

Ajan yapay zekâ: Otonom sistemler, sürdürülebilir etki ve akademik evrim için çok boyutlu bir çerçeve

Hüseyin PARMAKSIZ^{1,*}

¹İktisadi ve İdari Bilimler Fakültesi, Bilecik Şeyh Edebali Üniversitesi, 11000, Bilecik, Türkiye, ORCID: [0000-0001-8455-5625](https://orcid.org/0000-0001-8455-5625).

Özet

Ajan yapay zekâ (Ajan YZ), akıllı ajanlarda önemli bir ilerlemeyi ifade eder ve reaktif işlevlerden proaktif işlevlere geçiş yaparak bağlamsal muhakeme, öğrenme ve otonom eylemleri bir araya getirir. Bu çalışma, mevcut literatürdeki dört önemli eksikliği belirlemekte ve teorik temeller, operasyonel yapılar ve sürdürülebilirlik unsurlarını içeren kapsamlı bir çerçeve sunmaktadır. Amaçlar şunlardır: (1) özerklik, bellek, etkileşim ve öğrenme ile karakterize edilen ajan sistemlerinin bir listesini derlemek; (2) Sim ve n8n gibi düşük kodlu orkestrasyon platformlarının ajan tabanlı iş akışları üzerindeki etkisini incelemek; (3) Ajan YZ'nin belirli Birleşmiş Milletler Sürdürülebilir Kalkınma Hedeflerine (SDG'ler) nasıl katkıda bulunduğunu araştırmak; ve (4) Ajan YZ'deki araştırma temalarını belirlemek için 2023'ten 2025'e kadar Scopus indeksli 218 yayının hesaplamalı analizini yapmak. Araştırma, veri hazırlama için KNIME ve konu modelleme için Paralel LDA kullanarak, gelişmekte olan araştırma kümelerini ortaya koymakta ve etik yönetim, algoritmik hesap verebilirlik ve çevresel sürdürülebilirlik konularına yeterince odaklanılmadığını vurgulamaktadır. Bulgular, Ajan YZ'nin toplumsal faydalarını en üst düzeye çıkarmak için teknolojik ilerlemenin insan denetimi, şeffaflık çerçeveleri ve tasarım aşamasında etik ilkelere bağlılık ile entegre edilmesi gerektiğini göstermektedir.

Anahtar Kelimeler: Low-code Platformlar, Sürdürülebilir Kalkınma Hedefleri, Ajan Yapay Zekâ, İnsan-Yapay Zekâ İşbirliği, Yapay Zekâ Ajanları

Agentic artificial intelligence: A multidimensional framework for autonomous systems, sustainable impact, and scholarly evolution

Hüseyin PARMAKSIZ^{1,*}

¹Faculty of Economics and Administrative Sciences, Bilecik Seyh Edebali University, 11000, Bilecik, Türkiye, ORCID: [0000-0001-8455-5625](https://orcid.org/0000-0001-8455-5625).

Abstract

Agentic Artificial Intelligence (Agentic AI) represents a significant advancement in intelligent agents, transitioning from reactive to proactive capabilities, incorporating contextual reasoning, learning, and autonomous action. This study identifies four major deficiencies in the existing literature and presents a

*Sorumlu yazar (Corresponding author): huseyin.parmaksiz@bilecik.edu.tr

Başvuru / Received: 11.10.2025

Düzeltilme / Revised: 31.10.2025

Kabul / Accepted: 07.11.2025

comprehensive framework that includes theoretical foundations, operational structures, and sustainability aspects. The objectives are: (1) to compile a list of agentic systems characterized by autonomy, memory, interaction, and learning; (2) to examine the impact of low-code orchestration platforms such as Sim and n8n on agent-based workflows; (3) to investigate how Agentic AI contributes to specific United Nations Sustainable Development Goals (SDGs); and (4) to conduct a computational analysis of 218 Scopus-indexed publications from 2023 to 2025 to identify research themes in Agentic AI. The research utilizes KNIME for data preparation and Parallel LDA for topic modeling, revealing developing research clusters and highlighting the inadequate focus on ethical governance, algorithmic accountability, and environmental sustainability. The findings indicate that maximizing the societal benefits of Agentic AI necessitates the integration of technological advancement with human supervision, frameworks for transparency, and adherence to ethics-by-design principles.

Keywords: Low-code Platforms, Sustainable Development Goals, Agentic Artificial Intelligence, Human-AI Collaboration, AI Agents

1. Giriş

Yapay zekâ (YZ), birçok farklı çalışma alanını bir araya getirmektedir. Makinelere insanlarla aynı bilişsel yetenekleri kazandırma amacıyla başlamıştır. Alan Turing'in 1950 tarihli "Hesaplama Makineleri ve Zekâ" adlı makalesi, bu alanın temelini oluşturmuştur. Turing, "Makineler düşünebilir mi?" sorusunu sormuş ve makinelerin ne kadar akıllı olduğunu görmek için Turing Testi'ni geliştirmiştir. Bu test, bir makinenin farkı anlayamadan insanlarla ne kadar iyi konuşabildiğine dayanmaktadır (Turing, 1950). Joseph Weizenbaum, 1960'larda ELIZA programını geliştirmiştir. Doğal dil işlemeyi kullanan ilk diyalog sistemlerinden biri olduğu düşünülmektedir. ELIZA, özellikle "Rogerian psychotherapist" modelinde (Hatch, 2025), insanların söylediklerine basit, kural tabanlı yanıtlar vererek insanlarla etkileşim kurabilmiştir. Bu, YZ'nin insanlar ve makineler arasındaki iletişimde ne kadar faydalı olabileceğini göstermiştir. John Searle'ın Çin Odası Argümanı, 1980'lerde yapay zekâ felsefesine yönelik önemli bir eleştiriydi (Searle, 1999). Searle, bir sistemin sözdizimsel sembolleri işleme kapasitesinin gerçek bir anlam (semantik) anlayışına denk gelmediğini ileri sürerek makine zekâsının "bilinç" eşliğine ulaşamayacağını savundu. Bu argüman, bilişsel bilim ve sembolik yapay zekâ arasındaki tartışmalarda büyük bir felsefi değişime işaret ediyor. IBM Deep Blue'nun 1997'de satranç dünya şampiyonu Garry Kasparov'u yenmesi (Schaeffer ve Plaat, 1997), yapay zekânın ne kadar güçlü ve akıllı olduğunu göstermenin sembolik bir yoluydu. Arama algoritmalarının, minimax'ın ve sezgisel değerlendirme yöntemlerinin gelişmesi bu ilerlemeyi mümkün kıldı. Bu süre zarfında, yapay zekânın stratejik kararlar alma yeteneği insanüstü bir boyuta ulaşabildi.

21. yüzyılda makine öğrenimi (ML) ve derin öğrenme (DL) yöntemleri daha popüler hale geldikçe, yapay zekâ verilerden öğrenmek için istatistik kullanmaktan vazgeçti. Yapay sinir ağları (YSA), evrişimli sinir ağları (CNN), doğal dil işleme (NLP) ve GPT, BERT gibi büyük

veri tabanlı modeller, insanların öğrenme, görme ve yapma biçimlerini kopyalamaya başlayarak taklit edilebilirliğini göstermiştir. Robotik Süreç Otomasyonu (RPA), giderek daha fazla işletme süreçlerini otomatikleştirmek istedikçe yapay zekâ ekosisteminin önemli bir parçası haline geldi. RPA'nın amacı, tekrarlayan ve kurallara dayalı iş görevlerini insan etkileşimine ihtiyaç duymadan otomatikleştirmek için yazılım botlarını kullanmaktır. Bu sistemler eskiden yalnızca belirli kurallara göre çalışıyordu, ancak artık e-postalar, faturalar ve ses kayıtları gibi yapılandırılmamış verileri anlayabilen, kararlar alabilen ve yapay zekâ entegrasyonu yoluyla öğrenebilen akıllı otomasyon sistemleri (IPA) haline gelmektedir.

Son yıllarda, Agentic AI (Ajan Yapay zekâ) fikri de daha popüler hale geldi. Agentic AI, yalnızca araç değil, aynı zamanda hedef odaklı, bağımsız ve proaktif olan yapay zekâ sistemlerini ifade etmektedir. Bu tür sistemler, çevrelerinde olup biteni görebilir, kendi hedeflerini belirleyebilir veya genel talimatlara dayalı planlar oluşturabilir, birkaç adım gerektiren görevleri planlayabilir ve gerektiğinde insan geri bildirimlerini dikkate alarak davranışlarını değiştirebilir. Yapay zekânın (YZ) evrimi, kural tabanlı sistemlerden veri odaklı modellere ve günümüzde de yalnızca bilgiyi işlemekle kalmayıp aynı zamanda amaçlı hareket eden, dinamik ortamlara uyum sağlayan ve minimum insan müdahalesiyle hedefleri takip eden varlıklar olan aracı sistemlere doğru ilerlemektedir (Russell ve Norvig, 2021). Bu dönüşüm, yapay zekânın pasif bir analiz aracı olma rolünden çıkarak proaktif bir aktöre dönüşmesini ifade etmektedir. Bu sistemler, çevreden topladıkları verileri analiz eder; çok adımlı planlar oluşturur ve uygulama programlama arayüzleri (API), veritabanları ya da fiziksel eyleyiciler aracılığıyla eyleme geçerek “düşünme” ile “yapma” yetilerini bir araya getirir. Bu yeni nesil sistemler, planlama yapar ve bağlamsal farkındalığa sahiptir (Bandura, 2001). Ajan YZ, içerik oluşturma veya örüntü tanımanın ötesine geçer; amaçlılığı, stratejik planlamayı ve bağlamsal muhakemeyi bünyesinde barındırmaktadır. Kurumsal benimseme hızla artmakta; Forrester'in 2023 yılında yayınladığı "The Rise of AI Agents" raporuna göre, Agentic AI kullanımı son iki yılda %187 artış göstermiş ve MIT Technology Review'in yayınladığı "The State of AI Agents" araştırmasına göre karmaşık görevlerde statik sistemlere kıyasla %43 daha yüksek verimlilik sağladığı görülmektedir (Komtaş, 2025). Kuruluşlar müşteri hizmetleri, tedarik zinciri optimizasyonu ve yazılım geliştirme süreçleri için giderek daha fazla otonom aracı kullandıkça, bu paradigmanın titiz bir kavramsal ve ampirik anlayışı, yalnızca teknolojik ilerleme için değil, aynı zamanda küresel sürdürülebilirlik ve etik zorunluluklarla uyumu için de zorunlu hale gelmektedir.

