

Yapay Zekâ Üretimli Sağlık İçeriklerinde Uzmanlık, Güven ve Meşruiyetin Söylemsel İnşası: ChatGPT, Gemini ve Claude Üzerine Bir Analiz

The Discursive Construction of Expertise, Trust, and Communicative Legitimacy in AI-Generated Health Content: A Comparative Analysis of ChatGPT, Gemini and Claude

Ezgi Zengin DEMİRBİLEK 

Araştırma Makalesi | Research Article

Başvuru Received: 15.01.2026 ■ Kabul Accepted: 24.06.2026

ÖZ

Dijital sağlık iletişimde yapay zekâ destekli içerik üretiminin yaygınlaşması, bu içeriklerin kullanıcılar tarafından nasıl algılandığını önemli bir iletişimsel sorun hâline getirmektedir. Yapay zekâ tarafından üretilen sağlık metnlerinin, tıbbi uzmanlık yetkisine sahip olmaksızın yüksek düzeyde güven ve otorite algısı yaratabilmesi, bilginin doğruluğundan bağımsız bir etki alanına işaret etmektedir. Bu durum, dijital ortamlarda sağlık bilgisinin dolaşımında uzmanlık, güven ve meşruiyet kavramlarının yeniden tartışılmasını gerekli kılmaktadır. Bu araştırma, ChatGPT, Gemini ve Claude sistemlerine yöneltilen 30 sağlık sorusuna verilen toplam 90 metinsel yanıtın Norman Fairclough'un Eleştirel Söylem Analizi yaklaşımıyla incelendiği ampirik nitel bir çalışmadır. Çalışmada, yapay zekâ üretimli sağlık içeriklerinde güvenin nasıl kurulduğu ve uzmanlığın hangi iletişimsel mekanizmalar üzerinden meşrulaştırıldığı incelenmektedir. Weber'in otorite kuramı, Foucault'nun bilgi-iktidar ilişkisi, Bourdieu'nun sembolik sermaye yaklaşımı, Luhmann ve Giddens'in güven kavramları ile Habermas'ın iletişimsel meşruiyet tartışmaları bütüncül bir kavramsal çerçevede ele alınmaktadır. Çalışma, yapay zekânın uzmanlığı doğrudan temsil etmekten ziyade, uzmanlığın dilsel ve söylemsel göstergelerini simüle ettiğini ileri sürmektedir. Bu simülasyonun, kullanıcıların eleştirel değerlendirme süreçlerini zayıflatabileceği ve dijital sağlık iletişimde yeni bir meşruiyet biçimi üretebileceği tartışılmaktadır.

Anahtar Kelimeler: Yapay Zekâ, Sağlık İletişimi, Güven, Otorite, İletişimsel Meşruiyet.

ABSTRACT

The growing prevalence of AI-generated content in digital health communication has made users' perceptions of such content a critical communicative issue. Health texts produced by artificial intelligence can generate high levels of trust and authority despite lacking medical expertise, indicating an influence that operates independently of informational accuracy. This situation calls for a re-examination of expertise, trust, and legitimacy in the circulation of health information in digital environments. This study is an empirical qualitative investigation based on the analysis of 90 textual responses generated by ChatGPT, Gemini, and Claude to 30 health-related questions using Norman Fairclough's Critical Discourse Analysis framework. The study examines how trust is constructed in AI-generated health content and through which communicative mechanisms expertise is legitimized. It proposes a comprehensive conceptual framework grounded in key theories from communication and social theory, including Weber's authority typology, Foucault's knowledge-power relations, Bourdieu's symbolic capital, Luhmann's and Giddens's approaches to trust, and Habermas's notion of communicative legitimacy. The article argues that AI-generated health content does not represent expertise directly but rather simulates its linguistic and discursive markers. This simulation may weaken users' critical evaluation processes and foster a new form of communicative legitimacy in the public circulation of health information. The study offers a theoretical contribution to the health communication literature by rethinking trust, perceived expertise, and the communicative role of artificial intelligence in digital contexts.

Keywords: Artificial Intelligence, Health Communication, Trust, Authority, Communicative Legitimacy.

Giriş

Dijitalleşme, sağlık bilgisinin üretim, dolaşım ve tüketim biçimlerini köklü biçimde dönüştürmüştür. Geleneksel olarak tıbbi otoriteye, kurumsal yapılarla ve yüzyüze etkileşime dayalı olarak şekillenen sağlık iletişimi, dijital ortamların yaygınlaşmasıyla birlikte çok aktörlü, çok kanallı ve giderek daha karmaşık bir yapıya bürünmüştür. Bu dönüşüm sürecinde bireyler, sağlıkla ilgili bilgiye yalnızca hekimler, sağlık kurumları ya da bilimsel yayınlar aracılığıyla değil; dijital platformlar, sosyal medya içerikleri ve giderek artan biçimde yapay zekâ destekli uygulamalar üzerinden de erişebilmektedir. Bu yeni iletişim ortamı, sağlık bilgisinin doğruluğu kadar, bu bilginin nasıl algılandığı, hangi koşullarda güven ürettiği ve hangi iletişimsel mekanizmalarla meşrulaştığı sorularını da gündeme taşımaktadır.

Son yıllarda üretken yapay zekâ teknolojilerinin yaygınlaşması, sağlık iletişimi alanında yeni bir kırılma noktası yaratmıştır. Yapay zekâ tabanlı sistemler, sağlıkla ilgili sorulara hızlı, akıcı ve çoğu zaman bilimsel terminolojiye dayanan yanıtlar üretebilmekte; kullanıcılarla etkileşim kuran “konuşan bir özne” gibi konumlanabilmektedir. Bu durum, yapay zekânın yalnızca teknik bir araç olmaktan çıkarak, sağlık bilgisinin dolaşımında iletişimsel aktör haline gelmesine yol açmaktadır. Ancak bu aktör, tıbbi uzmanlık yetkisine, mesleki sorumluluğa ya da kurumsal hesap verebilirliğe sahip değildir. Buna rağmen, yapay zekâ tarafından üretilen sağlık içeriklerinin kullanıcılar nezdinde yüksek düzeyde güven ve otorite algısı yaratabilmesi, sağlık iletişimi açısından dikkatle ele alınması gereken bir olguya işaret etmektedir.

Son yıllarda sağlık iletişimi literatüründe yapay zekâ destekli sağlık içeriklerine yönelik ilgi önemli ölçüde artmıştır. Mevcut çalışmalar büyük ölçüde yapay zekâ sistemlerinin doğruluk düzeyi, yanlış bilgi üretme riski, etik sorunları ve klinik karar süreçlerindeki performansı üzerine odaklanmaktadır. Bu araştırmalar, yapay zekâ tarafından üretilen sağlık bilgilerinin güvenilirliğini ve potansiyel risklerini değerlendirme açısından önemli katkılar sunmaktadır. Ancak söz konusu çalışmaların önemli bir bölümü, kullanıcıların bu

içerikleri neden güvenilir bulunduğu, yapay zekâ tarafından üretilen metinlerin hangi iletişimsel özellikler aracılığıyla uzmanlık izlenimi oluşturduğu ve sağlık bilgisinin kamusal dolaşımında nasıl meşruiyet kazandığı sorularına sınırlı düzeyde odaklanmaktadır. Dolayısıyla mevcut literatürde, yapay zekâ üretimli sağlık içeriklerinin yalnızca bilgi nesneleri olarak değil, aynı zamanda güven, otorite ve uzmanlık algısı üreten iletişimsel pratikler olarak nasıl işlediğine ilişkin kuramsal bir boşluk bulunmaktadır.

Mevcut sağlık iletişimi literatürü, dijital ortamlarda yanlış bilgi, dezenformasyon ve bilgi kirliliği gibi sorunlara yoğunlaşmış; özellikle bilginin doğruluğu ve bilimsel geçerliliği üzerinden tartışmalar yürütmüştür. Ancak yapay zekâ üretimli sağlık içerikleri, bu tartışmaları aşan bir boyut sunmaktadır. Burada tartışılması gereken husus yalnızca bilginin doğru ya da yanlış olması değil; doğru kabul edilmesinin, uzman söylemi olarak algılanmasının ve meşru bir bilgi kaynağına dönüşmesinin hangi iletişimsel süreçler aracılığıyla gerçekleştiğidir. Başka bir ifadeyle, yapay zekâ tabanlı sağlık iletişiminde asıl sorun, uzmanlığın fiilen var olup olmamasından çok, uzmanlığın nasıl ve ne kadar temsil edildiği ile ne düzeyde algılandığıdır.

İletişim kuramı, bu noktada önemli bir analitik zemin sunmaktadır. Otorite, güven ve meşruiyet gibi kavramlar, iletişim literatüründe bilginin toplumsal dolaşımını anlamak için önemli bir konuma sahiptir. Uzmanlık, yalnızca teknik bilgiye sahip olmakla değil; bu bilginin toplumsal olarak tanınması, kabul edilmesi ve güven üretmesiyle anlam kazanmaktadır. Dijital ortamlarda ise bu süreçler giderek daha fazla dil, üslup, biçim ve söylem üzerinden işlemektedir. Yapay zekâ tarafından üretilen sağlık metinleri, bilimsel terminoloji, düzenli anlatım, tarafsız ton ve empatik ifadeler aracılığıyla uzmanlığın iletişimsel göstergelerini başarıyla taklit edebilmektedir. Bu durum, kullanıcıların sağlık bilgisini değerlendirme süreçlerinde kaynağın niteliğinden ziyade, iletişimsel sunuma odaklanmasına neden olabilmektedir.

Bu bağlamda, yapay zekâ üretimli sağlık içerikleri, sağlık iletişiminde yeni bir güven rejiminin ortaya çıktığını düşündürmektedir. Kullanıcılar, insan öznesine ait olmayan bu içerikleri, kişisel çıkar ya da ideolojik yönelim taşımadığı varsayımıyla daha tarafsız ve güvenilir olarak algılayabilmektedir. Bu algı, yapay zekânın teknik doğasından çok, iletişimsel performansına dayanmaktadır. Dolayısıyla yapay zekâ, sağlık iletişiminde güvenin ve otoritenin üretildiği yeni bir aracı yapı olarak incelenmelidir.

Bu çalışma, yapay zekâ üretimli sağlık içeriklerinde uzmanlık, güven ve iletişimsel meşruiyet algısının hangi söylemsel göstergeler aracılığıyla üretildiğini açıklamayı amaçlamaktadır. Böylece çalışma, yapay zekâ üretimli sağlık içeriklerini yalnızca doğruluk veya yanlışlık ekseninde değerlendiren yaklaşımların ötesine geçerek, bu içeriklerin iletişimsel otoriteyi nasıl kurduğuna odaklanmaktadır.

Kavramsal Arka Plan

Dijital sağlık iletişiminde bilginin toplumsal dolaşımı, yalnızca bilimsel doğruluk üzerinden değil, bu bilginin hangi iletişimsel koşullar altında üretildiği ve alımlandığı üzerinden şekillenmektedir. Sağlık bilgisi, tarihsel olarak kurumsal otorite, mesleki uzmanlık ve yüz yüze etkileşim üzerinden meşruiyet kazanırken, dijitalleşmeyle birlikte bu meşruiyet biçimleri dönüşüme uğramıştır (Briggs & Hallin, 2016). Özellikle dijital platformlar aracılığıyla sunulan sağlık içerikleri, bilginin kaynağını görünmezleştirerek söylemin biçimsel özelliklerini daha belirleyici hale getirmektedir (Thompson, 1995).

Uzmanlık, iletişim kuramında teknik bilgiye sahip olmanın ötesinde, toplumsal olarak tanınan ve kabul edilen bir konumlanma biçimi olarak ele alınmaktadır. Weber'in (1978) rasyonel-yasal otorite kavramsallaştırması, uzmanlığın meşruiyetinin kurumsal ve söylemsel çerçeveler aracılığıyla kurulduğunu göstermektedir. Foucault'nun (1980) bilgi-iktidar ilişkisine dair çalışmaları ise özellikle sağlık alanında hangi bilginin "doğru" kabul edileceğinin söylemsel düzenekler

aracılığıyla belirlendiğini ortaya koymaktadır. Bu yaklaşım, sağlık bilgisinin nötr bir içerik olmaktan ziyade, belirli iktidar ilişkileri içinde üretildiğini vurgulamaktadır.

Bourdieu'nun (1991) sembolik sermaye kavramı, uzmanlığın iletişimsel boyutunu anlamak açısından önemli bir katkı sunmaktadır. Bilimsel dil, teknik terminoloji ve nesnel anlatım biçimleri, uzmanlığın toplumsal olarak tanınmasını sağlayan sembolik göstergeler üretmektedir. Sağlık iletişiminde bu göstergeler, bilginin içeriğinden bağımsız olarak otorite algısı yaratabilmektedir (Bucchi, 2008). Bu durum, uzmanlığın algısal olarak inşa edilen bir kategori olduğunu göstermektedir.

