

Psikolojik Danışma ve Rehberlikte Konu Modellemesi

Nurten Karacan Ozdemir^{1*}



Kübra Atalay Kabasakal²



¹Hacettepe Üniversitesi, Rehberlik ve Psikolojik Danışmanlık Anabilim Dalı, Ankara, Türkiye
nurtenkaracan@hacettepe.edu.tr

²Hacettepe Üniversitesi, Eğitimde Ölçme ve Değerlendirme Anabilim Dalı, Ankara, Türkiye
katalay@hacettepe.edu.tr

Geliş tarihi:27.03.2024

Kabul tarihi:24.02.2025

Yayın tarihi:30.04.2025

Özet: Bu makalede, metin madenciliği yöntemlerinden biri olan konu modellemesi analizi tanıtılarak bunun psikolojik danışma ve rehberlik alanında kullanımı ele alınmıştır. Bu doğrultuda, bu derlemede (1) ruh sağlığı ile ilgili disiplinlerde (örn. danışmanlık, psikoloji, psikiyatri) trend olan konular, (2) belirli uzmanlık alanlarında (örn. kariyer danışmanlığı) veya konularda (örn. çok kültürlülük) öne çıkan temalar, (3) danışmanlık ve terapi süreciyle ilgili konular, (4) teşhis ve tanımlama ve (5) ölçme ve değerlendirme uygulamaları dahil olmak üzere alanın çeşitli yönlerini keşfetmek için konu modellemesi kullanan çalışmalar ele alınmıştır. Sonuç olarak, konu modelleme analizinin, psikolojik danışma ve rehberlik alanında ortaya çıkan eğilimleri, araştırma önceliklerini ve uygulamaya yönelik temaları belirlemede etkili bir yöntem olabileceği görülmüştür. Ancak, bu analiz yönteminde hem metodolojik sınırlılıkların (veri kalitesi, bağlam eksikliği ve yorum özneliği) hem de etik konuların (adil temsil, gizlilik ve şeffaflık) dikkate alınması önemli bir gereklilik olarak öne çıkmaktadır.

Anahtar Kelime: Psikolojik Danışma, Veri Madenciliği, Konu Modellemesi, Makine Öğrenmesi

GİRİŞ

Büyük nitel verilerin doğasında bulunan ayrıntılı bilgileri verimli bir şekilde kullanabilen potansiyel olarak umut verici veri analizi yöntemlerinden biri metin madenciliği yaklaşımıdır. Metin madenciliği, daha önce bilinmeyen ve ortaya çıkarılması kolay olmayan metin verilerinden bilgi üretme işlemi olarak tanımlanmaktadır (Buenano-Fernandez vd., 2020). Metin madenciliği, metinsel verilerden örüntüler, korelasyonlar ve içgörüler ortaya çıkarmaya yaramaktadır.

Teknolojinin, veri toplama sürecini hızlandırmasına ve daha büyük örneklem boyutlarından çok boyutlu bilgi toplanmasına olanak tanınmasına rağmen, büyük metin verilerini kullanan ve analiz eden çalışmaların eksikliği dikkat çekmektedir. Bu eksikliğin nedenlerinden biri, el ile nitel kodlamanın pahalı, zahmetli ve zaman alıcı olmasıdır (Feinerer & Wild, 2007). Bunun yanında mevcut nitel veri analizi araçlarının (ör., NVivo, MaxqDa) araştırmacıların büyük örneklemelerin metin verilerini kolayca kodlamasına olanak tanıyacak yeterli donanımına sahip olmaması da bir diğer eksiklik olarak gösterilebilir (Longo, 2020). Ayrıca, büyük metin verilerini analiz ederken her bir yorumu incelemek oldukça zorlayıcı olmaya devam etmektedir (Balahadia vd., 2016).

Son zamanlarda, bu sorunları aşmaya yardımcı olacak ve nitel veri analizi için gereken zaman ve işgücü talebini en aza indirecek şekilde, araştırmacıların bulgularında büyük ölçekli yazılı metin verilerini uygun maliyetli bir şekilde kullanmalarına olanak tanıyan makine öğrenmesine dayalı metin analitik teknikleri ortaya çıkmıştır. Bu yöntem, el ile kodlama kullanılarak analiz edilmesi mümkün olmayan büyük veri kümeleri ile başa çıkmanın bir yolu olarak geliştirilmiştir (Hopkins & King, 2010). Bu yöntemlerden biri konu modelleme analizidir.

Konu modelleme, metin belgesinin temel anlamsal yapısının belirlenmesine yönelik bir metin madenciliği araştırma alanıdır. Konu modelleme, büyük ölçekli belgelerin oluşturduğu yapılandırılmamış veri yapılarından anlamlı bilgiler çıkarmaya yarayan denetimsiz makine öğrenme algoritmalarının bir çerçevesidir (Kandula vd. 2011). Denetimli ve denetimsiz öğrenme yöntemleri makine öğrenmesindeki iki genel yaklaşımdır. Denetimsiz algoritmalar önceden herhangi bir bilgi sağlanmadan, verilerin keşfedilmesi yoluyla

kalıpları veya konuları belirleyebilir. Denetimsiz algoritmalar tematik analize benzer şekilde verilerden ortaya çıkan kelime kümelerini ve konuları belirler, ancak ortaya çıkan konuların yorumlanmasına yardımcı olmak için insan görüşü hala önemlidir. Konu modellemesi, tematik analizin (yani insan içgörüsü) ve makine öğreniminin (yani büyük miktarda metnin hızlı analizi) avantajlarından yararlanmaktadır (Roberts vd., 2014).

Konu modelleme yöntemleri kullanılarak, büyük miktarda, yapısal olmayan metin belgesi, otomatik olarak aranabilmekte, organize edilebilmekte ve özetlenebilmektedir. Konu modellemesi, belgeleri konu adı verilen çeşitli boyutlara indirgemesi bakımından faktör analizine benzemektedir. Konu modellemesinin farklı yöntemleri bulunmaktadır. Bunlardan biri olan yapısal konu modellemesi konu modellemesinde ortak değişkenlerin veya niteliklerin (örneğin, araştırmanın yapıldığı yıl, bölge, yazarın cinsiyeti, çalışma bağlamı vb.) dahil edilmesine izin vererek, bu ortak değişkenlerin belgedeki belirli konuların etkileyip etkilemediğini anlamak için bir regresyon çerçevesi kullanmaktadır (Roberts vd., 2014). Bu anlamda yapısal konu modellemesi, metnin ortak değişkenlere bağlı olarak nasıl değiştiğini hesaba katarak metinden elde edilen anlama daha fazla derinlik katmaktadır. Konu modellemesi ve yapısal konu modellemesinin geleneksel nitel metin analizine göre düşük yaygınlıktaki konuları yakalayamama gibi sınırlıkları bulunmaktadır (Baumer vd., 2017). Bu makalenin odağını veri madenciliği yöntemlerinden biri olan konu modellemesi analizi oluşturmaktadır. Bu doğrultuda konu modellemesi yöntemlerinden biri olan Gizil Dirichlet Tahsisi örnek olarak verilebilir.

Gizil Dirichlet Tahsisi

Konu modellemenin amacı, bir konuyu tanımlamak için sıkça birlikte geçen kelimeleri bir araya getirerek yazılı bir metindeki bir olgunun daha iyi anlaşılmasını sağlamaktır. Konu modellemesinin çeşitli varyantları mevcut olsa da orijinal konu modellerinden biri olan Gizil Dirichlet Tahsisi (GDT, Latent Dirichlet Allocation – LDA) ilk olarak Blei ve arkadaşları tarafından 2003 yılında literatüre tanıtılmıştır (Blei vd., 2003). Blei (2012) konu modellemesini orijinal metinlerdeki kelimeleri analiz ederek metinlerdeki temaları, bu temaların birbirleriyle nasıl bağlantılı olduklarını ve zaman içinde nasıl değiştiklerini keşfeden istatistiksel yöntemler olarak tanımlamaktadır.

GDT ile modellenecek olan belgeler kümesi külliyyat (corpus), külliyyat içindeki her bir öge belge (document) olarak adlandırılmaktadır. GDT yönteminde her belge için konu dağılımı ve her konu için de kelime dağılımı söz konusudur. Tamamen denetimsiz bir algoritma olan GDT herhangi bir ön bilgiye gerek kalmadan, kelimelerin belge içerisindeki bulunduğu yeri dikkate almadan kelimelerin birlikte bulunmasını inceleyen kelime torbası (bag of words) yaklaşımına dayalı olarak çalışmaktadır (Blei, 2012).

GDT bir konuya ait kelimelerin aynı belgede görülme olasılığının daha yüksek olduğu fikrine dayanan, bir dizi belgede gömülü bulunan gizil konuların belirlenmesinde kullanılan istatistiksel bir yöntemdir (Wang vd., 2018). Bir belgenin konusu doğal olarak o belgede geçen kelimelere ilişkilidir. Belgede geçen kelimeler gizil olarak dokümanın ana konu ya da konularını temsil eder. Bu kelimeler ilişkili oldukları konuların sözcük grubunu oluşturacaktır. Örneğin, psikoloji ile ilgili bir konuda duygular, bilişler, davranışlar gibi kelimeler ortaya çıkarken, ölçme ve değerlendirme ile ilgili bir konuda ölçek, değerlendirme ve geliştirme gibi kelimeler bulunabilecektir. Her iki konuda da ölçek gibi ortak bir kelime de olabilir. Belgeler kümesi (külliyyat) ne kadar büyük olursa hangi kelimelerin hangi konu grubuna dahil olduğunu belirlemek o kadar zorlaşmaktadır.

Modelin iki temel ilkesini basitçe ele alalım. Birinci ilke her belgenin konuların bir karışımı olduğu, ikinci ilke ise her konunun kelimelerin bir karışımı olduğudur. Her belge belirli oranlarda çeşitli konulardan kelimeler içerebilmektedir. Örneğin, bir sosyal medya platformunda yer alan ruh sağlığı ve ekonomi başlıklı paylaşımları inceleyen iki konu modelini ele alalım. Belge_1 %90 ruh sağlığı konusunu ve %10 ekonomi konusunu, Belge_2 ise %30 ruh sağlığı konusunu ve %70 ekonomi konusunu içermekte olsun. Ruh sağlığı başlığında psikoloji, terapi ve depresyon gibi kelimeler en yaygın olabilirken, ekonomi başlığı bütçe, enflasyon ve zam gibi kelimelerden oluşabilir. Daha da önemlisi, kelimeler konular arasında paylaşılabilir; bütçe gibi bir kelime her ikisinde de eşit olarak görülebilir.

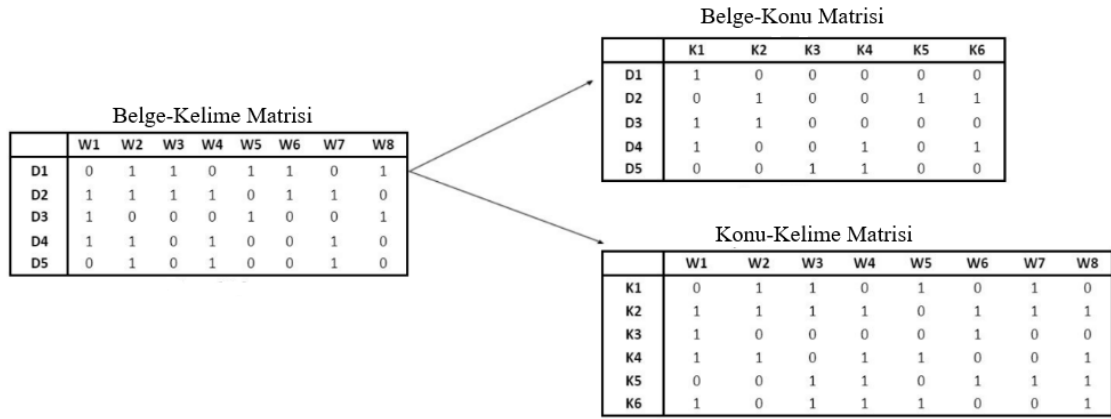
GDT algoritması birkaç kelimedenden (N) oluşan her belgenin (D), gizil konularda bir Dirichlet olasılık dağılımını temsil edebileceğini varsayar. GDT üretim sürecinde belgedeki kelimeleri farklı konulara tahsis etmek için kullandığı Dirichlet dağılımının adı ile anılır (Blei, 2012). GDT algoritmasında konular kelimelerin ve belgeler de konuların bir karışımını oluşturur, model bu bakımdan üç katmanlı hiyerarşik yapıya sahiptir.

GDT, bir konunun bir belgede olma olasılığını hesaplayarak bütün külliyat üzerindeki konuların saptanmasını sağlamaktadır.

GDT, belgeleri Şekil 1’de gösterildiği gibi belge-kelime matrisine dönüştürür. Şekil 1’de D1, D2, D3, D4, D5 belgeleri, K1, K2, K3, K4, K5 konuları ve W1, W2, W3, W4, W5, W6, W7, W8 kelimeleri göstermektedir. Belge-kelime matrisi bir külliyat içinde ortaya çıkan kelimelerin sıklığını tanımlayan istatistiksel bir gösterimdir. Bu matris belgelerdeki olası konuları içeren belge-konu matrisi, olası konuların içerdiği kelimeleri içeren konu-kelime matrisi olmak üzere iki alt matrise ayrılır. Bu matrisler zaten konu-kelime ve belge-konu dağılımlarını sağlamaktadır. Ancak, GDT algoritmasının temel amacı bu verilerin dağılımını iyileştirmektir. GDT, bu matrisleri iyileştirmek için örnekleme tekniklerini kullanır. Veriler matrislere dayalı olarak temsil edildikten sonra kelimeler vektörlere dönüştürülür (Blei, 2012).

Şekil 1

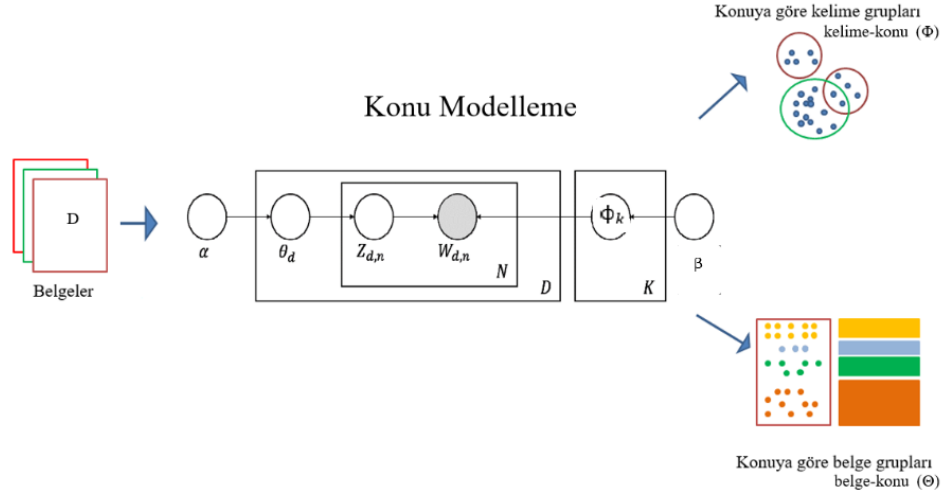
Belge-Kelime Matrisi



GDT belgelerin oluşturduğu külliyatı, konular olarak anlaşılabilen gizil (latent) olasılıksal değişkenler tarafından üretildiği varsayımına dayanmaktadır. GDT algoritmasında konuları belirlemek için kelimeler arasındaki eş oluşumdan yararlanan bir Bayesian hiyerarşik karışım modeli uygulanır. GDT üretken olasılıksal modelleme kullanmaktadır. Bu üretken süreç hem gözlenen hem de gizil rastgele değişkenler üzerinde ortak bir olasılık koşullu dağılımını tanımlar. Bu koşullu dağılıma sonsal dağılım da denir. GDT sonsal dağılım oluşturulması sürecinde Gibbs örnekleme algoritmasını kullanmaktadır (Blei, 2012). GDA modelinin daha iyi anlaşılabilmesi için Şekil 2’de grafiksel gösterimine yer verilmiştir.

Şekil 2

GDT Algoritmanın Şekilsel Gösterimi



Şekil 2’de α (belge-konu yoğunluğu) ve β (konu-kelime yoğunluğu) parametreleri işlemler sırasında bir kez örneklenen bağlantı parametreleridir. α parametresinin yüksek değerleri belgelerin daha fazla konunun karışımından; β parametresinin yüksek değerleri ise bir konun fazla sayıda kelimelerden oluştuğunu anlamına gelir. α değerinin düşük ya da yüksek olması belgelerin içerdiği konu sayısı ile ilgilidir. Özetle yüksek α parametresi belgelerin benzerliğini, yüksek β parametresi ise konuların benzerliğini gösterir.

Şekil 2’de $w_{d,n}$ d. belgedeki n. kelimeyi, $z_{d,n}$ d. belgedeki n’inci kelime için konu atamasını ve $\theta_{d,k}$ d. belgedeki k konu oranını göstermektedir. β_k , $\theta_{d,k}$, $z_{d,n}$ önceden bilinmeyen, $w_{d,n}$ bilinen değişkeninin işlenmesi sonucu elde edilen değerleri temsil etmektedir. $\theta_{d,k}$ α parametresi dağılımından çekilmektedir ($\theta_{d,k} \sim \text{dir}(\alpha)$). Konu sayısı K ve buna bağlı olarak α parametresinin dağılımı, konu değişkeni z’nin boyutluluğu önsel (a priori) olarak kabul edilir ve ayrıca sabit olduğu varsayılır. Her bir belge içindeki K konularının her biri için bir oranın tahmin edildiğine; ancak, bu tahmini oranlar oldukça küçük olabileceğine (örneğin, 0,001) dikkat etmek gerekir (Nowlin, 2015). Bu da o konunun o belge içinde yaygın olmadığını gösterir. α parametresinin standart Dirichlet önceliği 50/K’dır (Griffiths & Steyvers, 2004).

Şekil 2 incelendiğinde modelin bir dizi bağımlılık belirttiği anlaşılmaktadır. Örneğin, konu ataması $z_{d,n}$, belge başına konu oranlarına θ_d bağlıdır. Başka bir örnek olarak, gözlemlenen $w_{d,n}$ kelimesi $z_{d,n}$ konu atamasına ve $\beta_{1:k}$ konularının tümüne bağlıdır. (Operasyonel olarak, bu terim $z_{d,n}$ ’nin hangi konuya atıfta bulunduğu bakılarak ve $w_{d,n}$ kelimesinin bu konu içindeki olasılığına bakılarak tanımlanır). Bu bağımlılıklar GDT’yi tanımlar. Bunlar, üretim sürecinin arkasındaki istatistiksel varsayımlarda, ortak dağılımın belirli matematiksel biçiminde kodlanmıştır (Blei, 2012). Ortak olasılık dağılımı daha sonra Bayes kuralına göre, gözlenen değişkenler göz önüne alındığında gizli değişkenlerin koşullu veya sonsal dağılımını hesaplamak için kullanılır (Blei, 2012). Bu ortak dağılım, her bir belge için konu olasılıklarının sonsal dağılımını hesaplamak için aşağıda yer alan eşitliği kullanır:

$$p(\theta, z | w, \alpha, \beta) = \frac{p(\theta, z, w | \alpha, \beta)}{p(w | \alpha, \beta)}$$

Eşitlikte pay, tüm rastgele değişkenlerin ortak dağılımıdır ve payda, herhangi bir konu modeli altında gözlemlenen külliyatı elde etme olasılığıdır. Bu olasılıklar tüm olası konu yapıları üzerinden toplanabilir; ancak çok sayıda olası konu yapısı göz önüne alındığında bu hesaplanması zor hale gelmektedir. Bu nedenle, bu toplamın örneklemeye dayalı veya varyasyonel yaklaşımlar kullanılarak yaklaşık olarak hesaplanması gerekir (Blei, 2012). GDA algoritmasında, Gibbs örnekleme sonsal dağılımı tahmin etmek için kullanılır (Griffiths & Steyvers, 2004).

GDT algoritması daha önce bahsedildiği gibi ilk olarak belge-terim matrislerini oluşturur. Buna dayalı olarak, bir belgedeki her kelimeyi N kelimelik bir dizi olarak kavramsallaştırır (Blei, 2012). Daha sonra Dirichlet dağılımındaki k tane konudan birkaç tanesi seçilir. Belgeye bir konuya ait w_i kelimesi eklenir. w_i kelimesinin eklenmesinden sonra olasılık dağılımına göre bir konu seçilir. Bu şekilde bir dizi belgenin elde edildiğini varsayarak k tane konuyu ve bu konulara ait kelimeleri GDT algoritmasında Gibbs örnekleme ile saptamak için aşağıdaki adımlar izlenir.

1. Her bir belge için belgedeki her bir kelime rastgele K konularından birine atanır.

Her bir belge gözden geçirilir ve belgedeki her bir kelime K konudan birine rastgele atanır. Her d belgesi için, her w kelimesi gözden geçirilir. Belirli bir d belgesi için kaç kelimenin t konusuna ait olduğunu belirlemeye çalışılır. Eğer d 'deki birçok kelime t konusuna aitse, w kelimesinin t konusuna ait olma olasılığı daha yüksektir. w kelimesi nedeniyle t konusunda kaç belge olduğunu belirlenmeye çalışılır.

2. Bu rastgele atama hem tüm belgelerin hem de tüm konuların kelime dağılımlarını verir.

GDT, belgeleri konuların bir karışımı olarak temsil eder. Benzer şekilde, bir konu da kelimelerin bir karışımıdır. Bir kelimenin bir konuda olma olasılığı yüksekse, w kelimesine sahip tüm belgeler t konusu ile daha güçlü bir şekilde ilişkilendirilecektir. Benzer şekilde, w kelimesinin t konusunda olma olasılığı çok düşükse, w kelimesini içeren belgelerin t konusunda olma olasılığı çok düşük olacaktır, çünkü d belgesindeki kelimelerin geri kalanı başka bir konuya ait olacaktır ve dolayısıyla d bu konular için daha yüksek bir olasılığa sahip olacaktır. Dolayısıyla, w kelimesi t konusuna eklense bile, t konusuna bu tür çok sayıda belge getirmeyecektir.

3. Her bir belge için atamalar iyileştirilir.

Üçüncü aşamada her bir t konusu için iki oran hesaplanır.

$p(\text{konu } t | \text{belge } d) = d$ belgesindeki şu anda t konusuna atanmış kelimelerin oranı ve

$p(\text{kelime } w | \text{topic } t) = \text{bu } w \text{ kelimesinden gelen tüm belgeler üzerinden } t \text{ konusuna yapılan atamaların oranı.}$

Hesaplanan iki oran çarpımına yani t konusunun w kelimesini üretme olasılığına dayalı olarak w kelimesi yeni bir konuya atanır. Üçüncü adım defalarca tekrarlandığında konu atamalarının istikrarlı olduğu bir duruma ulaşılır. Böylece bu atamalar ile her bir belgenin içerdiği konu olasılıkları ve her konuya ait kelimeler, tekrarlanma sayıları ile elde edilir. GDT algoritmasını anlatırken işin matematiksel mantığı bu çalışmanın okuyucu kitlesine hitap etmeyeceği düşüncesi ile olasılık hesaplamalarında öncüllerin kullanımını gibi ayrıntılı hesaplamalara yer verilmeden açıklanmaya çalışılmıştır.