Bu gelişmeler paralelinde, Ajan YZ'nin potansiyeli hem fırsat hem de risk taşımaktadır. Örneğin, otonom sağlık ajanları kırsal bölgelerde erken teşhisi mümkün kılabileceği gibi, etik denetimden yoksun sistemler algoritmik adaletsizlikleri pekiştirebilmektedir (Guidance, 2021). Bu gerilim, Ajan YZ'nin yalnızca “daha akıllı” değil, aynı zamanda “daha adil ve sürdürülebilir” bir şekilde tasarlanması gerektiğini gündeme getirmektedir (Turan vd., 2025). Ancak mevcut literatür büyük ölçüde teknik performans, ölçeklenebilirlik ve sektörel uygulamalar üzerine odaklanmakta; sürdürülebilirlik, adalet ve etik yönetim gibi boyutlar ise kenara itilmektedir.

(Almulhim ve Yigitcanlar, 2025).

Bu bağlamda, bu çalışma mevcut literatürdeki birbiriyle ilişkili dört temel boşluğu ele almaktadır: (1) Birleşik bir ajan tipolojisinin (agentic AI kavramını kapsayan amaçlılık, etik otonomi ve bağlamsal farkındalık unsurlarını bütünleştiren) eksikliği, (2) Low-code orkestrasyon altyapılarına (sim, n8n v.d.) yönelik sınırlı ilgi, (3) Ajan YZ'nin Sürdürülebilir Kalkınma Hedefleri (SKH) üzerindeki rolünün sınırlı incelenmesi, (4) Ajan YZ literatürünün tematik haritalamasının eksikliği.

Bunları ele almak için, öncelikle Ajan YZ için teorik bir temel oluşturuyor, ardından operasyonel tezahürlerini (özellikle low-code platformlar aracılığıyla) analiz edilmektedir. Son olarak Scopus veri seti üzerinden konu modelleme tekniklerini (Newman vd., 2009) kullanarak akademik söylemin tematik evrimini haritalayan bir hesaplamalı araştırma tasarımı önerilmektedir. Böylece çalışma, Ajan YZ'nin yalnızca teknolojik bir ilerleme değil, aynı zamanda küresel sürdürülebilirlik gündeminin merkezine yerleştirilebilecek çok boyutlu bir sosyo-teknik dönüşüm aracı olarak değerlendirilmesini hedeflemektedir.

2. Literatür Taraması

Sürdürülebilir kalkınmanın tarihsel gelişimi, çevresel krizlerin küresel gündeme taşındığı Stockholm (1972) ve Rio (1992) Zirveleriyle ivme kazanmış; 2015 yılında kabul edilen Birleşmiş Milletler SKH ile küresel ölçekte somut hedeflere dönüştürülmüştür (Silah ve Eğilmez, 2025). Sürdürülebilirlik kavramı, günümüzde yalnızca çevresel dengenin korunması anlamına gelmemekte; aynı zamanda ekonomik adalet, sosyal kapsayıcılık ve kurumsal şeffaflığın bir sentezi olarak kavramsallaştırılmaktadır (Griggs vd., 2013; Van Niekerk, 2020; Leal-Arcas, 2025). Bu üçlü boyut, “insan, gezegen ve refah” ekseninde şekillenen Birleşmiş Milletler'in 2030 Gündemi'nde evrensel bir politika çerçevesiyle kodlanmıştır (Lee vd., 2016). SKH, 17 ana hedef ve 169 alt hedeften oluşan bu gündemin merkezinde (Okuyay, 2020) yer almakta ve küresel düzeyde eşitsizlik, iklim krizi, kaynak tükenmesi ve dijital uçurum gibi yapısal zorluklara karşı kolektif bir çözüm önermektedir (Nakicenovic vd., 2019). Ancak, SKH'lere ulaşım sürecinde karşılaşılan en büyük engellerden biri, teknolojik inovasyonun sürdürülebilirlik ilkeleriyle uyumlu bir şekilde tasarlanamamasıdır (Sachs, 2015; Sachs vd., 2019). Bu bağlamda, YZ hem bir risk hem de bir fırsat kaynağı olarak öne çıkmaktadır.

Geleneksel YZ sistemleri, çoğunlukla kural tabanlı veya istatistiksel öğrenmeye dayalı reaktif modellerden oluşmaktadır. Ancak son yıllarda büyük dil modellerinin (LLM) yükselişiyle birlikte, YZ sistemleri üretkenlik, bağlamsal anlama ve çok modlu entegrasyon kapasiteleri kazanmıştır. Bu dönüşümün bir sonraki aşaması olarak Ajan YZ literatürde yerini almaya başlamaktadır. Ajan YZ, yalnızca veri işlemekle kalmayıp, hedefler doğrultusunda plan yapabilen, çevresel geri bildirimlere göre stratejisini güncelleyebilen ve minimum insan müdahalesiyle karar verebilen otonom sistemleri tanımlamaktadır (Russell ve Norvig, 2021). Bu yaklaşım, klasik “akıllı ajan” teorisini (Wooldridge, 1999) dijital çağın dinamiklerine uyarlayarak, YZ'nin yalnızca araç değil, aynı zamanda aktör olarak düşünülmesini mümkün

kılmaktadır.

Ajan YZ'nin literatürdeki konumlandırılması, üç temel ayırım üzerinden yapılmaktadır. İlk olarak, Üretken YZ (Generative AI) ile karşılaştırıldığında, Ajan YZ'nin temel işlevi “ne”yi üretmek değil, “nasıl” hareket edeceğini belirlemektir (Schneider, 2025). Üretken modeller içerik üretirken, Ajan sistemler bu içerikleri bir hedefe ulaşmak için stratejik olarak kullanmaktadır. İkinci olarak, geleneksel otomasyon sistemleri sabit kurallara bağlı, esnek olmayan yapılar sunarken, Ajan sistemler belirsizlik altında karar verme, bağlamı anlama ve davranışlarını dinamik olarak ayarlama yeteneğine sahiptir (Bandura, 2008). Üçüncü olarak, Ajan YZ, insan–YZ işbirliği (Shneiderman, 2022) perspektifinden değerlendirildiğinde, insanı dışlayan bir otomasyon değil, insanı güçlendiren bir ortaklık modeli sunmaktadır (Floridi ve Cows, 2022).

Bu gelişmeler paralelinde, low-code ve no-code platformlar (n8n ve Sim Studio AI orkestrasyon odaklı, Metabase ve Apache Superset ise BI görselleştirme odaklı) Ajan YZ'nin yaygınlaşmasında kritik bir altyapı rolü üstlenmektedir (Viswanadhapalli, 2025). Sim ve n8n gibi platformlar (Barra vd., 2025; Yu vd., 2025), teknik uzmanlık gerektirmeden çok ajanlı sistemlerin tasarlanmasına, test edilmesine ve üretim ortamına entegre edilmesine olanak tanımaktadır (Rotar ve Zhang, 2025). Bu durum, özellikle kaynak kısıtlı kurumlar ve KOBİ'ler için dijital dönüşümü demokratikleştirmekte ve SKH 9 (Sanayi, İnovasyon ve Altyapı) kapsamında kapsayıcı inovasyonu mümkün kılmaktadır. Rotar ve Zhang (2025), low-code platformların “çoklu ajan orkestrasyonu” için görsel iş akışı dilleri sunarak, karmaşık YZ sistemlerinin kurumsal düzeyde benimsenmesini kolaylaştırdığını göstermiştir. Benzer şekilde, Jeong (2025), n8n, Langflow, Flowise gibi araçların, LLM'lerle entegre edildiğinde “otonom iş akışı ajanları” oluşturmanın mümkün olduğunu vurgulamıştır.

Ancak literatürde Ajan YZ'nin sürdürülebilirlik bağlamındaki etkileri yeterince derinlemesine incelenmemiştir. Mevcut çalışmalar genellikle teknik mimari, performans metrikleri veya sektörel uygulamalar üzerine odaklanırken, SKH'lerle kuramsal ve ampirik bağlantılar sınırlı kalmaktadır (Vinuesa vd., 2020; Cows vd., 2023). Örneğin, Ajan sistemlerin enerji tüketimi, karbon ayak izi, algoritmik adalet veya dijital eşitsizlik üzerindeki etkileri henüz sistematik olarak araştırılmamıştır. Bu durum, sürdürülebilir inovasyonun yalnızca teknolojik verimlilikle değil, aynı zamanda toplumsal ve ekolojik sorumlulukla tanımlanması gerekliliğini göz ardı eden bir dar görüşlülüğe işaret etmektedir.

Bu bağlamda, Vinuesa ve arkadaşları (2020), yapay zekânın Birleşmiş Milletler SKH'ye (169 alt hedeften oluşan) etkisini uzman görüşlerine dayalı bir değerlendirmeye haritalayarak, YZ'nin 134 alt hedefe olumlu katkı sağlayabileceğini, ancak 59 alt hedefi olumsuz etkileyebileceğini göstermektedir. Rolnick ve arkadaşları (2022), “iklim bilinci yapay zekâ” kavramını önererek, YZ sistemlerinin enerji verimliliği, emisyon izleme ve sürdürülebilir kaynak yönetimi gibi alanlarda nasıl yeniden tasarlanabileceğini tartışmıştır. Yine de bu öneriler, ajan sistemlerin otonom karar verme dinamiklerine uyarlanmamıştır. Angubasu (2022), Kenya bağlamında Ajan YZ'nin sağlık hizmetlerine entegrasyonunun SKH 3'e katkı

potansiyelini incelemiş ve yapay zekâ destekli otonom triyaj sistemlerinin kırsal kesimlerde erişilebilir ve güvenli sağlık hizmeti sunumuna destek olabileceğini öne sürmüştür. İlgili çalışma, nitel bir temelli kuram (grounded theory) yaklaşımıyla yürütülmüş ve katılımcı görüşlerine dayalı olarak AI tabanlı triyaj sistemlerinin uygulanması için politika, altyapı, düzenleyici mekanizmalar ve kullanıcı tercihlerini içeren kapsamlı bir çerçeve önermektedir.