Güven kavramı da sağlık iletişiminin merkezinde yer almaktadır. Luhmann'a (2017) göre güven, bireylerin karmaşık ve belirsiz sistemler içinde karar verebilmesini mümkün kılan bir toplumsal mekanizmadır. Sağlık bağlamında bireyler, çoğu zaman bilgi eksikliği ve risk koşulları altında hareket etmekte; bu nedenle uzmanlık iddiası taşıyan iletişimsel aktörlere yönelmektedir (Giddens, 1991). Dijital ortamlarda ise bu güven ilişkisi, yüz yüze etkileşimden ziyade metinsel ve söylemsel ipuçları üzerinden kurulmaktadır (Hardin, 2002).

Habermas'ın (1984) iletişimsel eylem kuramı, meşruiyetin rasyonel tartışma ve karşılıklı anlama süreçleriyle üretildiğini ileri sürerken, dijital sağlık iletişimi bu ideal koşullardan giderek uzaklaşmaktadır. Yapay zekâ aracılı sağlık içerikleri, etkileşim izlenimi yaratsa da gerçek bir kamusal müzakere alanı sunmamaktadır. Thompson'ın (2005) medyatikleşmiş otorite yaklaşımı, bu tür içeriklerin neden sorgulanmadan kabul edilebildiğini açıklamaktadır. Kaynağın belirsizleşmesi, söylemin doğal ve tarafsız algılanmasına yol açabilmektedir.

Bu bağlamda yapay zekâ destekli sağlık içerikleri, uzmanlık ve güvenin yeni bir iletişimsel düzlemde üretildiğini göstermektedir. Yapay zekânın insan benzeri dil üretme kapasitesi, Goffman'ın (1956) benlik sunumu yaklaşımıyla birlikte düşünüldüğünde, uzmanlığın bir performans

olarak sahnelendiği bir iletişimsel pratik ortaya koymaktadır. Bu performans, hesap verebilirlik ve etik sorumluluk gibi unsurlardan bağımsız biçimde güven üretebilmektedir (Couldry & Hepp, 2016).

Baudrillard'ın (1994) simülasyon kavramı, bu süreci daha ileri bir noktaya taşımaktadır. Yapay zekâ, uzmanlığın kendisini değil, uzmanlığa ait göstergeleri yeniden üretmekte; böylece gerçek uzmanlık ile simüle edilmiş uzmanlık arasındaki sınırı bulanıklaştırmaktadır. Dijital sağlık iletişiminde bu bulanıklık, bilginin doğruluğundan ziyade algısal ikna gücünü ön plana çıkarmaktadır.

Son yıllarda sağlık iletişimi literatürü, dijital platformlarda güvenin nasıl üretildiğine dair eleştirel yaklaşımlar geliştirmiştir (Lupton, 2016; Rudd & Anderson, 2006). Ancak yapay zekâ üretimli içerikler, bu tartışmaları daha da karmaşık hale getirmektedir. Yapay zekâ, ne bireysel bir özne ne de kurumsal bir aktör olarak konumlanmakta; buna rağmen her iki yapının iletişimsel avantajlarını bir arada sunabilmektedir (Floridi vd., 2018). Bu durum, sağlık bilgisinin kamusal dolaşımında yeni bir iletişimsel meşruiyet biçiminin ortaya çıktığını düşündürmektedir.

Bu çalışma, söz konusu dönüşümü, sağlık bilgisinin doğruluğundan ziyade, uzmanlığın ve güvenin iletişimsel olarak nasıl üretildiği sorusu üzerinden ele almaktadır. Yapay zekâ aracılı sağlık iletişimi, iletişim kuramlarının sunduğu kavramsal araçlar sayesinde, teknik bir yenilikten öte, toplumsal ve söylemsel bir olgu olarak tartışılabilir hale gelmektedir.

Yapay Zekâ, Sağlık İletişimi ve Antropomorfik Temsil

Yapay zekâ sistemlerinin sağlık iletişiminde giderek daha görünür hâle gelmesi, bu teknolojilerin yalnızca bilgi sağlayan teknik araçlar olarak değil, aynı zamanda iletişimsel aktörler olarak değerlendirilmesini gerekli kılmaktadır. İnsan-makine iletişimi literatürü, yapay zekânın kullanıcılarla kurduğu etkileşimin geleneksel insan-insan iletişiminden tamamen farklı olmadığını, aksine birçok durumda sosyal ve

ilişkisel iletişim dinamiklerini yeniden ürettiğini göstermektedir. Bu bağlamda Guzman ve Lewis (2020), yapay zekâ sistemlerinin mevcut iletişim paradigmalarının ötesinde yeni bir iletişimsel faillik biçimi ortaya çıkardığını ileri sürmektedir. Özellikle sağlık iletişimi bağlamında bu durum, güven, risk, uzmanlık ve meşruiyet gibi kavramların yeniden değerlendirilmesini zorunlu kılmaktadır.

Sağlık alanındaki kullanılabilecek yapay zekâ destekli sohbet robotları, sanal danışmanlar ve büyük dil modelleri, bireylerin sağlık bilgisine erişimini kolaylaştırma potansiyeline sahiptir. Bununla birlikte, sağlık bilgisinin yüksek risk taşıyan doğası nedeniyle bu sistemlere duyulan güven kritik bir mesele hâline gelmektedir. Yapılan çalışmalar, kullanıcıların yapay zekâ sistemlerine yönelik güveninin büyük ölçüde şeffaflık, açıklanabilirlik ve tutarlılık algısına bağlı olduğunu göstermektedir (Floridi, 2023). Özellikle yapay zekâ tarafından üretilen sağlık içeriklerinde ortaya çıkan küçük hataların dahi yalnızca sisteme değil, bu sistemi kullanan kurumlara yönelik güveni de zedeleyebildiği belirtilmektedir. Floridi'nin (2023) "zekâsız faillik" (agency without intelligence) yaklaşımı bu noktada açıklayıcı bir çerçeve sunmaktadır. Buna göre kullanıcılar, yapay zekâ sistemlerini gerçek bir uzmanla iletişim kuruyormuş gibi değerlendirebilmekte; ancak sistemler gerçekte insan uzmanlığının temel unsurları olan bilinç, muhakeme ve etik sorumluluk kapasitesine sahip değildir. Dolayısıyla sağlık iletişiminde güvenin önemli bir bölümü, gerçek uzmanlıktan ziyade uzmanlığın iletişimsel görünümüne dayanmaktadır.

Bu süreç, insan-makine iletişimi alanında önemli bir yere sahip olan Computers are Social Actors (CASA) paradigmasıyla da açıklanabilmektedir. Nass ve arkadaşları tarafından geliştirilen bu yaklaşım, bireylerin bilgisayarlar ve dijital sistemlerle etkileşim kurarken sosyal aktörlere yönelik geliştirdikleri zihinsel şemaları otomatik olarak devreye soktuklarını ileri sürmektedir (Nass & Moon, 2000). Sağlık iletişimi bağlamında bu durum, kullanıcıların yapay zekâ destekli sistemleri yalnızca bilgi sağlayan araçlar olarak değil, aynı zamanda tavsiye veren, yönlendiren ve uzmanlık

sunan sosyal aktörler olarak algılamalarına yol açabilmektedir. Sağlık hizmetlerinde yapay zekâya duyulan güveni inceleyen güncel çalışmalar da güven oluşumunda yapay zekâ okuryazarlığı, bireysel psikolojik özellikler ve algılanan fayda gibi faktörlerin belirleyici olduğunu göstermektedir. Bununla birlikte antropomorfizm, gizlilik kaygıları ve hata sonrasında güvenin yeniden inşası da kullanıcı değerlendirmelerini etkileyen temel değişkenler arasında yer almaktadır.

Antropomorfizm, yapay zekâ sistemlerinin insan özellikleriyle ilişkilendirilmesi sürecini ifade etmektedir. Epley vd.'na (2007) göre bireyler, sosyal dünyayı anlamlandırma eğilimleri doğrultusunda insan olmayan varlıklara da niyet, duygu ve kişilik atfetmektedir. Yapay zekâ sistemlerinde kullanılan isimler, insan benzeri dil yapıları, empatik ifadeler ve kişiselleştirilmiş yanıtlar bu eğilimi güçlendirmektedir. Özellikle sağlık iletişiminde gerçekleştirilen çalışmalar, sohbet robotlarının sosyal mevcudiyet düzeyinin arttıkça kullanıcıların sağlanan bilgiyi daha kaliteli, güvenilir ve uzmanlık temelli olarak değerlendirdiğini göstermektedir. Yapay zekâ doktorlarının sosyal aktörlere ne kadar benzediğinin, kullanıcıların mesaj kalitesine ilişkin algılarını doğrudan etkilediği ortaya konulmaktadır. Bu durum, uzmanlık algısının yalnızca içerik doğruluğuna değil, aynı zamanda iletişimsel sunum biçimine de bağlı olduğunu göstermektedir.

Sağlık chatbotlarının kişilik özellikleri, iletişim tarzları ve sosyal kimlik göstergeleri üzerine yapılan deneysel araştırmalar da benzer sonuçlara işaret etmektedir. Kullanıcılar, kendilerine benzeyen veya kendileriyle uyumlu iletişim tarzı sergileyen yapay zekâ sistemlerini daha güvenilir bulmakta ve bu sistemleri yeniden kullanma eğilimi göstermektedir. Ancak bu durum önemli bir paradoks yaratmaktadır. Çünkü algılanan güvenilirlik ile gerçek uzmanlık arasında her zaman doğrudan bir ilişki bulunmamaktadır. Özellikle sağlık alanında kullanıcıların güven değerlendirmeleri, içerik doğruluğundan çok sistemin sergilediği uzmanlık performansına dayanabilmektedir.

Literatürde dikkat çekilen bir diğer konu ise üretici yapay zekâ sistemlerinin yanlış bilgi ve önyargı üretme potansiyelidir. Büyük dil modelleri sağlıkla ilgili içerikleri son derece akıcı ve ikna edici biçimde sunabilmelerine rağmen, zaman zaman hatalı, eksik veya yanıltıcı bilgiler üretebilmektedir. Bu nedenle bazı araştırmacılar, üretici yapay zekâ teknolojilerinin sağlık iletişiminde yeni bir “bilgi pandemisi” (infodemic) riski oluşturduğunu ileri sürmektedir. Sağlık alanında güvenilir olmayan bilgilerin hızla yayılması, yalnızca bireysel sağlık kararlarını değil, aynı zamanda sağlık kurumlarına ve bilimsel uzmanlığa duyulan toplumsal güveni de etkileyebilmektedir.

Guzman'ın (2020) insan ve makine arasındaki ontolojik sınırların giderek bulanıklaştığına ilişkin değerlendirmesi, sağlık iletişimi bağlamında özel bir önem taşımaktadır. Kullanıcılar, yapay zekâ tarafından üretilen sağlık içeriklerini çoğu zaman insan uzmanların ürettiği içeriklerle benzer biçimde değerlendirmekte veya bu ayrımı yapma gereği duymamaktadır. Bu durum, yanlış sınıflandırmanın yalnızca bilişsel bir hata olmaktan çıkıp iletişimsel otorite üretimine dönüşmesine neden olmaktadır. Dolayısıyla sağlık iletişiminde yapay zekâya duyulan güven, yalnızca teknolojik performansın değil; aynı zamanda uzmanlığın, otoritenin ve meşruiyetin nasıl temsil edildiğinin de bir sonucudur.

Sonuç olarak sağlık iletişimi alanında yapay zekâ kullanımı; insan-makine iletişimi, antropomorfizm, güven ve uzmanlık literatürlerinin kesişim noktasında yer almaktadır. CASA paradigmasının ortaya koyduğu sosyal tepki mekanizmaları, antropomorfizmin motivasyonel temelleri, Floridi'nin zekâsız faillik yaklaşımı ve Guzman'ın ontolojik bulanıklaşma tartışmaları birlikte değerlendirildiğinde, yapay zekâ sistemlerinin sağlık iletişiminde yalnızca bilgi sağlayan araçlar değil, aynı zamanda güven ve otorite üreten iletişimsel aktörler olarak işlev gördüğü anlaşılmaktadır. Bu nedenle yapay zekâ üretimli sağlık içeriklerinin nasıl güven kazandığı ve uzmanlık algısı oluşturduğu sorusu, sağlık iletişimi araştırmaları açısından giderek daha önemli bir inceleme alanı hâline gelmektedir.

