birleştirerek grup profillerinin oluşturulmasını sağlayan birleştirme teknikleri ele alınmıştır. Mevcut teknikler, yapılan deneysel çalışmalar ile derinlemesine analiz edilmiş ve üç farklı birleştirme mekanizması önerilmiştir. Bu bağlamda ilk olarak gruptaki bireylerin tercihlerine ek olarak onların davranışsal özelliklerini de birleştirme sürecinde dikkate alan PwAvg tekniği geliştirilmiştir. Bu teknik, beş temel kişilik özelliğe (*açıklık, uyumluluk, duygusal istikrar, vicdan ve dışadönüklük*) göre gruptaki bireylerin grubun nihai kararı üzerindeki etkisini ölçmek ve bunları daha sonra birleştirme aşamasında derecelendirmeleri ağırlıklandırmak için kullanılmaktadır. Yapılan deneyler, PwAvg tekniğinin kıyaslama tekniği olarak seçilen üç farklı birleştirme tekniğinden daha iyi performans sergilediğini göstermiştir.

Bununla birlikte, grup üyelerinin sağladığı derecelendirmelerin farklı yönlerini dikkate alarak daha başarılı grup derecelendirmelerinin hesaplanmasını sağlayan ve AU_{AV} ve AV_{AU} olarak adlandırılan iki farklı melez birleştirme tekniği geliştirilmiştir. Bu teknikler, grup üyelerinin üzerinde fikir birliği sağladığı ve grup içerisinde popüler olan ürünleri birleştirme aşamasında öne çıkarmaya odaklanmıştır. Gerçekleştirilen deneysel çalışmalar, özellikle AU_{AV} tekniğinin grup öneri kalitesi açısından tüm temel birleştirme tekniklerinden daha iyi sonuçlar elde ettiğini göstermiştir.

Bu tez çalışmasının son ve en önemli katkılarından birisi, geliştirilen melez tekniklerin üzerine inşa edilen ve AwU olarak adlandırılan birleştirme tekniğidir. Spesifik olarak bu teknik, grup üyeleri tarafından sağlanan derecelendirmelerin dağılımını bilgi entropisi yardımıyla analiz eden bir tekniktir. Entropi hesaplaması ile AwU tekniği grup üyelerinin üzerinde uzlaşma sağlamadığı (entropisi yüksek) ürünleri birleştirme sürecinde göz ardı ederek yüksek kalitede grup önerilerinin üretilmesini sağlamaktadır. Deneysel çalışmalarda elde edilen sonuçlar, grup derecelendirmeleri hesaplanırken, AwU tekniğinin kullanılmasının, geliştirilen melez tekniklerden (AU_{AV} ve AV_{AU}) ve diğer tüm temel birleştirme tekniklerinden daha yüksek doğrulukta grup önerilerinin üretilmesini sağladığı göstermiştir.

KAYNAKÇA

- Adomavicius, G., & Tuzhilin, A. (2005). Toward the next generation of recommender systems: A survey of the state-of-the-art and possible extensions. *IEEE transactions on knowledge and data engineering*, 17(6), 734-749.
- Agarwal, A., Chakraborty, M., & Chowdary, C. R. (2017). Does order matter? Effect of order in group recommendation. *Expert Systems with Applications*, 82, 115-127.
- Ahmad, H. S., Nurjanah, D., & Rismala, R. (2017). A combination of individual model on memory-based group recommender system to the books domain. *2017 5th International Conference on Information and Communication Technology (ICoIC7)*, (s. 1-6).
- Amatriain, X., Jaimes, A., Oliver, N., & Pujol, J. M. (2011). Data mining methods for recommender systems. *Recommender systems handbook* (s. 39-71). içinde Springer.
- Amer-Yahia, S., Roy, S., Chawlat, A., Das, G., & Yu, C. (2009). Group Recommendation: Semantics and Efficiency. *Proc. VLDB Endow.*, 2(1), 754-765.
- Amirali, S.-A., & Boutilier, C. (2015). Preference-Oriented Social Networks: Group Recommendation and Inference. *Proceedings of the 9th ACM Conference on Recommender Systems* (s. 35-42). Vienna, Austria: Association for Computing Machinery.
- Ardissono, L., Goy, A., Petrone, G., Segnan, M., & Torasso, P. (2003). Intrigue: personalized recommendation of tourist attractions for desktop and hand held devices. *Applied artificial intelligence*, 17, 687-714.
- Baltrunas, L., Makcinskas, T., & Ricci, F. (2010). Group recommendations with rank aggregation and collaborative filtering}. *Proceedings of the fourth ACM conference on Recommender systems*, (s. 119-126).
- Barzegar Nozari, R., & Koohi, H. (2020). A novel group recommender system based on members' influence and leader impact. *Knowledge-Based Systems*, 205, 106296.