Birinci ve geçici olan aşamada kelimeler konulara atanır ancak konu sayısının araştırmacı tarafından belirlenmesi gerekmektedir. Kelimelerin konulara atanarak kelime-konu matrisin oluşturulmasında, kelimelerin farklı konular için kullanılmasına (çok anlamlılık) ve benzer kelimelerin kümelenmesine (eşanlamlılık) izin verilir. GDT algoritmasının kelime-konu matrisinde, konuları temsil eden kelimeler olasılıklara göre değerlendirilir ve her bir kelime birden fazla konu içerisinde bulunabilir. Ayrıca GDT modeli belgelerde geçen kelimelerin önemli semantik bilgileri taşıdığı ve benzer bir konudaki belge setinin aynı kelime setini içereceği varsayımı altında çalışır. Bu fikirden yola çıkarak, gizli konular, belgelerde sıklıkla geçen bir grup kelimeyi bulur. Bir konu içerisinde en yüksek olasılık değerine sahip kelimeler o konunun anlamının yorumlanmasını kolaylaştırır. Şekil 1'de GDT algoritması çeşitli belgelerden oluşan bir külliyatı, belge-konu (Θ) ve kelime-konu (Φ) matrisi oluşturmak suretiyle modellediği görülmektedir. Bu matrisler yardımıyla kelimeler ve ağırlıkları kullanılarak kullanışlı bilgiler ortaya çıkarılmakta ve bu kapsamda konular tespit edilebilmektedir. Modellemede kullanılmamış olan bir belgenin de hangi konuya ait olduğu yine bu matrisler yardımıyla atanabilir. GDT'nin çıktısı ise, her biri belirli bir konu kategorisine giren anahtar kelime gruplarından oluşur. Daha sonra bu anahtar kelime grupları benzersiz konu başlıkları ile etiketlenir. Bu nedenle çıktı, her biri 1 'den N 'ye kadar bir anahtar kelime grubunu ve 1 'den M 'ye kadar bir konu listesini temsil eder.

Yöntemsel Konular

Bir konu modellemesi çalışması çeşitli aşamalardan oluşmaktadır. Ancak bu aşamalar birbirinden bağımsız değildir, bazı aşamalar iç içe geçmiştir. Tüm araştırmalarda olduğu gibi konu modellemesi de hipotezler ve araştırma sorularının oluşturulması ile başlamalıdır. Konu modellemesi için en uygun olan bir metin kümesinin altında yatan gizli değişkenleri anlamamanın ilgi çekici olduğu sorulardır. İkinci aşama veri toplama aşamasıdır. Bu adım, hipotezleri test etmek veya araştırma sorularını yanıtlamak için gereken uygun metnin belirlenmesini ve bu metnin toplanması için bir yöntem geliştirilmesini içerir. Bir metin veri tabanını (külliyat) derlemek için veri türüne dayalı olarak kullanılabilir birçok veri toplama yöntemi vardır. Araştırmacı yeterli miktarda veri toplayabildiği sürece metin üreten herhangi bir tasarım kullanılabilir. Araştırmacılar, internet tabanlı metinlerden (örneğin, web sitesi metni, e-postalar veya Twitter, Facebook gibi sosyal medya verileri), mülakatlar ve odak grup görüşmelerinden elde ettikleri ses dosyalarını metne dönüştürerek veya katılımcılara anketler ve açık uçlu sorularla yanıt yazmalarını isteyerek, araştırma amacına uygun olarak farklı türde metinler kullanabilirler (Banks vd, 2018).

Konu modellemesinde dikkat edilmesi gerekenler örneklem büyüklüğü, veri ön işleme, konu sayısına karar verme ve model karmaşıklığını inceleme adımlarıdır. Bu adımlar sırasıyla ayrıntılı olarak açıklanmıştır.

Örneklem Büyüklüğü. Konu modellemesinde örneklem büyüklüğü analize dahil edilen belge sayısı ve her bir belgedeki kelime sayısı ile ilişkilidir. Konu modellemede gereken minimum kelime sayısı açısından genel bir kural bulunmamaktadır. Örneğin, sosyal medya paylaşımları sınırlı sayıda kelime içerdikleri için çok kısa olabilen belgelerdir. Metin türü bu şekilde kısa olduğunda, kullanıcı düzeyinde toplama işlemi yapılabilir ancak bu da örneklem büyüklüğünü (yani belge sayısını) azaltır. Metin verisi için minimum kelime sayısı konu bağlamına göre değişebilir. Özellikle anket çalışmalarında kullanıcılardan minimum bir kelime sayısında yazmaları istenebilir.

Metin uzunluğu kadar önemli olan noktalardan biri metnin kalitesidir. Yazım ve dilbilgisi açısından hatasız olan metinlerde konuların ortaya çıkması daha kolaydır. Aynı kelime bir belgede yanlış yazılmışsa veya kullanılmış ise, metindeki örüntüleri bulmak daha zordur. Anket çalışmalarında minimum bir kelime belirtilse de katılımcılar sadece kelime sayısını doldurmak için tekrarlayan ya da anlamsız kelimeler kullanabilirler. Bunun yanında metin türü sosyal medya gönderisi olduğunda kısaltmalar ve argo terimlerin varlığı araştırmacıları zorlayabilir. Örneklem büyüklüğünü etkileyen bir diğer husus belgenin düzeyidir. Bir cümle bir paragrafın içinde, bir paragraf bir sosyal medya gönderisinin içinde ve bir sosyal medya gönderisi de kullanıcının içinde yer alır. İlgilenilen hipotezlere veya araştırma sorularına bağlı olarak, araştırmacılar bir belgenin tanımlandığı düzeyi değiştirebilir bu da örneklem büyüklüğünü etkileyebilir (Tang vd, 2014).

Veri Ön İşleme. Nicel araştırmalarda analize başlamadan önce yapılması gerekli olan kayıp değer, uç değer incelemesi gibi metin verisi toplandıktan sonra ilk aşama veri ön incelemesi olmalıdır. Bu aşamada yapılan her işlem analiz sonucunu doğrudan etkilemektedir. Veri ön işleme aşaması ne kadar doğru ve iyi yapılırsa analiz sonucu da o kadar doğru ve güvenilir olmaktadır. Bir regresyon analizinde uç değerlerin ya da etkili değerlerin varlığı ve çıkarılması durumunda sonuçların değişebileceği gibi, metin ön işleme aşaması da potansiyel olarak sonuçları etkileyebilir. Ön işleme aşamasındaki adımlar, verileri daha sonraki analizler için hazırlamayı amaçlamaktadır. Bu adımlar sadece konu modellemesi yapmadan önce temiz ve uygun şekilde biçimlendirilmiş bir veri kümesi geliştirmek için değil, aynı zamanda konu modellemesi çıktısının yorumlanabilirliğini kolaylaştırmak için de gereklidir (Gencoglu, 2023). Bu bağlamda ön işleme keşifsel, yinelemeli bir aşamadır ve analizin ilerleyen aşamalarında bile bu aşamaya tekrar geri dönmek gerekebilir.

Ön işleme sürecinde, modelleme analizinin kalitesini etkileyebilecek tüm alakasız cevapların, çok kısa metinlerin veya geçersiz kayıtların kaldırılması ilk aşama olmalıdır. Devam eden süreç, parçalara ayırma (tokenizasyon/bölütleme) ve metnin temizlenmesini içerir. Bu aşama anlamlı olmayan karakterlerin ve noktalama işaretlerinin kaldırılmasını, tüm kelimelerin küçük harfe dönüştürülmesini ve boş cevapların ve boşluğa karşılık gelen cevapların hariç tutulmasını içermektedir (Gencoglu, 2023; Vijayarani vd. , 2015).

Parçalara ayırma (tokenization) işleminde tüm belge cümlelere, kelimelere, hecelere, harflere veya araştırma için belirlenen kurallara göre bölünmektedir. Her bir cümle, kelime, hece ve harf, noktalama işaretleri veya boşluklarla belirlenen araştırma için seçilmiş kurallara göre parçalara ayrılır. Bu aşamada dile özgü işlemler de gerekebilir. Örneğin Türkçe için Türkçe karakterlerin “ç” harfini “c” harfine, “ü” harfini “u”

harfine dönüştürmek gibi karakter dönüşümü yapılmalıdır. Veri setinde çok fazla Türkçe karakter bulunması işletim sistemindeki karakter kodlama problemi nedeniyle analiz sonucu negatif yönde etkilenebilir.

Ön işleme aşamasının en önemli adımlarından biri etkisiz kelimelerin (stopwords) değerlendirilmesi ve çıkarılmasını içerir. Etkisiz kelimeler bir araştırma bağlamında konuların belirlenmesine değer katmayacak kadar sık kullanılan sözcüklerdir. Etkisiz kelimeler metin içeriğiyle ilgili olmayan dile özgü sözcüklerdir, bu nedenle daha geçerli bir analiz için kaldırmaları gerekir. Etkisiz kelimeler kaldırıldıktan sonra kalan kelimeler genellikle isimler, fiiller ve sıfatlardır. “Ve, veya, ama, fakat, çünkü, üstelik, hatta” gibi bağlaçlar veya “ben, sen, o, şu” gibi zamirlerin tek başlarına anlamları olmadığı ve bu kelimelerin gruplandırma sürecine çok az değer kattığı veya hiç değer katmadığı için çıkarılması gerekir (Nowlin, 2015). Örneğin, külliyatta sadece tıbbi belgeler bulunuyorsa, insan, vücut, sağlık gibi kelimeler belgelerin çoğunda bulunabilir ve bu nedenle belgeyi öne çıkaracak herhangi bir özel bilgi eklemedikleri için çıkarılabilirler. Belgelerin en az %80 ~ %90'ında geçen bu etkisiz kelimeler herhangi bir bilgi kaybı olmadan sonuçların anlamlılığı için çıkarılmalıdır.

Analize başlamadan önce ilk olarak standart etkisiz kelimeler kaldırılmalıdır, ancak bulgular yorumlanmaya başlandığında ilave kelimeler de çıkarılabilir. GDT analizin yapıldığı açık kaynak kodlu yazılımlarda Türkçe etkisiz kelimeler kütüphaneleri istenilen seviyede olmayabilir. Bunun için Türkçe kelimelerden oluşan bir etkisiz kelimeler listesini çalışma kapsamında oluşturulması önerilebilir. Etkisiz kelimelerini bulmak için, kelime listeleri frekansa göre düzenlenebilir ve sık kullanılanlar semantik değer eksikliğine etkisiz kelime olarak seçebilir (Kadhim vd., 2014).

Etkisiz kelimelerin belirlenmesinde çalışma yapılan literatür hakkında bilgi sahibi olmak oldukça değerlidir. Bir analizde etkisiz görülebilecek bir kelime başka bir analizde önemli olabilir. Örneğin İngilizcedeki üçüncü tekil şahıs kelimesi “it” çoğu analiz için etkisiz kelimedir. Ancak bu kelime aynı zamanda bilgi teknolojisi kelimesinin de kısaltmasıdır ve bu alanda yapılacak bir çalışma için etkilidir. Etkisiz kelimelerin belirlenmesinde sonuçlardaki rastgele gürültüyü azaltmak için sıkça geçen kelimelerin çıkarılmasını öneren Zipf yasası (Newman, 2005) uygulanabilir. Zipf yasası, dağılımın kuyruğundaki seyrek kelimelerin çıkarılması için de geçerlidir. Bir kelimenin bir belgede geçmesi için minimum bir sayı belirlenebilir ve bunun üzerinde bir kelime seyrek terim olarak kabul edilmez. Ancak bazı kelimeler bazı belgelerde çok sayıda geçerken bazı belgelerde hiç geçmeyebilir. Bu nedenle, minimum sayıda belge belirlemek bazı kelimelerin analizden çıkarılmasına yol açacaktır.

Parçalara ayırma işleminden sonra, bir sonraki olası adım kelime köklerini belirleme (stemming) adımdır. Bu adım basitçe kök anlamı bozulmadan, eklerin çıkarılarak sözcüklerin kök veya temel biçimine indirgenmesi işlemidir (Jivani, 2011). Bu yöntemin amacı, çeşitli ekleri kaldırmak, kelime sayısını azaltmak, doğru eşleşen köklere sahip olmak, zamandan ve hafıza alanından tasarruf etmektir (Vijayarani vd., 2015). Ancak kök bulma sonucunda kalan sözcükler her zaman gerçek sözcükler olmayabilir; bunlar sadece anlamsal bir anlamı olmayan sözcük parçaları haline de gelebilir. Kök bulma, kelimeleri kök biçimlerine indirmek için daha temel ve kaba bir yöntem sağlarken, köklerini çözümleme (lemmatizasyon) yöntemi bir kelimenin amaçlanan anlamına göre temel veya sözlük biçimini belirlemeyi içeren daha karmaşık bir süreçtir. Bu işlemde ortaya çıkan kelimedeki her zaman dile ait olan özel bir kelimedir. Lemmatizasyon genellikle anlamsal doğruluğun çok önemli olduğu görevlerde tercih edilirken, kökten çıkarma basitlik ve hızın daha önemli olduğu görevlerde daha uygun olabilir. Lemmatization adımının gerekli olup olmadığı araştırmaya konu olan probleme bağlıdır (Weiss vd., 2010). Ne yazık ki, daha az sayıda yazılım paketi lemmatization yaklaşımını içermektedir. Hem kök belirleme hem de lemmatizasyon metin ön işlemede popüler seçenekler olmaya devam etse de, her bir yaklaşımın kullanılabilirliğine ilişkin araştırmalar karışıktır ve büyük ölçüde külliyata, yöntem ve araştırma diline bağlıdır (Banks vd, 2018). Özellikle, Türkçe gibi sondan eklemeli dillerde aynı kelime, metin içerisinde kullanımına göre farklı şekillerde görülebilmektedir. Kelimeler köklerine ayrılarak, bu sorunun üstesinden gelinebilir.

Kök bulma ve kök çözümleme adımlarından sonra kelime dizileri (n-gram) çalışmasının takip etmesi gerekmektedir. Kelime dizileri (n-gram) tekniğinde, bir dokümanda bitişik şekilde yer alan kelime ve kelime gruplarının kaç defa geçtiği hesaplanmaktadır. Tek bir kelimenin sıklık (frekans) sayısının hesaplanması tekli-gram (unigram), bitişik iki kelimenin sıklığının hesaplanması ikili-gram (bigram) ve bitişik üç kelimenin sıklık sayısının hesaplanması ise üçlü-gram (trigram) olarak adlandırılmaktadır (Kumar & Paul, 2016). Kelime dizileri çok faydalı ifadelerin belirlenmesinde anahtar olabilir. Örneğin, ikili-gramlarla analiz yapılırken Doğu Anadolu, Doğu ve Anadolu yerine ayrı ayrı iki kelime değil tek bir kelime olarak analize dahil edilecektir.

Sonuçlar tekli ve ikili gramların kullanımıyla değişecektir. Bununla birlikte, sonuçların her iki analizle de incelenmesi önerilir. Banks vd. (2018) genel olarak, araştırmacıların analizlere tekli-gramlarla başlamasını, ancak çoğu bağlamda ikili-gramların da dikkate alınmasını önermektedir.

GDT algoritmasının avantajlı ve dezavantajlı tarafları mevcuttur (Maier vd., 2018). GDT algoritmasının avantajları; büyük hacimdeki veri setinde yer alan konuların hızlı bir şekilde belirlenmesini sağlamaktadır. Gerçek hayata uygun olarak, bir belge birden fazla konuyu içirebilme olanağı mevcuttur. Dezavantajlı tarafları ise, sonuçların deterministik olmaması, yöntemsel konularda dikkat edilmesi gereken her noktanın, analiz sonuçları değiştirilmesidir. Konu sayısının araştırmacı tarafından belirlenmesi bir sınırlılık olarak sayılabilir.

Konu Sayısı Belirleme. K sayıda konunun belirlenmesi, konu modelleri kullanılırken karşılaşılan en büyük zorluklardan biridir (Blei & Lafferty, 2009). Denetimsiz bir algoritma olması nedeniyle, model tarafından türetilen kategorilerin araştırmacı tarafından gerekçelendirilebilir olmasını sağlama yükünü araştırmacıya yüklemektedir. GDT algoritması için uygun konu sayısını belirlemek araştırma sorusuna ve test edilen hipotezlere göre seçilmeli ve teori tarafından yönlendirilmelidir. Çok az sayıda konulu model, kelimelerin çakışmasına ve farklı konu başlıklarının belirlenememesine neden olabilir, ancak çok fazla konulu model ise anlamsal geçerliliği azaltabilir veya daha geniş ve kapsayıcı bir konuyu ortaya çıkarmakta başarısız olabilir. Bu nedenle konuların anlamsal geçerliliğe sahip olması, konu sayısı belirlemede en önemli gerekçe olmalıdır (Grimmer & Stewart, 2013). Anlamsal geçerlilik, her bir konunun, terimlerin o konuyla ilişkilendirilmesi ile fark edilebilecek açık ve tutarlı bir anlama sahip olması anlamına gelir (Krippendorff, 2003).

Araştırmacılar, değişen konu sayısı (K) değerleri ile yinelenmeli bir modelleme süreci kullanmalıdır. Konu sayısını azdan başlayıp artırarak veriyi keşfetmek, araştırmacıların verilerine daha aşina olmalarına yardımcı olur. Konu sayısını seçerken, araştırmacıların tutumluluk (parsimony) ilkesini uygulamaları önemlidir. Yani, metindeki konuların veya örtük yapıların sayısını ve dolayısıyla sonuçların karmaşıklığını artırmayı haklı çıkarmak için, daha fazla konu ile bulguları daha iyi yorumlayabileceğine dair bir gerekçe olmalıdır.

Araştırmacılar ortaya çıkan konuların neyi temsil ettiğini tartışmalı ve konu etiketleri geliştirmelidir. Ayrıca ortaya çıkan konuları mevcut literatürle ilişkilendirilmesini sağlamaya çalışmalıdır. Araştırmacılar konu etiketleri ve tanımlarından tatmin olduklarında ve orijinal metinden destekleyici örnekler belirlendiğinde analiz tamamlanmış olur. Konu modelleme aşamasında son adım olarak, araştırmacılar konuların ağ yapısını incelemelidir. Bir konu ağı, belirli konuların ne kadar ilişkili olduğunu görmeyi sağlar ve ortaya çıkan yapıların boyutluluğunu değerlendirmek için faydalıdır (Banks, 2018).

Model değerlendirme. Makine öğrenmesinde denetimli öğrenme algoritmalarında modeli değerlendirmek için hesaplanabilecek metrikler bulunmamaktadır. GDT denetimsiz bir öğrenme algoritması olduğundan hata metrikleri ile model performansının değerlendirilme olanağı yoktur. Ancak elde edilecek farklı modeller birbirleri ile karşılaştırılabilir. Dil modellemesinde sıklıkla kullanılan karmaşıklık (perplexity), cebirsel olarak kelime başına geometrik ortalama olasılığın tersine eşdeğerdir. Daha düşük bir karmaşıklık puanı daha iyi genelleme ve tahmin performansını gösterir (Blei vd., 2003).

PSİKOLOJİK DANIŞMA VE REHBERLİKTE KONU MODELLEMESİ

Bir veri analizi yöntemi olarak konu modellemesi; bazı araştırmacılar tarafından, büyük verinin makine öğrenimi yoluyla analiz edilmesini sağlayarak geleneksel araştırmaların tamamlayıcı bir işlevi olarak değerlendirilmektedir. Psikolojik danışma literatürünün konu modellemesine göre analiz edilmesinin, tarihsel süreç içinde alanın nasıl ve nereye doğru bir gelişim ve değişim gösterdiğini görmek ve alandaki değişim ve gelişmeleri izlemek açısından önemli içgörüler sağlayabileceğine işaret edilmektedir (Oh vd., 2017). Psikolojik danışma ve rehberlik alanında konu modellemesinin kullanımı kuram, araştırma ve uygulamaya yönelik önemli katkılar sağlama potansiyeli taşımaktadır.

Konu modellemesi ilk olarak, psikolojik danışma ve rehberlik alanyazınında öne çıkan konuların ortaya konulması kuramsal ve kavramsal temellerin test edilmesine, genişletilmesine ve geliştirilmesine (Chen & Wojcik, 2016) yönelik katkılar sağlayabilmektedir. Ayrıca, konu modellemesi, psikolojik danışma ve rehberlikte araştırma eğilimlerini belirlemek ve araştırmanın odağındaki konuları ortaya çıkarmak gibi önemli bir işlev görebilmektedir (Bittermann&Fischer, 2018). Bu noktada, geleneksel nitel veri analizi yöntemleri ile

kolay bir şekilde ortaya çıkarılamayabilecek gizil temaların görülmesine ve böylece araştırma açısından eksik kalan ya da gözden kaçan yönlerin keşfedilmesine katkı sağlayabilmektedir (Chen & Wojcik, 2016). Dahası, ortaya çıkan konular araştırma bulgularının sentezlenmesine ve bu sentezlere dayalı olarak etkili çıkarımların ortaya konulmasına ve mevcut problemlerin çözümüne ışık tutmaya yarayabilir. Ayrıca, genç araştırmacılara, alandaki araştırma eğilimlerini görmelerine ve bu doğrultuda kendi ilgi alanlarını keşfetmelerine yardımcı olabilir (Bittermann ve Fischer, 2018).

Konu modellemesinin, psikolojik danışma uygulamalarına yönelik olası katkılarının da altı çizilmektedir. Örneğin, anonimliği sağlanmış danışma transkriptlerinin konu modellemesi ile analiz edilmesinin danışma oturumlarının içeriği ve odağının ortaya konulmasına, danışanın yaşamında öne çıkan ve tekrarlayan temaların keşfedilmesine ve danışan deneyimlerinin çeşitli yönleri ile anlaşılmasına ve bu doğrultuda yardım sürecinin gözden geçirilmesine ve değerlendirilmesine katkı sağlayabilir (Chen & Wojcik, 2016). Ayrıca tekrarlayan temaların danışanlar ile oturumlarda şeffaf bir şekilde ele alınması, danışanların problemlerine ve bunların çözümlerine yönelik içgörüler kazanmasına katkı sağlayabilir. Farklı psikolojik danışma yöntem ve modellerinin etkililiğini anlamaya yönelik gerçekleştirilecek konu modellemesi araştırmaları (Liu & Gao, 2021) hem psikolojik danışma sürecinin geliştirilmesine hem de psikolojik danışman eğitimine önemli katkılar sağlayabilir. Psikolojik danışmada konu modellemesi ortaya koyduğu bulgular açısından toplumsal ihtiyaçların anlaşılmasına ve aynı zamanda toplumun ihtiyaçlarına cevap verecek şekilde uygulamalar geliştirilmesine katkı sağlayabilir (Bittermann & Fischer, 2018). Aynı zamanda, bu ihtiyaçlardan yola çıkarak politikalar geliştirilmesine yönelik anlamlı girdiler sağlayabilir. Bu işlevleri ile konu modellemesi, psikolojik danışma ve rehberlikte etkili, sistematik ve değerli bir araştırma yöntemi olarak değerlendirilmektedir (Bittermann & Fischer, 2018; Oh vd., 2017).

Psikolojik Danışma ve Rehberlik ve İlgili Alanyazında Konu Modellemesine Dair Örnekler

Büyük veri analizi ve makine öğrenimi konusundaki gelişmelerle birlikte konu modellemesinin bir veri analizi yöntemi olarak kullanımına dair araştırmaların ruh sağlığı ile ilgili alanyazında artmakta olduğu görülmektedir. Bu başlık altında, gizil konu modellemesi gibi yöntemlerle gerçekleştirilen konu modellemesi analizi kullanan araştırmalardan örnekler sunulmaktadır. Var olan çalışmalar, araştırma amaçlarına göre beş grup altında toparlanarak özetlenebilir: 1) Genel olarak ruh sağlığı ile ilgili alanlarda en çok araştırılan konuları ortaya çıkarmaya yönelik araştırmalar (ör., psikolojik danışma alanındaki trend konular) 2) Psikolojik danışma alanyazındaki belirli konulara ilişkin temaların araştırıldığı araştırmalar (ör., çok kültürlülük, kariyer danışmanlığı, bilinçli farkındalık) 3) Psikolojik danışma ve yardım süreci (ör., en etkili psikolojik danışma yöntemleri, terapötik ilişki, çevrimiçi danışma vb.) ile ilgili temaların araştırıldığı çalışmalar 4) Tanı ve tespiti yönelik gerçekleştirilen araştırmalar (ör., siber zorbalık, depresyon belirtilerinin tespiti gibi) 5) Ölçme, değerlendirme ve bireyi tanımayla yönelik çalışmalar kapsamında gerçekleştirilen araştırmalar (ör., anketlerde açık uçlu sorulara verilen yanıtların değerlendirilmesi).