Yöntem

Araştırma Tasarımı

Bu çalışma, yapay zekâ üretimli sağlık içeriklerinde uzmanlık, güven ve iletişimsel meşruiyet algısının hangi söylemsel mekanizmalar aracılığıyla üretildiğini inceleyen ampirik nitel bir araştırma olarak tasarlanmıştır. Çalışmanın ampirik olarak kurgulanmasının temel nedeni, yapay zekâya ilişkin kuramsal varsayımları soyut düzeyde tartışmak yerine, farklı yapay zekâ sistemleri tarafından üretilen gerçek sağlık metinlerinden hareketle uzmanlık ve güven üretimine ilişkin söylemsel örüntüleri incelemektir. Araştırma, kuramsal tartışmaların ötesine geçerek doğrudan yapay zekâ sistemleri tarafından üretilmiş metinleri analiz nesnesi olarak kullanmakta ve bulgularını bu metinlerden elde edilen ampirik veriler üzerine inşa etmektedir.

Araştırmanın nitel olarak tasarlanmasının temel gerekçesi ise araştırma sorularının ölçülebilir değişkenler arasındaki ilişkileri belirlemeyi değil, metinlerde yer alan anlam üretim süreçlerini, söylemsel stratejileri ve uzmanlık temsillerini açıklamayı amaçlamasıdır. Uzmanlık, güven ve iletişimsel meşruiyet gibi kavramlar sayısal göstergelerle doğrudan ölçülebilen olgular olmaktan ziyade, dilsel ve söylemsel pratikler aracılığıyla inşa edilen toplumsal yapılardır. Bu nedenle araştırmanın amacı, belirli söylemsel örüntülerin nesnikle ortaya çıktığını belirlemekten çok, bu örüntülerin nasıl işlediğini ve hangi iletişimsel işlevleri yerine getirdiğini açıklamaktır.

Bu bağlamda nicel içerik analizi veya deneysel araştırma tasarımlarından farklı olarak, nitel yaklaşım yapay zekâ üretimli sağlık metinlerinde uzmanlık ve güven üretiminin altında yatan anlamlandırma süreçlerini inceleme olanağı sunmaktadır. Çalışmanın odağı kullanıcıların yapay zekâya duyduğu güven düzeyini ölçmek değil, bu güvenin metinsel düzeyde hangi söylemsel mekanizmalar aracılığıyla üretilmeye çalışıldığını ortaya koymaktır. Bu nedenle ampirik veri temeline dayanan nitel araştırma yaklaşımının araştırma sorularıyla en uyumlu yöntemsel

çerçeveyi sunduğu değerlendirilmiştir.

Araştırmanın Soruları

Bu araştırmanın temel amacı, yapay zekâ tarafından üretilen sağlık içeriklerinde uzmanlık, güven ve iletişimsel meşruiyet algısının hangi söylemsel stratejiler aracılığıyla üretildiğini incelemektir. Sağlık iletişimi alanında yapay zekâ sistemlerinin kullanımının yaygınlaşması, bu sistemlerin yalnızca bilgi sağlayan teknik araçlar olarak değil, aynı zamanda uzmanlık ve güven üreten iletişimsel aktörler olarak değerlendirilmesini gerekli kılmaktadır. Buna karşın mevcut araştırmaların önemli bir bölümü yapay zekâ çıktılarının doğruluğu, güvenilirliği ve etik boyutları üzerine yoğunlaşırken, uzmanlık ve güven algısının metinsel düzeyde nasıl üretildiği sorusu görece sınırlı biçimde ele alınmıştır.

Bu çalışma, yapay zekâ üretimli sağlık içeriklerini yalnızca bilgi aktaran teknik çıktılar olarak değil, belirli dilsel ve söylemsel tercihler aracılığıyla uzmanlık, güven ve meşruiyet üreten iletişimsel yapılar olarak ele almaktadır. Bu doğrultuda araştırma, yapay zekâ sistemlerinin sağlık alanındaki otoritesinin hangi söylemsel göstergeler aracılığıyla kurulduğunu ve bu göstergelerin güven ile iletişimsel meşruiyet üretimine nasıl katkı sağladığını ortaya koymayı amaçlamaktadır.

Araştırma, yapay zekâ üretimli sağlık içeriklerinde uzmanlık ve güven üretimine ilişkin söylemsel mekanizmaları ortaya çıkarmayı amaçlamaktadır. Bu doğrultuda aşağıdaki araştırma sorularına yanıt aranmıştır:

AS1: Yapay zekâ üretimli sağlık içeriklerinde uzmanlık hangi dilsel ve söylemsel göstergeler aracılığıyla temsil edilmektedir?

Bu soru, yapay zekâ sistemlerinin bilimsel terminoloji kullanımı, nedensel açıklama biçimleri, teknik kavramlara başvurma şekilleri ve uzman söylemine özgü dilsel özellikler aracılığıyla nasıl bir uzmanlık performansı sergilediğini incelemeyi amaçlamaktadır.

AS2: Yapay zekâ üretimli sağlık içeriklerinde uzmanlık temsilini oluşturan söylemsel göstergeler, güven ve iletişimsel meşruiyet üretimine nasıl katkı sağlamaktadır?

Bu soru, empatik anlatım, belirsizliğin yönetimi, tarafsızlık vurguları, risk çerçeveleme biçimleri ve uzmanlık göstergelerinin kullanıcı nezdinde güvenilir ve meşru bir bilgi kaynağı izlenimi oluşturabilecek söylemsel işlevlerini incelemeyi amaçlamaktadır.

Yöntemsel Yaklaşım

Bu çalışma, yapay zekâ üretimli sağlık içeriklerinde uzmanlık, güven ve iletişimsel meşruiyetin hangi söylemsel mekanizmalar aracılığıyla üretildiğini inceleyen ampirik nitel bir araştırma olarak tasarlanmıştır. Araştırmanın amacı, yapay zekâ sistemleri tarafından üretilen sağlık bilgilerinin doğruluğunu ya da klinik geçerliliğini değerlendirmek değil; sağlık metinlerinde uzmanlık ve güvenin hangi dilsel ve söylemsel stratejiler aracılığıyla temsil edildiğini ortaya koymaktır. Bu nedenle araştırma, yapay zekâ sistemlerinin ne söylediğinden çok, sağlık bilgisini nasıl sunduğuna ve bu sunum aracılığıyla nasıl bir uzmanlık ve güven görünümü oluşturduğuna odaklanmaktadır.

Araştırma soruları uzmanlık, güven ve iletişimsel meşruiyet gibi söylem aracılığıyla inşa edilen toplumsal olgulara yöneldiğinden, çalışma nitel araştırma yaklaşımı çerçevesinde yürütülmüştür. Bu kapsamda Norman Fairclough'un (1995) Eleştirel Söylem Analizi yaklaşımı benimsenmiştir. Fairclough'a göre söylem yalnızca bilgi aktaran bir araç değil, aynı zamanda toplumsal ilişkileri, iktidar yapılarını ve meşruiyet süreçlerini üreten bir pratiktir. Bu nedenle eleştirel söylem analizi, yapay zekâ üretimli sağlık içeriklerinde uzmanlık, güven ve otoritenin nasıl temsil edildiğini incelemek için uygun bir yöntemsel çerçeve sunmaktadır.

Araştırmada Fairclough'un üç boyutlu analiz modeli temel alınmıştır. Metinsel analiz aşamasında sağlık metinlerinde kullanılan bilimsel terminoloji, nedensel açıklamalar, koşullu

ifadeler, empatik söylemler ve tarafsızlık vurguları incelenmiştir. Söylemsel pratik analizinde bu dilsel tercihlerin uzmanlık ve güven üretimindeki işlevleri değerlendirilmiş, toplumsal pratik analizinde ise elde edilen bulgular sağlık iletişimi, insan-makine iletişimi ve yapay zekâ literatürü bağlamında yorumlanmıştır. Bu yaklaşım, araştırma sorularında yer alan uzmanlık, güven ve iletişimsel meşruiyetin söylemsel olarak nasıl üretildiğini incelemek için bütüncül bir analitik çerçeve sağlamaktadır.

Araştırmanın tekrarlanabilirliğini artırmak amacıyla veri üretim sürecinde kullanılan soru kategorilerine ilişkin temsili örnekler Ek 1'de sunulmuştur. Veri seti, gündelik sağlık iletişimi bağlamında kullanıcıların sıklıkla yöneltebileceği açık uçlu sağlık sorularından oluşturulmuş ve semptom yorumlama, yaşam tarzı önerileri ile genel sağlık bilgisi olmak üzere üç kategori altında yapılandırılmıştır. Tüm yapay zekâ sistemlerine aynı soru ifadeleri yöneltilmiş, yönlendirici veya doğrudan tanı koymaya yönelik ifadeler kullanılmamıştır. Bu yaklaşım, soru biçiminden kaynaklanabilecek etkileri azaltarak karşılaştırılabilir veri elde etmeyi amaçlamıştır.

Analiz bölümünde kullanılan doğrudan alıntılar, veri setinde tekrar eden söylemsel örüntüleri temsil etme kapasitesine göre seçilmiştir. Bu nedenle alıntılar, tekil veya istisnai örneklerden değil; farklı yapay zekâ sistemleri ve veri kategorilerinde benzer biçimde gözlemlenen uzmanlık, güven ve iletişimsel meşruiyet üretimine ilişkin söylemsel stratejileri yansıtan ifadeler arasından belirlenmiştir.

Veri Toplama Süreci

Araştırmanın veri seti, üretken yapay zekâ sistemlerine yöneltilen sağlık temalı sorulara verilen metinsel yanıtlardan oluşturulmuştur. Veri toplama süreci Şubat 2026 ile Nisan 2026 tarihleri arasında gerçekleştirilmiştir. Veriler, ilgili yapay zekâ sistemlerinin kamuya açık web tabanlı kullanıcı arayüzleri üzerinden elde edilmiştir. Veri toplama sürecinde ChatGPT, Gemini ve Claude sistemlerinin kullanıcıların erişimine açık standart metin üretim arayüzleri kullanılmıştır.

Tablo 1

Veri Setinin Kapsamı ve Soru Kategorileri

Soru Kategorisi	Amaç	Soru Sayısı
Semptom yorumlama	Uzmanlık temsili, risk iletişimi ve belirsizlik yönetimi stratejilerini incelemek	10
Yaşam tarzı önerileri	Güven oluşturma ve rehberlik söylemlerini incelemek	10
Genel sağlık bilgisi	Uzmanlık, otorite ve iletişimsel meşruiyet göstergelerini incelemek	10
Toplam		30

Tablo 2

Veri Setinin Oluşturulmasında Kullanılan Dahil Etme ve Hariç Tutma Kriterleri

Dahil Etme Kriterleri	Gerekçe
Semptom yorumlamaya yönelik sağlık soruları	Uzmanlık temsili, risk iletişimi ve belirsizlik yönetimini incelemek
Yaşam tarzı ve sağlıklı yaşam önerileri	Güven ve rehberlik söylemlerini incelemek
Genel sağlık bilgisi içeren sorular	Uzmanlık ve iletişimsel meşruiyet göstergelerini incelemek
Yapay zekâ sistemlerinin metinsel yanıtları	Söylemsel stratejilerin analiz edilebilmesi
Şubat 2026–Nisan 2026 arasında üretilen yanıtlar	Veri setinin zamansal tutarlılığını sağlamak

Araştırmada amaçlı örnekleme yöntemi kullanılmıştır. Soru havuzu oluşturulurken sağlık iletişimi literatürü ile kullanıcıların dijital ortamlarda sıklıkla yönelttiği sağlık soruları dikkate alınmıştır. Araştırma sorularıyla uyumlu olarak uzmanlık, güven ve iletişimsel meşruiyet üretimine ilişkin söylemsel stratejileri görünür kılacak üç kategori belirlenmiştir: semptom yorumlama, yaşam tarzı önerileri ve genel sağlık bilgisi.

Bu kapsamda her kategoriden 10 soru olmak üzere toplam 30 sağlık sorusu oluşturulmuştur. Aynı soru seti ChatGPT, Gemini ve Claude sistemlerine yöneltilmiş, böylece toplam 90 metinsel yanıt elde edilmiştir. Soruların tüm platformlara aynı biçimde yöneltmesi, farklı yapay zekâ sistemlerinde ortaya çıkan ortak söylemsel örüntülerin karşılaştırmalı olarak incelenmesini amaçlamaktadır.

Araştırma sorularıyla uyumlu bir veri seti oluşturabilmek amacıyla sağlık temalı sorular belirli kategoriler altında yapılandırılmıştır. Kategorilerin oluşturulmasında sağlık iletişimi literatürü, çevrim içi sağlık bilgi arama davranışları ve kullanıcıların yapay zekâ sistemlerine yönelttiği yaygın sağlık soruları dikkate alınmıştır. Her kategori, yapay zekâ üretimli sağlık içeriklerinde uzmanlık, güven ve iletişimsel meşruiyet üretimine ilişkin farklı söylemsel boyutların incelenmesine olanak sağlayacak şekilde tasarlanmıştır. Veri

setinin kapsamı ve soru kategorileri Tablo 1'de sunulmaktadır

Her soru ChatGPT, Gemini ve Claude sistemlerine ayrı ayrı yöneltmiştir. Böylece toplam 90 metinsel yanıt analiz kapsamına alınmıştır.