Roll ve Wylie (2016), YZ Destekli Eğitim (AIED) alanının 25 yıllık gelişimini inceleyerek, gelecekte eğitimi dönüştürmede uyarlanabilir öğrenme sistemlerinin potansiyelini vurgulamaktadır. Çalışma, AIED'in mevcut sınıf uygulamalarını destekleyen “evrimsel” yaklaşımların yanı sıra, öğrencilerin günlük yaşamına, kültürlerine ve topluluklarına entegre olan “devrimci” yeniliklere odaklanması gerektiğini önermektedir. Ancak bu tür sistemlerin yaygınlaşması, etik yönetim, açıklanabilirlik ve insan-döngüsü-içi denetim gibi kritik boyutların dikkate alınmasını gerektirmektedir. Bu eksiklik, Floridi ve Cows (2022) tarafından öne sürülen bir paradoksa işaret etmektedir: yapay zekâ ne kadar gelişmiş olursa olsun, toplumsal değerlerle uyumsuz bir şekilde tasarlanırsa mevcut eşitsizlikleri derinleştirebilir veya yeni zararlar yaratabilir. Yazarlar, etik YZ için fayda, zararsızlık, özerklik, adalet ve açıklanabilirlik ilkelerinden oluşan bir çerçeve önermektedir. Bu nedenle, Ajan YZ'nin sürdürülebilir potansiyelinin açığa çıkarılması, yalnızca teknik gelişmişlikle değil, aynı zamanda tasarımdan itibaren etik ilkeler, katılımcı değerlendirme mekanizmaları ve disiplinlerarası araştırma yaklaşımları ile mümkün olabilecektir.

Bu çalışma hem kuramsal hem de ampirik düzeyde literatürdeki bu boşlukları doldurmayı hedeflemektedir. Kuramsal olarak, Ajan YZ için özerklik, bellek, etkileşim ve öğrenme boyutlarında sistematik bir tipoloji önermekte, ampirik olarak ise Scopus veri seti üzerinden tematik ve konu modelleme analizleriyle alanın tematik evrimini haritalamaktadır. Böylece, Ajan YZ'nin yalnızca “daha akıllı sistemler” değil, aynı zamanda “daha adil ve sürdürülebilir sistemler” inşa etme potansiyelini değerlendirmeyi amaçlamaktadır.

3. Ajan Yapay zekâ için Kavramsal ve Operasyonel Bir Çerçeve

Ajan YZ, geleneksel, reaktif YZ sistemlerinden hedef odaklı özerklik, bağlamsal uyarlanabilirlik ve proaktif muhakeme ile karakterize edilen çerçevelere dönüştürücü bir geçişi ifade etmektedir. Kavramsal temelde, Ajan YZ, akıllı ajanların karmaşık ortamlarla dinamik etkileşimini kolaylaştırmak üzere klasik “Algıla-Planla-Harekete Geç” (Sense-Plan-Act) çerçevesini (Russell ve Norvig, 2021) temel almakta, ancak bu çerçeveyi geri bildirimle öğrenme (learning via feedback) mekanizmalarıyla genişleterek özyinelemeli bir döngüye dönüştürmektedir.

Bu döngüde, ajanlar önce doğal dil ifadeleri, kullanıcı niyetleri ve gerçek zamanlı sensör akışları gibi çok modlu girdileri almaktadır. Bu, çalıştıkları bağlamı durumsal olarak anlamalarına yardımcı olmaktadır. Ardından, hedeflerine ulaşmalarına yardımcı olacak bir dizi eylem belirlemek için düşünsel veya sezgisel akıl yürütme kullanarak plan yapmaktadırlar. Ardından, eylem aşamasında, bu planlar API entegrasyonları, robotik aktüatörler veya dijital iş

akışı düzenleyicileri gibi farklı arayüz türleri kullanılarak hayata geçirilmektedir. Öğrenme kısmı çok önemlidir çünkü açık (kullanıcı düzeltmeleri) veya örtük (performans ölçütleri) geri bildirimleri alarak ve bunları zaman içinde iç modelleri, stratejileri ve davranış politikalarını iyileştirmek için kullanarak döngüyü tamamlamaktadır. Sadece sabit girdilere yanıt veren pasif YZ modellerinden farklı olarak, ajan sistemler doğası gereği teleolojik olmaktadır. Bu, düşünce ve eylemlerinin belirli hedeflere ulaşmaya yönelik olduğu anlamına gelir; bu da insanların davranışlarına benzer bir araçsal rasyonellik biçimini temsil etmektedir. Bu kavramsal çerçeve, Ajan YZ'nin işlevsel anatomisini tanımlamakla kalmaz, aynı zamanda giderek karmaşıklaşan sosyo-teknik ekosistemlerde özerk, uyarlanabilir ve amaç odaklı bir işbirlikçi olarak çalışma potansiyelini de vurgulamaktadır.

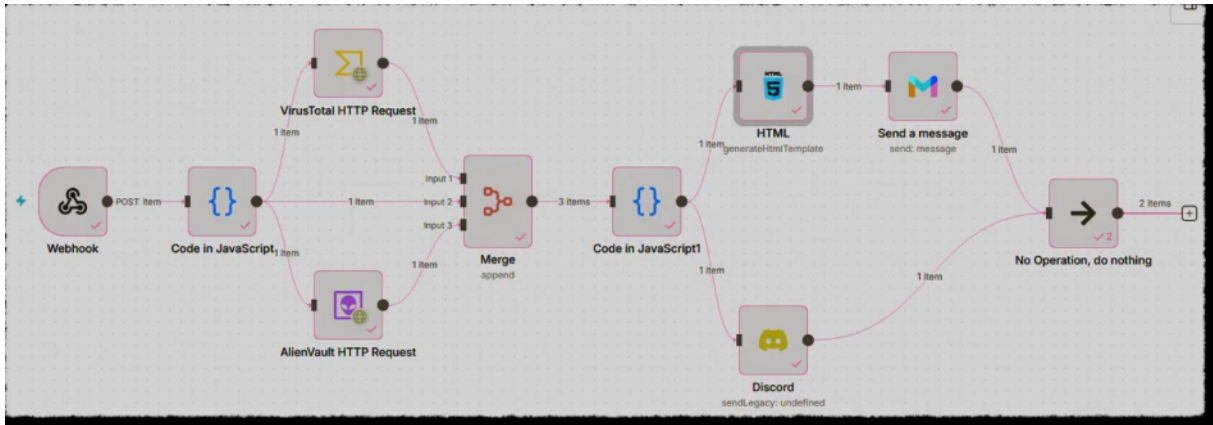
Ajan YZ, hem Üretken YZ hem de geleneksel otomasyondan kasıtlı olarak farklılık göstermektedir. Üretken YZ, metin, kod, görüntü veya çok modlu çıktılar gibi yeni şeyler bir araya getirmekte etkin rol oynamaktadır. Ancak, asıl görevi bağlamla ilgili eserler yaratarak ne olduğunu temsil etmektir. Diğer yandan, ajan YZ, üretken çıktıları, bir görevi yapmanın “nasıl”, “ne zaman” ve “neden” sorularını yanıtlayan amaçlı eylem dizilerini planlamak için girdi olarak kullanmaktadır. Bu, iki paradigmanın doğal olarak birbirini tamamladığı anlamına gelmektedir: üretken modeller ham bilişsel materyali sağlamakta, Ajan sistemler ise ona amaç ve hareket etme yeteneği kazandırmaktadır. Bağlamsal uyarlanabilirlikten yoksun, esnek olmayan if-then komut dosyalarıyla karakterize edilen kural tabanlı veya deterministik otomasyonun aksine, ajan sistemler bilişsel esneklik göstermektedir. Kullanıcıların gerçekten ne istediğini anlamakta, belirsizlikle başa çıkar ve ortamın nasıl değiştiğine göre planlarını anında değiştirmektedir. Bu yetenek, reaktif, önceden programlanmış davranıştan proaktif, hedef odaklı zekâya önemli bir geçişe olanak tanımaktadır. Ajan YZ, üretken yeteneklerin ve eski otomasyon sistemlerinin üzerinde yer alan, niteliksel olarak daha gelişmiş bir otonom muhakeme katmanı olarak yer almaktadır.

Ajan YZ, tasarımı, kullanımı ve etkileri çok boyutlu bir tipoloji ve destekleyici altyapı aracılığıyla sistematik olarak anlaşılabilen çok çeşitli akıllı sistemleri içermektedir. Ajan sistemleri, birbiriyle ilişkili dört (özerklik düzeyi, bellek kullanımı, çevreyle etkileşim ve öğrenme yeteneği) gruba ayrılmaktadır ve beş ana ajan türü ile sonuçlanmaktadır. Reaktif ajanlar bellek olmadan çalışır ve çevrelerindeki şeylere anında tepki vermektedir. Örnekler arasında kural tabanlı sohbet robotları ve alarm sistemleri bulunmaktadır. Otonom araçlar veya lojistik optimizasyon sistemleri gibi düşünsel ajanlar, içsel dünya modellerini kullanarak ne olacağını tahmin etmekte ve gelecek için plan yapabilmektedir. Sesli asistanlar (Alexa ve Siri gibi) ve işbirlikçi robotlar (cobotlar), kullanıcılar veya sistemlerle sürekli, iki yönlü iletişim kuran etkileşimli ajanlara örnek olabilmektedir. Uyarlanabilir ajanlar, geri bildirimleri kullanarak stratejilerini zaman içinde geliştirerek bir adım daha ileri gitmektedir. Bu, kişiselleştirilmiş öğrenme platformları ve dinamik fiyatlandırma motorları gibi uygulamaları mümkün kılmaktadır. Son olarak, Çoklu Ajan Sistemleri (Mahela vd., 2020), akıllı şebekeler, şehir trafiği kontrolü ve sürü robotikleri gibi zorlu ve yaygın sorunları çözmek için farklı

türdeki ajanları bir araya getirmektedir. Bu tipoloji, ajan becerilerini belirli bir alana özgü işlevsel, etik ve operasyonel ihtiyaçlarla eşleştirmek için kullanıcılara yapılandırılmış bir yol sunmaktadır.