İlk grupta, psikolojik danışma ve ilgili ruh sağlığı alanyazını ile bunların alt alanlarında (ör., kariyer psikolojik danışmanlığı, örgütsel psikoloji, evlilik ve aile psikolojik danışmanlığı gibi) tarihsel süreç içinde araştırmaların hangi konular etrafında toplandığını ortaya koymak, ayrıca araştırma süreçlerindeki güncel eğilimleri tespit etmek amacıyla gerçekleştirilen çalışmalar örnek olarak gösterilebilir. Örneğin, *Danışma Psikolojisi* (Journal of Counseling Psychology) dergisinde 1963'ten 2015 yılına kadar yayınlanan akademik makaleler konu modellemesi yöntemi ile analiz edilerek son 53 yılda araştırmaların hangi konularda yoğunlaştığı araştırılmıştır. Bulgular, 70 konu dağılımı ortaya koyarken bu konuların psikolojik danışma süreci ve sonuçları, çok kültürlülük, araştırma yöntemleri ve kariyer danışmanlığı olmak üzere dört kategoride toplandığı belirtilmiştir. Özellikle son yıllarda çeşitlilik ve çok kültürlülük konularına yönelik araştırma eğilimlerinde artış olduğu raporlanmıştır (Oh vd., 2017). Benzer şekilde, 1980'den 2016 yılına kadar Almanca konuşulan ülkelerde psikoloji alanında yapılan araştırmaların hangi konular etrafında toplandığı incelenmiştir. 314.573 yayının analiz edildiği araştırmada 500 konu ortaya konulmuş ve bunlar içinden hangilerinin trend konular olduğu araştırılmıştır. Buna göre artan bir araştırma eğilimi gösteren gündemdeki konular, nöropsikolojinin farklı yönleri, çevrimiçi psikolojik danışma, kültürler arası özellikler, travma ve görsel dikkat olarak belirlenmiştir (Bittermann & Fischer, 2018). Bu çalışmalara benzer şekilde sadece belli uzmanlık alanına odaklanan araştırmalar da bulunmaktadır. Örneğin, Liu vd. (2020), 2007 ile 2018 yılları arasında klinik psikoloji alanında yapılmış araştırmalarda öne çıkan konuları belirlemek için konu modellemesi analizi gerçekleştirmişler ve meta-analiz, psikoterapi, profesyonel gelişim ve depresyon konuların klinik psikolojide 12 yıl boyunca en öne çıkan konular olduğunu ortaya koymuşlardır. Bununla birlikte davranışsal müdahaleler

konusunun ise 2007 yılından bu yana araştırma konusu olarak artışta olduğunu rapor etmişlerdir. Bu araştırmalar, özellikle Psikolojik Danışma ve Rehberlik ve ilgili alanyazında araştırmaların odağının hangi konularda toplandığını ve araştırma eğilimlerini görmek açısından anlamlı sonuçlar üretmektedir. Bu açıdan bakıldığında bu araştırmalar, hangi konuların tarihsel süreç içinde önem kazandığını görmek kadar hangi konuların gözden kaçtığını görmek açısından da anlamlı çıkarımlar sağlayabilecektir.

İkinci grupta, psikolojik danışma ve ilgili alanyazında belirli sorun alanlarına ya da belirli müdahale biçimlerine yönelik gerçekleştirilen araştırmaların hangi konular etrafında odaklandığını ortaya çıkarmaya yönelik çalışmalar yer alabilir. Bu gruba, Kee vd. (2019) tarafından bilinçli farkındalık ile ilgili çalışmaların konu modellemesi analiz edildiği araştırma örnek olarak gösterilebilir. Araştırmacılar, Web of Science'da 20 Ekim 2017 yılına kadar konuyla ilgili yayınlanmış 5949 araştırmayı konu modellemesi ile analiz etmişlerdir. Bulgularda 106 konu 231 terim öne çıkarken bunlar sonrasında durum/sorun (kumar bağımlılığı, depresyon, kaygı bozuklukları, kanser gibi bilinçli farkındalık uygulamalarının kullanıldığı sorun ya da konu alanları), yapı/felsefe (Vipassana, Budizm, meditasyon gibi çalışmaların dayandırıldığı farklı felsefi kökenler), modalite/biçim (bilinçli farkındalık temelli bilişsel terapi, bilinçli farkındalık temelli stres azaltma programları, sanal gerçeklik uygulamaları gibi teknoloji tabanlı uygulamalar), hedef kitle/uygulama alanı (ergenler, çiftler, hamileler, sağlık çalışanları ve işyeri, okul gibi araştırmanın yapıldığı farklı kitleler ve yerler) ve araştırma yöntemi (kullanılan ölçme araçları, araştırmaların desenleri vb.) olmak üzere beş tema altında kategorize edilmiştir. Bu araştırma bilinçli farkındalık konusunda çalışan araştırmacılara şu ana kadar gerçekleştirilen araştırmaların hangi konularda toplandığını göstermesi açısından önemli girdiler sağlayabilir. İlginç bir diğer araştırma örgütsel psikolojide gerçekleştirilmiştir (Royston vd., 2022). Orduda farklı düzeylerde görev yapan subayların görevleri ile ilgili açık uçlu sorulara verdikleri yanıtlar konu modellemesi ile analiz edilmiş ve görev ve sorumluluk tanımlarının liderlik düzeyine ve deneyimine göre nasıl bir fark gösterdiği incelenmiştir. Bulgulara göre, alt düzey pozisyonlarda günlük görevler ve işle ilgili bilgi edinmeye odaklı konular öne çıkarken, liderlik seviyesi arttıkça denetim, ağ genişletme aktiviteleri ve süreç iyileştirme gibi konular öne çıkmıştır. Daha üst düzey liderlik seviyelerinde ise kullanılan dil, kurumsal komitelere katılım, stratejik liderlik, kurumsal operasyonları denetleme, mentörlük, stratejik planlama gibi konuları öne çıkarmıştır. Bu araştırmada konu modelleme analizinin iş analizi çalışmaları için önemine işaret edilirken, bu analiz yöntemi ile elde edilen bulguların belirli işlerin gerektirdiği nitelikler, çalışan özellikleri, iş bağlamı gibi konuların daha iyi anlaşılmasına katkı sağlayabileceğine vurgu yapılmıştır (Royston vd., 2022). Başka bir örnek ise Wang vd. (2016) tarafından ergenlerde madde kullanımı ve depresyon konusunda 2000 yılından 2014 yılına kadar PubMed'de yayınlanan toplam 17.723 araştırmanın özeti üzerinde gerçekleştirilen konu modellemesi analizidir. Analizler sonucunda beş, 20 ve 50 konulu modeller ortaya konulurken, örneğin beş konulu modelde bazı terimlerin bir arada yer almasına işaret edilmiştir. Örneğin, alkol ve beyin araştırmaları, alkol ve cinsel deneyimler, sigara ve alkol farklı konular altında ikili gruplar olarak bir arada yer alan terimler olmuştur. Benzer şekilde, 20 konu ortaya koyan modelde müdahale ve tedavi, aile etkisi ve akran sosyal ağı, diyet, fiziksel aktivite ve obezite, intihar ve beyin araştırmaları gibi ek konuların çıktığına işaret edilmiştir. 50 konu ortaya koyan modelde ise genetik ve çevresel etkiler, depresyon ve alkol kullanımı gibi konulara işaret edilmiştir. Bu üç modelde, ergenlerde madde kullanımı ile ilgili öne çıkan beş konunun; madde kullanımı, genel araştırma terimleri, sigara kullanımı, beyin araştırmaları ve ailesel ve arkadaş ağları olduğu rapor edilmiştir. Ergenlerde depresyon ile ilgili öne çıkan beş konunun ise; depresyon, psikiyatrik sorunlar, depresyonun tedavisi ve genel araştırma terimleri olduğu belirtilmiştir. Bulgularda, ergenlerde depresyon ve madde kullanımına dair öne çıkan tipik konuların yanı sıra ergenlerin madde kullanımı ile ilgili risk faktörleri, madde kullanımı ve depresyon arasındaki ilişki ve müdahale programı konuları öne çıktığı işaret edilmiştir. Ayrıca, ergenlerde madde kullanımı ve depresyon ile ilgili en öne çıkan temanın seks ve şiddet olduğu belirtilmiş; ergenlerde cinsel deneyim ve şiddetin madde kullanımı ve depresyonla ilişkisine işaret edilmiştir. Diğer öne çıkan ilginç konunun ise, çocukluktan yetişkinliğe gelişim olduğu belirtilmiştir. Bu çalışma ergenlerde depresyon ve madde kullanımı konusundaki risk faktörlerini, müdahaleye yönelik öne çıkan konuları ve bu alanlardaki güçlü yönler kadar ihtiyaçları belirlemede önemli sonuçlar ortaya koyduğu söylenebilir. Bu bulguların, konuyla ilgili gerçekleştirilecek müdahale programlarının nasıl şekillendirilmesi konusunda da önemli içgörüler sağlayabileceği söylenebilir.

Üçüncü grupta, psikolojik danışma ve psikolojik yardım süreçlerinin etkililiği ve sonuçları ile ilişkili olabilecek konuları ortaya çıkarmaya yönelik konu modelleme analizi ile gerçekleştirilen araştırmalar ele alınabilir. Bu grupta, Atzil-Slonim vd. (2019) tarafından psikoterapide danışanların işlevsellik düzeylerini ve terapötik ilişkideki kopuklukları belirlemek için konu modellemesi analizinin gerçekleştirildiği araştırma örnek olarak ele alınabilir. Bu çalışma kapsamında 52 terapist ve 58 danışan ile gerçekleşen 873 psikoterapi

oturumlarının transkriptleri analiz edilmiştir. Bu kapsamda, terapi süreçlerinde öne çıkan konuların neler olduğu; hangi konuların danışanların işlevselliğini ve terapist ile arasındaki terapötik ilişkideki kopukluğu en iyi şekilde tanımladığı ve analizlerde öne çıkan bu konulardaki değişimlere göre terapinin sonuçlarında bir değişiklik olup olmadığı araştırılmıştır. Bu incelemeleri desteklemek üzere, her oturumdan önce danışanlar tedavi sürecindeki gelişimin göstergesi olarak kabul edilen işlevselliklerinde, kişilerarası ilişkilerinde ve sosyal rollerdeki performanslarındaki değişimleri ve semptom düzeylerini bildirdikleri iki farklı ölçme aracı doldurmuşlardır. Ayrıca her oturumdan sonra terapistler de terapötik ilişkideki kopuklukları değerlendirdikleri bir soruya yanıt vermişlerdir. Çalışma kapsamında öncelikle transkriptler üzerinde konu modellemesi analizi yapılarak terapi sürecinde öne çıkan konular tespit edilmiştir. Buna göre pozitif deneyim, negatif deneyim, ilişkiler, tedavi, sağlık, günlük yaşam ve çeşitli konular öne çıkan konu alanları olmuştur. Ardından danışanların işlevsellik düzeylerini ve psikoterapi seanslarındaki ittifak kopmalarının hangi konularda ortaya çıktığı lojistik regresyon ile araştırılmıştır. Bulgulara göre, özellikle dört konunun (yalnızlık, ıstırap, fiziksel güçlükler ve öfke) danışanların düşük işlevsellikleri ile; eğlence, serbest zaman deneyimleri, kutlama ve seçim konularının ise yüksek işlevsellikleriyle ilişkili olduğu rapor edilmiştir. Ayrıca, analizlerde öne çıkan üç konu -iletişim, amaç belirleme, yardım ihtiyacı ve problemler-terapötik ittifaktaki kopuşla ilişkili bulunmuştur. Son olarak, terapi sürecindeki konular değiştikçe terapinin sonuçlarında da bir değişiklik olup olmadığı çok düzeyli büyüme modelleri ile analiz edilmiştir. Sonuçlara göre, terapi süreci boyunca yukarıda bahsedilen danışanların yüksek işlevsellikleri ile ilişkili bulunan konulardaki artışın semptomlarda azalma (düşük depresyon düzeyi gibi) ile ilişkili olduğu görülürken, düşük işlevsellikle ilişkili bulunan konulardaki değişimlere göre terapi sonuçlarında bir değişim olmadığı görülmüştür. Bu sonucun, olumsuz deneyimlere odaklanmaktansa olumlu deneyimlere odaklanmaya dayalı güncel psikoterapi yaklaşımlarını desteklediğine işaret edilmiştir. Benzer şekilde, Howes vd. (2013) tarafından gerçekleştirilen çalışmada terapi sürecinde öne çıkan konuların danışanların semptomlarını ve terapi sonuçlarını yordayıp yordamadığı, eğer böyle bir yordama gerçekleşiyorsa, terapistler tarafından gerçekleştirilen açıklamalar yerine konu modellemesi gibi makine öğrenmesine dayalı otomatik veri analizi yöntemlerinin kullanılıp kullanılamayacağı ve konu modelleme analizinin yorumlanabilir ve/veya terapistler tarafından alınan kararlarla karşılaştırılabilir nitelikte konular üretip üretmediği gibi sorulara yanıt aranmıştır. Bu doğrultuda, 2006-2008 yılları arasında 31 psikiyatrist ve bilgilendirilmiş onam alınan 138 danışan arasında gerçekleşen en kısıtı 5 dk. en uzununu 1 saat süren psikolojik danışma oturumlarının transkriptleri kullanılmıştır. Tedavi sürecinin bir parçası olmayan araştırmacılar danışanların semptomlarını değerlendirmiş, danışanlardan tedavi oturumlarına ilişkin doyum düzeylerini değerlendirmek üzere bir ölçme aracına yanıt vermeleri istenmiş ve hem danışanlardan hem de doktorlardan bu süreçteki terapötik ilişkiyi değerlendirmeye yönelik bir ölçme aracına yanıt vermeleri istenmiştir. Toplamda 138 danışma transkripti içinden 12'si iki uzman tarafından kodlanmış ve bu oturumlarda öne çıkan 20 konu belirlenmiştir. Aynı zamanda, 138 transkript konu sayısı 20'ye sabitlenerek konu modellemesi yoluyla ayrıca analiz edilmiştir. Böylece uzman insan kodlayıcılarla makine öğrenmesine dayalı konu modellemesinin öne çıkardığı konular karşılaştırılmıştır. Hem makine öğrenmesine dayalı konu modellemesi analizi ile hem de uzmanlar tarafından elle yapılan tematik kodlamanın benzer yordama gücüne sahip olduğu görülmüştür. Ancak, makine öğrenmesine dayalı otomatik kodlama ile ortaya çıkarılan konuların semptomları yordamazken uzmanlar tarafından yapılan kodlamaların semptomları yordadığı görülmüştür. Öte yandan konu modellemesi analizinin uzmanlar tarafından yapılan kodlamaya kıyasla terapötik ilişkiyi daha iyi yordayan konular/temalar ortaya koyduğu rapor edilmiştir. Bu araştırmalar, oturum transkriptleri üzerinde gerçekleştirilen konu modellemesi analizlerinin psikolojik yardım süreci ve sonuçlarını ve danışan ile psikolojik danışman/terapist arasındaki terapötik ilişkiyi ve ittifakı değerlendirmede kullanılabileceğine işaret eden öncül bulgular üretmeleri açısından ilgili alan yazına önemli katkılar sağlamaktadır.

Dördüncü grup, konu modellemesi analizinin tanı ve tespit gibi amaçlarla kullanımına yönelik örnekleri kapsamaktadır. Buna örnek olarak, Franz vd. (2020) tarafından ergenler tarafından kullanılan bir internet destek grubu olan TeenHelp.org platformunda kendine zarar verme ve intihar ile ilgili alt gruplarda gerçekleştirilen paylaşım ve gönderilerin konu analizi ile incelendiği araştırma verilebilir. Aynı içerikler aynı zamanda üç insan kodlayıcı tarafından da kodlanarak intihar ve kendine zarar vermeye yönelik tespitleri içeren temalar (intihar içermeyen kendine zarar verme, intihar düşüncesi, intihar girişimi, pasif intihar düşüncesi, intiharı planlama, depresyon, sosyal kaygılar, istismar, ilaç kullanımı, madde kullanımı ve diğer şeklinde) ortaya konulmuştur. Benzer şekilde konu modellemesi analizinin kendine zarar vermeye yönelik düşünce ve davranışları tespit edip edemediği değerlendirilmiştir. Bulgulara göre, 15 konu ortaya koyan analiz sonuçlarına göre altı tanesinin (intihar, arkadaşlar, aile, depresyon, intihar içermeyen kendine zarar davranışları ve sosyal olmak üzere) bu tespiti içerdiği ifade edilmiştir. Dahası, makine öğrenmesinin intihar ve depresyonu ayrı konular olarak tespit edebildiğine işaret edilmiştir. Araştırmada aynı zamanda, makine öğrenmesine dayalı

konu modellemesi analizinin insan kodlayıcılar tarafından ortaya konulan temalarla ne kadar örtüştüğü de araştırılmıştır. Bu amaçla yapılan lojistik regresyon analizi bulgularında makine öğrenmesiyle ortaya çıkarılan ve hipotezleri karşılayan altı konunun insan kodlayıcılar tarafından belirlenen dört tema ile örtüştüğü ifade edilmiştir. Makine öğrenmesi ile tespit edilen konuların cinsiyete ve yaşa göre fark gösterip göstermediği t testi ile analiz edilmiş ve “intihar” konusunun erkeklerde daha yüksekken, “intihar içermeyen kendine zarar verme” konusunun kadınlarda daha yüksek görüldüğü raporlanmıştır. Bu araştırma, makine öğrenmesine dayalı konu analizinin ergenlerin paylaşımlarından intihar ve kendine zarar verme ile ilgili temaları ortaya çıkararak tespitlerde bulunmaya yönelik nasıl kullanılabileceğini örneklemesi açısından önemlidir. Bu gruptaki çalışmaların serbest/denetimsiz konu modellemesi analizlerinden ziyade denetimli konu modellemesi analizleri ile daha sofistike bir şekilde ele alındığı söylenebilir (ör., Mobin vd., 2024). Dolayısıyla sadece serbest konu modellemesi analizlerinin, incelenen sorunla ilgili temel konuları ortaya çıkarması tek başına tespit ve tanı amaçlı kullanılması pek mümkün görünmemekle birlikte üzerinde değerlendirmeler yapmak ve olası hipotezleri üretmek açısından bir fikir vereceği söylenebilir. Ancak denetimli konu modellemesi analizleri ve diğer derin öğrenme analizleri ile tanı ve tespiti yönelik makine öğrenmesi analizlerinin gerçekleştirildiği görülmektedir.

Son grup, bireyi tanıma ve ölçme değerlendirme süreçlerinde konu modellemesi analizinin nasıl kullanılabileceğine yönelik örnekleri içermektedir. Örneğin Finch vd. (2021), üstün yetenekli çocukların belirlenmesine yönelik geliştirmek istedikleri bir ölçek için, öncelikle üstün yetenekli çocukların ebeveynleri tarafından yazılı olarak verilen üstün yetenekli olmanın farklı yönlerine dair tanımları konu modellemesi yoluyla analiz etmişlerdir. Bu analizler sonucunda ortaya çıkan konular ölçek maddelerinin geliştirilmesinde kullanılmıştır. Başka bir araştırmada da (Rohrer vd., 2017), konu modellemesi analizi nicel araştırma içine entegre edilmiş, katılımcılara uygulanan ölçme araçlarının yanı sıra açık uçlu olarak sorulan “Sizi endişelendiren diğer şeyler neler?” sorusuna katılımcıların verdiği yanıtlar konu modellemesi ile analiz edilmiştir. Bu araştırma kapsamında, Almanya’da Sosyo-Ekonomik Forum kapsamında, 2000-2011 yılları arasında katılımcıların başka hangi konular hakkında endişe duyduklarına dair açık uçlu soruya verdikleri yanıtları içeren 35.000 yazılı yanıt incelenmiştir. Bulgularda öne çıkan konular; çocukların geleceği, çocuklar, gençler, okul, aile sağlığı, yükselen fiyatlar, zengin-fakir, Almanya’daki yabancılar, işsizlik, iş bulma, emeklilik ve ekonomik güvence, Almanya’nın gelişimi, kişilerarası konular, ahlaki değerlerdeki düşüş, politikacılar, politikadaki yozlaşma ve savaş ve terör olarak rapor edilmiştir. Konuların zaman içindeki değişimi incelendiğinde ise savaş ve terör ile fiyatlardaki artış konularının özellikle uluslararası savaşların yaşandığı yıllarda (ör. Suriye’deki iç savaş gibi) ve ekonomik kriz zamanlarında yükselişe geçen konular olduğu belirtilmiştir. Bu örnek, nicel ve nitel veri analizlerinin bir arada kullanımına bir örnek olduğu gibi, geniş katılımlı araştırmalarda açık uçlu sorulara verilen yanıtların analizi için elde yapılan kodlamaların sınırlılıklarının ötesine geçebilecek bir alternatif sunması açısından önemli görülebilir. Başka bir ilginç araştırmada ise (Jayaratne & Jayatilleke, 2020) iş başvurusu aşamasında adayların açık uçlu sorulara verdikleri yanıtlar analiz edilerek HEXACO kişilik modeline göre hangi tipe girdikleri tespit edilmeye çalışılmıştır. Çalışma kapsamında 46.000 iş başvurusunda adayların çevrimiçi yazışma botu kullanarak açık uçlu sorulara verdikleri yanıtların yanı sıra HEXACO kişilik modeline yönelik yanıtladıkları bir ölçme aracı sonuçları kullanılmıştır. İş mülakatları şu gibi açık uçlu soruları içermiştir: “Sizi ne motive eder? Hangi konuda tutkulusunuz? Herkes her zaman aynı fikirde olmayabilir. Sizinle aynı fikirde olmayan bir akranınız, takım arkadaşınız veya arkadaşınız oldu mu? Siz ne yaptınız? Bir şeyi başarmak için elinizden gelenin daha fazlasını yaptığınız bir duruma örnek verin. Bunu başarmak sizin için neden önemliydi? Bazen işler her zaman planlandığı gibi gitmez. Bir teslim tarihini veya kişisel taahhüdünüzü yerine getiremediğiniz bir zamanı anlatın. Ne yaptınız? Bu size kendinizi nasıl hissettirdi? Satışta hızlı düşünmek kritik önem taşır. Sizi bu konuda nitelikli kılan nedir? Bir örnek verin.” Adayların bu sorulara verdikleri yanıtlar makine öğrenmesine dayalı konu modelleme analizi (gizil dirichlet tahsisi) ile incelenmiş ve modeller arasında ortalama düzeyde ($r = 0.39$) bir korelasyon bulunmuştur. Farklı kişilik modellerini tanımlayan terimlerin uygunluğuna dair 117 gönüllü katılımcının geribildirimi ile teyit alınmıştır. Bu çalışmada, iş mülakatları esnasında sosyal beğenirlik gibi unsurlardan etkilenen öz bildirime dayalı bireyi tanıma teknikleri kullanmak yerine açık uçlu mülakat sorularına verilen yanıtların konu modelleme analizleri yoluyla adayların kişilik tiplerini belirleme ve işe uygunluğunu değerlendirmede kullanımı açısından önemli bir ilerleme sağlayabileceğine vurgu yapılmıştır.