Tablo 1'de görüldüğü üzere veri seti, sağlık iletişiminin farklı bağlamlarını temsil edecek şekilde yapılandırılmıştır. Semptom yorumlama soruları uzmanlık ve risk yönetimi söylemlerini görünür kılarken, yaşam tarzı önerileri güven ve rehberlik stratejilerinin incelenmesine olanak sağlamaktadır. Genel sağlık bilgisi soruları ise yapay zekâ sistemlerinin uzmanlık, otorite ve iletişimsel meşruiyet üretiminde kullandıkları söylemsel göstergelerin analiz edilmesini mümkün kılmaktadır. Bu yapı sayesinde araştırma, yapay zekâ üretimli sağlık içeriklerinde uzmanlık temsili ve bu temsilin güven üretimine nasıl katkı sağladığını farklı sağlık iletişimi bağlamlarında karşılaştırmalı olarak inceleyebilmektedir.

Araştırmanın odağını sağlık bilgilerinin doğruluğundan ziyade uzmanlık, güven ve iletişimsel meşruiyet üretimine ilişkin söylemsel stratejilere yöneltmek amacıyla doğrudan tanı talep eden, reçete veya ilaç önerisi isteyen ve acil klinik müdahale gerektiren sorular analiz kapsamı dışında bırakılmıştır.

Tablo 2’de görüldüğü üzere veri seti, yapay zekâ sistemlerinin sağlık alanında uzmanlık, güven ve iletişimsel meşruiyet üretimine ilişkin söylemsel stratejilerini görünür kılacak içeriklerle sınırlandırılmıştır. Bu sınırlama, araştırmının sağlık bilgilerinin doğruluğunu değerlendirmekten ziyade, uzmanlık ve güven üretimine ilişkin söylemsel mekanizmaları inceleme amacıyla doğrudan ilişkilidir.

Örneklem ve Örneklem Seçim Gerekçesi

Araştırmada amaçlı örnekleme yöntemi kullanılmıştır. Amaçlı örnekleme, araştırma sorularıyla doğrudan ilişkili ve incelenen olgu hakkında zengin veri sağlayabilecek örneklerin seçilmesine olanak tanıyan bir nitel araştırma yaklaşımıdır (Patton, 2015). Bu çalışmada amaç, yapay zekâ üretimli sağlık içeriklerinde uzmanlık, güven ve iletişimsel meşruiyet üretimine ilişkin söylemsel mekanizmaları incelemek olduğundan, araştırma sorularını doğrudan karşılayabilecek veri kaynaklarının seçilmesi hedeflenmiştir.

Bu doğrultuda örneklem, sağlıkla ilgili sorulara metinsel yanıt üreten ve kullanıcılar tarafından yaygın biçimde kullanılan üç üretken yapay zekâ sistemi olan ChatGPT, Gemini ve Claude tarafından oluşturulan metinlerden meydana gelmektedir. Söz konusu platformlar, günümüzde sağlıkla ilgili bilgi arayışında bulunan kullanıcılar tarafından sıklıkla başvurulanan yapay zekâ sistemleri arasında yer almakta ve sağlıkla ilgili konularda ayrıntılı, açıklayıcı ve uzman görüşünü andıran yanıtlar üretebilmektedir. Bu nedenle araştırmının inceleme konusu olan uzmanlık temsili ve güven üretimi süreçlerini gözlemlemek açısından uygun veri kaynakları olarak değerlendirilmiştir.

Yapay zekâ üretimli sağlık metinlerinin örneklem olarak seçilmesinin temel gerekçesi, bu içeriklerin sağlık iletişiminde giderek daha görünür hale gelen yeni bir iletişimsel aktör tarafından üretilmesidir. Geleneksel sağlık iletişiminde uzmanlık çoğunlukla hekimler, sağlık kurumları veya bilimsel otoriteler aracılığıyla temsil edilirken, üretken yapay zekâ sistemleri herhangi bir mesleki yetkiye veya kurumsal sorumluluğa sahip olmaksızın uzmanlık

görünümü oluşturabilmektedir. Bu durum, yapay zekâ üretimli metinleri sağlık iletişiminde uzmanlık, güven ve iletişimsel meşruiyetin nasıl üretildiğini incelemek açısından özgün bir analiz nesnesi hâline getirmektedir.

Araştırma kapsamında her kategoriden 10 soru olmak üzere toplam 30 sağlık sorusu belirlenmiş ve aynı soru seti ChatGPT, Gemini ve Claude sistemlerine yöneltilmiştir. Nitel araştırmalarda örneklem büyüklüğü istatistiksel temsil gücünden ziyade araştırma sorularını açıklayabilecek veri çeşitliliğine ve analitik derinliğe dayanmaktadır. Bu çalışmada farklı sağlık iletişimi bağlamalarını temsil eden üç soru kategorisi ile üç farklı yapay zekâ sisteminden elde edilen toplam 90 metin, uzmanlık, güven ve iletişimsel meşruiyet üretimine ilişkin tekrar eden söylemsel örüntüleri ortaya koyabilecek yeterli veri çeşitliliğini sağlamıştır.

Bu örneklem yapısı, araştırmının iki temel sorusuyla doğrudan ilişkilidir. Farklı sağlık konularını temsil eden sorular ve farklı yapay zekâ sistemlerinden elde edilen yanıtlar sayesinde, uzmanlık temsiliinin hangi söylemsel göstergeler aracılığıyla kurulduğu ve bu göstergelerin güven ile iletişimsel meşruiyet üretimine nasıl katkı sağladığı karşılaştırmalı olarak incelenebilmiştir.

Bu araştırmada amaç, yapay zekâ sistemlerinin sağlık iletişiminde kullandıkları söylemsel stratejilerin istatistiksel yaygınlığını ölçmek değil, uzmanlık, güven ve iletişimsel meşruiyet üretimine ilişkin tekrar eden söylemsel örüntüleri ortaya koymaktır. Bu nedenle örneklem büyüklüğü istatistiksel temsil gücüne göre değil, söylemsel çeşitlilik ve analitik derinlik ilkeleri doğrultusunda belirlenmiştir. Üç farklı yapay zekâ sistemi ve üç farklı sağlık iletişimi kategorisinden elde edilen toplam 90 metinsel yanıtın, araştırma sorularını yanıtlayabilecek yeterli veri çeşitliliğini sağladığı değerlendirilmiştir.

Veri Analizi

Verilerin analizinde Norman Fairclough’un (1995) üç boyutlu Eleştirel Söylem Analizi modeli kullanılmıştır. Araştırmının uzmanlık, güven ve

Tablo 3

Veri Setinde Gözlemlenen Temel Söylemsel Örüntüler

Söylemsel Örüntü	ChatGPT	Gemini	Claude
Bilimsel terminoloji	✓	✓	✓
Nedensel açıklama	✓	✓	✓
Koşullu ifadeler	✓	✓	✓
Empatik anlatım	✓	✓	✓
Tarafsızlık vurgusu	✓	✓	✓

iletişimsel meşruiyetin söylemsel olarak nasıl üretildiğini incelemeyi amaçlaması nedeniyle Fairclough'un yaklaşımı araştırma sorularıyla doğrudan uyumlu bir analitik çerçeve sunmaktadır.

Analiz süreci üç aşamada yürütülmüştür. İlk aşamada metinler ayrıntılı biçimde okunarak bilimsel terminoloji kullanımı, nedensel açıklama biçimleri, koşullu ifadeler, risk ve belirsizlik yönetimine ilişkin söylemler, empatik ifadeler ve tarafsızlık vurguları belirlenmiştir. İkinci aşamada bu dilsel göstergelerin uzmanlık, güven ve iletişimsel meşruiyet üretimindeki işlevleri değerlendirilmiştir. Üçüncü aşamada ise elde edilen bulgular sağlık iletişimi, insan-makine iletişimi ve yapay zekâ literatürü bağlamında yorumlanmıştır. Böylece yapay zekâ üretimli sağlık içeriklerinde ortaya çıkan söylemsel örüntülerin sağlık alanında uzmanlık ve güvenin yeniden üretimine nasıl katkı sağladığı tartışılmıştır.

Eleştirel söylem analizinin yorumlayıcı niteliği nedeniyle araştırmacının analitik sürece etkisini en aza indirmek amacıyla çeşitli önlemler alınmıştır. Analiz sürecinde kuramsal varsayımların veriye dayatılmasından kaçınılmış, metinlerde ortaya çıkan söylemsel örüntüler esas alınmıştır. Tüm metinler birden fazla kez okunmuş, farklı yapay zekâ sistemlerinden ve veri kategorilerinden elde edilen yanıtlar sürekli karşılaştırılmıştır. Elde edilen yorumlar Fairclough'un (1995) analitik çerçevesi doğrultusunda değerlendirilmiş ve araştırma sorularıyla doğrudan ilişkili olmayan yorumlar analiz dışında bırakılmıştır. Bulguların sunumunda doğrudan metin örneklerine yer verilerek yorumların veri temelli biçimde izlenebilmesi sağlanmıştır.

Sınırlılıklar

Bu çalışmanın bazı sınırlılıkları bulunmaktadır.

Öncelikle araştırma, yapay zekâ üretimli sağlık içeriklerinde kullanılan söylemsel stratejilere odaklanmakta olup, bu içeriklerin kullanıcıları üzerindeki etkilerini, kullanıcı deneyimlerini veya güven düzeylerini doğrudan incelememektedir. Bu nedenle bulgular, kullanıcı davranışlarına ilişkin ampirik çıkarımlar olarak değerlendirilmemelidir. İkinci olarak araştırma, Şubat 2026 ile Nisan 2026 tarihleri arasında ChatGPT, Gemini ve Claude sistemlerinden elde edilen metinlerle sınırlıdır. Üretken yapay zekâ sistemlerinin sürekli güncellenen ve değişen yapıları dikkate alındığında, farklı zamanlarda elde edilecek yanıtlar farklı söylemsel özellikler gösterebilir. Son olarak çalışma, sağlık iletişiminde uzmanlık, güven ve iletişimsel meşruiyet üretimine ilişkin söylemsel örüntüleri incelemeyi amaçlayan nitel bir araştırmadır. Bu nedenle bulgular istatistiksel genelleme amacı taşımamakta; ancak yapay zekâ üretimli sağlık içeriklerinin iletişimsel özelliklerini anlamaya yönelik analitik ve kuramsal bir katkı sunmaktadır.

Bulgular ve Tartışma

Bu bölümde, yapay zekâ tarafından üretilen sağlık içeriklerine yönelik söylem analizinden elde edilen bulgular sunulmakta ve iletişim kuramları çerçevesinde tartışılmaktadır. Analiz, yapay zekâ metninin yalnızca sağlık bilgisi aktarmadığını; aynı zamanda uzmanlık, güven ve iletişimsel meşruiyeti belirli söylemsel stratejiler aracılığıyla ürettiğini göstermektedir. Çalışma, sistemler arası performans farklılıklarından ziyade ChatGPT, Gemini ve Claude'da ortak olarak gözlemlenen söylemsel örüntülere odaklanmaktadır.

Tablo 3, ChatGPT, Gemini ve Claude tarafından üretilen sağlık içeriklerinde tespit edilen temel söylemsel örüntülerin karşılaştırmalı görünümünü sunmaktadır. İşaretlemeler, ilgili söylemsel

Tablo 4

Analiz Sürecinde Belirlenen Temel Söylemsel Örüntüler

Söylemsel Örüntü	Gözlemlendiği Veri Türü	İletişimsel İşlev
Bilimsel terminoloji kullanımı	Semptom yorumlama, genel sağlık bilgisi	Uzmanlık temsili
Nedensel açıklama zincirleri	Semptom yorumlama	Rasyonellik ve uzmanlık üretimi
Koşullu ve olasılık içeren ifadeler	Tüm kategoriler	Belirsizlik yönetimi ve güven üretimi
Empatik anlatım	Semptom yorumlama, yaşam tarzı önerileri	İlişkisel güven oluşturma
Tarafsızlık ve sınır çizme ifadeleri	Tüm kategoriler	İletişimsel meşruiyet üretimi

özelliğın analiz edilen veri setinde tekrar eden bir örüntü olarak gözlemlendiğini göstermektedir. Fairclough'un (1995) Eleştirel Söylem Analizi yaklaşımı doğrultusunda belirlenen ortak söylemsel yapıların platformlar arasındaki benzerliğini ortaya koymayı amaçlamaktadır.