Altyapıyı basitleştirirken ifade gücünü koruyan düşük kodlu orkestrasyon platformları, bu tür ajan mimarilerini gerçek hayatta oluşturmayı kolaylaştırmaktadır. Sim (Sim Studio AI, 2025) ve n8n (n8n.io, 2025), bu alanda öne çıkan iki örnek olarak yaygın biçimde tercih edilmektedir. Sim, kullanıcıların doğal dil kullanarak ajan davranışlarını ayarlamasına ve birden fazla ajanın etkileşimini test etmesine olanak tanıyarak karmaşık ajan mantığını tasarlamayı kolaylaştırmaktadır. n8n ise ajanları CRM'ler, veritabanları ve LLM API'leri gibi çok sayıda harici hizmete bağlayan açık kaynaklı bir iş akışı otomasyon aracı olarak sistemde görev almaktadır. Bu yapı, ajan ekosistemlerinin adeta bir “duyusal sinir sistemi” gibi çalışmasını sağlamaktadır.

Bu entegrasyonun pratik bir örneği, bir n8n otomasyon akışında görülmektedir: Wazuh ile tespit edilen kötü amaçlı dosyalar, VirusTotal ve AlienVault OTX gibi tehdit istihbaratı kaynakları da dahil edilerek ajan tabanlı analizle incelenmiş ve otomatik uyarı mekanizması devreye alınmıştır. Bu sayede sosyal mühendislik saldırılarından korunulmuş ve farkındalık davranışa dönüşmüştür. Bu süreç, Şekil 1’de görselleştirilen ajan orkestrasyonu şemasıyla desteklenmektedir. Böylelikle, Sim’in yüksek düzeyli muhakeme ve davranış koordinasyonu yeteneği ile n8n’in veri alımı ve eylem yürütme gücü sorunsuz biçimde bütünleşmekte, gelişmiş otomasyonun teknik olmayan ekipler ve KOBİ’ler için erişilebilir hale gelmesini sağlamaktadır. Bu sinerji, Ajan YZ’nin yalnızca kavramsal değil, aynı zamanda uygulanabilir bir dijital dönüşüm aracına dönüşmesini desteklemektedir.



Şekil 1. Ajan YZ destekli kötü amaçlı dosya tespiti otomasyonu: n8n iş akışı

Ajan YZ, operasyonları daha verimli hale getirmenin yanı sıra, Birleşmiş Milletler SKH’nin gerçekleştirilme şeklini dönüştürme potansiyeline sahiptir. Bu teknoloji, sağlık, eğitim, istihdam, üretim ve iklim eylemi gibi farklı alanlarda somut katkılar sunabilmektedir.

Ajan sistemlerin bu hedeflerle olan ilişkisi Tablo 1’de özetlenmiştir. Özellikle SKH 3 kapsamında otonom sağlık ajanları erken teşhisi kolaylaştırabilir; SKH 4 için uyarlanabilir öğrenme sistemleri fırsat eşitliğini artırabilir; SKH 8 kapsamında rutin görevlerin otomasyonu

Tablo 1. Ajan YZ ve SKH Arasındaki Kavramsal Eşleşme

SKH Kodu	Hedef Başlığı	Ajan YZ ile İlişkili Uygulama Alanı	Potansiyel Katkı Türü
SKH 3	İyi Sağlık ve Refah	Otonom sağlık ajanları, erken teşhis sistemleri, uzaktan izleme.	Sağlık erişimini artırma, teşhis süresini kısaltma.
SKH 4	Nitelikli Eğitim	Uyarlanabilir öğrenme platformları, bireyselleştirilmiş eğitim ajanları.	Eğitimde kapsayıcılık ve kişiselleştirme.
SKH 8	İnsana Yakışır İş ve Ekonomik Büyüme	Rutin görev otomasyonu, insan-ajan işbirliği modelleri.	Verimlilik artışı, yaratıcı işlere odaklanma.
SKH 9	Sanayi, Yenilikçilik ve Altyapı	Low-code/No-code ajan orkestrasyon platformları (ör. Sim, n8n).	KOBİ'ler için dijital dönüşümün demokratikleşmesi.
SKH 12	Sorumlu Tüketim ve Üretim	Tedarik zinciri optimizasyonu, kaynak verimliliği ajanları.	Atık azaltımı, sürdürülebilir üretim planlaması.
SKH 13	İklim Eylemi	Karbon ayak izi izleme, enerji yönetimi ajanları.	Emisyon azaltımı ve sürdürülebilir enerji kullanımı.

yaratıcılığa alan açabilir (Ergani, 2024); SKH 9 için düşük kodlu platformlar (Sim Studio AI, 2025; n8n.io, 2025) kapsayıcı dijitalleşmeyi destekleyebilir; SKH 12 ve SKH 13 kapsamında tedarik zincirleri ve enerji yönetimi ajanları sürdürülebilir üretim ve karbon azaltımına katkı sağlayabilir. Ancak bu potansiyel, yalnızca etik tasarım ilkeleri, algoritmik adalet ve kapsayıcı yönetim çerçeveleri benimsendiğinde sürdürülebilir faydaya dönüşebilmektedir (Floridi ve Cows, 2022; Vinuesa vd., 2020).

Bu nedenle, ajan sistemleri daha bağımsız hale geldikçe güçlü etik ve düzenleyici korumalara ihtiyaç duymaktadırlar. En büyük sorunlardan bazıları, kararların hesap verebilir bir şekilde alınmasını, karmaşık ajan davranışlarının açıklanabilmesini ve tarihsel veya sistemik önyargıların daha da kötüleşmemesini sağlamaktır. Bu riskleri azaltmak için, (1) sağlık hizmetleri veya ceza adaleti gibi yüksek riskli alanlarda insan denetimli (HITL) gözetim; (2) değiştirilemeyen ve daha sonra incelenen karar günlükleri aracılığıyla denetlenebilirlik; ve (3) AB AI Yasası, GDPR ve ISO/IEC 24027 gibi düzenleyici gereklilikleri doğrudan ajan mimarilerine yerleştiren tasarımdan itibaren uyumluluk çerçeveleri olmak üzere üç temel fikir önerilmektedir (Cantero Gamito ve Marsden, 2024). Proaktif yönetim olmadan, ajan YZ dijital eşitsizlikleri daha da kötüleştirebilir veya büyük ölçekte zararı otomatikleştirebilir. Bu nedenle, sorumlu inovasyon, teknolojiye ilerlemenin etik öngörü ve katılımcı denetimle el ele gitmesi gerektiği anlamına gelmektedir.

4. Araştırma Metodolojisi: Ajan Yapay zekâ Haritalandırılması

Çalışmada, Ajan YZ alanının bilimsel gelişimi ampirik olarak incelenmiştir. Tematik analiz (Braun ve Clarke, 2019) ile denetimsiz konu modellemesi gibi hesaplamalı ve manuel içerik analizi tekniklerini bir araya getiren hibrit bir içerik analizi yaklaşımı uygulanmıştır

(Krippendorff, 2018; Nelson, 2020). Bu yaklaşım, karma yöntem araştırmalarının “convergent parallel design” (Creswell ve Plano Clark, 2023; Delaney vd., 2017) desenine benzer şekilde, nitel ve nicel veri kaynaklarının aynı araştırma sorusuna eş zamanlı olarak katkı sağlaması prensibine dayanmaktadır. Ancak bu çalışmada, geleneksel karma yöntem desenlerinden ziyade hesaplamalı içerik analizi çerçevesinde geliştirilen hibrit analiz modelleri (Nelson, 2020) temel alınmıştır. Bu çerçevede, büyük metin veri setlerinin otomatik işlenmesiyle elde edilen nicel örüntüler, araştırmacıların manuel kodlama ve yorumlama süreçleriyle zenginleştirilerek bütüncül bir analiz sunulmuştur.

Veri toplama sürecinde, Scopus veritabanı kullanılmıştır. Sorgu ifadesi olarak *TITLE-ABS-KEY(“agentic ai” OR “agentic artificial intelligence”)* anahtar kelimeleri seçilmiştir. Bu sorgu ile 218 adet akademik yayın belirlenmiş ve analiz veri seti oluşturulmuştur. Veri seti, 2023–2025 yıllarını kapsamaktadır. Bu tarih aralığının seçilmesinin temel gerekçesi, “Ajan Yapay Zekâ” kavramının hem kavramsal hem de teknik düzeyde literatürde net bir çerçeve kazanmaya başladığı dönemin 2023 sonrası olduğunun gözlemlenmesidir. Özellikle 2023 yılında büyük dil modellerinin (LLM) gelişmesiyle birlikte, sistemlerin yalnızca pasif araçlar olmaktan çıkıp hedef odaklı, planlama yapabilen, araçları kullanabilen ve çevresiyle etkileşim kurabilen “ajanlar” haline gelmesi, bu paradigmanın bilimsel tartışmalara konu olmasına yol açmıştır (Shinn vd., 2023). Bu dönüm noktası, ajan YZ’nin sadece bir metafor değil, aynı zamanda teknik bir mimari ve araştırma ajanlığı olarak literatürde yer edinmesini sağlamıştır. Bu nedenle, bu çalışma, LLM devriminin ardından şekillenen yeni ajan paradigmalarını yakalamayı hedeflemektedir.