Psikolojik danışma ve ilgili alanyazındaki konu modellemesinin kullanıldığı araştırmaları sınıflarken araştırmanın amacı kadar araştırmada kullanılan veri kaynağına göre de bir sınıflamaya gidilebilir. Buna göre, (1) yayınlanmış akademik dokümanların, (2) sosyal medya ve/veya internetteki (bloglar gibi) etkileşim içeren site ve platformlardaki yazışmaların, yorumların ve paylaşımların, (3) psikolojik danışma ve terapi

transkriptlerinin ve (4) veri toplama araçlarındaki açık uçlu sorulara verilen yazılı yanıtların kullanıldığı araştırmalar olarak bir sınıflama yapmak mümkündür. Birinci gruptaki araştırmalara yukarıda çok sayıda örnek verilmiştir. Araştırma makaleleri dışındaki belgelerin nasıl analiz edilebileceğine dair başka bir örnek Ni vd. (2023) tarafından gerçekleştirilen psikolojik ilk yardım konusunda yazılmış el kitaplarının analiz edildiği çalışma verilebilir. Bu doğrultuda, Google’da “psikolojik ilk yardım,” “el kitabı,” “rehber,” “kılavuz” gibi anahtar kelimelerle ulaşılan farklı ülkelerdeki farklı kurum kuruluşlar tarafından İngilizce dilinde yazılmış 14 el kitabı konu modellemesi yoluyla analiz edilmiştir. Bulgularda, mülteciler, oryantasyon aktiviteleri, toplum temelli uygulamalar, post-travmatik stres bozukluğu ve diğer psikolojik konular, eğitim materyalleri, hizmet verenlere özel yönergeler, psikolojik ilk yardım konusundaki burslar, ruh sağlığı ve psikososyal destek ve Avustralya’ya özgü hizmet modeli gibi konular raporlanmıştır. Bu tarz analizler, benzer konularda el kitapları, uygulama pratikleri ya da modeller geliştirmede ya da var olanları değerlendirmede ülkelere, kurum ya da kuruluşlara yol gösterici olabilir. Ayrıca, konu ile ilgili çalışanlara uygulamalarını gözden geçirme fırsatı sunabilir.

İkinci gruptaki araştırmalara örnek olarak Seah ve Shim (2018) tarafından, konu alanlarına göre çeşitli alt gruplarda bireylerin gönderiler paylaşp yorumlar yaptığı bir sosyal platform olan Reddit isimli web sayfasının, intihar ve depresyon ile ilgili Singapur alt grubundaki gönderi ve yorumları konu modellemesi ile analiz edilen araştırma ele alınabilir. Çalışmada 2010-2018 yılları arasında depresyon ve intiharla ilgili konularda yapılan 385 gönderi ve 21.000 yorum üzerinde analizler gerçekleştirilmiştir. Sonuçlara göre yaşam ve ölüme dair karar ve ötanazi, arkadaşlık, sosyal destek ve ilişkiler, okul yaşamı, beklentiler ve stres, toplumsal görüşler, kariyer, stres ve zorluklar, daha geniş bir düzeyde toplumdaki tutumlar, polise rapor edilen vakalar, deneyimleri paylaşma ve yorum yapma konusunda teşvikler, LGBT bireylerin kabulü ve inançlar, intihar girişimleri, ailevi sorunlar, geçmiş yaşantılar ve anılar ve tedavi ve yardım arayışı olmak üzere 14 konu öne çıkmıştır. Bu araştırma, görece daha rahat bir şekilde görüşlerini paylaşabilecekleri varsayılan bir sosyal platformda yardıma ihtiyacı olduğu varsayılacak bireylerin depresyon ve intihar konusundaki paylaşımlarında öne çıkan konuların belirlenmesi yoluyla bu konuda güçlük yaşayan bireyler için gerçekleştirilecek yardım çalışmalarının hangi alanları kapsayabileceği konusunda da fikir verebileceği gibi depresyon ve intiharı önleme konusunda gerçekleştirilebilecek çalışmalara da ışık tutabilir. Benzer şekilde, Zhang vd. (2021) COVID-19 döneminde evden çalışma tutum ve deneyimleri daha iyi anlayabilmek için 2020 yılında Mart ve Haziran ayları arasında Twitter’da “evden çalışma”, “uzaktan çalışma”, “telework” gibi anahtar kelimeleri kullanarak taradıkları bir milyondan fazla tivitisi konu modellemesi yoluyla analiz etmişlerdir. Sonuçlarda öne çıkan konular, evden çalışma, siber güvenlik, ruh sağlığı, iş-yaşam dengesi, takım çalışması ve liderlik gibi temaları içermiştir. Bu sonuçlar, çalışanların evden çalışma ile ilgili deneyimlerde güçlük yaşadıkları noktaları anlamak ve ihtiyaçlarını belirlemek ayrıca evden çalışmaya yönelik genel eğilimi ve tutuma dair fikir edinmek açısından önemli görülebilir. Bu grupta yer alabilecek ilgi çekici diğer bir araştırmada (Carron-Arthur vd., 2016) ise ruh sağlığı ile ilgili internetteki destek grubunda paylaşılan gönderiler konu modellemesi ile analiz edilmiştir. Ancak bu araştırmada bu tarz platformlarda %75 oranındaki gönderilerin aynı kişiler tarafından yapıldığına işaret edilerek bu kullanıcılar süper kullanıcılar olarak adlandırılmış ve konu modellemesi analizinde öne çıkan konuların süper kullanıcılar ile diğer kullanıcılara göre değişim gösterip göstermediği kay kare testi ile incelenmiştir. Buna göre süper kullanıcıların ebeveynlik rolleri, ruh sağlığı ve pozitif değişim gibi konularda gönderi yapma olasılığı yüksek bulunurken, depresyon, ilaçla tedavi, terapi ve anksiyete konularında daha az olduğu görülmüş ancak öne çıkan yedi konudan beşi hakkında gönderide bulunma olasılıklarının da yüksek olduğuna işaret edilmiştir. Bu araştırma özellikle sosyal medya ya da internet ortamında ruh sağlığı ile ilgili konularda gerçekleştirilen gönderi ve paylaşımlar üzerinde gerçekleştirilen konu analizlerinde, gönderi sayısı ya da sıklığı yerine sürekli gönderide bulunan %75’lik grup ile diğer gruplar arasındaki konu farklılıklarını incelemenin önemine işaret etmektedir. Bu bağlamda aynı ya da benzer yorumların sürekli aynı kişiler tarafından yazılması durumunun öne çıkan konu alanlarını şekillendirmedeki belirleyici rolünün farkında olarak daha detaylı incelemeler yapma gereğine işaret etmektedir.

Üçüncü gruptaki çalışmalara da benzer şekilde yukarıdaki sınıflama kapsamında örnekler verilmiştir (Bkz. Atzil-Slonim vd., 2019). Ek olarak, Atkins vd. (2012) diğerlerinin, çift ve aile terapisinde konu modellemesi analizinin kullanımını önerdikleri makale incelenebilir. Yazarlar, bu makalede, çiftlerle gerçekleştirilen deneysel bir çalışma kapsamında elde edilen terapi transkriptlerinin konu modellemesi ile analizini örnekler üzerinden tanıtmışlardır. Örneğin tüm transkriptler üzerinde gerçekleştirdikleri konu modellemesi analizi ile terapi süreci boyunca çiftler arasında öne çıkan konuları üç farklı modelde sunmuşlardır: içerik odaklı konular (aile, iş, para, cinsellik gibi), duygu odaklı konular (öfke, üzüntü,

incinmişlik, çaba gibi) ve terapi ile ilgili konular (iletişim, problem çözme, beyin fırtınası gibi). Konu modellemesi analizinin ayrıca çiftlerle gerçekleştirilen terapi süreci boyunca oturumlarda öne çıkan konulardaki değişimleri incelemek açısından nasıl kullanılabileceğini de örneklendirmişlerdir. Örneğin, transkriptlerini inceledikleri bir çiftin terapi sürecinde iletişim odaklı konuların bir blok şeklinde öne çıktığı görülürken dört ile sekizinci oturumlar bu konuların yaygınlığının arttığı ve bu oturumlardan sonra konu odağının problem çözmeye doğru değiştiği aktarılmıştır. Son olarak dördüncü grupta, yukarıda örnekleri verilen, ölçme değerlendirme ve bireyi tanıma kapsamında açık uçlu soruların konu modellemesi ile analizini içeren uygulamalar örnek gösterilebilir (Bkz. Finch vd., 2021; Rohrer vd., 2017).

Kullanılan konu modellemesi analizi yöntemine göre de 1) eğitilmiş (supervised) konu modelleme, (2) eğitilmemiş/serbest (unsupervised) konu modelleme ve (3) geniş dil modellerine dayalı üretken yapay zekâ kullanan araştırmalar olarak da bir sınıflama yapılabilir. Bu makalede farklı konu modelleme analizi yöntemlerinden ziyade kullanım amacının öne çıkarılması benimsendiği için konu modellemesi analizlerinin farklı amaçlar için nasıl kullanıldığı esas alınarak bir sınıflama yapılmaya çalışılmıştır. Bu noktada, konu modellemesi analizinin psikolojik danışma ve ilgili ruh sağlığı alanyazınında ilgi çekici ve gün geçtikçe gelişen bir kullanımı söz konusu olsa da bu süreçte bazı sınırlılıkların ve etik konuların öne çıktığı söylenebilir.

KONU MODELLEMESİ ANALİZİNİN KULLANIMINDA SINIRLILIKLAR VE ETİK HUSUSLAR

İlgili alanyazın incelendiğinde konu modellemesi analizinin kullanımında göz önünde bulundurulması gereken sınırlılıklara işaret edildiği görülmektedir (Banks vd., 2018; Bittermann & Fischer, 2018; Kobayashi vd., 2018; Liu & Gao, 2021; Oh vd., 2017; Otsuka vd., 2021). İlk olarak, konu modellemesinin sonuçları konu modellemesi analizine tabi tutulan veri setine oldukça bağlıdır. Bu nedenle, psikolojik danışma alanında araştırma yaparken literatürün yeterince taranmış olup olmaması, dahil edilme ve hariç tutulma gibi kriterler belirlendiyse bunların neye göre belirlendiği ve hangi çalışmaların analize dahil edildiği, internet ortamındaki paylaşım ve gönderilerin sıklığı veya doğru olmama olasılığı ya da yetersiz transkriptler gibi durumlardan oldukça etkilenme olasılığı vardır. Verilerin kalitesinin sonuçları önemli bir şekilde etkileyebileceği göz önünde bulundurulmalıdır. Konu modelleme analizi, doğru konu ve tema dağılımları ortaya koyabilmek için geniş bir veri setini gerekli kılmaktadır. Bu yüzden küçük veri seti (ör., birkaç danışma transkripti gibi) analiz için uygun olmayacaktır. Özellikle yayınlanan makalelerin sadece özetlerini analiz etme yoluna giden konu modellemesinde, bulguların ortaya koyduğu konular özetlerin sağladığı bilgilerle sınırlı olacaktır.

Ayrıca konu modellemesi analizinde sistemi eğitmek için kullanılan mevcut sınıflandırma sistemlerinin kullanılması, araştırılan alandaki son konuları ya da öne çıkan konuları belirlemede yetersiz ve eksik kalabilir. Veri setinin büyüklüğünden bağımsız olarak, konu modellemesi analizi sonuçlarında bazen bir konunun kapsadığı terimler birden fazla temayı yansıtabilir ve bu durum tutarlılık ve bütünlük açısından soru işaretleri oluşturabilir. Konu modellemesi analizinden elde edilen bulguların geçerliğini sağlamak için belirli ölçütlerin belirlenmesi, veri analizinde çeşitlemeye gitme, uzman görüşüne başvurma gibi çeşitli stratejilerin belirlenmesi önerilmektedir. Bu noktada bu veri analizi yönteminin kendisinden kaynaklı bir diğer sınırlılık ise sürecin oldukça teknolojiye bağımlı olmasıdır. Araştırmacılar sadece makine öğrenmesinin ortaya koyduğu bulgulara bağlı kaldıklarında, daha geniş bir perspektiften ve eleştirel bir şekilde konuya yaklaşımdan uzaklaşabilirler. Ayrıca, teknolojiye bağlı bir veri analizi yöntemi olması veri madenciliği analizlerine dair bir uzmanlığa sahip olmayı gerekli kılmaktadır, bu durum da psikolojik danışma oturumlarına dair transkriptlerin yorumlanarak psikolojik danışma sürecini ve sonuçlarını değerlendirme ya da bir kurumda çalışan ihtiyaçlarını ya da geri bildirimlerini değerlendirme gibi amaçlarla konu modellemesinin kullanımında bu yöntemlere hâkim olmayan uygulayıcılar açısından bir sınırlılık getirmektedir.

Bulgularla ilgili olası diğer bir sınırlılık ise, konu modellemesi analizi ile, özellikle psikolojik danışma transkriptleri ya da sosyal paylaşım sitelerindeki gönderiler analiz edildiğinde, bağlamı değerlendirme, nüansları fark etme ve bireysel deneyimlerdeki karmaşıklığı tam olarak kapsayamama olasılığıdır. Dolayısıyla konu modellemesi analizi, metin içeren verideki temaları ve örüntüleri ortaya çıkarabilirken bu verinin elde edildiği bağlamı ve psikolojik olguların karmaşıklığını tam yansıtmayabilir. Bir başka sınırlılık bulguların yorumlanması ile ilgilidir. Bulguların ortaya koyduğu konular ve konular altında temsil edilen terimlerin yorumlanması araştırmacılar tarafından yapılmaktadır. Bu durum ise araştırmacıların kişisel görüşleri, yargıları ya da önyargıları gibi öznel süreçlerden etkilenme riski taşımaktadır. Bu nedenle, bulgular yorumlanırken dikkatli bir inceleme ve değerlendirme gerektirmektedir. Son olarak, özellikle

psikolojik danışma transkriptlerinin kullanıldığı konu modellemesi analizlerinde bulguların sadece incelenen bağlamla ilgili konular ve örüntüler ortaya koyma olasılığı vardır. Bu nedenle bu bulguların daha geniş bir psikolojik danışma bağlamına ya da diğer gruplara genellenebilirliğinde sınırlılık içerecektir.

Bu sınırlılıklara yönelik bir önlem olarak, psikoloji, psikolojik danışmanlık ve rehberlik gibi insan odaklı araştırmalar içeren alanlarda konu modellemesi analizine dayalı uygulamaların ve araştırmaların planlanması sürecinde etik konuların hem alanın kendi etik ilkeleri çerçevesinde hem de yayın etiği çerçevesinde düşünülmesi önerilebilir. Bu doğrultuda, araştırmacılar, ilk olarak, kullanılan veri kaynakları ister akademik makaleler, raporlar ya da dokümanlar gibi çeşitli yayınlar olsun ister sosyal medya ve internetteki çeşitli paylaşım platformlarından elde edilen paylaşımlar ve gönderiler olsun, toplanan içeriklerin önyargısız, tarafsız, adil bir temsil sağlayacak şekilde bir araya getirildiğinden emin olmalıdırlar. Bu adım, analizler sonucunda ortaya çıkacak konuların içeriklerini, sıklıklarını, dağılımlarını vb. şekillendirecek bir öneme sahiptir. Dolayısıyla bu işlemde etkilenecek bulguların, konuya ilişkin tespit ya da saptamalarda ve uygulamaya yönelik içeriklerin oluşturulmasında kritik öneme sahip olacağı hatırlanmalıdır. Akabinde, bulguların yorumlanmasında şeffaflık esas alınmalı; veri analizi yöntemindeki olası sınırlılıkların ve analiz sürecinde araştırmacının rolünün ve duruşunun (konuya olan ilgisi, olası önyargıları gibi) açık bir şekilde aktarılması göz önünde bulundurulmalıdır. Ayrıca, araştırmanın araştırılmaya değer olup olmadığı, araştırma bulgularının belirli gruplar, bireyler ya da toplumlar üzerindeki etkileri göz önünde bulundurulmalıdır. Örneğin, spesifik olarak bir grubu (ör., LGBTI bireyleri) ya da toplumu etiketleme potansiyeli içeren konuların (ör., taşıdıkları hastalıklar ya da ruh sağlığı problemleri gibi) ortaya çıkarılmasının ya da araştırmanın bu gibi yorumlara neden olabilecek şekilde tasarlanmasının toplumsal yararlılık ilkesi ile örtüşmeyeceği hatırlanmalıdır. Danışanlara ait daha özel veri kaynaklarının (transkriptler gibi) kullanılması durumunda bilgilendirilmiş onam alınması, katılımcıların gizliliğinin sağlanması, danışanları ya da bir grubu yansıtacak tanımlayıcı bilgilerin transkriptlerden çıkarıldığından emin olunması ve mahremiyetin sağlanması ve son olarak tüm veri kaynakları için veri güvenliğinin sağlanması gibi sorumlu bir kullanım son derece önem arz etmektedir.

SONUÇ VE GELECEK ARAŞTIRMALAR İÇİN ÖNERİLER

Psikolojik danışma alanında veri analizi yöntemi olarak konu modellemesi kullanan araştırmalar, bulguları açısından kuramsal temellere, araştırma ve uygulamaya yönelik önemli katkılar sunma potansiyeli barındırmaktadır. Öncelikle, yukarıda verilen örneklerde olduğu gibi, psikolojik danışma alanındaki önemli gelişmeler, yükselen eğilimler, araştırma ve uygulamanın yoğunlaştığı alanlar konu modellemesi yoluyla analiz edilebilir. Bu inceleme, konu dağılımına bakılarak durum saptaması yapılmasına, incelenen konunun çeşitli boyutlarının daha iyi anlaşılmasına, uygulamalarda hangi bileşenlere yer verilmeli konusunun değerlendirilmesine ve var olan uygulamaların ve politikaların gözden geçirilmesine katkı sağlayabilir. Bu konuda, şu ana kadar yapılan araştırmalara bakıldığında okul psikolojik danışmanlığı, kariyer psikolojik danışmanlığı ve aile ve evlilik danışmanlığı alanında bu araştırmaların oldukça az olduğu görülmektedir. Örneğin evlilik doyumu ile ilgili araştırmalardaki trend konuların ortaya çıkarılması, konu dağılımlarının yıllar içindeki değişimlerine bakılması, boşanma oranları ile karşılaştırılması gibi araştırmalar bu konuda yapılacak önleyici ve destekleyici çalışmalara anlamlı girdiler sağlayabilecektir. Veri madenciliğine dayanan bu analiz tüm dünyadaki araştırmaları ele alarak yapılabileceği gibi sadece ülkemizdeki araştırmalar özelinde incelenerek daha yerel ve kültüre duyarlı bulgularla uygulamaya yönelik daha spesifik sonuçlara ulaşılabilir. Benzer şekilde, iş doyumu, örgütsel bağlılık, mesleki tükenmişlik gibi konularda yapılacak araştırmalar, kariyer psikolojik danışmanlığı alanına önemli katkılar sağlayabilir. Bu araştırmalardan elde edilen bulgular iş verimliliğini değerlendirmek ve ekonomik gelişimi desteklemek açısından mevcut işleyişlerin gözden geçirilmesine ve politikalar geliştirilmesine katkı sağlayabilir. Daha spesifik konulara odaklanılarak da konu modellemesi analizi gerçekleştirilebilir. Örneğin, Karacan-Ozdemir vd. (incelemede) STEM kariyerlerinde başarılı olabilmek için gerekli olan becerileri konu modellemesi analizi ile incelemişlerdir. Bu araştırmanın bulguları, öğrencilerin STEM kariyer gelişimlerini desteklemeye yönelik programların tasarlanmasına ve var olan programların değerlendirilmesine yönelik katkılar sunabilecektir. Benzer şekilde, ülkemizde okul psikolojik danışmanlığı alanında günümüze kadar yapılan tüm araştırmalar konu modellemesi ile incelenebilir ve yıllar içindeki konu dağılımlarındaki değişimler, toplumsal krizler, gelişmeler ve değişimler ışığında yorumlanarak geleceğe yönelik çıkarımlar yapılabilir ve bu doğrultuda kapsamlı gelişimsel rehberlik programlarının tasarımına yönelik önemli girdiler sağlanabilir. Daha spesifik bir öneri olarak, okullarda krize müdahale ile ilgili hazırlanmış el kitapları ve uygulayıcı kılavuzları konu modellemesi ile analiz edilebilir ve öne çıkan konu dağılımlarına göre ülkemizdeki yaygın olarak kullanılan krize müdahale programlarının içerikleri gözden geçirilebilir.

Konu modellemesi aynı zamanda aşağıdan yukarıya doğru bir yaklaşımla psikolojik danışma alanındaki araştırma eğilimlerini ortaya koymak ve gelişmekte olan psikolojik araştırmalara dair gündemi yakalamak için kullanılabilir (Bittermann & Fisher, 2018). Ayrıca, sistematik literatür taraması çalışmalarına entegre edilebilir ve bu yolla literatür tarama çalışmaları genişletilebilir (Bittermann ve Fisher, 2018; Chen & Wojcik, 2016). Bulgular, yeni araştırmaların tasarlanmasında araştırmacılara yol gösterebilir. Ayrıca, hangi konu alanlarının daha gündemde olduğu ve öne çıktığı, hangi konu alanlarının ise ihmal edildiği gibi noktalara dair farkındalıklar sağlanabilir. Dahası, yeni araştırma sorularının ve hipotezlerin geliştirilmesine katkı sağlayabilir (Chen & Wojcik, 2016).

Konu modellemesi analizinin potansiyel kullanım alanlarına yönelik gelecek araştırmalar için bir diğer öneri, yukarıda verilen örneklerde olduğu gibi, Facebook paylaşımları, tivitler gibi sosyal medya ve internetteki gönderi, paylaşım ve geribildirimlerin incelenmesine yöneliktir (Banks vd., 2018). Bu doğrultuda yapılacak araştırmalarda da psikolojik danışma ve rehberlik alanı açısından oldukça yüksek potansiyeller bulunmaktadır. Örneğin, bilgisayar oyunları ile ilgili spesifik internet sitelerinde ve platformlardaki yorumları ve gönderileri konu analizi ile incelemek ve dahası yaş gruplarına göre konulardaki değişimleri ortaya çıkarmak bilgisayar oyunlarına yönelik bağımlılığı anlamak, açıklamak ve farklı yönleri ile değerlendirmek açısından anlamlı katkılar sağlayabilir. Bulgular, öngörülemeyen risk faktörlerini değerlendirmede, (çocuklar, ergenler ve yetişkinler gibi) farklı gelişim dönemlerine özgü bilgisayar bağımlılığı eğilimlerini belirlemede ve bu doğrultuda yardım süreçlerini planlamada ya da var olanları değerlendirmede önemli girdiler sunabilecektir. Benzer şekilde, yükselen bir sorun olan kumar bağımlılığını anlamaya yönelik iddia sitelerindeki kullanıcı yorumlarını ve gönderilerini konu analizi ile incelemek bu konuda yapılacak çalışmalar açısından önemli katkılar sağlayabilecektir.