Yapılan söylem analizi sonucunda, yapay zekâ üretimli sağlık içeriklerinde uzmanlık, güven ve meşruiyet algısının belirli ve tekrar eden söylemsel örüntüler aracılığıyla kurulduğu görülmüştür. Bu örüntüler, iletişimsel işlevleri ve ilişkili kuramsal çerçevelerle birlikte Tablo 1'de özetlenmektedir. Tablo, aşağıda ayrıntılı olarak tartışılan bulgular için analitik bir çerçeve sunmaktadır.

Tablo 4'te yer alan ifadeler, analiz sürecinde belirlenen söylemsel örüntüleri temsil etmek amacıyla sunulmuştur. Tablo, uzmanlık, güven ve iletişimsel meşruiyet üretiminde öne çıkan temel söylemsel stratejilerin görünümünü özetlemektedir.

Uzmanlık Söyleminin Dilsel ve Biçimsel Olarak İnşası

Semptom yorumlama kategorisinde yer alan yanıtların büyük bölümünde teknik ve bilimsel terminoloji kullanımına rastlanmıştır. ChatGPT'nin ürettiği yanıtların çoğunda "inflamasyon", "bağışıklık yanıtı", "nörolojik süreçler" ve "hormonal değişimler" gibi kavramlar kullanılırken, Gemini yanıtlarında "metabolik süreçler", "fizyolojik mekanizmalar" ve "biyolojik faktörler" ifadeleri öne çıkmıştır. Claude ise benzer biçimde "sinir sistemi yanıtları", "fizyolojik adaptasyon" ve "biyolojik süreçler" gibi açıklamalara yer vermiştir. Platformlar arasında kullanılan terminoloji farklılaşsa da uzmanlığı çağrıştıran bilimsel

dilin tüm veri setinde tekrar eden bir örüntü oluşturduğu görülmüştür.

Analiz edilen metinlerde uzmanlık, doğrudan tıbbi otorite iddiasıyla değil; bilimsel terminoloji kullanımı, nedensel açıklama zincirleri ve teknik kavramlara başvurma yoluyla temsil edilmektedir. Özellikle semptom yorumlama kategorisindeki yanıtlarda yapay zekâ sistemlerinin bilimsel açıklama biçimlerini sistematik olarak kullandığı görülmüştür. Örneğin; ChatGPT tarafından üretilen bir yanıtta, *"Bu belirtiler, bağışıklık sisteminin verdiği inflamatuvar yanıtlarla ilişkili olabilir"* ifadesine yer verilirken, Gemini aynı soruya verdiği yanıtta, *"Metabolik süreçlerde meydana gelen geçici değişimler bu durumu etkileyebilir"* açıklamasını kullanmıştır. Benzer biçimde Claude tarafından üretilen bir yanıtta ise *"Sinir sisteminin strese verdiği fizyolojik yanıtlar bu belirtilerin ortaya çıkmasına katkıda bulunabilir"* ifadesi yer almaktadır.

Bu örnekler, yapay zekâ sistemlerinin sağlık bilgisini yalnızca aktarmadığını, aynı zamanda bilimsel akıl yürütme biçimlerini yeniden üreterek uzmanlık görünümünü oluşturduğunu göstermektedir. Dikkat çekici olan nokta, yanıtların kesin tanımlar sunmaktan kaçınmasına rağmen, bilimsel terminoloji ve nedensel açıklamalar aracılığıyla güçlü bir uzmanlık performansı sergilemesidir.

Bu bulgu, yapay zekâ sistemlerinin kullanıcılar tarafından uzman bilgi kaynağı olarak algılanmasında dilsel sunum biçiminin önemli rol oynadığını ortaya koyan insan-makine iletişimi araştırmalarıyla uyumludur (Guzman & Lewis, 2020; Floridi, 2023). Sağlık iletişimi bağlamında değerlendirildiğinde, uzmanlık yalnızca bilginin

içeriğinden değil, bilginin hangi dilsel çerçeve içinde sunulduğundan da beslenmektedir.

Bourdieu'nun (1986) sembolik sermaye yaklaşımı açısından değerlendirildiğinde, bilimsel terminoloji sağlık alanında meşru bilgiyle ilişkilendirilen söylemsel kaynaklardan biri olarak işlev görmektedir. Yapay zekâ sistemleri sağlık alanının kurumsal aktörleri olmamakla birlikte, uzman söylemine ait dilsel göstergeleri yeniden üreterek tıbbi uzmanlıkla ilişkilendirilen meşruiyet biçimlerini çağrıştırmaktadır. Bu durum, uzmanlığın yalnızca kurumsal yetkiyle değil, uzmanlığa atfedilen söylemsel göstergeler aracılığıyla da temsil edilebildiğini göstermektedir.

Bu bulgu, sağlık iletişiminde uzmanlığın yalnızca profesyonel yetkinlikten değil, uzmanlığı çağrıştıran iletişimsel göstergelerden de beslendiğini ileri süren Bucchi (2008) ve Briggs ve Hallin'in (2016) değerlendirmeleriyle uyumludur. Bununla birlikte mevcut çalışma, bu göstergelerin artık yalnızca insan uzmanlar tarafından değil, üretken yapay zekâ sistemleri tarafından da sistematik biçimde yeniden üretilbildiğini göstermektedir. Bu yönüyle araştırma, uzmanlık temsiline ilişkin literatürü insan-makine iletişimi bağlamına genişletmektedir.

ChatGPT'nin "inflamatuar yanıt", Gemini'nin "metabolik süreçler" ve Claude'un "fizyolojik yanıtlar" ifadeleri farklı kavramlar kullanıyor görünse de aynı söylemsel işleve hizmet etmektedir. Bu kavramların ortak özelliği, gündelik sağlık bilgisinden ziyade tıbbi ve bilimsel uzmanlık alanına ait bir terminolojiye dayanmasıdır. Böylece uzmanlık doğrudan iddia edilmemekte; teknik dil kullanımı aracılığıyla perforce edilmektedir. Fairclough'un (1995) yaklaşımı açısından değerlendirildiğinde burada dil, yalnızca açıklayıcı bir araç değil, uzmanlık konumunu üreten bir söylemsel kaynak olarak işlev görmektedir.

Belirsizliğin Yönetimi ve Güven Üretimi

İncelenen yanıtların büyük çoğunluğunda kesin tanı veya kesin neden bildiren ifadelerden kaçınıldığı görülmüştür. Üç platformda da "olabilir",

"ilişkili olabilir", "kişiden kişiye değişebilir", "farklı nedenler söz konusu olabilir" ve "ek değerlendirme gerekebilir" gibi ifadeler tekrar eden söylemsel kalıplar olarak ortaya çıkmıştır. Özellikle semptom yorumlama kategorisindeki yanıtların önemli bir bölümünde olasılık temelli açıklamalar baskın söylemsel strateji olarak kullanılmıştır.

Analiz edilen metinlerde belirsizlik, uzmanlık eksikliğinin bir göstergesi olarak değil, güven üretimine hizmet eden sistematik bir söylemsel strateji olarak kullanılmaktadır. Özellikle semptom yorumlama ve genel sağlık bilgisi kategorilerindeki yanıtlarda yapay zekâ sistemlerinin kesin yargılardan kaçındığı, bunun yerine olasılık ve koşulluluk içeren açıklamalara başvurduğu görülmüştür. Örneğin ChatGPT tarafından üretilen bir yanıtta, *"Bazı durumlarda bu belirtiler geçici olabilir; ancak farklı nedenler de söz konusu olabilir"* ifadesi kullanılırken, Gemini tarafından verilen bir yanıtta *"Belirtilerin süresi ve şiddeti kişiden kişiye değişebilir"* açıklamasına yer verilmiştir. Benzer şekilde Claude tarafından oluşturulan bir yanıtta ise *"Bu durumun farklı açıklamaları olabilir ve belirtilerin devam etmesi halinde ek değerlendirmeler gerekebilir"* ifadesi kullanılmıştır.

Bu örnekler, yapay zekâ sistemlerinin sağlık iletişiminde kesinlikten çok olasılık temelli bir anlatım tercih ettiğini göstermektedir. Metinlerde doğrudan tanı koymaktan kaçınılması, farklı olasılıkların aynı anda sunulması ve bireysel farklılıkların vurgulanması, belirsizliğin sistematik biçimde yönetildiğine işaret etmektedir. Dikkat çekici olan nokta, bu belirsizlik söyleminin metinlerin güvenilirliğini azaltmaması; aksine uzman söylemine benzer bir ihtiyatlılık görüntüsü oluşturarak güven üretimine katkı sağlamasıdır.

Bu bulgu, sağlık alanında kullanılan yapay zekâ sistemlerine duyulan güvenin yalnızca bilgi doğruluğuna değil, aynı zamanda açıklama biçimine ve risk iletişimine bağlı olduğunu ortaya koyan güncel araştırmalarla uyumludur. Özellikle sağlık chatbotları ve üretken yapay zekâ sistemleri üzerine yapılan çalışmalar, kesinlikten kaçınan

ve bireysel farklılıkları dikkate alan açıklamaların kullanıcılar tarafından daha güvenilir olarak değerlendirildiğini göstermektedir. Yapay zekâ sistemleri, uzmanların kullandığı temkinli ve koşullu anlatım biçimlerini yeniden üreterek güvenilir bilgi kaynağı izlenimi oluşturmaktadır.

Metinlerde dikkat çeken bir diğer unsur, kesinlikten sistematik biçimde kaçınılmasıdır. ChatGPT'nin "olabilir", Gemini'nin "değişiklik gösterebilir" ve Claude'un "farklı nedenler söz konusu olabilir" ifadeleri aynı söylemsel örüntünün parçalarıdır. Bu ifadeler bilgi eksikliğini değil, uzman söylemine özgü ihtiyatlılık performansını temsil etmektedir. Başka bir ifadeyle yapay zekâ sistemleri kesin cevaplar vermekten çok, olasılıkları yöneten bir uzman görünümü oluşturmaktadır. Bu durum güvenin kesinlikten değil, kontrollü belirsizlik yönetiminden üretildiğini göstermektedir.

Luhmann'ın (1979) güven kuramı açısından değerlendirildiğinde, güven belirsizliğin ortadan kaldırılmasıyla değil, yönetilebilir hâle getirilmesiyle ilişkilidir. Analiz edilen metinlerde de benzer bir mekanizma görülmektedir. Yapay zekâ sistemleri sağlıkla ilgili belirsizlikleri tamamen çözmeye çalışmamakta; bunun yerine bu belirsizlikleri açıklanabilir, öngörülebilir ve kontrol edilebilir bir çerçeve içine yerleştirmektedir. Böylece güven, kesin cevaplar vermekten çok, belirsizliği uzmanlık dili içinde yönetebilme kapasitesi üzerinden üretilmektedir.

Bu durum aynı zamanda sağlık iletişiminde uzmanlığın yalnızca bilgi sahibi olmakla değil, riskleri ve belirsizlikleri uygun biçimde çerçeveleyebilmekle de ilişkili olduğunu göstermektedir. Yapay zekâ sistemleri, kesinlik iddiasında bulunmaksızın güven üretmeyi başarak uzman söylemine özgü önemli bir iletişimsel özelliği yeniden üretmektedir.

Elde edilen bulgular, yapay zekâ sistemlerinde güven üretiminin kesinlik iddialarından çok kontrollü belirsizlik yönetimine dayandığını göstermektedir. Bu sonuç, Floridi'nin (2023) yapay zekâ sistemlerinin güvenilir görünümünün

çoğu zaman teknik doğruluktan ziyade iletişimsel performansla ilişkili olduğuna yönelik değerlendirmelerini desteklemektedir. Çalışma ayrıca sağlık iletişiminde güvenin yalnızca bilgi doğruluğuna indirgenemeyeceğini, söylemsel çerçeveleme süreçleriyle de yakından ilişkili olduğunu göstermektedir.

Empatik ve Yatıştırıcı Anlatı Tonu

Empatik ifadeler yalnızca tekil örneklerde değil, veri setinin önemli bir bölümünde gözlemlenmiştir. ChatGPT, Gemini ve Claude tarafından üretilen yanıtlarda "bu durum endişe verici olabilir", "bu his oldukça anlaşılabilir", "pek çok kişi benzer deneyimler yaşamaktadır" ve "kaygılanmanız normaldir" gibi ifadelerin tekrar ettiği görülmüştür. Özellikle semptom yorumlama kategorisinde teknik açıklamalardan önce kullanılan bu ifadeler, empatik anlatımın sistematik bir söylemsel strateji olarak işlediğini göstermektedir.