4.1. Veri İşleme ve Analiz Süreci

Literatürde NLP ve makine öğrenme çalışmalarında analiz yapılırken farklı ve yeni araçlar bulunmaktadır. R ve Python’un öncülüğünde başlayan bu trend, Weka, KNIME ve Orange gibi programlama tabanlı olmayan diğer ücretsiz analitik araçlarına ve platformlarına da yayılmaktadır (Delen, 2024). KNIME (Konstanz Information Miner) (Berthold vd., 2008), Mobyly (Néron vd., 2009), Galaxy (Goecks vd., 2010), Taverna (Hull vd., 2006), Kepler (Ludäscher vd., 2006), geWorkbench (Floratos vd., 2010), Conveyor (Linke vd., 2011) gibi iş akışı yönetim sistemlerinden ayrılmaktadır. KNIME, belirli bir alana özgü bir çözüm sunmaktan ziyade, güçlü veri ön işleme ve veri analizi yetenekleriyle öne çıkan esnek bir entegrasyon omurgası niteliğindedir (Jagla vd., 2011). Bu yönüyle farklı veri kaynaklarının ve analiz yöntemlerinin bir araya getirilmesine olanak tanıyarak, çok disiplinli araştırmalarda etkin bir şekilde kullanılabilir. KNIME Analytics Platform, başlıklar, özetler, anahtar kelimeler ve alıntılar gibi ham bibliyografik verileri işlemektedir. Bu, metin ön işleme (tokenizasyon, durdurma kelimesi kaldırma, lemmatizasyon) ve özellik mühendisliği için tekrarlanabilir, görsel iş akışları oluşturulmasına olanak tanımaktadır. Ardından, büyük veri kümelerinde daha iyi çalışan bir Latent Dirichlet Allocation (LDA) türü olan Paralel LDA kullanarak gizli tematik kümeler belirlenmektedir. Paralel LDA, dağıtılmış bilgi işlem

kaynaklarından sonuçlar çıkarmayı kolaylaştırmakta, bu da onu dinamik, yüksek hacimli akademik veri kümeleri için iyi bir seçim haline getirmektedir.

Konu algılama sürecimiz, Şekil 2’de gösterildiği gibi, birbirini takip eden ve birbiriyle bağlantılı olmayan dört adımdan oluşmaktadır. Konu modelleme analizimizin güvenilirliğini ve geçerliliğini sağlamak için, KNIME Analytics Platform (v5.5.0) kullanarak titiz ve tekrarlanabilir bir metin ön işleme süreci uygulanmıştır. Birincil veri kaynağı, çalışmanın odak noktası ve metodolojisinin özetlenmiş semantik temsilini içeren her akademik çalışmanın “Özet/Abstract” alanı olarak seçilmiştir.

Veri Okuma, CSV Reader düğümü (1. adım) CSV formatındaki dosyadan özetleri okumaktadır. Sonrasında, Ön İşleme (2. adım) gelmektedir; bu adımda metin, noktalama işaretleri, sayılar, genel ve alana özgü durdurma kelimeleri temizlenerek ve büyük/küçük harfler normalize edilerek sistematik bir şekilde temizlenmektedir. Bu adım, dilin aynı olmasını ve analizin daha doğru olmasını sağlamak için yapılmaktadır. Konu Algılama, Konu Çıkarıcı (Paralel LDA) düğümünü (3. adım) kullanarak, terimlerin belge kümesinde ne sıklıkla birlikte görüldüğüne bakarak 15 gizli konu belirlenmektedir. Her konuyu en iyi temsil eden 8 anahtar kelime, onu benzersiz kılan unsurları ifade etmektedir.

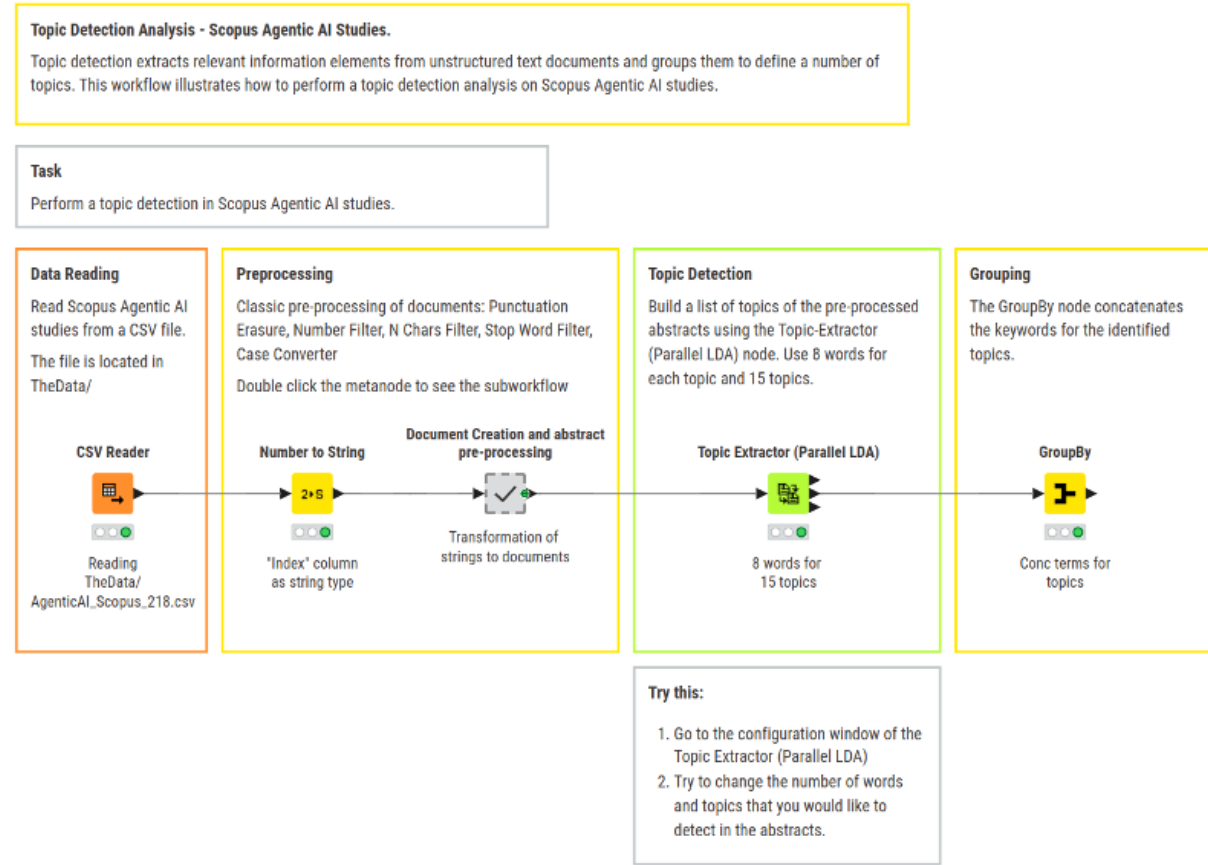
Gruplama aşaması, son olarak GroupBy düğümünü (4. adım) kullanarak her konu için ana anahtar kelime kümelerini bir araya getirmekte ve listelemektedir (Şekil 3). Bu, temaların daha kolay anlaşılmasını sağlamaktadır. Bu iş akışı, metin madenciliğine yeniden üretilebilir, ölçeklenebilir ve erişilebilir bir yaklaşım sunarak, sosyal bilimlerde araştırmacıların ileri düzey programlama uzmanlığı gerektirmeden, Ajan YZ literatüründe ortaya çıkan tematik kümeleri sistematik olarak tanımlamasına ve incelemesine olanak tanımaktadır.

Şekil 3’te gösterilen KNIME platformunda ön işleme dizisi, ham metin verilerini LDA gibi doğal dil işleme (NLP) görevleri için uygun, standartlaştırılmış, gürültü azaltılmış bir biçime dönüştürmek için tasarlanmış yedi ardışık işlemten oluşmaktadır.

Dizgiyi Belgeye Dönüştürme (Strings to Document) düğümü (1. adım), özet dizgilerini belge nesnelere dönüştürmektedir. Bu dönüşüm, KNIME’nin metin madenciliği uzantılarının girişi sadece karakter dizileri olarak değil, dilbilimsel birimler olarak tanınması ve işlemlerinden dolayı kritiktir. Her özet, ilişkili meta verilerle (ör. makale kimliği, yayın yılı) ayrı bir belge nesnesi haline gelmekte ve analiz sırasında izlenebilirliği sağlanmaktadır.

Sütun Filtreleme (Column Filter) düğümü (2. adım), yeni oluşturulan belge sütununu korumak ve diğer tüm meta veri sütunlarını (ör. başlık, yazarlar, DOI, atıf sayısı) atmak için kullanılmaktadır. Bu, sonraki metin işleme adımlarının yalnızca özet içeriği üzerinde çalışmasını sağlayarak dilbilimsel olmayan özelliklerin etkisini en aza indirmektedir.

Noktalama İşaretlerinin Silinmesi (Punctuation Erasure) düğümü (3. adım), virgül, nokta, tırnak işareti, parantez ve tire dahil olmak üzere tüm alfasayısal olmayan karakterleri metinden kaldırmaktadır. Bu adım, kelime sınırlarını korurken noktalama işaretlerinin neden olduğu sözcüksel varyasyonu azaltmaktadır. Örneğin, “agent-based systems.” “agent based systems” haline gelmekte ve tutarlı tokenleştirmeyi kolaylaştırmaktadır.