Psikolojik danışma ve rehberlikte konu modellemesi analizi, kültürlerarası araştırmalar açısından da oldukça zengin ve içgörü sağlayıcı bulgular üretebilme potansiyeli taşımaktadır. Bu araştırmalar, farklı kültürlerde hangi konuların öne çıktığının ve kültüre göre ne gibi farklılıkların oluştuğunun anlaşılmasına katkı sağlaması ve kültüre duyarlı uygulamalar geliştirilmesi açısından anlamlı katkılar sağlayabilir (Liu ve Gao, 2021). Örneğin, Otsuka vd. (2021), psikolojik araştırmalarda hangi konuların evrensel olduğunu, hangi konuların ise farklı coğrafi bölgelere ve tarihsel süreç içinde nasıl değiştiğini incelemişlerdir. Ana bulgular arasında psikolojideki araştırma faaliyetlerinin ve eğilimlerinin hem coğrafi bölgelerden hem de zaman eğilimlerinden güçlü bir şekilde etkilendiğini rapor etmişlerdir. Buna göre, klinik araştırmalarla ilgili konuların Avrupa'da Kuzey-Orta Amerika'dan daha yüksek oranda olduğunu ortaya koymuşlardır. Kültürler arası araştırmaya bir başka örnek, bu makalenin yazarının da içinde bulunduğu sosyal duygusal öğrenme becerileri konusunda araştırma yapan çok kültürlü bir araştırma grubunun çalışmaları verilebilir. Araştırma ekibi, karar verme becerilerinde hem evrensel olan temaları hem de kültüre göre değişen temaları ortaya çıkarmak için, eğitimcilerden toplanan nitel verileri konu modellemesi ile analiz etmişlerdir. Öncül bulgular arasında, Türkiye'de, öne çıkan temalar ve terimler arasında diğer kültürlerden farklı olarak duygu ifadelerinin (öfke, baskı hissetme gibi) öne çıktığı görülmüştür (Kanzaki vd., incelemede). Bu örneklerde olduğu gibi, psikolojik danışma ve rehberlik alanında bulguları anlamlı katkılar sağlayacak pek çok konu, konu modellemesi analizi ile incelenebilir. Örneğin, psikolojik danışma sürecinin etkililiğini belirleyen unsurlar, terapötik ilişkide öne çıkan konular, etkili bulunan psikolojik danışma yöntem ve yaklaşımları, danışanların psikolojik danışma sürecini yarıda bırakma ve terk etme durumlarında öne çıkan temalar gibi konular, konu modellemesi ile incelenebilir ve hem evrensel temalar hem de kültüre özgü konular ortaya konabilir. Böylece, Batı odaklı psikolojik danışma kuram ve yaklaşımlarının kültüre özgü uygulanmasına yönelik anlamlı içgörüler kazanılabilir. Bu örnekleri çoğaltmak mümkündür.

Konu modellemesinin psikolojik danışmada bir diğer kullanım alanı, ölçme ve değerlendirmeye yönelik uygulamalar dahilinde olabilir. Konu modellemesi, öntest ve sontest arasında konularda çıkan farklılıkları incelemek yoluyla psikolojik danışma ve rehberlik uygulamalarının etkililiğini değerlendirmede kullanılabilir (Liu & Gao, 2021). Örneğin, okullarda uygulanan gelişimsel sınıf rehberliği programları, akran zorbalığı programları, psikolojik sağlamlıkla ilgili programlar gibi uygulamaların etkililiğini değerlendirmede uygulayıcılardan ve/veya yararlanıcılardan alınacak geribildirimlerde öne çıkan konu dağılımları incelenebilir. Bu değerlendirmeler, mevcut programların etkililiğine dair sağlayacağı bilgilerin yanı sıra programların iyileştirilmesinde ve gelecek programların geliştirilmesinde anlamlı girdiler sağlayabilir. Benzer şekilde, örneğin ülke çapında uygulanan okul temelli programların değerlendirilmesinden elde edilen bulgular, bu konularla ilgili politikaların geliştirilmesinde önemli katkılar sağlayabilir (Oh vd., 2017). Bir başka örnek olarak, okul psikolojik danışmanlığı çalışmaları kapsamında, öğrencilerin gizlilik ve mahremiyet açısından

anonimliği sağlanmış otobiyografilerinin ya da uygulanan çeşitli anketlerde yer alan açık uçlu sorulara verdikleri yanıtların konu modellemesi ile analiz edilmesi olabilir. Böylece, okuldaki ihtiyaçlar ve sorun alanlarının belirlenmesine yönelik önemli bulgulara erişilebilir. Ayrıca, öğrencilerin otobiyografilerindeki temaların yıllara göre değişimini incelemeye yönelik boylamsal araştırmalar, gelişimsel ihtiyaçlar ve farklılıkların anlaşılmasına katkı sağlayabilir. Benzer şekilde, geriye dönük olarak, aynı düzeydeki (örn., 7. sınıf düzeyinde) öğrenci otobiyografileri analiz edilerek yıllar ilerledikçe aynı yaş grubunda öne çıkan temalarda değişimler var mı, varsa ne yöne doğru gibi incelemeler yapılabilir ve geleceğe yönelik kestirimlerde bulunulabilir. Kariyer psikolojik danışmanlığı ve örgütsel psikoloji alanlarında kullanımına örnek olarak, bir kurumdaki çalışan memnuniyeti ya da geribildirimlerinin, yıllık raporlar, stratejik planlar, toplantı tutanakları gibi kurumsal dokümanların konu analizi ile incelenmesi verilebilir. Buradan elde edilecek sonuçlar, kurumlara, çalışan ihtiyaçlarını değerlendirmede ve gelişim alanlarını belirlemede olduğu kadar kurumun güçlü ve zayıf yönlerini değerlendirmede de fayda sunabilir (Banks vd., 2018).

Konu modellemesi psikolojik danışman eğitiminin geliştirilmesinde de kullanılabilir. Yukarıdaki örneklerde ve önerilerde detaylandırıldığı üzere, psikolojik danışma ve rehberlik alanındaki araştırmaların ve konuların konu modellemesi analizi ile incelenmesi sonucu öne çıkan konular, psikolojik danışman eğitiminde kullanılan müfredatın gözden geçirilmesinde ve ele alınmayan, gözden kaçan konuların olması durumunda programların bu doğrultuda güncellenmesinde kullanılabilir. Bu, psikolojik danışman eğitiminin güncel araştırma eğilimleri ve iyi uygulamalarla uyumlu olacak şekilde yapılandırılmasına yardımcı olabilir (Oh vd., 2017). Örneğin, etkili psikolojik danışma yaklaşımları ve uygulamalarında öne çıkan konuların belirlenmesi, terapötik ilişkinin kültüre özgü belirleyicilerinin ortaya konması gibi araştırma bulguları psikolojik danışma uygulamaları derslerinde ele alınabilir ve öğrencilere bu ders kapsamında verilen süpervizyon süreçlerinde göz önünde bulundurulabilir. Öte yandan, gelecek araştırmalarda, su ana kadar verilen tüm önerilerde öne çıkabilecek etik konular ve gereklilikler de göz önünde bulundurularak, konu modellemesinin psikolojik danışma ve rehberlikte kullanımına dair etik ilkeler ve standartlar oluşturulabilir (Liu & Gao, 2021).

Konu modellemesinin, öncesinde değinildiği gibi, bir veri analizi yöntemi olarak sınırlılıkları göz önünde bulundurulduğunda diğer veri analizi yöntemleriyle entegre edilerek kullanılması ya da farklı konu modellemesi analizlerinin harmanlanarak kullanılması yönünde önerilerde bulunulabilir. İlk olarak, psikolojik danışma ve rehberlik alanında konu modellemesi analizi, nitel araştırma desenlerinden gömülü desen kapsamında gerçekleştirilebilir (Banks vd., 2018) ve/veya nitel veri analizi yöntemlerinden içerik analizi ya da duygu analizi gibi yöntemlerle (Kobayashi vd., 2018) bir arada kullanılarak gerçekleştirilebilir.

Psikolojik danışma ve rehberlikte öne çıkan konuları belirlemeye yönelik araştırmalarda (örn., Oh vd., 2017) olduğu gibi, tarihsel süreç içinde konuların değişiminin incelenmesi ve bu incelemelerde konu yapılarının değişken doğasını yakalanması ve daha incelikli bir anlayış kazanılması için dinamik konu modellemesi gibi daha gelişmiş konu modellemesi analizleri yapılabilir (Kobayashi vd., 2018). Ayrıca, konu modellemesinin boylamsal veriler ve ağ verileri gibi daha karmaşık verilerin işlenebileceği şekilde geliştirilmesi yönünde çalışmalar yapılabilir (Otsuka vd., 2021). Dahası, kelime gömme ve derin öğrenme gibi nöral olasılık yöntemlerinin geliştirilmesi yönünde çalışmalar hızlandırılabilir (Banks vd., 2018). Gelecekteki araştırmalar, metin analizinin ötesinde, konu modellemesinin uygulanabilirliğini genişletmek ve analitik derinliği artırmak için ses ve video gibi alternatif veri kaynaklarıyla entegrasyonunu keşfedebilir (Kobayashi ve ark., 2018). Son olarak, sınırlılıklar başlığı altında belirtildiği üzere, konu modellemesi analizi ile ortaya konan sonuçların geçerliğinin değerlendirilmesinde doğrulanmış konu modellerinin geliştirilmesi yönünde çalışmalar yapılabilir; konu modellerinin kalitesini değerlendirmek için belirlenmiş kriterlerin kullanılarak ortaya konan konuların doğrulanması yönünde standartlaştırılmış prosedürler geliştirilebilir (Otsuka vd., 2021).

Sonuç

Bu makalede, hızlı teknolojik gelişmeler altında makine öğrenmesine dayalı metin madenciliği veri analizi yöntemlerinden biri olan konu modellemesinin psikolojik danışma ve rehberlik alanında kullanımı ele alınmıştır. Konuyla ilgili yapılan çalışmaların sınırlı olduğu ancak psikolojik danışma ve rehberlikte kullanım alanı açısından oldukça zengin potansiyeller taşıdığı söylenebilir. Özellikle, psikolojik danışma ve rehberliğin farklı uzmanlık alanları altında gerçekleştirilecek nitel araştırmalar kapsamında gerçekleştirilecek konu modellemesi analizleri, bulguları açısından kuramsal bilgi birikimine, araştırmaya, uygulamaya ve ilgili konular yönünde politikalar geliştirilmesine yönelik sağlayacağı girdilerle alanın gelişimine önemli katkılar

sağlayabilir. Bu bağlamda, veri analizi yöntemin kendisinden kaynaklanan sınırlılıklar kadar araştırma sürecinin doğasına yönelik sınırlılıklar ve konuya ilişkin etik hususlar göz önünde bulundurularak, psikolojik danışma ve rehberlik alanının gelişimine yönelik katkı sağlayabilecek araştırmaların yapılmasına ihtiyaç duyulmaktadır.

KAYNAKÇA

- Atkins, D. C., Rubin, T. N., Steyvers, M., Doeden, M. A., Baucom, B. R., & Christensen, A. (2012). Topic models: A novel method for modeling couple and family text data. *Journal of Family Psychology*, 26, 816–827. <https://doi.org/10.1037/a0029607>
- Atzil-Slonim, D., Juravski, D., Bar-Kalifa, E., Gilboa-Schechtman, E., Tuval-Mashiach, R., Shapira, N., & Goldberg, Y. (2021). Using topic models to identify clients' functioning levels and alliance ruptures in psychotherapy. *Psychotherapy (Chicago, Ill.)*, 58(2), 324–339. <https://doi.org/10.1037/pst0000362>
- Banks, G. C., Woznyj, H. M., Wesslen, R. S., & Ross, R. L. (2018). A review of best practice recommendations for text analysis in R (and a user-friendly app). *Journal of Business and Psychology*, 33(4), 445–459. <https://doi.org/10.1007/s10869-017-9528-3>
- Balahadia, F. F., Fernando, M. C. G., & Juanatas, I. C. (2016). Teacher's performance evaluation tool using opinion mining with sentiment analysis. In IEEE region symposium (TENSYP) (pp. 95–98). <https://doi.org/10.1109/TENCONSpring.2016.7519384>
- Baumer, E. P., Mimno, D., Guha, S., Quan, E., & Gay, G. K. (2017). Comparing grounded theory and topic modeling: Extreme divergence or unlikely convergence? *Journal of the Association for Information Science and Technology*, 68, 1397–1410. <https://doi.org/10.1002/asi.23786>
- Buenano-Fernandez, D., Gonzalez, M., Gil, D., & Lujan-Mora, S. (2020). Text mining of open-ended questions in self-assessment of university teachers: An LDA topic modeling approach. *IEEE Access*, 8, 35318–35330. <https://doi.org/10.1109/ACCESS.2020.2974983>
- Blei, D. M. (2012). Probabilistic topic models. *Communications of the ACM*, 55(4), 77-84. <https://doi.org/10.1145/2133806.2133826>
- Blei, D. M., & Jordan, M. I. (2003, July). Modeling annotated data. In *Proceedings of the 26th annual international ACM SIGIR conference on Research and development in informaion retrieval* (pp. 127-134).
- Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent dirichlet allocation. *Journal of machine Learning research*, 3(Jan), 993-1022.
- Blei, D., & Lafferty, J. (2009). Topic Models. In A. Srivastava & M. Sahami (Eds.), *Text Mining: Classification, Clustering, and Applications* (pp. 71-94). Boca Raton, FL: Chapman & Hall/CRC.
- Bittermann, A., & Fischer, A. (2018). How to identify hot topics in psychology using topic modeling. *Zeitschrift für Psychologie*, 226(1), 3–13. <https://doi.org/10.1027/2151-2604/a000318>
- Carron-Arthur, B., Reynolds, J., Bennett, K., Bennett, A., & Griffiths, K. M. (2016). What's all the talk about? Topic modelling in a mental health Internet support group. *BMC Psychiatry*, 16(1), 367. <https://doi.org/10.1186/s12888-016-1073-5>
- Chen, E. E., & Wojcik, S. P. (2016). A practical guide to big data research in psychology. *Psychological methods*, 21(4), 458–474. <https://doi.org/10.1037/met0000111>
- Feinerer, I., & Wild, F. (2007). Automated coding of qualitative interviews with latent semantic analysis. In Heinrich C. Mayr, Dimitris Karagiannis (Ed.), *"Lecture notes in informatics" (LNI): P-107, Information Systems Technology and its Applications, May 23-25, 2007, Kharkiv, Ukraine* (pp. 66 - 78). Köllen Druck+Verlag GmbH.
- Finch, W. H., Hernández Finch, M.E.H., French, B. F., & Claire, B. (2021). Latent Dirichlet Analysis to Aid Equity in the Identification for Gifted Education. *Psychological Test and Assessment Modeling* 63(3),

- 305-36. https://www.psychologie-aktuell.com/fileadmin/Redaktion/Journale/ptam-2021-3/PTAM_3-2021_2_kor.pdf
- Franz, P. J., Nook, E. C., Mair, P., & Nock, M. K. (2020). Using Topic Modeling to Detect and Describe Self-Injurious and Related Content on a Large-Scale Digital Platform. *Suicide & life-threatening behavior*, 50(1), 5–18. <https://doi.org/10.1111/sltb.12569>
- Gencoglu, B., Helms-Lorenz, M., Maulana, R., Jansen, E. P., & Gencoglu, O. (2023). Machine and expert judgments of student perceptions of teaching behavior in secondary education: Added value of topic modeling with big data. *Computers & Education*, 193, 104682. <https://doi.org/10.1016/j.compedu.2022.104682>.
- Griffiths, T. L., & Steyvers, M. (2004). Finding scientific topics. *Proceedings of the National Academy of Sciences of the United States of America*, 101(1 Supplement), 5228–5235.
- Grimmer, J., & Stewart, B. M. (2013). Text as data: The promise and pitfalls of automatic content analysis methods for political texts. *Political Analysis*, 21(3), 267–297. <https://doi.org/10.1093/pan/mps028>
- Hopkins, D. J., & King, G. (2010). A method of automated nonparametric content analysis for social science. *American Journal of Political Science*, 54(1), 229–247. <https://doi.org/10.1111/j.1540-5907.2009.00428.x>
- Howes, C., Purver, M., & McCabe, R. (2013). Howes, C., Purver, M., & McCabe, R. (2013). Investigating Topic Modelling for Therapy Dialogue Analysis. Proceedings of IWCS 2013 Workshop on Computational Semantics in Clinical Text. <https://qmro.qmul.ac.uk/xmlui/handle/123456789/5300>
- Jayarathne M., & Jayatilake B. (2020). Predicting personality using answers to open-ended interview questions. *IEEE Access*, 8, 115345–115355. 10.1109/ACCESS.2020.3004002
- Jivani, A. G. (2011). A comparative study of stemming algorithms. *Journal of Computer Applications in Technology*, 2(6), 1930-1938.
- Kanzaki, M. Suzuki, H., Abdoulaye, O., Agnieszka O., Agnieszka, S., Andrei, A., Bacanlı, F., Donnelly, H., Ferrari, L., Karacan-Ozdemir, N., Kourmoussi, N., Kounenou, K., Marsay, G., Park, C., Scoda, A.D., Solberg, V. S. H., & Zuzanna Z. (incelemede). Cultural translation of career decision-making in adolescents. *International Journal for Educational and Vocational Guidance*.
- Karacan-Ozdemir, N., Park, C., & Solberg, S. (incelemede). STEM career competencies. A qualitative framework. *Journal of Career Development*.
- Kadhim, A. I. (2018). An evaluation of preprocessing techniques for text classification. *International Journal of Computer Science and Information Security (IJCSIS)*, 16(6), 22-32.
- Kandula, S., Curtis, D., Hill, B., & Zeng-Treitler, Q. (2011). *Use of topic modeling for commending relevant education material to diabetic patients*. In Annual symposium proceedings/AMIA symposium (pp. 674–682).
- Kee, Y. H., Li, C., Kong, L. C., Tang, C. J., & Chuang, K. (2019). Scoping review of mindfulness research: A topic modelling approach. *Mindfulness*, 10, 1474–1488. <https://doi.org/10.1007/s12671-019-01136-4>
- Kobayashi, V. B., Mol, S. T., Berkers, H. A., Kismihók, G., & Den Hartog, D. N. (2018). Text Mining in Organizational Research. *Organizational Research Methods*, 21(3), 733-765. <https://doi.org/10.1177/1094428117722619>
- Krippendorff, K. (2003). *Content Analysis: An Introduction to Its Methodology* (2nd ed.). Thousand Oaks, CA: Sage Publications, Inc.
- Kumar, A. & Paul, A. (2016). *Mastering Text Mining with R*. Birmingham: Packt Publishing.

- Liu, J., & Gao, L. (2021). Analysis of topics and characteristics of user reviews on different online psychological counseling methods. *International journal of medical informatics*, 147, 104367. <https://doi.org/10.1016/j.ijmedinf.2020.104367>
- Liu, S., Zhang, R.Y., & Kishimoto, T. (2020). Analysis and prospect of clinical psychology based on topic models: hot research topics and scientific trends in the latest decades. *Psychol Health Med*, 26(4), 395-407. <https://doi.org/10.1080/13548506.2020.1738019>
- Longo, L. (2020). Empowering qualitative research methods in education with artificial intelligence. In: Costa, A., Reis, L., Moreira, A. (Eds) Computer supported qualitative research. WCQR 2019. Advances in Intelligent Systems and Computing, vol 1068. Springer, Cham. https://doi.org/10.1007/978-3-030-31787-4_1
- Maier, D., Waldherr, A., Miltner, P., Wiedemann, G., Niekler, A., Keinert, A., Pfetsch, B., Heyer, G., Reber, U., Häussler, T., Schmid-Petri, H. ve Adam, S. (2018). Applying lda topic modeling in communication research: toward a valid and reliable methodology. *Communication Methods and Measures*, 12(3), 93–118. <https://doi.org/10.1080/19312458.2018.1430754>
- Mobin, I., Mridha, M. F., & Mahmud, S. M. H. (2024). Exploratory analysis of suicidal tendency in depression investigation social media post. Research Square. <https://doi.org/10.21203/rs.3.rs-3823798/v1>
- Newman, M. E. (2005). Power laws, Pareto distributions and Zipf's law. *Contemporary Physics*, 46, 323–351.
- Ni, C. F., Lundblad, R., Dykeman, C., Bolante, R., & Łabuński, W. (2023). Content analysis of psychological first aid training manuals via topic modelling. *European journal of psychotraumatology*, 14(2), 2230110. <https://doi.org/10.1080/20008066.2023.2230110>
- Nowlin, M.C. (2015). Modeling issue definitions using quantitative text analysis. *Policy Studies Journal*, 44(3), 309-331. <https://doi.org/10.1111/psj.12110>
- Roberts, M. E., Stewart, B. M., Tingley, D., Lucas, C., Leder-Luis, J., Gadarian, S. K., ... & Rand, D. G. (2014). Structural topic models for open-ended survey responses. *American journal of political science*, 58(4), 1064-1082. <https://doi.org/10.1111/ajps.12103>
- Rohrer, J. M., Brümmer, M., Schmukle, S. C., Goebel, J., & Wagner, G. G. (2017). What else are you worried about? – Integrating textual responses into quantitative social science research. PLoS ONE 12(7), e0182156. <https://doi.org/10.1371/journal.pone.0182156>
- Royston, R. P., Lin, N., & Amey R. C. (2022). Applying Topic Modeling to narrative statements to determine how job responsibility descriptions differ by leadership level and demographic factors. Society for Industrial and Organizational Psychology Annual Conference APR 27-30
- Seah, J. H., K., & Shim, K. J. (2018). Data mining approach to the detection of suicide in social media: A case study of Singapore. 2018 IEEE International Conference on Big Data (Big Data): Seattle, WA, December 10-13: Proceedings. 5442-5444. https://ink.library.smu.edu.sg/sis_research/4340
- Tang, J., Meng, Z., Nguyen, X., Mei, Q., & Zang, M. (2014). Understanding the limiting factors of topic modeling via posterior contraction analysis. *Proceedings of the 31st International Conference on Machine Learning*, 337–345.
- Oh, J., Stewart, A., & Phelps, R. (2017). Topics in the journal of counseling psychology, 1963–2015.. *Journal of Counseling Psychology*, 64(6), 604-615. <https://doi.org/10.1037/cou0000218>
- Otsuka, S., Ueda, Y., & Saiki, J. (2021). Diversity in psychological research activities: quantitative approach with topic modeling. *Front Psychol*, 16(12), 773-916. <https://doi.org/10.3389/fpsyg.2021>
- Vijayarani, S., Ilamathi, M. J., & Nithya, M. (2015). Preprocessing techniques for text mining-an overview. *International Journal of Computer Science & Communication Networks*, 5(1), 7-16.

- Wang, S. H., Ding, Y., Zhao, W., Huang, Y. H., Perkins, R., Zou, W., & Chen, J. J. (2016). Text mining for identifying topics in the literatures about adolescent substance use and depression. *BMC public health*, 16, 279. <https://doi.org/10.1186/s12889-016-2932-1>
- Wang, W., Feng, Y. ve Dai, W. (2018). Topic analysis of online reviews for two competitive products using Latent Dirichlet Allocation. *Electronic Commerce Research and Applications*, 29, 142–156. <https://doi.org/10.1016/J.ELERAP.2018.04.003>
- Weiss, S.M., Indurkha, N. ve Zhang, T. (2010). Fundamentals of Predictive Text Mining. London: Springer-Verlag London Limited.
- Zhang, C., Yu, M. C., & Marin, S. (2021). Exploring public sentiment on enforced remote work during COVID-19. *The Journal of Applied Psychology*, 106(6), 797–810. <https://doi.org/10.1037/apl0000933>

Topic Modeling in Psychological Counseling and Guidance

Nurten Karacan Ozdemir^{1*} 
Kübra Atalay Kabasakal² 

¹Hacettepe University, Department of
Guidance and Psychological Counseling,
Ankara, Türkiye
nurtenkaracan@hacettepe.edu.tr

²Hacettepe University, Department of
Measurement and Assessment, Ankara,
Türkiye
katalay@hacettepe.edu.tr

*Corresponding Author

Received: 27.03.2024
Accepted: 24.02.2025
Available Online: 30.04.2025

Abstract: This study introduces topic modeling analysis, a text mining method, and examines its applications in the field of counseling and guidance. Accordingly, this paper reviews studies that use topic modeling to explore various aspects of the field, including (1) trending topics in mental health-related disciplines (e.g., counseling, psychology, psychiatry), (2) prominent themes in specific specialties (e.g., career counseling) or issues (e.g., multiculturalism), (3) issues related to the counseling and therapy process, (4) diagnosis and identification, and (5) assessment and evaluation practices. In conclusion, it was concluded that topic modeling analysis can be an effective method for identifying emerging trends, research priorities and practical themes in the field of counseling and guidance. However, it is important to consider both methodological limitations (data quality, lack of context and subjectivity of interpretation) and ethical issues (fair representation, confidentiality and transparency) in this analysis method.