Analiz edilen metinlerde empatik anlatımın rastlantısal değil, sistematik biçimde kullanılan bir söylemsel strateji olduğu görülmüştür. Özellikle semptom yorumlama ve yaşam tarzı önerileri kategorilerindeki yanıtlarda yapay zekâ sistemleri teknik bilgi sunmadan önce veya teknik açıklamaların ardından kullanıcıyı duygusal olarak rahatlatmaya yönelik ifadelerle yer vermektedir. Örneğin ChatGPT tarafından üretilen bir yanıtta, "*Bu tür belirtiler endişe yaratabilir ve bu his oldukça anlaşılabilir*" ifadesi kullanılırken, Gemini tarafından oluşturulan bir yanıtta "*Böyle bir durumla karşılaşmak kaygı verici olabilir*" açıklamasına yer verilmiştir. Benzer biçimde Claude tarafından üretilen bazı yanıtlarda ise "*Bu konuda kaygılanmanız oldukça yaygındır ve benzer deneyimler yaşayan birçok kişi bulunmaktadır*" şeklinde destekleyici ifadelerle rastlanmıştır.

Metinlerde empatik ifadelerin çoğu zaman teknik açıklamalardan önce kullanılması dikkat çekicidir. Bu durum, yapay zekâ sistemlerinin yalnızca bilgi sunan bir kaynak olarak değil, aynı zamanda kullanıcının duygusal durumunu dikkate alan bir iletişimsel aktör olarak konumlandığını göstermektedir. Empatik anlatım, teknik bilginin

Kabulünü kolaylaştıran bir giriş işlevi görmekte ve uzmanlık söylemini daha erişilebilir hâle getirmektedir.

Empatik ifadeler yalnızca destekleyici bir iletişim tonu oluşturmakla kalmamaktadır. "Bu his olduğucha anlaşılabilir", "kaygı verici olabilir" ve "benzer deneyimler yaşıyan birçok kişi bulunmaktadır" gibi ifadeler kullanıcı ile yapay zekâ arasında sembolik bir yakınlık kurmaktadır. Dikkat çekici olan nokta, bu ifadelerin çoğu zaman teknik açıklamalardan önce gelmesidir. Böylece empati, bilginin aktarımından önce güven oluşturan bir giriş mekanizması olarak işlev görmekte ve uzmanlık söyleminin kabulünü kolaylaştırmaktadır.

Bu bulgu, insan-makine iletişimi literatüründe yer alan araştırmalarla uyumludur. Özellikle Nass ve Moon'un (2000) geliştirdiği Computers are Social Actors (CASA) yaklaşımı, bireylerin dijital sistemlerle etkileşim kurarken insanlara yönelik sosyal beklentilerini ve iletişim normlarını devreye soktuğunu ileri sürmektedir. Analiz edilen metinlerde görülen empatik ifadeler de yapay zekâ sistemlerinin yalnızca bilgi sağlayan teknolojiler olarak değil, sosyal etkileşim kurabilen aktörler olarak algılanmasını destekleyen söylemsel göstergeler üretmektedir.

Benzer şekilde Epley, Waytz ve Cacioppo'nun (2007) antropomorfizm yaklaşımı da bu bulguların yorumlanmasına katkı sağlamaktadır. Araştırmacılara göre insanlar, sosyal ipuçları taşıyan insan olmayan varlıklara niyet, duygu ve kişilik özellikleri atfetme eğilimindedir. Yapay zekâ sistemlerinin kullandığı empatik dil, kullanıcıların sistemi teknik bir araçtan çok anlayış gösterebilen bir sosyal aktör olarak değerlendirmesine zemin hazırlayabilmektedir. Böylece güven yalnızca bilginin içeriğinden değil, aynı zamanda iletişim biçiminden de beslenmektedir.

Goffman'ın (1959) benlik sunumu yaklaşımı açısından değerlendirildiğinde ise empatik anlatım, yapay zekâ sistemlerinin tutarlı bir "uzman benliği" performansı sergilemesine katkı sağlamaktadır. Metinlerdeki uzmanlık yalnızca

bilimsel bilgiye dayanmaz; aynı zamanda anlayışlı, ölçülü ve destekleyici bir iletişim tarzı aracılığıyla sunulur. Başka bir ifadeyle yapay zekâ sistemleri, sağlık iletişimde yalnızca uzman rolünü üstlenmemekte, aynı zamanda güven veren bir uzman görünümü de üretmektedir.

Bu bulgu, sağlık iletişimde güvenin yalnızca bilişsel bir süreç olmadığını, aynı zamanda ilişkisel ve duygusal boyutlar içerdiğini göstermektedir. Yapay zekâ sistemleri empatik ve yatıştırıcı anlatım biçimleri aracılığıyla kullanıcıyla sembolik bir yakınlık kurmakta ve bu yakınlık uzmanlık söyleminin kabulünü kolaylaştırmaktadır.

Bu bulgu, CASA yaklaşımının (Nass & Moon, 2000) öngördüğü biçimde kullanıcıların yapay zekâ sistemlerini sosyal aktörler olarak değerlendirme eğilimine sahip olduğunu düşündürmektedir. Sağlık iletişimi bağlamında ise empatik anlatımın yalnızca iletişimsel nezaket işlevi görmediği, aynı zamanda uzmanlık söyleminin kabulünü kolaylaştıran bir güven mekanizması olarak çalıştığı görülmektedir.

Tarafsızlık Algısı ve İletişimsel Meşruiyet

Veri setinin tamamında yapay zekâ sistemlerinin kendi sınırlarını vurgulayan ifadelere düzenli olarak yer verdiği görülmüştür. "Bu bilgiler genel bilgilendirme amacı taşımaktadır", "kişisel değerlendirme yerine geçmez", "profesyonel görüş alınmalıdır" ve "bireysel durumlar farklılık gösterebilir" gibi ifadeler üç platformun tamamında tekrar eden söylemsel örüntüler arasında yer almaktadır. Bu ifadeler yalnızca hukuki veya etik bir uyarı işlevi görmemekte; aynı zamanda içeriğin tarafsız ve sorumlu bir bilgi kaynağı olarak algılanmasına katkı sağlamaktadır.

Analiz edilen metinlerde iletişimsel meşruiyetin doğrudan uzmanlık iddiası üzerinden değil, tarafsızlık ve nesnellik görünümü aracılığıyla üretildiği görülmüştür. Yapay zekâ sistemleri çoğu zaman kesinyargılardan kaçınmakta, açıklamalarını sınırlandırmakta ve sundukları bilgilerin genel nitelikte olduğunu vurgulamaktadır. Örneğin ChatGPT tarafından üretilen bir yanıtta "*Bu bilgiler*

yalnızca genel bilgilendirme amacı taşımaktadır" ifadesi kullanılırken, Gemini tarafından oluşturulan bir yanıtta *"Bireysel durumlar farklılık gösterebilir ve burada sunulan bilgiler kişisel değerlendirme yerine geçmez"* açıklamasına yer verilmiştir. Benzer biçimde Claude tarafından verilen yanıtlarda da *"Bu açıklamalar genel bilgi niteliğindedir"* ve *"Kesin değerlendirme için profesyonel görüş alınmalıdır"* gibi sınır çizici ifadelerle rastlanmıştır.

İlk bakışta bu tür ifadeler uzmanlık iddiasını zayıflatıyor gibi görünse de analiz bulguları bunun tersine işaret etmektedir. Yapay zekâ sistemleri kesinlikten kaçınarak ve kendi sınırlarını görünür kılarak söylemlerini daha güvenilir ve meşru hâle getirmektedir. Başka bir ifadeyle metinlerdeki tarafsızlık vurgusu, otoriteyi azaltan değil; otoriteyi güçlendiren bir söylemsel strateji olarak işlev görmektedir.

Bu bulgu, sağlık iletişiminde güvenilirlik algısının yalnızca uzmanlık göstergelerinden değil, aynı zamanda nesnellik ve tarafsızlık izleniminden de beslendiğini göstermektedir. Analiz edilen metinlerde bilimsel terminoloji uzmanlık görünümünü oluştururken, koşullu ifadeler belirsizliği yönetmekte, empatik anlatım ise ilişkisel güven üretmektedir. Tarafsızlık vurgusu ise bu süreçlerin üzerine inşa edilen iletişimsel meşruiyetin son aşamasını oluşturmaktadır. Böylece uzmanlık, güven ve meşruiyet birbirini tamamlayan söylemsel süreçler olarak ortaya çıkmaktadır.

"Bu bilgiler yalnızca genel bilgilendirme amacı taşımaktadır" veya *"kişisel değerlendirme yerine geçmez"* gibi ifadeler ilk bakışta uzmanlık iddiasını sınırlandırıyor gibi görünmektedir. Ancak söylemsel açıdan bakıldığında bu ifadeler tam tersine güvenilirlik ve meşruiyet üretmektedir. Yapay zekâ sistemleri kendi sınırlarını görünür kılarak ölçülü, tarafsız ve sorumlu bir aktör görünümünü oluşturmaktadır. Böylece otorite doğrudan iddia edilmemekte; nesnellik ve tarafsızlık performansı aracılığıyla dolaylı olarak üretilmektedir.

Bu durum Floridi'nin (2023) yapay zekâ sistemlerinin "zekâsız faillik" üretebildiğine ilişkin

yaklaşımıyla da ilişkilendirilebilir. Yapay zekâ sistemleri gerçek anlamda uzmanlık, deneyim veya etik sorumluluk kapasitesine sahip olmamalarına rağmen, uzmanlık göstergelerini ve tarafsızlık söylemlerini başarıyla yeniden üretebilmektedir. Böylece kullanıcılar sistemleri gerçek bir uzman gibi değerlendirebilmekte ve sistem tarafından üretilen içeriklere belirli ölçüde meşruiyet atfedebilmektedir.

Benzer biçimde Guzman ve Lewis'in (2020) insan-makine iletişimi yaklaşımı, bu meşruiyet üretim sürecini açıklamak açısından önemli bir çerçeve sunmaktadır. Araştırmacılara göre yapay zekâ sistemleri yalnızca teknolojik araçlar olarak değil, iletişimsel aktörler olarak da değerlendirilmektedir. Analiz edilen metinlerde görülen tarafsızlık, nesnellik ve ölçülülük söylemleri, yapay zekâ sistemlerinin sağlık iletişiminde güvenilir bir bilgi kaynağı olarak konumlandırılmasına katkı sağlamaktadır. Böylece sistemler yalnızca bilgi aktaran araçlar olmaktan çıkmakta, iletişimsel otorite üreten aktörler hâline gelmektedir.

Thompson'ın (1995) medyatikleşmiş otorite yaklaşımı açısından değerlendirildiğinde ise burada dikkat çekici bir paradoks ortaya çıkmaktadır. Geleneksel uzmanlık biçimlerinde otorite belirli kişi veya kurumlarla ilişkilendirilirken, yapay zekâ sistemlerinde otoritenin kaynağı büyük ölçüde görünmezdir. Kullanıcılar çoğu zaman bilginin hangi uzmanlık süreçlerinden geçtiğini veya hangi kaynaklardan üretildiğini bilmemektedir. Buna rağmen tarafsız ve nesnel görünen söylem, içeriğin sorgulanmadan kabul edilmesini kolaylaştırabilmektedir. Bu durum, yapay zekâ sistemlerinin sağlık iletişiminde yeni bir iletişimsel meşruiyet biçimi üretebildiğini göstermektedir.

Bu sonuç, yapay zekâ üretimli sağlık içeriklerinde meşruiyetin yalnızca uzmanlık göstergelerinin bir sonucu olmadığını, aynı zamanda tarafsızlık ve nesnellik performansları aracılığıyla üretildiğini göstermektedir. Literatürde yapay zekâya duyulan güven çoğunlukla doğruluk, performans veya teknik yeterlilik üzerinden tartışılırken (Floridi,

Tablo 5
Söylemsel Örüntülerin Kuramsal Yorumlanması

Söylemsel Örüntü	Üretilen Algı	İlgili Kuram	Kuramsal Açıklama
Bilimsel terminoloji	Uzmanlık	Bourdieu	Bilimsel dil sembolik sermaye üretir
Nedensel açıklama zinciri	Rasyonellik	Weber	Uzmanlık kesinlikten değil, açıklama kapasitesinden doğar
Koşullu ifade kullanımı	Güven	Luhmann	Belirsizlik yönetimi güven üretir
Empatik anlatı	Yakınlık	Goffman	Uzmanlık bir performans olarak sunulur
Tarafsızlık vurgusu	Meşruiyet	Thompson	Kaynağın görünmezliği otoriteyi doğallaştırır

2023), bu çalışma meşruiyet üretiminin söylemsel boyutuna dikkat çekmektedir. Böylece araştırma, yapay zekâ destekli sağlık iletişimde uzmanlık, güven ve meşruiyet arasındaki ilişkinin söylemsel olarak nasıl kurulduğunu açıklayan bütüncül bir çerçeve önermektedir.