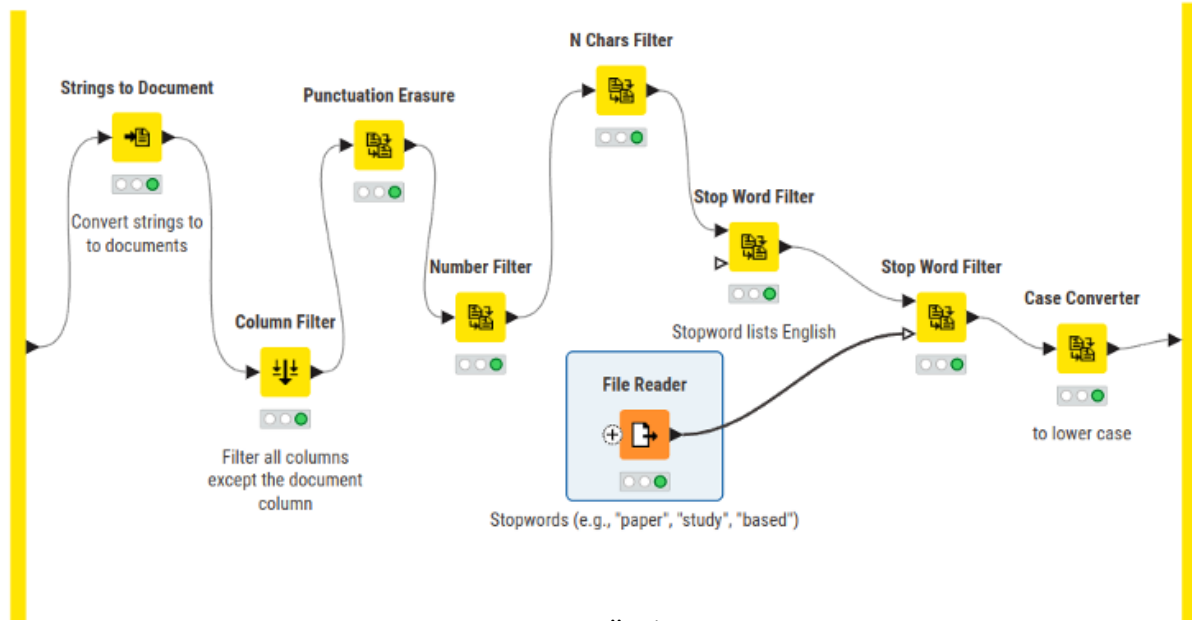


Şekil 2. KNIME Konu Belirleme İş Akışı

Sayı Filtresi (Number Filter) düğümü (4. adım) tüm sayısal belirteçleri (ör. “2023”, “3.14”, “Şekil 11”) ortadan kaldırmaktadır. Sayılar belirli bağlamlarda (ör. istatistiksel sonuçlar) anlamsal bir anlam taşıyabilmekte, ancak konu modellemesine dahil edilmeleri genellikle tematik tutarlılığa katkıda bulunmadan gürültü yaratmaktadır. Bunların kaldırılması, gizli konuların yorumlanabilirliğini artırmaktadır.

N Karakter Filtresi (N Chars Filter) düğümü (5. adım), üç karakterden kısa belirteçleri (ör. “a”, “an”, “it”, “of”) atmaktadır. Bu kısa belirteçler genellikle işlev kelimeleri veya kökleme/lemmatizasyonun artefaktlarıdır ve nadiren konu farklılaşmasına katkıda bulunmaktadır. Yalnızca ≥ 3 karakterli terimleri tutmak, son terim-belge matrisinde sinyal-gürültü oranını iyileştirmektedir.

Yüksek sıklıkta kullanılan, anlamsal değeri düşük kelimeleri kaldırmak için iki paralel Durdurma Kelimesi Filtresi (Stop Word Filter ve File Reader) düğümleri (6. adım) uygulanmaktadır. İlki, yaygın bağlaçlar, edatlar ve yardımcı fiiller (ör. “the”, “and”, “is”, “were”) içeren KNIME’nin yerleşik İngilizce durdurma kelimesi listesini kullanmaktadır. İkincisi, akademik özetlerde sıklıkla bulunan ancak tematik özgüllükten yoksun, alan bağımsız terimler içeren, bilimsel literatür için özel olarak derlenen özel bir durdurma kelimesi listesi kullanmaktadır, örneğin “makale”, “çalışma”, “dayalı”, “kullanarak”, “önerilen”, “sunulan”, “tartışılan”. Bu özel liste, konu bütünlüğünü zayıflatan “dolgu” terimlerini ortadan kaldırmak için ön LDA çıktılarının tekrarlı manuel incelemesi yoluyla derlenerek yazar tarafından



Şekil 3. KNIME ile Ön İşleme Adımları

oluşturulmaktadır.

Son olarak, bir Büyük/Küçük Harf Dönüştürücü (Case converter) düğümü (7. adım) kalan tüm belirteçleri küçük harfe dönüştürmektedir. Bu standardizasyon, aynı kelimenin büyük/küçük harf kullanımı nedeniyle farklı varlıklar olarak değerlendirilmesini önlemekte (ör. “Agent” ve “agent”), böylece terim sıklığı toplama işlemini iyileştirmekte ve konu modelinin istikrarını artırmaktadır.

Bu iş akışını uyguladıktan sonra, orijinal 218 özet, yaklaşık 12.500 benzersiz terimden (durdurma kelimeleri çıkarıldıktan sonra) oluşan, temizlenmiş, normalleştirilmiş ve anlamsal olarak zenginleştirilmiş bir metin kümesine dönüştürülmektedir. Ortaya çıkan belge-terim matrisi, Paralel LDA modellemesi için girdi olarak kullanılmakta ve keşfedilen konuların biçimlendirme, noktalama veya sözcüksel fazlalıkların bir sonucu değil, anlamlı kavramsal kümeleri yansıttığından emin olunmasını sağlanmaktadır. Bu ön işleme stratejisi, hesaplamalı sosyal bilimler ve tematik analizdeki en iyi uygulamalarla uyumlu olmaktadır (Grimmer ve Stewart, 2013; Blei vd., 2003) ve bulgularımızın tekrarlanabilirliğini ve şeffaflığını desteklemektedir; bu da YZ etiği ve politikası alanındaki ampirik araştırmalar için kritik unsurlar olmaktadır.

5. Scopus Ajan Yapay zekâ Araştırmalarında Tematik Kümeler: LDA ile Elde Edilen Bulgular

Kavramsal analizi ampirik olarak temellendirmek amacıyla, Bölüm 4’te belirtilen sorgu dizisi kullanılarak elde edilen veri seti üzerinde konu modelleme analizi gerçekleştirilmiştir. Metinler, KNIME platformu aracılığıyla ön işleme tabi tutulmuş; bu süreçte tokenleştirme, lemmatizasyon ve alana ilişkin olmayan terimlerin (stopwords ve domain dışı kelimeler) filtrelenmesi uygulanmıştır. Konu sayısı, koherens (tutarlılık) puanları ve anlamsal

yorumlanabilirlik dikkate alınarak $k = 15$ olarak belirlenmiş ve Paralel LDA yöntemi uygulanmıştır.

Yao ve arkadaşları (2009), SparseLDA'nın geleneksel LDA'ya (Parmaksız ve Akarsu, 2025) kıyasla yaklaşık 20 kat daha hızlı çalışabildiğini, mevcut hızlı örnekleme yöntemlerinden ise iki kat daha verimli olduğunu ve aynı zamanda önemli ölçüde daha az bellek tükettiğini göstermişlerdir. Bu çalışmada da işlem süresini kısaltmak ve bellek kullanımını optimize etmek amacıyla paralel LDA algoritması tercih edilmiştir. Analiz sonuçları Şekil 4'te verilmiştir. LDA ile belirlenen konulara ait SKH ilişkileri ise Şekil 5'te Digraph yapısıyla

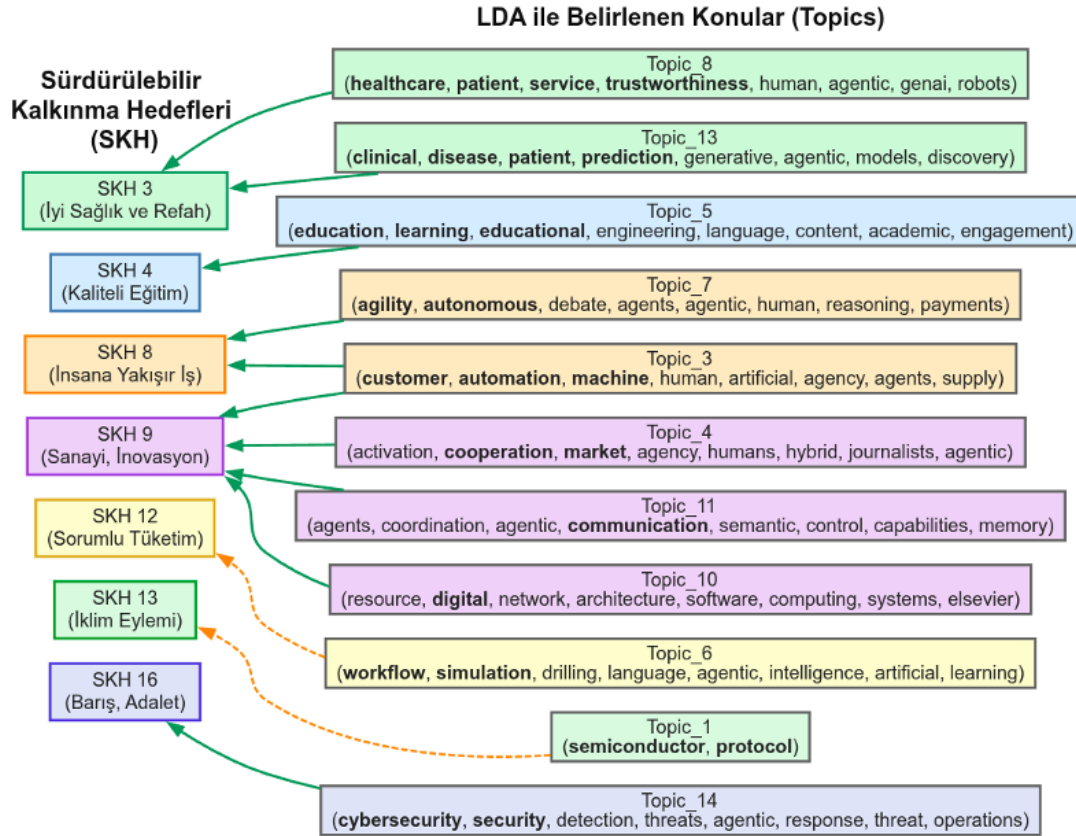
Topic id ↑ String	Concatenate(Term) String
topic_0	agentic, models, elsevier, learning, language, potential, real-time, challenges
topic_1	language, semiconductor, protocol, model, archival, languages, agency, quality
topic_2	clinical, learning, medication, language, machine, prince, utility, simulation
topic_3	human, artificial, customer, agency, agents, supply, automation, machine
topic_4	activation, agentic, market, agency, cooperation, humans, hybrid, journalists
topic_5	education, learning, educational, engineering, language, content, academic, engagement
topic_6	drilling, workflow, language, simulation, agentic, intelligence, artificial, learning
topic_7	agility, debate, agents, agentic, payments, human, reasoning, autonomous
topic_8	healthcare, human, service, patient, agentic, trustworthiness, genai, robots
topic_9	agents, agentic, software, voice, elsevier, models, directions, architecture
topic_10	resource, digital, network, architecture, software, computing, elsevier, systems
topic_11	agents, coordination, agentic, communication, semantic, control, capabilities, memory
topic_12	agentic, intelligence, elsevier, artificial, challenges, agents, ethical, generative
topic_13	models, clinical, disease, generative, patient, agentic, prediction, discovery
topic_14	cybersecurity, security, detection, threats, agentic, response, threat, operations