Keywords: Psychological Counseling, Text Mining, Topic Modelling, Machine Learning

INTRODUCTION

Text mining is a promising analytical approach that can effectively leverage the detailed information inherent in large qualitative datasets. It is defined as the process of extracting knowledge from previously unknown and inaccessible textual data (Buenano-Fernandez et al., 2020). This technique enables the identification of patterns, correlations, and insights from textual sources.

Despite technology's ability to accelerate data collection and enable the gathering of multidimensional information from larger sample sizes, there is a noticeable lack of studies that utilize and analyze large textual datasets. One of the reasons for this deficiency is that manual qualitative coding is expensive, laborious, and time-consuming (Feinerer & Wild, 2007). Additionally, the current qualitative data analysis tools (e.g., NVivo, MaxqDa) do not possess sufficient capabilities to allow researchers to easily code text data from large samples (Longo, 2020). Furthermore, examining each individual interpretation when analyzing large textual datasets continues to be a considerable challenge (Balahadia et al. 2016).

Recently, machine learning-based text analytics techniques have emerged that can help overcome these issues and minimize the time and labor demands of qualitative data analysis, allowing researchers to cost-effectively utilize large-scale written text data in their findings. This method has been developed to manage large datasets that are impractical to analyze through manual coding (Hopkins & King, 2010). One such approach is topic modeling analysis.

Topic modeling is a text mining research field focused on determining the fundamental semantic structure of textual documents. Topic modeling provides a framework of unsupervised machine learning algorithms that can extract meaningful information from the unstructured data structures created by large-scale documents (Kandula et al. 2011). Supervised and unsupervised learning methods are two general approaches

in machine learning. Unsupervised algorithms can identify patterns or topics without any prior information, by exploring the data. Like thematic analysis, unsupervised algorithms identify word clusters and topics emerging from the data, but human insight is still crucial to aid the interpretation of the resulting topics. Topic modeling leverages the advantages of both thematic analysis and machine learning.

By utilizing topic modeling techniques, large volumes of unstructured textual documents can be automatically searched, organized, and summarized. Topic modeling is like factor analysis in the way it reduces documents into various dimensions referred to as topics. There are different methods of topic modeling, including structural topic modeling, which allows for the inclusion of covariates or attributes in the topic modeling process. Structural topic modeling employs a regression framework to understand whether these shared variables influence the specific topics present in the documents (Roberts et al., 2014). In this sense, structural topic modeling adds further depth to the insights derived from the text by accounting for how the text varies in relation to the included covariates. While topic modeling and structural topic modeling have limitations in capturing low-prevalence topics compared to traditional qualitative text analysis, this paper focuses on topic modeling analysis as one of the data mining techniques (Baumer et al., 2017). For instance, Latent Dirichlet Allocation can be provided as an example of a topic modeling method.

Latent Dirichlet Allocation

The goal of topic modeling is to facilitate a better understanding of phenomena in written texts by grouping together frequently co-occurring words that define a topic. While there are various topic modeling approaches, one of the original topic models, Latent Dirichlet Allocation (LDA), was first introduced to the literature by Blei and colleagues in 2003 (Blei et al., 2003). Blei (2012) defines topic modeling as statistical methods that analyze the words in original texts to discover the themes within the texts, how these themes are interrelated, and how they evolve over time.

The dataset to be modeled using LDA is referred to as a corpus, with each element within the corpus considered a document. The LDA method involves determining the topic distribution for each document and the word distribution for each topic. As a completely unsupervised algorithm, LDA operates based on the bag-of-words approach, which examines the co-occurrence of words without considering their position within the document, and does not require any prior knowledge (Blei, 2012).

LDA is a statistical technique that identifies latent topics embedded within a collection of documents, based on the premise that words belonging to the same topic are more likely to co-occur within the same document (Wang et al., 2018). The topic of a document is naturally related to the words it contains. The words in a document implicitly represent the main topic(s) of the document. These words will form the word group associated with their respective topics. For instance, in a psychology-related topic, words such as emotions, cognitions, and behaviors may emerge, while in a measurement and evaluation-related topic, words like scale, assessment, and development may be found. There may even be a common word like scale across both topics. The larger the set of documents, the more challenging it becomes to determine which words belong to which topic group.

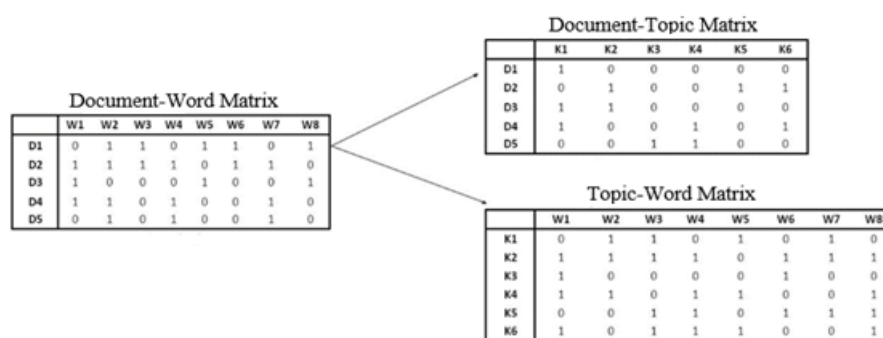
Let us briefly review the two fundamental principles of the model. The first principle is that each document is a mixture of topics, and the second principle is that each topic is a mixture of words. Each document can contain words from diverse topics in different proportions. For example, consider two topic models examining health-related and economy-related posts on a social media platform. Document_1 may consist of 90% health-related content and 10% economy-related content, while Document_2 may contain 30% health-related content and 70% economy-related content. In the health-related topic, the most familiar words may be psychology, therapy, and depression, while the economy-related topic may comprise words like budget, inflation, and raise. Importantly, words can be shared across topics; a word like 'budget' may be equally prevalent in both topics.

LDA is a statistical technique that assumes each document (D) consisting of a few words (N) can represent a Dirichlet probability distribution of latent topics. The Dirichlet distribution used to allocate words in the document to different topics is named after the LDA generative process (Blei, 2012). In the LDA algorithm, topics are a mixture of words, and documents are a mixture of topics, resulting in a three-layered hierarchical structure. LDA calculates the probability of a topic occurring in a document, allowing the identification of topics across the entire corpus.

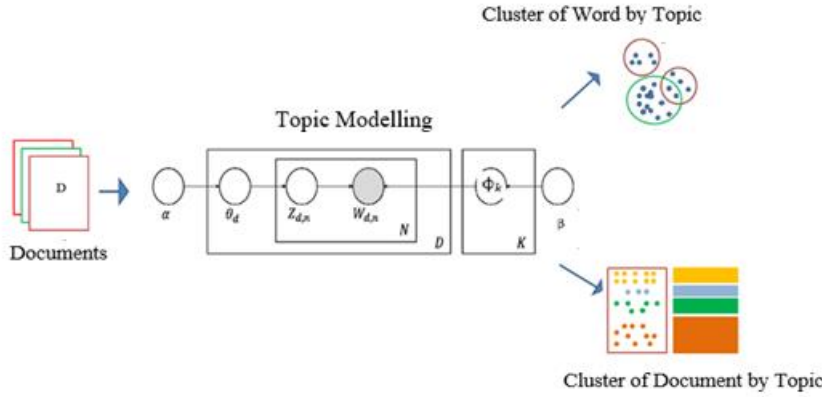
LDA converts documents into a document-word matrix as shown in Figure I. Figure I shows documents D1, D2, D3, D4, D5, topics K1, K2, K3, K4, K5 and words W1, W2, W3, W4, W5, W6, W7, W8. The document-word matrix is a statistical representation describing the frequency of words occurring in a corpus. This matrix is divided into two sub-matrices: the document-topic matrix, which contains the possible topics in the documents, and the topic-word matrix, which contains the words contained in the possible topics. These matrices already provide the topic-word and document-topic distributions. However, the main goal of the LDA algorithm is to improve the distribution of this data. LDA uses sampling techniques to improve these matrices. Once the data is represented based on matrices, words are converted into vectors (Blei, 2012).

Figure 1

Document-Word Matrix



LDA assumes that the corpus of documents is generated by latent probabilistic variables that can be understood as topics. The LDA algorithm applies a Bayesian hierarchical mixture model that exploits co-occurrence between words to identify topics. LDA uses generative probabilistic modeling. This generative process defines a joint probability conditional distribution over both observed and latent random variables. This conditional distribution is also called the posterior distribution. LDA uses the Gibbs sampling algorithm to generate the posterior distribution (Blei, 2012). For a better understanding of the GDA model, a graphical representation is given in Figure 2.

Figure 2*A Graphical Representation of LDA*

In Figure 2, the parameters α (document-topic density) and β (topic-word density) are link parameters that are sampled once during processing. Higher values of α mean that documents are composed of a mixture of more topics, while higher values of β mean that a topic is composed of more words. A low or high value of α is related to the number of topics the documents contain. In summary, a high α parameter indicates similarity of documents, while a high β parameter indicates similarity of topics.

In Figure 2, $W_{d,n}$ represents the n th word in document d , $z_{d,n}$ represents the topic assignment for the n th word in document d , and $\theta_{d,k}$ represents the proportion of k topics in document d . β_k , $\theta_{d,k}$, $Z_{d,n}$ represent the values obtained as a result of processing the previously unknown and $w_{d,n}$ the known variable. $\theta_{d,k}$ is drawn from the distribution of the parameter α ($\theta_{d,k} \sim \text{dir}(\alpha)$). The number of topics K and the associated distribution of the parameter α and the dimensionality of the topic variable z are taken as priori and assumed to be constant. Note that a proportion is estimated for each of the K topics within each document; however, these estimated proportions can be quite small (e.g., 0.001) (Nowlin, 2015). This indicates that the topic is not common within that document. The standard Dirichlet prior to the parameter α is $50/K$ (Griffiths & Steyvers, 2004).

Figure 2 shows that the model specifies a series of dependencies. For example, the topic assignment $Z_{d,n}$ depends on the rate of topics per document θ_d . As another example, the observed word $W_{d,n}$ depends on the topic assignment $Z_{d,n}$ and all topics $\beta_{1:k}$. (Operationally, this term is defined by looking at which topic $Z_{d,n}$ refers to and the probability of $W_{d,n}$ within that topic). These dependencies define the LDA. They are encoded in the statistical assumptions behind the production process in the specific mathematical form of the joint distribution (Blei, 2012). The joint probability distribution is then used to calculate the conditional or posterior distribution of latent variables given the observed variables, according to Bayes' rule (Blei, 2012). This joint distribution uses the following equation to calculate the posterior distribution of topic probabilities for each document:

$$p(\theta, z | w, \alpha, \beta) = \frac{p(\theta, z, w | \alpha, \beta)}{p(w | \alpha, \beta)}$$

In the equation, the numerator is the joint distribution of all random variables, and the denominator is the probability of obtaining the observed corpus under any topic model. These probabilities can be summed over all topic structures, but this becomes difficult to calculate given the considerable number of topic structures. Therefore, this sum must be approximated using sampling-based or variational approaches (Blei, 2012). In the LDA algorithm, Gibbs sampling is used to estimate the posterior distribution (Griffiths & Steyvers, 2004).

As mentioned before, the LDA algorithm first constructs document-word matrices. Based on this, it conceptualizes each word in a document as a sequence of N words (Blei, 2012). Then a few of the k topics in the Dirichlet distribution are selected. The word ' w_i belonging to a topic' is added to the document. After the addition of w_i , a topic is selected according to the probability distribution. Assuming a set of documents is obtained this way, the following steps are followed to detect k topics and their words with Gibbs sampling in the LDA algorithm.

1. For each document, each word in the document is randomly assigned to one of the K topics.

Each document is reviewed and each word in the document is randomly assigned to one of the K topics. For each document d , each word w is reviewed. For a given document d , we try to determine how many words belong to topic t . If many words in d belong to topic t , there is a higher probability that word w belongs to topic t . Attempt to determine how many documents are on topic t because of word w .

2. This random assignment yields word distributions of both all documents and all topics.

LDA represents documents as a mixture of topics. Similarly, a topic is a mixture of words. If the probability of a word being in a topic is high, all documents with the word w will be more strongly associated with topic t . Similarly, if the probability of a word w being in topic t is very low, then documents containing the word w will have a very low probability of being in topic t , because the rest of the words in document d will belong to another topic and hence d will have a higher probability for those topics. So, even if the word w is added to topic t , it will not bring many such documents to topic t .

3. Assignments for each document are optimized.

In the third stage, two proportions are calculated for each topic t .

$p(\text{topic } t \mid \text{document } d)$ = the proportion of words in document d that are currently assigned to topic t , and

$p(\text{word } w \mid \text{topic } t)$ = proportion of assignments to topic t over all documents from this word w .

Based on the product of the two calculated rates, i.e., the probability that topic t produces the word w , the word w is assigned to a new topic. When the third step is repeated many times, a state is reached where the topic assignments are stable. Thus, with these assignments, the topic probabilities contained in each document and the words belonging to each topic are obtained with the number of repetitions. While explaining the LDA algorithm, the work's mathematical logic is explained without detailed calculations such as the use of priors in probability calculations, as it is thought that this work will not appeal to this study's readership.

In the initial and provisional stage, words are assigned to topics, but the number of topics must be determined by the researcher. In the creation of the word-topic matrix by assigning words to topics, it is allowed for words to be used for different topics and for similar words to be clustered. In the word-topic matrix of the LDA algorithm, the words representing the topics are evaluated according to probabilities, and each word can be present in multiple topics. Also, the LDA model assumes that the words occurring in the documents carry important semantic information and that a set of documents on a similar topic will contain the same set of words. Based on this idea, latent topics find a group of words that frequently appear in the documents. The words with the highest probability within a topic facilitate interpretation of its meaning. As shown in Figure 1, the LDA algorithm models a corpus consisting of various documents by creating a document-topic and word-topic matrix. With the help of these matrices, useful information is extracted using the words and their weights, and in this context, topics can be identified. The topic assignment of a document not used in the modeling can also be made using these matrices. The output of LDA consists of keyword groups, each belonging to a specific topic category. These keyword groups are then labeled with unique topic titles. Therefore, the output represents a list of topics from 1 to M , each with a keyword group from 1 to N .

Methodological Issues

A topic modeling study consists of several stages. However, these stages are not independent of each other; some stages are intertwined. As with all research, topic modeling should begin with the formulation of hypotheses and research questions. Questions that are most appropriate for topic modeling are those in which it is interesting to understand the hidden variables underlying a set of texts. The second stage is data collection. This involves identifying the appropriate text needed to test hypotheses or answer research questions and developing a method for collecting this text. There are many data collection methods that can be used to compile a text database (corpus), depending on the type of data. Any text-generating design can be used if the researcher can collect enough data. Researchers can use different types of texts depending on the research purpose, ranging from internet-based texts (e.g. website text, emails or social media data such as Twitter, Facebook), transcribing audio files from interviews and focus groups, or asking participants to write responses to questionnaires and open-ended questions (Banks et al., 2018).

The steps to be considered in topic modeling are sample size, data preprocessing, deciding on the number of topics, and examining model complexity. These steps are explained in detail, respectively.

Sample Size in Topic Modeling: The sample size in topic modeling is related to the number of documents included in the analysis and the number of words in each document. There is no universal rule regarding the minimum number of words required for topic modeling. For example, social media posts can be noticeably short as they contain a limited number of words. When the text data is this brief, aggregation at the user level can be performed, but this reduces the sample size. The minimum word count for text data can vary depending on the context. Particularly in survey studies, participants may be asked to write a minimum number of words.

As important as text length is the quality of the text. It is easier for topics to emerge in texts that are spelled and grammatically correct. If the same word is misspelled or used in a document, it is more difficult to find patterns in the text. Although surveys specify a minimum word count, respondents may use repetitive or meaningless words just to fill the word count. In addition, when the text type is a social media post, the presence of abbreviations and slang terms can challenge researchers. Another aspect that affects the sample size is the level of the document. A sentence is contained within a paragraph, a paragraph within a social media post and a social media post within a user. Depending on the hypotheses or research questions of interest, researchers can change the level at which a document is described, which can affect the sample size (Tang et al., 2014).

Data Preprocessing. In quantitative research, the first step after collecting text data, such as missing value and outlier analysis, which must be done before starting the analysis, should be data pre-processing. Every action taken at this stage directly affects the result of the analysis. The more accurate and better the data preprocessing stage is, the more accurate and reliable the analysis result is. Just as the presence and removal of outliers or influential values in a regression analysis can change the results, the text preprocessing stage can potentially affect the results. The steps in the preprocessing phase aim to prepare the data for further analysis. These steps are necessary not only to develop a clean and appropriately formatted dataset before conducting topic modeling, but also to facilitate the interpretability of the topic modeling output (Gencoglu, 2023). In this context, preprocessing is an exploratory, iterative stage, and it may be necessary to return to it even later in the analysis.

The preprocessing process should be the first step to remove all irrelevant answers, noticeably short text or invalid records that could affect the quality of the modeling analysis. The subsequent process involves tokenization and cleaning of the text. This stage includes removing meaningless characters and punctuation marks, converting all words to lower case, and excluding blank answers and answers corresponding to blanks (Gencoglu, 2023; Vijayarani et al., 2015).

In tokenization, the entire document is divided into sentences, words, syllables, letters or according to the rules set for the research. Each sentence, word, syllable, and letter are tokenized according to the rules

chosen for the research, which are determined by punctuation marks or spaces. Language-specific processing may also be required at this stage. For example, for Turkish, character conversions such as converting "ç" to "c" and "ü" to "u" are required. If there are too many Turkish characters in the data set, the analysis result may be negatively affected due to the character encoding problem in the operating system.

One of the most important steps of the preprocessing stage involves the assessment and removal of stop words. Stop words are words that are used so frequently that they do not add value to the identification of topics in a research context. Stop words are language-specific words that are not relevant to the text content, so they should be removed for a more valid analysis. The words that remain after removing ineffective words are usually nouns, verbs, and adjectives. Conjunctions such as "and, or, but, but, but, because, moreover, even" or pronouns such as "I, you, he, she, that" need to be removed as they have no meaning on their own and add little or no value to the grouping process (Nowlin, 2015). For example, if the corpus contains only medical documents, words such as human, body, health can be found in most of the documents and can be omitted as they do not add specific information to make the document stand out. These ineffective words, which occur in at least 80% ~ 90% of the documents, should be removed for the meaningfulness of the results without any loss of information.

Standard stop words should be removed first before starting the analysis, but additional words can be removed when the findings are interpreted. Turkish stop words libraries may not be at the desired level in open-source software where LDA analysis is performed. For this reason, it may be suggested to create a list of ineffective words consisting of Turkish words within the scope of the study. To find ineffective words, word lists can be organized by frequency and frequently used words can be selected as ineffective words due to their lack of semantic value (Kadhim et al., 2014).

It is valuable to have information about the literature studied in determining stop words. A word considered a stop word in one analysis may be important in another. For example, the third person singular word "it" in English is a stop word for most analyses. However, this word is also an abbreviation for the word information technology and is effective for a study in this field. To identify stop words, Zipf's law (Newman, 2005) can be applied, which suggests removing frequently occurring words to reduce random noise in the results. Zipf's law also applies to the removal of infrequent words in the tail of the distribution. A minimum number of occurrences of a word in a document can be set, above which a word is not considered a sparse term. However, some words may occur many times in some documents and not at all in others. Therefore, setting a minimum number of documents will lead to the exclusion of some words from the analysis.

After tokenization, the next possible step is stemming. This step is simply the process of reducing words to their root or basic form by removing affixes, while keeping the root meaning intact (Jivani, 2011). The aim of this method is to remove various affixes, reduce the number of words, have correctly matched stems, and save time and memory space (Vijayarani et al., 2015). However, the words that remain after stemming may not always be real words; they may be just fragments of words with no semantic meaning. While stemming provides a more basic and coarse method for reducing words to their root forms, lemmatization is a more complex process that involves determining the basic or lexical form of a word according to its intended meaning. The word that emerges from this process is always a special word that belongs to the language. Lemmatization is preferred for tasks where semantic accuracy is crucial, while stemming may be more appropriate for tasks where simplicity and speed are more important. Whether the lemmatization step is necessary depends on the problem under investigation (Weiss et al., 2010). Unfortunately, fewer software packages include the lemmatization approach. While both stemming and lemmatization remain popular choices for text preprocessing, research on the usefulness of each approach is mixed and highly dependent on the corpus, method, and research language (Banks et al., 2018). In suffixal languages such as Turkish, the same word may appear in different forms depending on its use in the text. By separating words into their roots, this problem can be overcome.

Stemming and lemmatization steps should be followed by word sequences (n-grams). In the n-gram technique, the number of occurrences of adjacent words and phrases in a document is calculated. Calculating

the frequency of a single word is called unigram, the frequency of two adjacent words is called bigram, and the frequency of three adjacent words is called trigram (Kumar & Paul, 2016). Word sequences can be key to identifying extremely useful expressions. For example, when analyzing with bigrams, Eastern Anatolia will be included in the analysis as a single word, rather than two separate words, Eastern and Anatolia. The results will vary with the use of unigrams and bigrams. It is therefore recommended to examine the results with both analyses. Banks et al. (2018) recommend that researchers start their analyses with unigrams, but in most contexts, bigrams should also be considered.

The LDA algorithm has advantages and disadvantages (Maier et al., 2018). The advantages of the LDA algorithm are that it enables the rapid identification of topics in a large data set. In accordance with real life, a document can contain more than one topic. The disadvantages are that the results are not deterministic, and every point that needs to be considered in methodological issues changes the results of the analysis. A limitation is that the number of topics can be determined by the researcher.

Determining the Number of Topics: Determining the number of topics is one of the biggest challenges when using topic models (Blei & Lafferty, 2009). As it is an unsupervised algorithm, it places the burden on the researcher to ensure that the categories derived by the model are justifiable by the researcher. Determining the appropriate number of topics for the LDA algorithm should be chosen according to the research question and the hypotheses being tested and should be guided by theory. Too few topic models can lead to word overlap and failure to identify different topics, but too many topic models can reduce semantic validity or fail to reveal a broader and more inclusive topic. Therefore, the semantic validity of the topics should be the most important reason for determining the number of topics (Grimmer & Stewart, 2013). Semantic validity means that each topic has a clear and coherent meaning that can be recognized by associating terms with that topic (Krippendorff, 2003).

Researchers should use an iterative modeling process with varying number of K values. Exploring the data by starting with a small number of subjects and increasing the number of subjects helps researchers to become more familiar with their data. When choosing the number of subjects, it is important for researchers to apply the principle of parsimony. That is, to justify increasing the number of topics or latent constructs in the text and thus the complexity of the results, there should be a rationale that with more topics one can better interpret the findings.

Researchers should discuss what the emerging topics represent and develop topic labels. They should also try to relate the emerging issues to the existing literature. The analysis is complete when the researchers are satisfied with the topic labels and definitions, and when supporting examples from the original text are identified. As an ultimate step in the topic modeling phase, researchers should examine the network structure of the topics. A topic network allows to see how certain topics are related and is useful for assessing the dimensionality of emerging structures (Banks, 2018).