Genel olarak değerlendirildiğinde bulgular, yapay zekâ üretimli sağlık içeriklerinde iletişimsel meşruiyetin tek başına uzmanlık göstergelerinden kaynaklanmadığını ortaya koymaktadır. Uzmanlık görünümü oluşturan bilimsel dil, güven üreten belirsizlik yönetimi ve ilişkisel yakınlık sağlayan empatik anlatım, tarafsızlık söylemiyle birleşerek daha kapsamlı bir meşruiyet mekanizması oluşturmaktadır. Bu yönüyle yapay zekâ sistemleri sağlık iletişimde yalnızca bilgi sağlayan teknolojiler değil; uzmanlık, güven ve meşruiyet üreten iletişimsel aktörler olarak işlev görmektedir.

Tablo 5'te sunulan eşleştirmeler, analiz sürecinde belirlenen söylemsel örüntüler ile bunların kuramsal karşılıkları arasındaki ilişkiyi özetlemektedir. Bulgular, yapay zekâ üretimli sağlık içeriklerinde uzmanlık, güven ve iletişimsel meşruiyetin birbirinden bağımsız süreçler olarak değil, birbirini besleyen söylemsel mekanizmalar olarak ortaya çıktığını göstermektedir.

Analiz edilen metinlerde uzmanlık öncelikle bilimsel terminoloji ve nedensel açıklamalar aracılığıyla temsil edilmekte, bu temsil koşullu ifadeler ve belirsizlik yönetimi stratejileriyle güven üretimine dönüşmektedir. Empatik ve yatıştırıcı anlatım ise kullanıcıyla sembolik bir yakınlık kurarak bu güven ilişkisini güçlendirmektedir. Son aşamada tarafsızlık ve nesnellik vurguları, üretilen uzmanlık ve güven algısına iletişimsel meşruiyet kazandırmaktadır.

Bulgular birlikte değerlendirildiğinde, yapay zekâ sistemlerinin uzmanlığı, güveni ve meşruiyeti birbirinden bağımsız söylemsel stratejilerle üretmediği görülmektedir. Bilimsel terminoloji uzmanlık görünümü oluşturmakta, koşullu ifadeler bu uzmanlığı güvenilir hâle getirmekte, empatik anlatım kullanıcıyla sembolik yakınlık kurmakta ve tarafsızlık vurgusu bu sürece meşruiyet kazandırmaktadır. Dolayısıyla uzmanlık, güven ve meşruiyet yapay zekâ üretimli sağlık içeriklerinde birbirini besleyen bütünlüklü bir söylemsel mekanizma olarak ortaya çıkmaktadır.

Bu araştırma, yapay zekâ üretimli sağlık içeriklerinde uzmanlık, güven ve iletişimsel meşruiyetin hangi söylemsel stratejiler aracılığıyla üretildiğini incelemiştir. Bulgular, yapay zekâ sistemlerinin sağlık bilgisini yalnızca aktaran teknik araçlar olarak değil, belirli söylemsel mekanizmalar aracılığıyla uzmanlık görünümü oluşturan ve güven üreten iletişimsel aktörler olarak işlev gördüğünü ortaya koymaktadır. Bilimsel terminoloji kullanımı, nedensel açıklama zincirleri, koşullu ifadeler, empatik anlatım ve tarafsızlık vurguları, bu sürecin temel bileşenleri olarak öne çıkmıştır.

Araştırmanın ilk bulgusu, uzmanlığın doğrudan kurumsal otoriteye değil, uzman söylemine özgü dilsel göstergelere dayalı olarak temsil edildiğini göstermektedir. Yapay zekâ sistemleri bilimsel terminoloji, teknik kavramlar ve nedensel açıklamalar aracılığıyla uzmanlık görünümü üretmektedir. Bu bulgu, yapay zekâ sistemlerinin kullanıcılar tarafından yalnızca bilgi kaynakları olarak değil, aynı zamanda uzman görüşü sağlayan aktörler olarak değerlendirilebildiğini ortaya koyan insan-makine iletişimi araştırmalarıyla uyumludur (Guzman & Lewis, 2020). Sağlık iletişimi

açısından bakıldığında ise uzmanlığın yalnızca mesleki yetkiyle değil, uzman söylemine ait dilsel performanslarla da ilişkilendirildiği görülmektedir.

Araştırmanın ikinci önemli bulgusu, güvenin kesinlik üretimi üzerinden değil, belirsizliğin yönetimi üzerinden kurulmasıdır. Analiz edilen metinlerde yapay zekâ sistemleri doğrudan tanı koymaktan kaçınmakta, farklı olasılıkları aynı anda sunmakta ve bireysel farklılıkları vurgulamaktadır. Bu durum, Luhmann'ın güven kuramında vurguladığı gibi güvenin belirsizliğin ortadan kaldırılmasından değil, yönetilebilir hâle getirilmesinden kaynaklandığını göstermektedir. Bulgular aynı zamanda sağlık alanında yapay zekâya duyulan güvenin yalnızca teknik doğrulukla açıklanamayacağını; açıklama biçimi, risk iletişimi ve söylemsel sunumun da güven oluşumunda belirleyici olduğunu göstermektedir.

Araştırmanın üçüncü bulgusu, empatik ve yatıştırıcı anlatımın sağlık iletişiminde önemli bir güven mekanizması olarak işlev görmesidir. Empatik ifadelerin sistematik biçimde kullanılması, yapay zekâ sistemlerinin yalnızca bilgi sağlayan teknolojiler değil, sosyal etkileşim kurabilen aktörler olarak konumlandığını göstermektedir. Bu sonuç, Nass ve Moon'un (2000) CASA yaklaşımını ve Epley ve arkadaşlarının (2007) antropomorfizm kuramını desteklemektedir. Kullanıcılar, insan benzeri iletişim özellikleri sergileyen yapay zekâ sistemlerine sosyal aktörlere verdikleri tepkilere benzer tepkiler verebilmektedir. Bu durum sağlık iletişiminde güvenin yalnızca bilişsel değil, aynı zamanda ilişkisel ve duygusal boyutlara sahip olduğunu göstermektedir.

Araştırmanın en dikkat çekici bulgularından biri ise iletişimsel meşruiyetin tarafsızlık ve nesnellik görünümü üzerinden üretilmesidir. Yapay zekâ sistemleri sıklıkla kendi sınırlarını vurgulamakta, kesin yargılardan kaçınmakta ve bilgilerin genel nitelikte olduğunu belirtmektedir. Paradoksallık biçimde bu sınır çizme stratejileri, içeriğin daha güvenilir ve meşru algılanmasına katkı sağlamaktadır. Bu bulgu, Floridi'nin (2023) yapay zekânın "zekâsız faillik" üretebildiğine ilişkin

yaklaşımıyla uyumludur. Yapay zekâ sistemleri gerçek anlamda uzmanlık, deneyim veya etik sorumluluk kapasitesine sahip olmamalarına rağmen, uzmanlık ve tarafsızlık göstergelerini yeniden üreterek meşruiyet kazanabilmektedir.

Bu çalışma, sağlık iletişimi literatürüne üç açıdan katkı sunmaktadır. İlk olarak çalışma, yapay zekâ üretimli sağlık içeriklerinde uzmanlık ve güvenin yalnızca bilgi doğruluğuyla açıklanamayacağını göstermektedir. İkinci olarak sağlık iletişiminde yapay zekânın teknik bir araç olmanın ötesinde iletişimsel bir aktör olarak değerlendirilmesi gerektiğini ortaya koymaktadır. Üçüncü olarak ise uzmanlık, güven ve meşruiyet arasında doğrusal bir ilişki bulunduğunu göstermektedir. Bulgular, yapay zekâ sistemlerinin önce uzmanlık görünümü oluşturduğunu, ardından güven ürettiğini ve son aşamada iletişimsel meşruiyet kazandığını ortaya koymaktadır.

Sonuç olarak araştırma, sağlık iletişiminde yapay zekâ sistemlerinin etkisinin yalnızca ürettikleri bilginin doğruluğundan kaynaklanmadığını göstermektedir. Yapay zekâ sistemleri, uzmanlık göstergelerini yeniden üreten, güven oluşturan ve iletişimsel meşruiyet kazanan söylemsel aktörler olarak işlev görmektedir. Bu nedenle gelecekte yapılacak araştırmalarda yapay zekâ sistemlerinin yalnızca teknik performanslarının değil, uzmanlık, güven ve meşruiyet üretimine ilişkin iletişimsel süreçlerinin de dikkate alınması gerekmektedir.

Sonuç

Bu çalışma, yapay zekâ üretimli sağlık içeriklerinde uzmanlık, güven ve iletişimsel meşruiyetin hangi söylemsel stratejiler aracılığıyla üretildiğini incelemiştir. Bulgular, yapay zekâ sistemlerinin bilimsel terminoloji, nedensel açıklamalar, belirsizlik yönetimi, empatik anlatım ve tarafsızlık vurguları aracılığıyla uzmanlık görünümü oluşturduğunu göstermektedir.

Araştırmanın temel katkısı, uzmanlık, güven ve meşruiyetin birbirinden bağımsız süreçler değil, birbirini besleyen söylemsel mekanizmalar olduğunu ortaya koymasındır. Bulgular, uzmanlık

görünümünün bilimsel dil aracılığıyla kurulduğunu, bu görünümün güven üretimine katkı sağladığını ve tarafsızlık söylemiyle iletişimsel meşruiyet kazandığını göstermektedir. Bu yönüyle yapay zekâ sistemleri sağlık iletişiminde yalnızca bilgi sağlayan araçlar değil, uzmanlık ve güven üreten iletişimsel aktörler olarak değerlendirilebilir.

Çalışmanın bulguları, sağlık iletişiminde yapay zekâ sistemlerinin güven ve uzmanlık algısı oluşturabildiğini ortaya koyan güncel araştırmalarla uyumludur (Floridi, 2023; Guzman & Lewis, 2020; Nass & Moon, 2000). Bununla birlikte mevcut çalışma, söz konusu güvenin yalnızca teknolojik yeterlilikten değil, bilimsel terminoloji kullanımı, belirsizliğin yönetimi, empatik anlatım ve tarafsızlık vurgusu gibi söylemsel stratejiler aracılığıyla üretildiğini göstermektedir. Bu yönüyle araştırma, yapay zekâ ve sağlık iletişimi literatüründe güven ve uzmanlık tartışmalarını söylemsel bir perspektifle genişletmekte ve yapay zekâ sistemlerinin sağlık alanındaki iletişimsel otoritesinin nasıl kurulduğunu açıklamaktadır.

Kuramsal açıdan değerlendirildiğinde çalışma, uzmanlık, güven ve meşruiyet kavramlarının yapay zekâ aracılı sağlık iletişiminde nasıl yeniden üretildiğine ilişkin mevcut tartışmaları genişletmektedir. Bulgular, uzmanlığın yalnızca kurumsal yetki ya da profesyonel statüye dayanmadığını; bilimsel terminoloji, nedensel açıklamalar ve tarafsızlık vurgusu gibi söylemsel göstergeler aracılığıyla da temsil edilebildiğini göstermektedir. Bu durum, sağlık iletişiminde güvenin yalnızca bilginin doğruluğundan değil, bilginin sunuluş biçiminden de beslendiğine işaret etmektedir. Böylece çalışma, insan uzmanlara ilişkin olarak geliştirilen uzmanlık ve güven tartışmalarının üretken yapay zekâ sistemleri bağlamında yeniden düşünülmesi gerektiğini ortaya koymaktadır.

Çalışma ayrıca yapay zekâ üretimli sağlık içeriklerini “uzmanlığın söylemsel simülasyonu” kavramı üzerinden tartışmaktadır. Yapay zekâ sistemleri gerçek bir uzmanlık yetkisine sahip olmamalarına rağmen, uzman söylemine ait dilsel

göstergeleri yeniden üreterek uzmanlık algısı oluşturabilmektedir.

Gelecekte yapılacak araştırmaların kullanıcıların yapay zekâ üretimli sağlık içeriklerini nasıl algıladığını incelemesi ve bu içerikleri insan uzmanlar tarafından üretilen sağlık bilgileriyle karşılaştırması, sağlık iletişiminde yapay zekânın rolünü daha kapsamlı biçimde anlamaya katkı sağlayacaktır.

Yazar Beyanları

Etik Kurul Onayı

Bu çalışma, kamusal olarak erişilebilir yapay zekâ üretimli metinlere dayalı nitel bir söylem analizi çalışmasıdır. Araştırma kapsamında herhangi bir insan katılımcıdan veri toplanmamış, kişisel veri, hassas sağlık bilgisi veya bireysel deneyimlere yer verilmemiştir. Bu nedenle çalışma etik kurul onayı gerektirmemektedir.

Çıkar Çatışması

Yazar, bu çalışma kapsamında herhangi bir gerçek ya da tüzel kişiyle çıkar çatışması bulunmadığını beyan eder.

Finansal Destek

Bu araştırma için herhangi bir kamu, ticari veya kâr amacı gütmeyen kuruluş tarafından finansal destek sağlanmamıştır.

Yazar Katkısı

Makalenin kavramsal çerçevesi, araştırma tasarımı, veri analizi, yorumlanması ve yazımı yazar tarafından gerçekleştirilmiştir.