Şekil 4. KNIME ile Elde Edilen Paralel LDA Çıktıları

görselleştirilerek verilmiştir. SKH'lerle ilişkilendirilen topic'lerdeki term'lerin net belirginleştirilmesi için kalın fontta vurgulanmıştır. Ayrıca şekilde SKH'lere açık olarak katkı verdiği düşünülen Topic'ler yeşil düz çizgi, gizli katkılar ise turuncu kesikli çizgilerle görselleştirilmiştir. LDA analiziyle belirlenen tematik yapıların, SKH'lere yönelik bilimsel katkı alanları, Tablo 2'de detaylı olarak sunulmuştur.

Elde edilen analiz sonuçları incelendiğinde, topic_3 ("human, artificial, customer, agency, agents, supply, automation, machine") ve topic_7 ("agility, debate, agents, agentic, payments, human, reasoning, autonomous") gibi konular, SKH 8 (İnsana Yakışır İş ve Ekonomik Büyüme) ve SKH 9'un (Sanayi, İnovasyon ve Altyapı) merkezinde yer alan yapay zekâ özerkliği, insan denetimi ve sosyo-teknik entegrasyonun sınırlarını sorgulamaktadır. Bu kümeler, otomasyonlu müşteri hizmetlerinde, finansal sistemlerde ve iş gücünü artıran teknolojilerde insan kontrolünü vurgulayarak, otomasyonun insan onurunu ve ekonomik katılımı ortadan kaldırmak yerine artırmasını sağlama hedefini dolaylı olarak desteklemektedir. topic_11 ("agents, coordination, agentic, communication, semantic, control, capabilities, memory"), dayanıklı ve uyarlanabilir dijital ekosistemleri mümkün kılan, birlikte çalışabilir,

akıllı altyapıların tasarımını ilerletmek suretiyle SKH 9’a daha da katkıda bulunmaktadır.



Şekil 5. Ajan YZ Çalışmalarının LDA Konuları ve SKH İlişkileri

Topic_9 (“agents, agentic, software, voice, elsevier, models, directions, architecture”) ve topic_12 (“agentic, intelligence, elsevier, artificial, challenges, agents, ethical, generative”) kapsamındaki araştırmalar, ölçeklenebilir ve akıllı otomasyon için temel oluşturan LLM ile entegre etken çerçevelerine ve üretken akıl yürütmeye odaklanmaktadır. Bu gelişmeler, birden fazla SKH genelinde ilerlemenin temelini oluştursa da, mevcut çerçevelerde açık bir sürdürülebilirlik veya kapsayıcılık kriteri bulunmamaktadır. Endüstriyel ve operasyonel iş akışlarını inceleyen topic_6 (“drilling, workflow, language, simulation, agentic, intelligence, artificial, learning”), süreç optimizasyonu ve kaynak verimliliği yoluyla SDG 12 (Sorumlu Tüketim ve Üretim) için gizli bir potansiyel barındırmaktadır, ancak bu bağlantı literatürde henüz dile getirilmemiştir. Benzer şekilde, topic_1’deki “semiconductor/yarı iletken” ve “protocol/protokol” ifadeleri, SKH 13 (İklim Eylemi) ile uyumlu, enerji tasarruflu bilgi işlemi destekleyebilecek donanım-yazılım ortak tasarımına işaret etmektedir, ancak sürdürülebilirlik ölçütleri söylemde yer almamaktadır.

Birkaç küme, temel SKH’lerle güçlü ve açık bir uyum sergilemektedir. Topic_8 (“healthcare, human, service, patient, agentic, trustworthiness, genai, robots”) ve topic_13 (“models, clinical, disease, generative, patient, agentic, prediction, discovery”), yapay zekâ destekli klinik karar desteği, hastalık tahmini ve hasta merkezli bakımı mümkün kılarak, özellikle kaynak kısıtlı ortamlarda değerli olan SKH 3’ü (İyi Sağlık ve Refah) doğrudan

Tablo 2. Ajan Yapay zekâ Çalışmalarının LDA ile Belirlenen Tematik Yapıları ve Bilimsel Etkileri

SKH	İlgili Topic	LDA Anahtar Kelimeler	Katkı Alanları
SKH 3 – İyi Sağlık ve Refah	Topic_8, Topic_13	healthcare, patient, service, trustworthiness, clinical, disease, prediction, discovery	YZ destekli klinik karar desteği, hastalık tahmini ve hasta merkezli bakım uygulamalarıyla özellikle kaynak kısıtlı ortamlarda sağlık hizmetlerinin erişimini ve kalitesini artırmaktadır.
SKH 4 – Kaliteli Eğitim	Topic_5	education, learning, educational, engagement	Kişiselleştirilmiş ve uyarlanabilir öğrenme sistemleri aracılığıyla küresel ölçekte yüksek kaliteli eğitime erişimi desteklemektedir.
SKH 8 – İnsana Yakışır İş ve Ekonomik Büyüme	Topic_3, Topic_7	automation, customer, machine, agility, autonomous, human	Otomasyonlu müşteri hizmetleri, finansal sistemler ve iş gücünü artıran teknolojilerde insan denetimini vurgulayarak, otomasyonun insan onurunu ve ekonomik katılımı artırmasını amaçlamaktadır.
SKH 9 – Sanayi, İnovasyon ve Altyapı	Topic_3, Topic_4, Topic_10, Topic_11	market, cooperation, communication, digital, systems, architecture	Akıllı, dayanıklı dijital altyapılar; birlikte çalışabilir sistemler ve ölçeklenebilir bilgi işlem ağları aracılığıyla inovasyonu ve kapsayıcı sanayileşmeyi desteklemektedir.
SKH 12 – Sorumlu Tüketim ve Üretim	Topic_6 (potansiyel)	workflow, simulation, intelligence, learning	Süreç optimizasyonu ve kaynak verimliliği yoluyla sürdürülebilir üretim için potansiyel taşımaktadır, ancak literatürde açık bağlantı henüz kurulmamıştır.
SKH 13 – İklim Eylemi	Topic_1 (potansiyel)	semiconductor, protocol	Enerji verimli bilgi işlem ve donanım-yazılım ortak tasarımı yoluyla düşük karbonlu dijital altyapıların geliştirilmesine katkı potansiyeli taşımaktadır.
SKH 16 – Barış, Adalet ve Güçlü Kurumlar	Topic_14	cybersecurity, security, detection, response, operations	Gerçek zamanlı tehdit tespiti ve otonom müdahale sistemleriyle dijital güvenliği güçlendirerek demokratik süreçleri ve kurumsal dayanıklılığı korumaktadır.

ilerletmektedir. Topic_5 (“education, learning, educational, engineering, language, content, academic, engagement”), küresel olarak yüksek kaliteli eğitime erişimi genişletebilen kişiselleştirilmiş, uyarlanabilir öğrenme sistemleri aracılığıyla SKH 4’e (Kaliteli Eğitim) katkıda bulunmaktadır. Bu arada, topic_4 (“activation, agentic, market, agency, cooperation, humans, hybrid, journalists”) ve topic_3’ün ”supply/tedarik” ve ”customer/müşteri” sistemlerine odaklanması, medya, lojistik ve dijital pazarlarda çevik, insan-yapay zekâ iş birliği modellerini teşvik ederek, inovasyonu ve kapsayıcı sanayileşmeyi geliştirerek SKH 9’u desteklemektedir.

Topic_14 (“cybersecurity, security, detection, threats, agentic, response, threat, operations”), aracı yapay zekâyı, giderek dijitalleşen bir dünyada demokratik süreçleri ve ekonomik istikrarı korumak için gerekli olan gerçek zamanlı tehdit tespiti, otonom olay müdahalesi ve dayanıklı dijital altyapıyı mümkün kılarak SKH 16 (Barış, Adalet ve Güçlü Kurumlar) için kritik bir araç olarak konumlandırıyor. Topic_10 (“resource, digital,

network, architecture, software, computing, elsevier, systems”) ise kapsayıcı dijital dönüşümün temelini oluşturan ölçeklenebilir, ağ tabanlı bilgi işlem mimarileri üzerine araştırmalarla SKH 9’u daha da desteklemektedir.

Bu umut verici bağlantılara rağmen, dikkat çekici bir eksiklik devam ediyor: ”önyargı”, ”hesap verebilirlik”, ”adalet”, ”çevresel etki”, ”karbon ayak izi” gibi temel sürdürülebilirlik ve etik kavramlar veya SKH’ye açıkça atıfta bulunma gibi kavramlar, baskın araştırma konularında büyük ölçüde yer almamaktadır. ”Etik” terimini içeren topic_12 bile, işlevselleştirilebilir etik, algoritmik adalet veya gezegensel sınırlarla özlü bir şekilde ilgilenmemektedir. Bu kopukluk, önemli bir teori-pratik boşluğunu ortaya koymaktadır: Politika çerçeveleri ve üst düzey stratejiler sürdürülebilir kalkınma için sorumlu yapay zekâyı savunurken, deneysel ve teknik araştırma topluluğu, eşitlik, çevresel sürdürülebilirlik veya kapsayıcı erişim gibi SKH ile uyumlu kriterleri henüz sistematik olarak sistem tasarımı, değerlendirmesi veya dağıtımına entegre etmemektedir.