- Bilge, A., & Polat, H. (2011). An improved profile-based CF scheme with privacy. *Proceedings - 5th IEEE International Conference on Semantic Computing, ICSC 2011*.
- Bilge, A., & Polat, H. (2013). A scalable privacy-preserving recommendation scheme via bisecting k-means clustering. *Information Processing & Management*, 49(4), 912-927.
- Blondel, V. D., Guillaume, J.-L., Lambiotte, R., & Lefebvre, E. (2008). Fast unfolding of communities in large networks. *Journal of Statistical Mechanics: Theory and Experiment*, 2008(10), P10008.
- Bobadilla, J., Serradilla, F., & Bernal, J. (2010). A new collaborative filtering metric that improves the behavior of recommender systems. *Knowledge-Based Systems*, 23(6), 520-528.
- Boratto, L., & Carta, S. (2010). State-of-the-art in group recommendation and new approaches for automatic identification of groups. *Information retrieval and mining in distributed environments* (s. 1-20). içinde Springer.
- Boratto, L., & Carta, S. (2014). Using Collaborative Filtering to Overcome the Curse of Dimensionality When Clustering Users in a Group Recommender System. *Proceedings of the 16th International Conference on Enterprise Information Systems - Volume 2* (s. 564-572). Lisbon, Portugal: SCITEPRESS - Science and Technology Publications, Lda.
- Boratto, L., & Carta, S. (2015a). Art: Group recommendation approaches for automatically detected groups. *International Journal of Machine Learning and Cybernetics*, 6(9), 953-980.
- Boratto, L., & Carta, S. (2015b). The rating prediction task in a group recommender system that automatically detects groups: architectures, algorithms, and performance evaluation. *Journal of Intelligent Information Systems*, 45(2), 221-245.
- Boratto, L., Carta, S., & Fenu, G. (2016). Discovery and representation of the preferences of automatically detected groups: Exploiting the link between group modeling and clustering. *Future Generation Computer Systems*, 64, 165-174.

- Boratto, L., Carta, S., & Fenu, G. (2017). Investigating the role of the rating prediction task in granularity-based group recommender systems and big data scenarios. *Information Sciences*, 378, 424-443.
- Boratto, L., Carta, S., & Satta, M. (2010). Groups Identification and Individual Recommendations in Group Recommendation Algorithms. *PRSAT@ recsys*, (s. 27-34).
- Boratto, L., Carta, S., Chessa, A., Agelli, M., & Clemente, M. L. (2009). Group recommendation with automatic identification of users communities. *2009 IEEE/WIC/ACM International Joint Conference on Web Intelligence and Intelligent Agent Technology*. 3, s. 547-550. IEEE.
- Cantador, I., & Castells, P. (2011). Extracting multilayered Communities of Interest from semantic user profiles: Application to group modeling and hybrid recommendations. *Computers in Human Behavior*, 27(4), 1321-1336.
- Chao, D. L., Balthrop, J., & Forrest, S. (2005). Adaptive Radio: Achieving Consensus Using Negative Preferences. *Proceedings of the 2005 International ACM SIGGROUP Conference on Supporting Group Work* (s. 120-123). Sanibel Island, Florida, USA: Association for Computing Machinery.
- Choi, K., & Suh, Y. (2013). A New Similarity Function for Selecting Neighbors for Each Target Item in Collaborative Filtering. *Knowledge-Based Systems*, 37, 146–153.
- Christensen, I. A., & Schiaffino, S. (2011). Entertainment recommender systems for group of users. *Expert Systems with Applications*, 38(11), 14127-14135.
- Costa Jr, P. T., & McCrae, R. R. (1992). Four ways five factors are basic. *Personality and individual differences*, 13(6), 653-665.
- Crossen, A., Budzik, J., & Hammond, K. J. (2002). Flytrap: intelligent group music recommendation. *Proceedings of the 7th international conference on Intelligent user interfaces*, (s. 184-185).
- Fatemi, M., & Tokarchuk, L. (2012). An Empirical Study on IMDb and Its Communities Based on the Network of Co-Reviewers. *Proceedings of the First Workshop on*

Measurement, Privacy, and Mobility. Bern, Switzerland: Association for Computing Machinery.