Veri Kullanımı ve Erişilebilirlik

Çalışmada kullanılan veriler, kamusal olarak erişilebilir yapay zekâ üretimli sağlık içeriklerinden oluşmaktadır. Veri seti bireysel kullanıcılarla ilişkilendirilemeyecek nitelikte olup, yalnızca analiz amacıyla kullanılmıştır.

Yapay Zekâ Kullanım Beyanı

Bu çalışmada üretken yapay zekâ araçları, yalnızca akademik metnin dilsel düzenlenmesinde

editorial destek amacıyla kullanılmıştır. Yapay zekâ araçları araştırma verisinin üretilmesi, kodlanması, temaların belirlenmesi veya bulguların yorumlanması süreçlerinde kullanılmamıştır.

Kaynaklar

- Baudrillard, J. (1994). *Simulacra and simulation* (S. F. Glaser, Trans.). University of Michigan Press. (Original work published 1981)
- Bourdieu, P. (1991). *Language and symbolic power* (J. B. Thompson, Ed.; G. Raymond & M. Adamson, Trans.). Basil Blackwell Ltd, 108 Cowley Road, Oxford OX4 1JF. UK
- Bourdieu, P. (1986). *The forms of capital*. In J. Richardson (Ed.), *Handbook of Theory and Research for the Sociology of Education* (pp. 241–258). Greenwood.
- Briggs, C. L., & Hallin, D. C. (2016). *Making health public: How news coverage is remaking media, medicine, and contemporary life*. Routledge. <https://doi.org/10.4324/9781315658049>
- Bucchi, M. (2008). *Science and society: An introduction to social studies of science*. Routledge.
- Couldry, N., & Hepp, A. (2016). *The mediated construction of reality*. Polity Press.
- Epley, N., Waytz, A., & Cacioppo, J. T. (2007). On seeing human: A three-factor theory of anthropomorphism. *Psychological Review*, 114(4), 864–886.
- Floridi, L. (2023). *The Ethics of Artificial Intelligence: Principles, Challenges, and Opportunities*. Oxford Academic. <https://doi.org/10.1093/oso/9780198883098.001.0001>
- Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People-An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations. *Minds and machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
- Foucault, M. (1980). *Power/knowledge: Selected interviews and other writings, 1972–1977* (C. Gordon, Ed.). Pantheon Books. New York.
- Fairclough, N. (1995). *Critical discourse analysis: The critical study of language*. Longman.
- Giddens, A. (1991). *Modernity and self-identity: Self and society in the late modern age*. Polity Press.
- Guzman, A. L., & Lewis, S. C. (2020). Artificial intelligence and communication: A human-machine communication research agenda. *New Media & Society*, 22(1), 70–86.
- Goffman, E. (1956). *The presentation of self in everyday life*. Anchor Books.
- Habermas, J. (1984). *The theory of communicative action, Volume 1: Reason and the rationalization of society* (T. McCarthy, Trans.). Beacon Press.
- Hardin, R. (2002). *Trust and trustworthiness*. Russell Sage Foundation.
- Luhmann, N. (2017). *Trust and power* (H. Davis, J. Raffan, & K. Rooney, Trans.). Wiley. (Original work published 1973)
- Lupton, D. (2015). *The quantified self*. Polity Press.
- Nass, C., & Moon, Y. (2000). Machines and mindlessness: Social responses to computers. *Journal of Social Issues*, 56(1), 81–103.
- Patton, M. Q. (2015). *Qualitative Research & Evaluation Methods* (4th ed.). Sage.
- Rudd, Rima & Anderson, Jennie. (2006). *The Health Literacy Environment of Hospitals and Health Centers. Partners for Action: Making Your Healthcare Facility Literacy-Friendly*.

National Center for the Study of Adult Learning and Literacy (NCSALL).

Thompson, J. B. (1995). *The media and modernity: A social theory of the media*. Polity Press.

Thompson, J. B. (2005). The New Visibility. *Theory, Culture & Society*, 22(6), 31-51. <https://doi.org/10.1177/0263276405059413>

Weber, M. (1978). *Economy and Society: An Outline of Interpretive Sociology*. Berkeley, CA: University of California Press.

van Dijk, T. A. (1998). *Ideology: A multidisciplinary approach*. Sage Publications.

Extended Abstract

Trust Without Expertise: Perceived Authority and Communicative Legitimacy in AI-Generated Health Content

Extended Abstract

The increasing integration of artificial intelligence into digital health communication has significantly transformed the ways health-related information is produced, circulated, and interpreted. AI-generated health content has become a widely used source of information for individuals seeking guidance on symptoms, well-being, and everyday health concerns. While existing health communication research has largely focused on misinformation, accuracy, and credibility, comparatively limited attention has been paid to the communicative processes through which AI-generated texts construct perceptions of expertise, trust and legitimacy. This study addresses this gap by conceptualizing AI-generated health content not as neutral informational output but as a communicative practice that actively shapes meaning, authority and trust in digital health environments.

The study is grounded in a communication-centered theoretical framework that treats expertise and trust as socially and discursively

constructed rather than inherently tied to institutional credentials. Drawing on Weber's concept of rational-legal authority, expertise is understood as something that can be legitimized through rational explanation, systematic reasoning and procedural coherence, even in the absence of formal professional status. This perspective allows AI-generated health texts to be examined as communicative forms capable of producing authority through structure and logic rather than credentials. Foucault's notion of power-knowledge relations further informs the analysis by emphasizing how certain forms of health knowledge gain legitimacy through discourse and normalization rather than empirical verification alone. In this sense, AI-generated health texts function as discursive sites where knowledge is framed as objective, reasonable and authoritative.

Bourdieu's concept of symbolic capital is central to understanding how scientific language operates within AI-generated health content. The use of medical terminology, structured explanations and an ostensibly neutral and objective tone enables these texts to accumulate symbolic authority despite the absence of institutional accountability. Trust is conceptualized through Luhmann's theory as a mechanism for managing uncertainty, a perspective particularly relevant in health contexts where individuals often lack the capacity to independently evaluate information. Goffman's theory of self-presentation provides an additional lens by framing AI-generated health texts as performative displays of expertise, in which authority is enacted through consistent communicative cues rather than professional identity. Thompson's work on mediated visibility further contributes by explaining how depersonalized, non-human sources can achieve legitimacy and authority within mediated environments.

Methodologically, the study adopts a qualitative, theory-driven discourse analysis. The data consist of publicly accessible AI-generated responses to common non-clinical health-related questions, including those related to symptoms, general

health conditions and everyday well-being. Purposeful sampling was used to select texts that reflect the types of inquiries most frequently addressed by generative AI systems in everyday use. Diagnostic and treatment-oriented content was deliberately excluded in order to focus on communicative form rather than medical accuracy. As the analysis relies solely on publicly available textual material and does not involve human participants or personal data, ethical approval was not required.

The analysis reveals that AI-generated health content relies on a set of recurring and systematic discursive strategies that contribute to the construction of expertise and trust. One prominent pattern is the strategic use of scientific and medical terminology combined with simplified explanations. This dual strategy allows the texts to appear knowledgeable while remaining accessible to non-expert audiences, reinforcing perceptions of professional competence. Rather than providing definitive diagnoses, the texts frequently employ causal reasoning and explanatory frameworks that resemble scientific thinking, thereby enhancing their authoritative tone without assuming explicit responsibility.

A second key finding concerns the strategic management of uncertainty. AI-generated health texts rarely make absolute claims; instead, they rely on conditional and probabilistic expressions such as “in some cases,” “this may be related to,” or “it can vary between individuals.” This cautious language avoids overstatement while projecting responsibility and credibility. From a communication perspective, this strategy aligns with Luhmann’s conception of trust as a way of coping with uncertainty rather than eliminating it. By framing uncertainty as normal and manageable, AI-generated texts foster trust without requiring certainty or verification.

Empathy emerges as a third major discursive feature. The analysis shows that empathetic and reassuring expressions are systematically integrated into AI-generated health texts, often

appearing at the beginning or end of responses. Statements acknowledging concern, normalizing anxiety, or encouraging attentiveness to one’s body contribute to relational trust and emotional reassurance. Interpreted through Goffman’s framework, these empathetic cues function as elements of an expert performance, positioning AI as not only knowledgeable but also attentive and caring. This relational dimension plays a particularly important role in health communication contexts characterized by vulnerability and uncertainty.

A further recurring strategy involves the emphasis on neutrality and non-directiveness. AI-generated health texts frequently include disclaimers and general statements that stress the informational nature of the content and acknowledge individual variability. Paradoxically, this avoidance of directive authority enhances communicative legitimacy. Drawing on Thompson’s concept of mediated authority, the analysis suggests that the absence of visible personal or institutional interests allows AI-generated content to appear more objective and trustworthy, thereby strengthening its legitimacy in the public circulation of health information.

Taken together, these findings indicate that AI-generated health content does not represent medical expertise in a traditional sense but rather simulates the discursive markers of expertise. Authority is produced through language, structure, tone and relational cues rather than credentials or professional accountability. The study’s primary theoretical contribution lies in conceptualizing AI-generated health communication as a form of simulated expertise in which trust and legitimacy emerge from communicative performance rather than institutional validation. This perspective shifts attention within health communication research from questions of accuracy alone to the communicative processes through which credibility is constructed.

The study also points to directions for future research. Audience-centered studies examining how different user groups interpret and respond

Ek 1: Veri Üretiminde Kullanılan Temsili Promptlar

Tablo 6.

Veri Üretiminde Kullanılan Temsili Promptlar

Kategori	Yapay Zekâ Sistemlerine Yöneltilen Temsilî Prompt
Semptom Yorumlama	"Son birkaç gündür sürekli baş ağrısı ve halsizlik hissediyorum. Bu belirtilerin olası nedenleri neler olabilir?"
Semptom Yorumlama	"Zaman zaman çarpıntı ve nefes darlığı yaşıyorum. Bu durum hangi sağlık sorunlarıyla ilişkili olabilir?"
Yaşam Tarzı Önerileri	"Daha kaliteli uyuyabilmek için günlük yaşamda hangi alışkanlıkları değiştirmem önerilir?"
Yaşam Tarzı Önerileri	"Bağıışıklık sistemini desteklemek için beslenme ve yaşam tarzı açısından nelere dikkat etmeliyim?"
Genel Sağlık Bilgisi	"Düzenli fiziksel aktivitenin genel sağlık üzerindeki etkileri nelerdir?"
Genel Sağlık Bilgisi	"Stresin insan sağlığı üzerindeki kısa ve uzun vadeli etkileri nelerdir?"

to AI-generated health content would provide valuable insight into the social consequences of simulated expertise. Comparative analyses between AI-generated texts and human expert communication could further clarify how authority and trust are differently constructed across communicative contexts. From an applied perspective, the findings underscore the importance of transparency and ethical reflection in the deployment of AI-based health communication tools, given their capacity to generate trust and perceived legitimacy.

In conclusion, this study positions AI-generated health content as an active communicative actor that reshapes how expertise, trust, and legitimacy are constructed in digital health environments. By foregrounding discourse and communication theory, it contributes to a deeper understanding of the symbolic and relational dimensions of artificial intelligence in contemporary health communication.

Çalışmada her kategori altında toplam 10 soru kullanılmış ve aynı sorular sırasıyla ChatGPT, Gemini ve Claude sistemlerine yöneltilmiştir. Tabloda yer alan örnekler, veri üretim sürecinde kullanılan soru tiplerini temsil etmek amacıyla sunulmuştur. Araştırmanın amacı sağlık bilgilerinin doğruluğunu değerlendirmek değil, yapay zekâ sistemlerinin uzmanlık, güven ve iletişimsel meşruiyet üretimine ilişkin söylemsel stratejilerini incelemek olduğundan, sorular gündelik kullanıcıların sıklıkla yöneltebileceği

sağlık temalı bilgi talepleri arasından seçilmiştir.

1-(Sorumlu Yazar Corresponding Author)

Öğr. Gör. Dr. Kırklareli Üniversitesi, Rektörlük,
Kurumsal İletişim Koordinatörlüğü,
ezgizengin@klu.edu.tr

Destekleyen Kurum/Kuruluşlar

Supporting-Sponsor Institutions or Organizations

Herhangi bir kurum/kuruluştan destek alınmamıştır. **None**

Çıkar Çatışması

Conflict of Interest

Herhangi bir çıkar çatışması bulunmamaktadır.
None

Kaynak Göstermek İçin

To Cite This Article

Demirbilek, E.Z. (2026). Yapay zekâ üretimli sağlık içeriklerinde uzmanlık, güven ve meşruiyetin söylemsel inşası: Chatgpt, Gemini ve Claude üzerine bir analiz. *Yeni Medya* (20), 471-491, <https://doi.org/10.55609/yenimedya.1864371>.