Model evaluation: Supervised learning algorithms in machine learning do not have metrics that can be calculated to evaluate the model. Since LDA is an unsupervised learning algorithm, it is not possible to evaluate model performance with error metrics. However, different models can be compared with each other. Perplexity, often used in language modeling, is algebraically equivalent to the inverse of the geometric mean probability per word. A lower perplexity score indicates better generalization and prediction performance (Blei et al., 2003).

TOPIC MODELING IN PSYCHOLOGICAL COUNSELING AND GUIDANCE

Topic modeling as a data analysis method is considered by some researchers as a complementary function to traditional research by enabling the analysis of big data through machine learning. It is pointed out that analyzing the counseling literature according to topic modeling can provide important insights to see how and where the field has developed and changed in the historical process and to monitor changes and

developments in the field (Oh et al., 2017). Using topic modeling in counseling and guidance can make significant contributions to theory, research, and practice.

First, topic modeling can contribute to the testing, extending, and developing the theoretical and conceptual foundations (Chen & Wojcik, 2016) by revealing the prominent topics in counseling and guidance literature. In addition, topic modeling can serve an important function such as identifying research trends in counseling and guidance and revealing the topics in the focus of the research (Bittermann & Fischer, 2018). At this point, it can contribute to the uncovering of latent themes that may not be easily identified by traditional qualitative data analysis methods and thus contribute to the unveiling of missing or overlooked aspects of the research (Chen & Wojcik, 2016). Moreover, emerging themes can help to synthesize research findings and draw effective inferences based on these syntheses and shed light on solving existing problems. It can also help young researchers to see the research trends in the field and explore their own areas of interest (Bittermann & Fischer, 2018).

The possible contributions of topic modeling to counseling practices are also underlined. For example, analyzing anonymized counseling transcripts with topic modeling can contribute to revealing the content and focus of counseling sessions, discovering prominent and recurrent themes in the client's life, understanding various aspects of client experiences, and reviewing and evaluating the helping process accordingly (Chen & Wojcik, 2016). In addition, transparently addressing recurrent themes in sessions with clients can contribute to gaining insights into clients' problems and their solutions. Topic modeling research to understand the effectiveness of different counseling methods and models (Liu & Gao, 2021) can make significant contributions to both the development of the counseling process and counselor education. The findings of topic modeling in the field of counseling psychology can contribute to the understanding of social needs and to the development of practices to respond to the needs of society (Bittermann & Fischer, 2018). At the same time, it can provide meaningful inputs for developing policies based on these needs. With these functions, topic modeling is considered as an effective, systematic, and valuable data analysis method in counseling and guidance (Bittermann & Fischer, 2018; Oh et al., 2017).

Examples of Topic Modeling in Psychological Counseling and Guidance and Related Literature

With the developments in big data analysis and machine learning, it is seen that research on the use of topic modeling as a data analysis method is increasing in the literature on mental health. Under this heading, examples of studies using topic modeling analysis with methods such as latent topic modeling are presented. Existing studies can be summarized under five groups according to their research objectives: 1) Research to identify the most researched topics in mental health-related fields in general (e.g., trending topics in counseling) 2) Research exploring topics related to specific themes in the counseling literature (e.g., multiculturalism, career counseling, mindfulness) 3) Studies investigating topics related to the counseling and helping process (e.g., most effective counseling methods, therapeutic relationship, online counseling, etc.) 4) Research on diagnosis and detection (e.g., cyberbullying, detection of depression symptoms) 5) Research conducted within the scope of measurement, evaluation and individual assessment studies (e.g. evaluation of responses to open-ended questions in questionnaires).

In the first group, studies that aim to uncover the thematic focus of research in counseling and related mental health fields (such as career counseling, organizational psychology, and marriage and family counseling) over time, as well as to identify current trends in research, serve as illustrative examples. For example, academic articles published in the *Journal of Counseling Psychology* from 1963 to 2015 were analyzed with the topic modeling method to investigate which topics the research has been concentrated on in the last 53 years. The findings revealed a distribution of 70 topics, which were grouped into four categories: counseling process and outcomes, multiculturalism, research methods, and career counseling. Especially in recent years, it has been reported that there has been an increase in research trends on diversity and multiculturalism (Oh et al., 2017). Similarly, research in psychology conducted in German-speaking countries from 1980 to 2016 was analyzed to identify emerging topics. In this study, which examined 314,573 publications, 500 distinct topics were identified, and the analysis focused on determining which of these were

trending. The findings revealed that topics with an increasing research trend included various aspects of neuropsychology, online counseling, cross-cultural characteristics, trauma, and visual attention (Bittermann & Fischer, 2018). In line with such studies, there are also investigations that focus on more specialized areas. For instance, Liu et al. (2020) performed a topic modeling analysis to determine the prominent themes in clinical psychology research between 2007 and 2018. They found that meta-analysis, psychotherapy, professional development, and depression were the most frequently studied topics during this 12-year period. Additionally, their study indicated that behavioral interventions have been a steadily rising research focus since 2007. These studies produce meaningful results in terms of identifying the focal points of research and emerging trends, especially in Psychological Counseling and Guidance and related literature. From this perspective, they offer valuable insights into which topics have gained importance over time and which areas may have been overlooked in the historical research process.

In the second group, there may be studies aimed at revealing the topics around which the research on specific problem areas or specific forms of intervention in counseling and related literature focuses. Kee et al. (2019) analyzed the topic modeling of studies on mindfulness in this group. The researchers analyzed 5949 studies published in Web of Science until October 20, 2017, with topic modeling. In the analysis, 106 topics and 231 terms emerged, which were systematically categorized into five distinct themes. The first theme, target group/application area, encompassed various populations and settings where mindfulness research was conducted, including adolescents, couples, pregnant women, health professionals, as well as workplace and school environments. The second theme, situation/problem, focused on specific conditions and challenges where mindfulness practices were implemented, such as gambling addiction, depression, anxiety disorders, and cancer. Structure/philosophy emerged as the third theme, addressing different philosophical foundations underlying mindfulness practices, particularly Vipassana meditation and Buddhist principles. The fourth theme, modality/form, highlighted various therapeutic approaches, with mindfulness-based cognitive therapy being a notable example. Finally, the research method theme incorporated methodological aspects, including measurement tools, research designs, and related procedural elements used in the studies. This research can provide important inputs to researchers working on mindfulness in terms of showing the topics on which the research conducted so far has been gathered. Another interesting study was conducted in organizational psychology (Royston et al., 2022). The responses of officers serving at distinct levels in the army to open-ended questions about their duties were analyzed with topic modeling and it was examined how the definitions of duties and responsibilities differed according to leadership level and experience. According to the findings, in lower-level positions, topics focused on daily tasks and obtaining information about the job, while topics such as supervision, network expansion activities and process improvement came to the fore as the leadership level increased. At higher levels of leadership, the language used emphasized topics such as participation in corporate committees, strategic leadership, overseeing corporate operations, mentoring, and strategic planning. In this study, while pointing out the importance of topic modeling analysis for job analysis studies, it was emphasized that the findings obtained with this analysis method can contribute to a better understanding of issues such as the qualifications required by certain jobs, employee characteristics, and job context (Royston et al., 2022). Another example is the topic modeling analysis conducted by Wang et al. (2016) on the abstracts of 17,723 studies on substance use and depression in adolescents published in PubMed from 2000 to 2014. As a result of the analysis, five-, 20- and 50-topic models were put forward, for example, it was pointed out that some terms were included together in the five-topic model. For example, alcohol and brain research, alcohol and sexual experiences, smoking and alcohol were grouped in pairs under different topics. Similarly, in the model with 20 topics, it was pointed out that additional topics such as intervention and treatment, family influence and peer social network, diet, physical activity and obesity, suicide and brain research emerged. In the 50-issue model, issues such as genetic and environmental influences, depression and alcohol use were pointed out. In these three models, it was reported that the five prominent topics related to substance use in adolescents were substance use, general research terms, smoking, brain research, and family and friend networks. The five prominent topics related to depression in adolescents were depression, psychiatric problems, treatment of depression and general research terms. In the findings, it was pointed out that in addition to the typical topics related to depression and substance use in adolescents, risk factors related to substance use in adolescents, the relationship between substance use and depression, and intervention programs were the most prominent topics. In addition, the most prominent topic related to substance use and depression in

adolescents was sex and violence, and the relationship between sexual experience and violence in adolescents and substance use and depression was pointed out. The other interesting topic was development from childhood to adulthood. It can be said that this study has revealed important results in identifying risk factors for depression and substance use in adolescents, prominent issues for intervention, and needs as well as strengths in these areas. It can be said that these findings can also provide important insights on how to shape the intervention programs to be carried out on the subject.

Studies employing topic modeling analysis to examine the effectiveness and outcomes of psychological counseling and psychological helping processes constitute the third group of research. A significant example within this category is the research conducted by Atzil-Slonim et al. (2019), who utilized topic modeling analysis to assess client functionality in psychotherapy and identify disruptions in therapeutic relationships. Their study involved a comprehensive analysis of transcripts from 873 psychotherapy sessions, encompassing interactions between 52 therapists and 58 clients. In this context, it was investigated which topics were prominent in the therapy processes, which topics best described the clients' functioning and the disconnection in the therapeutic relationship between the clients and the therapist, and whether there was a change in the outcomes of the therapy according to the changes in these topics that were prominent in the analyses. To support these investigations, before each session, clients completed two different instruments in which they reported changes in their functioning, interpersonal relationships, performance in social roles, and symptom levels, which were considered as indicators of progress in the treatment process. In addition, after each session, the therapists responded to a question in which they evaluated the disconnections in the therapeutic relationship. Within the study's scope, firstly, topic modeling analysis was conducted on the transcripts and the prominent topics in the therapy process were identified. Accordingly, positive experiences, negative experiences, relationships, treatment, health, daily life, and diverse topics were the most prominent topics. Then, logistic regression was used to investigate the functioning levels of the clients and the topics on which the alliance breaks in psychotherapy sessions occur. According to the findings, it was reported that four topics (loneliness, suffering, physical difficulties, and anger) were associated with low functioning of the clients, while the topics of fun, leisure experiences, celebration and choice were associated with high functioning. Moreover, three key topics that emerged from the analyses—communication, goal setting, and the need for help and problems—were found to be closely related to disruptions in the therapeutic alliance. To further investigate this relationship, multilevel growth models were employed to examine whether therapy outcomes changed as the focus of topics shifted throughout the therapeutic process. The results indicated that an increase in topics associated with high-functioning client behaviors (such as communication and goal setting) was linked to a reduction in symptoms, including lower levels of depression. However, changes in topics related to low-functioning client behaviors did not produce significant differences in therapy outcomes. It was pointed out that this result supports current psychotherapy approaches based on focusing on positive experiences rather than focusing on negative experiences. Similarly, Howes et al. (2013) sought answers to questions such as whether the topics that are prominent in the therapy process predict clients' symptoms and therapy outcomes, if such prediction occurs, whether machine learning-based automated data analysis methods such as topic modeling can be used instead of explanations made by therapists, and whether topic modeling analysis produces topics that can be interpreted and/or compared with the decisions made by therapists. Accordingly, transcripts of counseling sessions between 31 psychiatrists and 138 clients who gave informed consent between 2006 and 2008, the shortest of which lasted 5 minutes and the longest of which lasted 1 hour, were used. The researchers, who were not part of the treatment process, assessed the clients' symptoms, asked the clients to respond to a questionnaire to assess their satisfaction with the treatment sessions, and asked both the clients and the doctors to respond to a questionnaire to assess the therapeutic relationship during this process. In total, 12 out of 138 counseling transcripts were coded by two experts and 20 prominent topics were identified in these sessions. At the same time, 138 transcripts were further analyzed through topic modeling by fixing the number of topics to 20. Thus, the topics highlighted by expert human coders and machine learning based topic modeling were compared. It was found that both machine learning-based topic modeling analysis and manual thematic coding by experts had similar predictive power. However, it was observed that the topics revealed by machine learning-based automatic coding did not predict symptoms, while coding by experts predicted symptoms. On the other hand, it was reported that topic modeling analysis revealed topics/themes that better predicted the therapeutic relationship compared to expert coding. These studies make important

contributions to the relevant literature in terms of producing preliminary findings indicating that topic modeling analyses conducted on session transcripts can be used to evaluate the process and outcomes of psychological help and the therapeutic relationship and alliance between the client and the counselor/therapist.

The fourth group includes examples of topic modeling analysis for diagnostic and detection. For instance, Franz et al. (2020) conducted a study utilizing topic analysis to examine posts and shares made in subgroups related to self-harm and suicide on TeenHelp.org, an internet support group frequented by adolescents. The same content was independently coded by three human coders, revealing themes such as non-suicidal self-harm, suicidal ideation, suicide attempts, passive suicidal ideation, suicide planning, depression, social concerns, abuse, drug use, substance use, and other related issues. Similarly, the effectiveness of topic modeling analysis in identifying self-harm thoughts and behaviors was evaluated. The findings revealed 15 distinct topics, six of which (suicide, friends, family, depression, non-suicidal self-harm behaviors, and social) were considered relevant for detection purposes. Notably, the analysis showed that machine learning distinguished suicide and depression as separate topics. The study also explored the degree of overlap between the machine learning-based topic modeling analysis and the themes identified by human coders. Logistic regression analysis indicated that the six machine learning-identified topics that met the research hypotheses corresponded with four of the themes identified by human coders, highlighting the potential for automated methods to effectively complement human analysis in this domain. The t-test was used to analyze whether the topics identified by machine learning differed according to gender and age, and it was reported that the topic of “suicide” was higher in males, while the topic of “non-suicidal self-harm” was higher in females. This study is important in terms of exemplifying how machine learning-based topic analysis can be used to identify themes related to suicide and self-harm in adolescents' posts. Studies in this group are more sophisticated with supervised topic modeling analyses than free/unsupervised ones (e.g., Mobin et al., 2024).

The final group of studies illustrates how topic modeling analysis can be applied in assessment processes. For instance, Finch et al. (2021) employed topic modeling to analyze written definitions of various aspects of giftedness provided by parents of gifted children. This analysis informed the development of a scale designed to identify gifted children by generating relevant scale items based on the extracted topics. Similarly, Rohrer et al. (2017) integrated topic modeling into a quantitative research framework. In addition to administering standardized measurement tools, the researchers analyzed participants' responses to the open-ended question, “What other things worry you?” using topic modeling. This approach enabled a deeper exploration of participants' concerns beyond predefined survey items. This study analyzed 35,000 written responses to an open-ended question on what other issues respondents were concerned about in the context of the Socio-Economic Forum in Germany between 2000 and 2011. The most prominent issues reported in the findings were the future of children, children, young people, school, family health, rising prices, rich and poor, foreigners in Germany, unemployment, finding a job, retirement and economic security, Germany's development, interpersonal issues, decline in moral values, politicians, corruption in politics and war and terrorism. An analysis of topic changes over time indicates that themes related to war, terrorism, and rising prices become more prominent, particularly during periods of international conflict (e.g., the Syrian Civil War) and economic crises. This example demonstrates the combined use of quantitative and qualitative data analysis, offering an alternative to traditional hand coding for analyzing responses to open-ended questions in large-scale surveys. By leveraging topic modeling, researchers can overcome the limitations of manual coding and gain deeper insights into evolving societal concerns. In another interesting study (Jayaratne & Jayatilleke, 2020), the responses of candidates to open-ended questions during the job application phase were analyzed to determine which type they fall into according to the HEXACO personality model. Within the scope of the study, the results of a measurement tool that candidates answered for the HEXACO personality model as well as their answers to open-ended questions using an online chatbot were used in 46,000 job applications. The job interviews included open-ended questions such as: "What motivates you? What are you passionate about? Not everyone will always agree with you. Have you ever had a peer, teammate or friend disagree with you? Give an example of a situation where you went beyond achieving something. Why was it important for you to achieve this? Sometimes things do not always go as planned. Describe a time when you failed to meet a deadline or personal commitment. What did you do and how did that make you feel? Thinking fast is critical in sales. What makes you qualified to do this? Give an example." Candidates' responses to these questions

were analyzed using machine learning based topic modeling analysis (latent Dirichlet allocation) and a moderate correlation ($r= 0.39$) was found between the models. The appropriateness of the terms describing the different personality models was confirmed by feedback from 117 volunteer participants. This study highlights that instead of using self-report-based individual recognition techniques that can be influenced by factors such as social desirability during job interviews, responses to open-ended interview questions can provide an important advance in the use of topic modeling analyses to identify personality types of candidates and assess their suitability for the job.

When classifying studies in counseling and related literature that utilize topic modeling, a classification can be made based on both the data source used and the research purpose. Accordingly, studies can be categorized as follows: (1) published academic documents, (2) posts, comments and shares on social media and/or interactive sites and platforms on the internet (such as blogs), (3) counseling and therapy transcripts, and (4) written responses to open-ended questions in data collection instruments. Numerous studies falling into the first category have been discussed above. Another example of how non-research documents can be analyzed is the study by Ni et al. (2023), which examined handbooks on psychological first aid. In this study, 14 handbooks written in English by various institutions and organizations across different countries were analyzed using topic modeling. The researchers identified relevant documents through Google searches with keywords such as “psychological first aid,” “handbook,” “guide,” and “manual.” The findings revealed key topics, including refugees, orientation activities, community-based practices, post-traumatic stress disorder and other psychological issues, training materials, specific guidelines for service providers, scholarships on psychological first aid, mental health and psychosocial support, and an Australia-specific service model. Such analyses can assist countries, institutions, and organizations in developing or evaluating handbooks, codes of practice, or models on related topics. Additionally, they provide an opportunity for professionals in the field to review and refine their practices.

As an example of the second group of studies, Seah and Shim (2018) analyzed the posts and comments in the Singapore subgroup related to suicide and depression on Reddit, a social platform where individuals share posts and comments in various subgroups according to subject areas, using topic modeling. The study analyzed 385 posts and 21,000 comments on topics related to depression and suicide between 2010-2018. According to the results, 14 topics emerged, including life and death decisions and euthanasia, friendship, social support and relationships, school life, expectations and stress, social views, career, stress and challenges, attitudes in society at a broader level, cases reported to the police, incentives to share experiences and comment, acceptance of LGBT people and beliefs, suicide attempts, family problems, past experiences and memories, and seeking treatment and help. By identifying the prominent topics in the posts on depression and suicide of individuals who may be assumed to be in need of help on a social platform where they are assumed to be able to share their views relatively more easily, this research may also shed light on the areas that can be covered by the help studies to be carried out for individuals who have difficulties in this regard, as well as the studies that can be carried out on the prevention of depression and suicide. Similarly, Zhang et al. (2021) analyzed more than one million tweets on Twitter between March and June 2020 using keywords such as “working from home”, ‘telecommuting’ and “telework” through topic modeling to better understand attitudes and experiences of working from home during COVID-19. Prominent topics in the results included themes such as telecommuting, cybersecurity, mental health, work-life balance, teamwork, and leadership. These results can be seen as important in understanding the difficulties and needs of employees in their experiences with telecommuting and gaining insight into the general trend and attitude towards telecommuting. In another interesting study (Carron-Arthur et al., 2016), the posts shared in an online support group related to mental health were analyzed by topic modeling. However, in this study, it was pointed out that 75% of the posts on such platforms were made by the same people, and these users were called super users, and it was examined with the chi-square test whether the topics that stood out in the topic modeling analysis varied between super users and other users. Accordingly, super users were found to be more likely to post on topics such as parenting roles, mental health, and positive change, and less likely to post on depression, medication, therapy, and anxiety, but they were also found to be more likely to post on five of the seven prominent topics. This research points to the importance of examining the differences in topics between the 75% group who post continuously and the other groups, rather than the number or frequency of posts, especially in topic analyses conducted on

posts and shares on mental health issues on social media or the internet. In this context, the fact that the same or similar comments are constantly written by the same people points to the need to conduct more detailed examinations by being aware of the determining role of the situation in shaping the prominent topics.

Similarly, examples of studies in the third category have been provided within the above classification (see Atzil-Slonim et al., 2019). Additionally, Atkins et al. (2012) proposed the use of topic modeling analysis in couple and family therapy. In their study, the authors applied topic modeling to analyze therapy transcripts obtained from an experimental study involving couples. Their analysis identified key themes in therapy discussions across three distinct models: content-oriented topics (e.g., family, work, money, sexuality), emotion-oriented topics (e.g., anger, sadness, hurt, effort), and therapy-related topics (e.g., communication, problem-solving, brainstorming). Furthermore, they showed how topic modeling can be used to track changes in discussion topics over therapy. For instance, in the case of one couple, communication-oriented topics dominated discussions in a particular phase, increased in prevalence between the fourth and eighth sessions, and subsequently shifted toward problem-solving. Finally, in the fourth category, studies on the application of topic modeling to analyze responses to open-ended questions in assessment and evaluation processes, and in understanding individuals, serve as relevant examples.

According to the topic modeling analysis method used, a classification can be made as (1) supervised topic modeling, (2) unsupervised topic modeling, and (3) research using generative artificial intelligence based on large language models. In this article, since the purpose of use rather than the different methods of topic modeling analysis is emphasized, we have tried to make a classification based on how topic modeling analysis is used for different purposes. At this point, it can be said that although there is an interesting and developing use of topic modeling analysis in counseling and related mental health literature, some limitations and ethical issues come to the fore in this process.

Limitations and Ethical Considerations in the Use of Topic Modeling Analysis

When the related literature is reviewed, it is seen that there are limitations that should be taken into consideration when using topic modeling analysis (Banks et al., 2018; Bittermann & Fischer, 2018; Kobayashi et al., 2018; Liu & Gao, 2021; Oh et al., 2017; Otsuka et al., 2021). First, the results of topic modeling are highly dependent on the data set subjected to topic modeling analysis. Therefore, when conducting research in the field of counseling, it is likely to be highly influenced by whether the literature has been adequately reviewed, if criteria such as inclusion and exclusion have been determined, how they have been determined and which studies have been included in the analysis, the frequency or inaccuracy of posts on the internet, or inadequate transcripts. Note that the data quality can significantly impact the results. Topic modeling analysis requires a large data set to produce accurate topic and theme distributions. Therefore, a small data set (e.g., a few counseling transcripts) would not be suitable for the analysis. Especially in topic modeling that only analyzes abstracts of published articles, the topics revealed by the findings will be limited to the information provided by the abstracts.

Moreover, using existing classification systems to train topic modeling algorithms may be inadequate or incomplete in identifying emerging or dominant topics within a given research area. Regardless of the dataset size, topic modeling results may sometimes encompass multiple themes within a single topic, raising concerns about consistency and comprehensiveness. To enhance the validity of findings derived from topic modeling, researchers are advised to employ various strategies, such as establishing clear criteria, applying triangulation in data analysis, and consulting expert opinions. Another limitation of this analytical method is its strong reliance on technology. When researchers depend exclusively on machine learning-generated findings, they may risk losing a broader and more critical perspective on the subject. Additionally, the technology-dependent nature of topic modeling necessitates expertise in data mining analyses. This presents a challenge for practitioners unfamiliar with such methods, particularly when applying topic modeling for purposes such as evaluating psychological counseling processes through session transcripts or assessing employee needs and feedback within an organization.