- Fatemi, M., & Tokarchuk, L. (2013). A Community Based Social Recommender System for Individuals & Groups. *Proceedings of the 2013 International Conference on Social Computing* (s. 351-356). Washington, DC, USA: IEEE Computer Society.
- Felfernig, A., Boratto, L., Stettinger, M., & Tkalčič, M. (2018). Evaluating Group Recommender Systems. *Group Recommender Systems : An Introduction* (s. 59-71). içinde Springer International Publishing.
- Fortunato, S., & Castellano, C. (2012). Community Structure in Graphs. *Computational Complexity: Theory, Techniques, and Applications* (s. 490-512). içinde New York, NY: Springer New York.
- Goren-Bar, D., & Glinansky, O. (2004). FIT-recommending TV programs to family members. *Computers & Graphics*, 28(2), 149-156.
- Gorla, J., Lathia, N., Robertson, S., & Wang, J. (2013). Probabilistic Group Recommendation via Information Matching. *Proceedings of the 22Nd International Conference on World Wide Web* (s. 495-504). Rio de Janeiro, Brazil: ACM.
- Harper, F. M., & Konstan, J. A. (2015). The MovieLens Datasets: History and Context. *ACM Trans. Interact. Intell. Syst.*, 5(4), 19.
- Herlocker, J. L., Konstan, J. A., Borchers, A., & Riedl, J. (2017). An Algorithmic Framework for Performing Collaborative Filtering. *SIGIR Forum*, 51(2), 227-234.
- Hu, R., & Pu, P. (2010). A study on user perception of personality-based recommender systems. *International conference on user modeling, adaptation, and personalization* (s. 291-302). Springer.
- Jameson, A. (2004). More than the sum of its members: challenges for group recommender systems. *Proceedings of the working conference on Advanced visual interfaces*, (s. 48-54).

- Jameson, A., & Smyth, B. (2007). Recommendation to groups. *The adaptive web* (s. 596-627). içinde Springer.
- Jeong, H. J., & Kim, M. H. (2019). HGGC: A hybrid group recommendation model considering group cohesion. *Expert Systems with Applications*, 136, 73-82.
- Kaleli, C. (2014). An entropy-based neighbor selection approach for collaborative filtering. *Knowledge-Based Systems*, 56, 273-280.
- Kassak, O., Kompan, M., & Bielikova, M. (2016). Personalized hybrid recommendation for group of users: Top-N multimedia recommender. *Information Processing & Management*, 52(3), 459-477.
- Khazaei, E., & Alimohammadi, A. (2018). An automatic user grouping model for a group recommender system in location-based social networks. *ISPRS International Journal of Geo-Information*, 7(2), 67.
- Kim, J., Kim, H., Oh, H., & Ryu, Y. U. (2010). A group recommendation system for online communities. *International journal of information management*, 30(3), 212--219.
- Koren, Y. (2008). Factorization meets the neighborhood: a multifaceted collaborative filtering model. *Proceedings of the 14th ACM SIGKDD international conference on Knowledge discovery and data mining*, (s. 426-434).
- Koren, Y. (2010). Factor in the neighbors: Scalable and accurate collaborative filtering. *ACM Transactions on Knowledge Discovery from Data (TKDD)*, 4(1), 1-24.
- Lancichinetti, A., Fortunato, S., & Kertesz, J. (2009). Detecting the overlapping and hierarchical community structure in complex networks. *New journal of physics*, 11(3), 033015.
- Lemire, D., & Maclachlan, A. (2005). Slope one predictors for online rating-based collaborative filtering. *Proceedings of the 2005 SIAM International Conference on Data Mining* (s. 471-475). SIAM.
- Lieberman, H., Van Dyke, N., & Vivacqua, A. (1999). Let's browse: a collaborative browsing agent. *Knowledge-Based Systems*, 12(8), 427-431.

- Liu, Y., Wang, B., Wu, B., Zeng, X., Shi, J., & Zhang, Y. (2016). Cogrec: A community-oriented group recommendation framework. *International Conference of Pioneering Computer Scientists, Engineers and Educators* (s. 258-271). Springer.
- Mahyar, H., Ghalebi, K. E., Morshedi, S. M., Khalili, S., Grosu, R., & Movaghar, A. (2017). Centrality-based group formation in group recommender systems. *International World Wide Web Conferences Steering Committee*, (s. 1187-1196).
- Marquez, J., & Ziegler, J. (2016). Hootle+: A group recommender system supporting preference negotiation. *CYTED-RITOS international workshop on groupware* (s. 151-166). Springer.
- Masthoff, J. (2011). Group recommender systems: Combining individual models. *Recommender systems handbook* (s. 677-702). içinde Springer.
- McCarthy, J. F. (2002). Pocket restaurantfinder: A situated recommender system for groups. *Workshop on Mobile Ad-Hoc Communication at the 2002 ACM Conference on Human Factors in Computer Systems*, 8.
- McCarthy, J. F., & Anagnost, T. D. (1998). MusicFX: an arbiter of group preferences for computer supported collaborative workouts. *Proceedings of the 1998 ACM conference on Computer supported cooperative work*, (s. 363-372).
- McCarthy, K., Salamo, M., Coyle, L., McGinty, L., Smyth, B., & Nixon, P. (2006). Cats: A synchronous approach to collaborative group recommendation. *Florida Artificial Intelligence Research Society Conference (FLAIRS)*, (s. 86-91).
- Mnih, A., & Salakhutdinov, R. R. (2008). Probabilistic matrix factorization. *Advances in neural information processing systems*, (s. 1257-1264).
- Newman, M., & Girvan, M. (2004). Finding and evaluating community structure in networks. *Physical Review E*, 69(2).
- Nguyen, T., Harper, F., Terveen, L., & Konstan, J. (2017). User Personality and User Satisfaction with Recommender Systems. *Information Systems Frontiers*, 20, 1-17.

- Ntoutsi, I., Stefanidis, K., Norvag, K., & Kriegel, H.-P. (2012). greco: A group recommendation system based on user clustering. *International Conference on Database Systems for Advanced Applications* (s. 299-303). Springer.
- O'Connor, M., Cosley, D., Konstan, J. A., & Riedl, J. (2001). PolyLens: a recommender system for groups of users. *ECSCW 2001* (s. 199-218). Springer.
- Quijano-Sanchez, L., Recio-Garcia, J. A., & Diaz-Agudo, B. (2010). Personality and social trust in group recommendations. *2010 22Nd IEEE international conference on tools with artificial intelligence*. 2, s. 121-126. IEEE.
- Quijano-Sanchez, L., Recio-Garcia, J. A., & Diaz-Agudo, B. (2011). Happymovie: A facebook application for recommending movies to groups. *2011 IEEE 23rd International Conference on Tools with Artificial Intelligence* (s. 239-244). IEEE.
- Recio-Garcia, J. A., Jimenez-Diaz, G., Sanchez-Ruiz, A. A., & Diaz-Agudo, B. (2009). Personality aware recommendations to groups. *Proceedings of the third ACM conference on Recommender systems*, (s. 325-328).
- Ricci, F., Rokach, L., & Shapira, B. (2011). Introduction to recommender systems handbook. *Recommender systems handbook* (s. 1-35). içinde Springer.
- Rossi, S., & Cervone, F. (2016). Social Utilities and Personality Traits for Group Recommendation: A Pilot User Study. *ICAART (1)*, (s. 38-46).
- Sacharidis, D. (2019). Top-N group recommendations with fairness. *Proceedings of the 34th ACM/SIGAPP Symposium on Applied Computing*, (s. 1663-1670).
- Sarwar, B., Karypis, G., Konstan, J., & Riedl, J. (2002). Recommender systems for large-scale e-commerce: Scalable neighborhood formation using clustering. *Proceedings of the fifth international conference on computer and information technology*, (s. 291-324).
- Schafer, J. B., Frankowski, D., Herlocker, J., & Sen, S. (2007). Collaborative Filtering Recommender Systems. *The Adaptive Web: Methods and Strategies of Web Personalization* (s. 291-324). içinde Springer-Verlag.