Another limitation of the findings is that topic modeling analysis, especially when analyzing counseling transcripts or posts on social networking sites, may not be able to assess context, recognize nuances, and fully capture the complexity of individual experiences. Thus, while topic modeling analysis may reveal themes and patterns in textual data, it may not fully reflect the context in which this data was obtained and the complexity of psychological phenomena. Another limitation concerns the interpretation of the findings. The interpretation of the topics revealed by the findings and the terms represented under the topics are made by the researchers. This carries the risk of being influenced by subjective processes such as subjective opinions, judgments, or biases of the researchers. Therefore, interpretation of the findings requires scrutiny and evaluation. Finally, especially in topic modeling analyses using counseling transcripts, the findings may only reveal topics and patterns related to the context under study. Therefore, there will be limitations in the generalizability of these findings to a broader counseling context or to other groups.

As a precautionary measure against these limitations, it is recommended that researchers planning applications and studies based on topic modeling analysis in human-centered fields such as psychology, psychological counseling, and guidance consider ethical issues both within the framework of their discipline's ethical principles and publication ethics. Researchers should first ensure that the data sources used—whether academic articles, reports, documents, or posts from social media and various online platforms—are compiled in a manner that ensures an unbiased, fair, and representative sample of the collected content. This step is crucial, as it influences the content, frequency, and distribution of topics that emerge from the analysis. The findings generated from this process play a critical role in shaping conclusions and informing practice. Furthermore, the interpretation of the findings should be grounded in transparency. Researchers should explicitly acknowledge the limitations of the data analysis method, as well as their own role and potential biases in the analytical process, such as their interests in the subject matter. Additionally, they should consider whether the study warrants further investigation and assess the broader implications of its findings for specific groups, individuals, or societies. For example, researchers should be mindful that identifying issues that could stigmatize a particular group (e.g., LGBTI individuals) or society (e.g., associations with diseases or mental health conditions) or designing a study in a way that could lead to such interpretations would not align with the principle of social good. When utilizing more sensitive data sources, such as therapy transcripts, responsible data use is essential. This includes obtaining informed consent, ensuring participant confidentiality, removing any identifying information related to clients or groups from transcripts, maintaining privacy, and safeguarding data security for all data sources.

CONCLUSION AND SUGGESTIONS FOR FUTURE RESEARCH

Studies using topic modeling as a data analysis method in counseling can make significant contributions to theoretical foundations, research, and practice in terms of their findings. First, as in the examples above, important developments, rising trends, and areas where research and practice are concentrated in counseling can be analyzed through topic modeling. This analysis can contribute to the determination of the situation based on the distribution of topics, to a better understanding of the various aspects of the examined topic, to the evaluation of which components should be included in practices, and to the review of existing practices and policies. When we look at the research on this subject so far, there are few studies in school counseling, career counseling and family and marriage counseling. For example, studies such as revealing the trend topics in research on marital satisfaction, looking at the changes in the distribution of topics over the years, and comparing them with divorce rates can provide meaningful inputs to preventive and supportive studies to be conducted on this subject. This analysis based on data mining can be carried out by analyzing research all over the world, or by examining only the research in our country. More specific results can be reached for practice with more local and culture-sensitive findings. Similarly, research on job satisfaction, organizational commitment, and burnout can make important contributions to career counseling. Findings from these studies can contribute to the review of current processes and the development of policies in terms of evaluating work productivity and supporting economic development. Topic modeling analysis can also be conducted by focusing on more specific topics. For example, Karacan-Ozdemir et al. (In review) examined the skills necessary to be successful in STEM careers through topic modeling analysis. The findings of this study may contribute to the design of programs to support students' STEM career development and the evaluation of

existing programs. Similarly, all research studies conducted in the field of school counseling in Türkiye to date can be examined with topic modeling and changes in the distribution of topics over the years can be interpreted in the light of social crises, developments and changes, and inferences can be made for the future and important inputs can be provided for the design of comprehensive developmental guidance programs in this direction. As a more specific suggestion, handbooks, and practitioner guides on crisis intervention in schools can be analyzed through topic modeling and the contents of commonly used crisis intervention programs in our country can be reviewed according to the prominent topic distributions.

Topic modeling can also be used to reveal research trends in the field of counseling with a bottom-up approach and to capture the agenda of emerging psychological research (Bittermann & Fisher, 2018). In addition, it can be integrated into systematic literature review studies and in this way, literature review studies can be expanded (Bittermann & Fisher, 2018; Chen & Wojcik, 2016). The findings can guide researchers in designing novel studies. In addition, awareness can be raised about which research topics are on the agenda and more prominent, and which research topics are neglected. Moreover, it can contribute to the development of new research questions and hypotheses (Chen & Wojcik, 2016).

Another suggestion for future research on the potential uses of topic modeling analysis, as in the examples given above, is to examine posts, shares, and feedback on social media and the internet, such as Facebook posts, tweets (Banks et al., 2018). Research can be conducted in this direction in psychological counseling and guidance. For example, examining comments and posts on specific websites and platforms related to computer games through topic analysis and furthermore, revealing the changes in topics according to age groups can make significant contributions to understanding, explaining, and evaluating addiction to computer games in various aspects. The findings may provide important inputs in assessing unforeseen risk factors, identifying computer addiction tendencies specific to different developmental stages (such as children, adolescents, and adults) and planning or evaluating existing help processes accordingly. Similarly, analyzing user comments and posts on betting sites with topic analysis to understand gambling addiction, which is a rising problem, may provide important contributions to future studies on this subject.

Topic modeling analysis in psychological counseling and guidance can produce rich and insightful findings in cross-cultural research. These studies can make meaningful contributions in terms of understanding which topics are prominent in distinct cultures and what differences occur according to culture and developing culturally sensitive practices (Liu & Gao, 2021). For example, Otsuka et al. (2021) examined which topics in psychological research are universal and which topics vary across geographical regions and over time. Among the main findings, they report that research activities and trends in psychology are strongly influenced by both geographical regions and time trends. Accordingly, they found that issues related to clinical research were higher in Europe than in North-Central America. Another example of cross-cultural research is the work of a multicultural research group, including the author of this article, who conducted research on social emotional learning skills. The research team analyzed qualitative data collected from educators using topic modeling to uncover both universal topics and themes that varied by culture in decision-making skills. Preliminary findings indicate that, unlike in other cultures, expressions of emotions (e.g., anger, feeling pressure) emerged as prominent topics and terms in Türkiye (Kanzaki et al., in review). As demonstrated in these examples, numerous issues in psychological counseling and guidance can be explored through topic modeling analysis, with findings that can contribute meaningfully to the field. For instance, topic modeling can be used to examine factors influencing the effectiveness of the counseling process, key themes in the therapeutic relationship, counseling methods and approaches perceived as effective, and prominent reasons for client dropout or discontinuation of counseling. Through such analyses, both universal and culture-specific themes can be identified. This, in turn, can provide valuable insights into the culturally specific application of Western-oriented counseling theories and approaches. These examples can be multiplied.

Another use of topic modeling in counseling can be in measurement and evaluation applications. Topic modeling can be used to evaluate the effectiveness of counseling and guidance practices by examining the differences in topics between pretest and posttest (Liu & Gao, 2021). For example, in evaluating the effectiveness of practices such as developmental classroom guidance programs, peer bullying programs, and

resiliency-related interventions implemented in schools, the distribution of prominent topics in the feedback from practitioners and/or participants can be examined. In addition to providing information on the effectiveness of existing programs, these evaluations can provide meaningful inputs for the improvement of programs and the development of future programs. Similarly, for example, findings from the evaluation of school-based programs implemented across the country can provide important contributions to the development of policies related to these issues (Oh et al., 2017). As another example, within the scope of school counseling studies, students' autobiographies, which are anonymized in terms of confidentiality and privacy, or their responses to open-ended questions in various questionnaires can be analyzed with topic modeling. Thus, important findings can be accessed to identify needs and problem areas in the school. In addition, longitudinal studies examining the changes in topics in students' autobiographies over the years can contribute to the understanding of developmental needs and differences. Similarly, by retrospectively analyzing student autobiographies at the same level (e.g., 7th grade level), it is possible to examine whether there are changes in the prominent topics in the same age group over the years, and if so, in which direction, and to make predictions for the future. An example of its use in the fields of career counseling and organizational psychology can be the analysis of employee satisfaction or feedback in an institution through topic analysis of institutional documents such as annual reports, strategic plans, and meeting minutes. The results obtained here can be useful for organizations in assessing their strengths and weaknesses, employee needs, and identifying development areas (Banks et al., 2018).

Topic modeling can also be used in the development of counselor education. As detailed in the examples and suggestions above, the topics that come to the forefront as a result of examining the research and issues in the field of counseling and guidance through topic modeling analysis can be used in reviewing the curriculum used in counselor education and updating the programs accordingly in case there are issues that are not addressed or overlooked. This can help to structure counselor education in line with current research trends and best practices (Oh et al., 2017). For example, research findings such as identifying the prominent topics in effective counseling approaches and practices and revealing the culture-specific determinants of the therapeutic relationship can be addressed in counseling practice courses and can be taken into consideration in the supervision processes given to students within the scope of this course. On the other hand, in future research, ethical principles and standards can be established for the use of topic modeling in counseling and guidance, considering the ethical issues and requirements that may come to the fore in all the suggestions given so far (Liu & Gao, 2021).

Considering the limitations of topic modeling as a data analysis method, as mentioned before, suggestions can be made to integrate it with other data analysis methods or to use it by blending different topic modeling analyses. First, topic modeling analysis in the field of counseling and guidance can be conducted within the scope of grounded theory approach, one of the qualitative research designs (Banks et al., 2018), and/or can be used in combination with qualitative data analysis methods such as content analysis or sentiment analysis (Kobayashi et al., 2018).

As in studies aimed at identifying prominent topics in psychological counseling and guidance (e.g., Oh et al., 2017), more advanced topic modeling techniques, such as dynamic topic modeling, can be employed to analyze the evolution of topics over time. This approach enables researchers to capture the shifting nature of topic structures and gain a more nuanced understanding of their historical development (Kobayashi et al., 2018). Additionally, efforts can be directed toward enhancing topic modeling methods to accommodate more complex data types, such as longitudinal and network data (Otsuka et al., 2021). Furthermore, advancements in neural probabilistic approaches, including word embedding and deep learning, can be accelerated to refine topic modeling techniques (Banks et al., 2018). Beyond text analysis, future research can explore the integration of topic modeling with alternative data sources, such as audio and video, to expand its applicability and improve analytical depth (Kobayashi et al., 2018). Finally, as mentioned under the limitations, studies can be conducted to develop validated topic models to assess the validity of the results obtained through topic modeling analysis; standardized procedures can be developed to validate the topics presented using established criteria to assess the quality of topic models (Otsuka et al., 2021).

Conclusion

This article examines the use of topic modeling, a text mining data analysis method based on machine learning, within the field of psychological counseling and guidance amid rapid technological advancements. While research on this subject remains limited, topic modeling holds significant potential for applications in psychological counseling and guidance. Conducting topic modeling analyses within qualitative research across various domains of psychological counseling and guidance can provide valuable insights. These findings can contribute to the advancement of the field by informing theoretical knowledge, research, practice, and policy development. However, further studies are needed to explore its potential while also addressing the inherent limitations of the data analysis method, the challenges associated with the research process, and the ethical considerations involved.

REFERENCES

- Atkins, D. C., Rubin, T. N., Steyvers, M., Doeden, M. A., Baucom, B. R., & Christensen, A. (2012). Topic models: A novel method for modeling couple and family text data. *Journal of Family Psychology*, 26, 816–827. <https://doi.org/10.1037/a0029607>
- Atzil-Slonim, D., Juravski, D., Bar-Kalifa, E., Gilboa-Schechtman, E., Tuval-Mashiach, R., Shapira, N., & Goldberg, Y. (2021). Using topic models to identify clients' functioning levels and alliance ruptures in psychotherapy. *Psychotherapy (Chicago, Ill.)*, 58(2), 324–339. <https://doi.org/10.1037/pst0000362>
- Banks, G. C., Woznyj, H. M., Wesslen, R. S., & Ross, R. L. (2018). A review of best practice recommendations for text analysis in R (and a user-friendly app). *Journal of Business and Psychology*, 33(4), 445–459. <https://doi.org/10.1007/s10869-017-9528-3>
- Balahadia, F. F., Fernando, M. C. G., & Juanatas, I. C. (2016). Teacher's performance evaluation tool using opinion mining with sentiment analysis. In IEEE region symposium (TENSYP) (pp. 95–98). <https://doi.org/10.1109/TENCONSpring.2016.7519384>
- Baumer, E. P., Mimno, D., Guha, S., Quan, E., & Gay, G. K. (2017). Comparing grounded theory and topic modeling: Extreme divergence or unlikely convergence? *Journal of the Association for Information Science and Technology*, 68, 1397–1410. <https://doi.org/10.1002/asi.23786>
- Buenano-Fernandez, D., Gonzalez, M., Gil, D., & Lujan-Mora, S. (2020). Text mining of open-ended questions in self-assessment of university teachers: An LDA topic modeling approach. *IEEE Access*, 8, 35318–35330. <https://doi.org/10.1109/ACCESS.2020.2974983>
- Blei, D. M. (2012). Probabilistic topic models. *Communications of the ACM*, 55(4), 77-84. <https://doi.org/10.1145/2133806.2133826>
- Blei, D. M., & Jordan, M. I. (2003, July). Modeling annotated data. In *Proceedings of the 26th annual international ACM SIGIR Conference on Research and Development In Information Retrieval* (pp. 127-134).
- Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent dirichlet allocation. *Journal of Machine Learning Research*, 3(Jan), 993-1022.
- Blei, D., & Lafferty, J. (2009). Topic Models. In A. Srivastava & M. Sahami (Eds.), *Text Mining: Classification, Clustering, and Applications* (pp. 71-94). Chapman & Hall/CRC.
- Bittermann, A., & Fischer, A. (2018). How to identify hot topics in psychology using topic modeling. *Zeitschrift für Psychologie*, 226(1), 3–13. <https://doi.org/10.1027/2151-2604/a000318>
- Carron-Arthur, B., Reynolds, J., Bennett, K., Bennett, A., & Griffiths, K. M. (2016). What is all the talk about? Topic modelling in a mental health Internet support group. *BMC Psychiatry*, 16(1), 367. <https://doi.org/10.1186/s12888-016-1073-5>
- Chen, E. E., & Wojcik, S. P. (2016). A practical guide to big data research in psychology. *Psychological Methods*, 21(4), 458–474. <https://doi.org/10.1037/met0000111>
- Feinerer, I., & Wild, F. (2007). Automated coding of qualitative interviews with latent semantic analysis. In Heinrich C. Mayr, Dimitris Karagiannis (Ed.), *"Lecture notes in informatics" (LNI): P-107, Information Systems Technology and its Applications, May 23-25, 2007, Kharkiv, Ukraine* (pp. 66 - 78). Köllen Druck+Verlag GmbH.

- Finch, W. H., Hernández Finch, M.E.H., French, B. F., & Claire, B. (2021). Latent Dirichlet Analysis to Aid Equity in the Identification for Gifted Education. *Psychological Test and Assessment Modeling* 63(3), 305-36. https://www.psychologie-aktuell.com/fileadmin/Redaktion/Journale/ptam-2021-3/PTAM_3-2021_2_kor.pdf
- Franz, P. J., Nook, E. C., Mair, P., & Nock, M. K. (2020). Using Topic Modeling to Detect and Describe Self-Injurious and Related Content on a Large-Scale Digital Platform. *Suicide & Life-threatening Behavior*, 50(1), 5–18. <https://doi.org/10.1111/sltb.12569>
- Gencoglu, B., Helms-Lorenz, M., Maulana, R., Jansen, E. P., & Gencoglu, O. (2023). Machine and expert judgments of student perceptions of teaching behavior in secondary education: Added value of topic modeling with big data. *Computers & Education*, 193, 104682. <https://doi.org/10.1016/j.compedu.2022.104682>.
- Griffiths, T. L., & Steyvers, M. (2004). Finding scientific topics. *Proceedings of the National Academy of Sciences of the United States of America*, 101(1 Supplement), 5228–5235.
- Grimmer, J., & Stewart, B. M. (2013). Text as data: The promise and pitfalls of automatic content analysis methods for political texts. *Political Analysis*, 21(3), 267–297. <https://doi.org/10.1093/pan/mps028>
- Hopkins, D. J., & King, G. (2010). A method of automated nonparametric content analysis for social science. *American Journal of Political Science*, 54(1), 229–247. <https://doi.org/10.1111/j.1540-5907.2009.00428.x>
- Howes, C., Purver, M., & McCabe, R. (2013). Howes, C., Purver, M., & McCabe, R. (2013). Investigating Topic Modelling for Therapy Dialogue Analysis. Proceedings of IWCS 2013 Workshop on Computational Semantics in Clinical Text. <https://qmro.qmul.ac.uk/xmlui/handle/123456789/5300>
- Jayarathne M., & Jayatilake B. (2020). Predicting personality using answers to open-ended interview questions. *IEEE Access*, 8, 115345–115355. 10.1109/ACCESS.2020.3004002
- Jivani, A. G. (2011). A comparative study of stemming algorithms. *Journal of Computer Applications in Technology*, 2(6), 1930-1938.
- Kanzaki, M. Suzuki, H., Abdoulaye, O., Agnieszka O., Agnieszka, S., Andrei, A., Bacanlı, F., Donnelly, H., Ferrari, L., Karacan-Ozdemir, N., Kourmousi, N., Kounenou, K., Marsay, G., Park, C., Scoda, A.D., Solberg, V. S. H., & Zuzanna Z. (in review). Cultural translation of career decision-making in adolescents. *International Journal for Educational and Vocational Guidance*.
- Karacan-Ozdemir, N., Park, C., & Solberg, S. (in review). STEM career competencies. A qualitative framework. *Journal of Career Development*.
- Kadhim, A. I. (2018). An evaluation of preprocessing techniques for text classification. *International Journal of Computer Science and Information Security (IJCSIS)*, 16(6), 22-32.
- Kandula, S., Curtis, D., Hill, B., & Zeng-Treitler, Q. (2011). *Use of topic modeling for commending relevant education material to diabetic patients*. In Annual symposium proceedings/AMIA symposium (pp. 674–682).
- Kee, Y. H., Li, C., Kong, L. C., Tang, C. J., & Chuang, K. (2019). Scoping review of mindfulness research: A topic modelling approach. *Mindfulness*, 10, 1474–1488. <https://doi.org/10.1007/s12671-019-01136-4>

- Kobayashi, V. B., Mol, S. T., Berkers, H. A., Kismihók, G., & Den Hartog, D. N. (2018). Text Mining in Organizational Research. *Organizational Research Methods*, 21(3), 733-765. <https://doi.org/10.1177/1094428117722619>
- Krippendorff, K. (2003). *Content Analysis: An Introduction to Its Methodology* (2nd ed.). Sage Publications, Inc.
- Kumar, A. & Paul, A. (2016). *Mastering Text Mining with R*. Packt Publishing.
- Liu, J., & Gao, L. (2021). Analysis of topics and characteristics of user reviews on different online psychological counseling methods. *International Journal of Medical Informatics*, 147, 104367. <https://doi.org/10.1016/j.ijmedinf.2020.104367>
- Liu, S., Zhang, R.Y., & Kishimoto, T. (2020). Analysis and prospect of clinical psychology based on topic models: hot research topics and scientific trends in the latest decades. *Psychol Health Med*, 26(4), 395-407. <https://doi.org/10.1080/13548506.2020.1738019>
- Longo, L. (2020). Empowering qualitative research methods in education with artificial intelligence. In: Costa, A., Reis, L., Moreira, A. (Eds) Computer supported qualitative research. WCQR 2019. Advances in Intelligent Systems and Computing, vol 1068. Springer, Cham. https://doi.org/10.1007/978-3-030-31787-4_1
- Maier, D., Waldherr, A., Miltner, P., Wiedemann, G., Niekler, A., Keinert, A., Pfetsch, B., Heyer, G., Reber, U., Häussler, T., Schmid-Petri, H. ve Adam, S. (2018). Applying lda topic modeling in communication research: toward a valid and reliable methodology. *Communication Methods and Measures*, 12(3), 93–118. <https://doi.org/10.1080/19312458.2018.1430754>
- Mobin, I., Mridha, M. F., & Mahmud, S. M. H. (2024). Exploratory analysis of suicidal tendency in depression investigation social media post. Research Square. <https://doi.org/10.21203/rs.3.rs-3823798/v1>
- Newman, M. E. (2005). Power laws, Pareto distributions and Zipf's law. *Contemporary Physics*, 46, 323–351.
- Ni, C. F., Lundblad, R., Dykeman, C., Bolante, R., & Łabuński, W. (2023). Content analysis of psychological first aid training manuals via topic modelling. *European Journal of Psychotraumatology*, 14(2), 2230110. <https://doi.org/10.1080/20008066.2023.2230110>
- Nowlin, M.C. (2015). Modeling issue definitions using quantitative text analysis. *Policy Studies Journal*, 44(3), 309-331. <https://doi.org/10.1111/psj.12110>
- Roberts, M. E., Stewart, B. M., Tingley, D., Lucas, C., Leder-Luis, J., Gadarian, S. K., ... & Rand, D. G. (2014). Structural topic models for open-ended survey responses. *American Journal of Political Science*, 58(4), 1064-1082. <https://doi.org/10.1111/ajps.12103>
- Rohrer, J. M., Brümmer, M., Schmukle, S. C., Goebel, J., & Wagner, G. G. (2017). What else are you worried about? – Integrating textual responses into quantitative social science research. PLoS ONE, 12(7), e0182156. <https://doi.org/10.1371/journal.pone.0182156>
- Royston, R. P., Lin, N., & Amey R. C. (2022). Applying Topic Modeling to narrative statements to determine how job responsibility descriptions differ by leadership level and demographic factors. Society for Industrial and Organizational Psychology Annual Conference APR 27-30
- Seah, J. H., K., & Shim, K. J. (2018). Data mining approach to the detection of suicide in social media: A case study of Singapore. 2018 IEEE International Conference on Big Data (Big Data): Seattle, WA, December 10-13: Proceedings. 5442-5444. https://ink.library.smu.edu.sg/sis_research/4340

- Tang, J., Meng, Z., Nguyen, X., Mei, Q., & Zang, M. (2014). Understanding the limiting factors of topic modeling via posterior contraction analysis. *Proceedings of the 31st International Conference on Machine Learning*, 337–345.
- Oh, J., Stewart, A., & Phelps, R. (2017). Topics in the journal of counseling psychology, 1963–2015. *Journal of Counseling Psychology*, 64(6), 604-615. <https://doi.org/10.1037/cou0000218>
- Otsuka, S., Ueda, Y., & Saiki, J. (2021). Diversity in psychological research activities: quantitative approach with topic modeling. *Front Psychol*, 16(12), 773-916. <https://doi.org/10.3389/fpsyg.2021>
- Vijayarani, S., Ilamathi, M. J., & Nithya, M. (2015). Preprocessing techniques for text mining-an overview. *International Journal of Computer Science & Communication Networks*, 5(1), 7-16.
- Wang, S. H., Ding, Y., Zhao, W., Huang, Y. H., Perkins, R., Zou, W., & Chen, J. J. (2016). Text mining for identifying topics in literature about adolescent substance use and depression. *BMC Public Health*, 16, 279. <https://doi.org/10.1186/s12889-016-2932-1>
- Wang, W., Feng, Y. ve Dai, W. (2018). Topic analysis of online reviews for two competitive products using Latent Dirichlet Allocation. *Electronic Commerce Research and Applications*, 29, 142–156. <https://doi.org/10.1016/J.ELERAP.2018.04.003>
- Weiss, S.M., Indurkha, N. ve Zhang, T. (2010). Fundamentals of Predictive Text Mining. Springer-Verlag London Limited.
- Zhang, C., Yu, M. C., & Marin, S. (2021). Exploring public sentiment on enforced remote work during COVID-19. *The Journal of Applied Psychology*, 106(6), 797–810. <https://doi.org/10.1037/apl0000933>