- Seo, Y.-D., Kim, Y.-G., Lee, E., Seol, K.-S., & Baik, D.-K. (2018). An enhanced aggregation method considering deviations for a group recommendation. *Expert Systems with Applications*, 93, 299-312.
- Serbos, D., Qi, S., Mamoulis, N., Pitoura, E., & Tsaparas, P. (2017). Fairness in Package-to-Group Recommendations. *Proceedings of the 26th International Conference on World Wide Web* (s. 371-379). Perth, Australia: International World Wide Web Conferences Steering Committee.
- Shi, J., Wu, B., & Lin, X. (2015). A Latent Group Model for Group Recommendation. *2015 IEEE International Conference on Mobile Services*, (s. 233-238).
- Sridevi, M., & Rao, R. R. (2017). DECORS: A Simple and Efficient Demographic Collaborative Recommender System for Movie Recommendation. *Advances in Computational Sciences and Technology*, 10(7), 1969-1979.
- Vozalis, M., & Margaritis, K. G. (2004). Collaborative filtering enhanced by demographic correlation. *AIAI symposium on professional practice in AI, of the 18th world computer congress*.
- Wang, Q., & Fleury, E. (2011). Uncovering overlapping community structure. *Complex networks* (s. 176-186). içinde Springer.
- Yalçın, E., Bilge, A., & Yüksek, A. G. (2019). An empirical evaluation of aggregation techniques used in group recommendation. *Proceedings of the 8th International Conference on Advanced Technologies (ICAT 2019)*, (s. 186-191).
- Yargıç, A., & Bilge, A. (2019). Privacy-preserving multi-criteria collaborative filtering. *Information Processing & Management*, 56(3), 994-1009.
- Zhao, C., Li, X., & Cang, Y. (2015). Bisecting k-means clustering based face recognition using block-based bag of words model. *Optik - International Journal for Light and Electron Optics*, 126(19), 1761-1766.
- Zhao, W.-L., Deng, C.-H., & Ngo, C.-W. (2018). k-means: A revisit. *Neurocomputing*, 291, 195-206.

Zhiwen, Y., Xingshe, Z., & Daqing, Z. (2005). An adaptive in-vehicle multimedia recommender for group users. *2005 IEEE 61st Vehicular technology conference* (s. 2800-2804). IEEE.



ÖZGEÇMİŞ

ORCID NO: 0000-0003-3818-6712

Adı Soyadı : Emre YALÇIN
Yabancı Dil : İngilizce
Doğum Yeri ve Yılı : Sivas / 03.09.1990
E-Posta : eyalcin@cumhuriyet.edu.tr

Eğitim Geçmişi:

- Doktora (2018-) : Eskişehir Teknik Üniversitesi, Fen Bil. Ens., Bilgisayar Mühendisliği A.B.D.
- Doktora (2016-2018) : Anadolu Üniversitesi, Fen Bil. Ens., Bilgisayar Mühendisliği A.B.D.
- Yüksek Lisans (2014-2016) : Anadolu Üniversitesi, Fen Bil. Ens., Bilgisayar Mühendisliği A.B.D.
- Lisans (2007-2012) : İstanbul Üniversitesi, Bilgisayar Mühendisliği Bölümü

Meslek Geçmişi:

- Araştırma Görevlisi (2018-) : Cumhuriyet Üniversitesi, Bilgisayar Mühendisliği Bölümü
- Araştırma Görevlisi (2014-2018) : Anadolu Üniversitesi, Bilgisayar Mühendisliği Bölümü
- Araştırma Görevlisi (2013-2014) : Cumhuriyet Üniversitesi, Bilgisayar Mühendisliği Bölümü
- Araştırma Görevlisi (2012-2013) : Bozok Üniversitesi, Bilgisayar Mühendisliği Bölümü