Ajan YZ’nin tüm toplumsal potansiyelini gerçekleştirmek, yalnızca teknik gelişmişlik değil, aynı zamanda tasarıma dayalı etik, sürdürülebilirlik bilincine sahip değerlendirme ölçütleri ve kapsayıcı ortak yaratım süreçleri aracılığıyla SKH ile bilinçli bir uyum da gerektirmektedir. Bu tür bir entegrasyon olmadan, ajan sistemler, eşitlik, dayanıklılık ve uzun vadeli insan ve gezegen refahı pahasına verimlilik veya performans için optimizasyon yapma riskiyle karşı karşıya kalmaktadır.

6. Sonuç ve Öneriler

Bu çalışma, Ajan YZ’nin yalnızca bir teknolojik ilerleme değil, aynı zamanda sağlık, eğitim, iklim eylemi ve ekonomik büyüme gibi kritik alanlarda SKH’lere somut katkı sunabilecek sosyo-teknik bir dönüşüm aracı olarak değerlendirilmesi gerektiğini savunmaktadır. Bulgular, özellikle SKH 3 kapsamında YZ destekli triyaj sistemlerinin kaynak kısıtlı bölgelerde erken teşhis ve müdahale imkânı yaratabileceğini, SKH 4 kapsamında uyarlanabilir öğrenme sistemlerinin bireysel öğrenme stillerine göre içerik sunarak eğitimde eşitliği artırma potansiyeline sahip olduğunu göstermektedir. Benzer şekilde, tedarik zinciri optimizasyonu ve enerji tüketiminin gerçek zamanlı izlenmesi yoluyla SKH 12 (Sorumlu Tüketim ve Üretim) ve SKH 13 (İklim Eylemi) hedeflerine doğrudan katkı sağlanabilmektedir. Ancak bu potansiyelin tam anlamıyla hayata geçirilebilmesi, teknik gelişmişliğin ötesinde etik, kurumsal ve epistemolojik bir dönüşümü zorunlu kılmaktadır. LDA analizi sonucunda ortaya çıkan en çarpıcı bulgu, ”karbon ayak izi”, ”algoritmik adalet” ve ”kapsayıcılık” gibi sürdürülebilirlik ve etik kavramlarının akademik söylemde marjinal konumda kalmasıdır. Bu durum, Ajan YZ’nin geliştirilmesinde ”performans” ve ”verimlilik” paradigmasının, ”adalet” ve ”sürdürülebilirlik” ilkelerine kıyasla aşırı ölçüde önceliklendirildiğine işaret etmektedir. Söz konusu kopukluk, Floridi’nin (2013) ”ahlaki bozulma” (moral disruption) olarak tanımladığı bir durumu yansıtmaktadır: teknolojik ilerleme, toplumsal değerler ve etik normlarla uyumsuz bir şekilde ilerlediğinde, zararlı veya adaletsiz sonuçlar doğurabilir. Bu durum, teknolojinin asla bir amaç değil; aksine, insanlığın refahı ve gezegenin ekolojik bütünlüğü için bir araç olması gerektiği

gerçeğini bir kez daha gözler önüne sermektedir.

Bu riskleri azaltmak ve Ajan YZ'nin toplumsal faydasını maksimize etmek adına üç temel öneri öne çıkarılmaktadır. İlk olarak, tasarımdan itibaren etik ve sürdürülebilirlik entegrasyonu ilkesi benimsenmelidir: Ajan sistemlerin mimari tasarım aşamasından itibaren çevresel etki analizleri, adalet odaklı değerlendirme metrikleri ve açıklanabilirlik protokolleri sistematik olarak entegre edilmelidir; bu süreci AB Yapay Zekâ Yasası ve ISO/IEC 24027 gibi düzenleyici çerçeveler rehberlik edebilir. İkinci olarak, özellikle sağlık, adalet ve güvenlik gibi yüksek riskli alanlarda insan-döngüsü-içi denetim mekanizmaları (human-in-the-loop) zorunlu kılınmalıdır. Bu yaklaşım yalnızca teknik hataların düzeltilmesini değil, toplumsal değerlerin karar mekanizmalarına aktif olarak yerleştirilmesini de mümkün kılmaktadır. Üçüncü olarak, Ajan YZ'nin sürdürülebilirlik potansiyelinin tam anlamıyla değerlendirilebilmesi için disiplinlerarası araştırma ve ortak yaratım teşvik edilmelidir; bilgisayar mühendisliği, sosyal bilimler, çevre bilimleri ve etik alanlarındaki uzmanların iş birliği, teknik çözümlerin toplumsal ve çevresel bağlamlarla uyumlu olmasını sağlayacaktır. Gelecekteki araştırmaların özellikle hibrit zekâ mimarileri (üretken, duygusal ve ahlaki muhakemeyi birleştiren sistemler), kültürlerarası etik değerlendirmeler ve gezegensel sınırlar içinde YZ (enerji tüketimi, e-atık, su kullanımı gibi çevresel etkilerin niceliksel analizi) gibi alanlarda yoğunlaşması, Ajan YZ'nin yalnızca "daha akıllı" değil, aynı zamanda "daha adil ve sürdürülebilir" bir teknoloji haline gelmesini mümkün kılacaktır.

CRediT Yazar Katkı Beyanı

Kavramsallaştırma: Hüseyin Parmaksız, **Metodoloji:** Hüseyin Parmaksız, **Yazılım:** Hüseyin Parmaksız, **Araştırma:** Hüseyin Parmaksız, **Veri Düzenleme:** Hüseyin Parmaksız, **Yazım - İlk Taslak:** Hüseyin Parmaksız, **Yazım - İnceleme & Düzenleme:** Hüseyin Parmaksız, **Görselleştirme:** Hüseyin Parmaksız, **Denetim:** Hüseyin Parmaksız.

Çıkar Çatışması Beyanı

Yazarlar, herhangi bir çıkar çatışması olmadığını beyan etmektedir.

Referanslar

- Almulhim, A. I., ve Yigitcanlar, T. (2025). Understanding Smart Governance of Sustainable Cities: A Review and Multidimensional Framework. *Smart Cities*, 8(4), 113.
- Angubasu, J. (2022). The Implementation of Ai Self-triage Systems as a Digital Health Solution for Primary Healthcare in Kenya: Challenges and Prospects (Doctoral dissertation, University of Nairobi).
- Bandura, A. (2001). Social Cognitive Theory: An Agentic Perspective. *Annual Review of Psychology*, 52(1), 1-26.
- Bandura, A. (2008). Toward an Agentic Theory of The Self. *Advances in Self Research*, 3, 15-49.

- Barra, F. L., Rodella, G., Costa, A., Scalogna, A., Carenzo, L., Monzani, A., ve Corte, F. D. (2025). From Prompt to Platform: An Agentic AI Workflow for Healthcare Simulation Scenario Design. *Advances in Simulation*, 10(1), 29.
- Berthold, M., Cebon, N., Dill, F., Gabriel, T., Kötter, T., Meinl, T., ... ve Wiswedel, B. (2008). Data Analysis, Machine Learning and Applications SE-38, *Studies in Classification, Data Analysis, and Knowledge Organization*.
- Blei, D. M., Ng, A. Y., ve Jordan, M. I. (2003). Latent Dirichlet Allocation. *Journal of Machine Learning Research*, 3(Jan), 993-1022.
- Braun, V., ve Clarke, V. (2019). Psikolojide Tematik Analizin Kullanımı. *Journal of Qualitative Research in Education*, 7(2).
- Cantero Gamito, M. ve Marsden, C. T. (2024). Artificial Intelligence Co-regulation? The Role of Standards in The EU AI Act. *International Journal of Law and Information Technology*, 32, eaae011.
- Cowls, J., Tsamados, A., Taddeo, M. ve Floridi, L. (2023). The AI Gambit: Leveraging Artificial Intelligence to Combat Climate Change-Opportunities, Challenges, and Recommendations. *Ai & Society*, 38(1), 283-307.
- Creswell, J. W. ve Plano Clark, V. L. (2023). Revisiting Mixed Methods Research Designs Twenty Years Later. *Handbook of Mixed Methods Research Designs*, 1(1), 21-36.
- Delaney, Y., McCarthy, J. ve Beecham, S. (2017, June). Convergent Parallel Design Mixed Methods Case Study in Problem-Based Learning. In ECRM 2017 16th European Conference on Research Methods in Business and Management (p. 408). Academic Conferences and publishing limited.
- Delen, D. (2024). Landscape of Tools for Business Analytics and Data Science—A Tutorial on KNIME.
- Ergani, G. Ç. (2024). Sürdürülebilirlik Raporlaması İçin Standartlara Uygun Veri Toplama ve Analitik Yaklaşımlar ile Karar Verme: Simülasyon Tabanlı bir Uygulama (Master's thesis, Marmara Üniversitesi (Turkey)).
- Floratos, A., Smith, K., Ji, Z., Watkinson, J. ve Califano, A. (2010). geWorkbench: An Open Source Platform for Integrative Genomics. *Bioinformatics*, 26(14), 1779-1780.
- Floridi, L. (2013). The ethics of Information. Oxford University Press (UK).
- Floridi, L. ve Cowls, J. (2022). A Unified Framework of Five Principles for AI in Society. *Machine Learning and The City: Applications in Architecture and Urban Design*, 535-545.
- Goecks, J., Nekrutenko, A., Taylor, J. ve Galaxy Team Team@ Galaxyproject. org. (2010). Galaxy: a Comprehensive Approach for Supporting Accessible, Reproducible, and Transparent Computational Research in The Life Sciences. *Genome Biology*, 11(8), R86.
- Griggs, D., Stafford-Smith, M., Gaffney, O., Rockström, J., Öhman, M. C., Shyamsundar, P., ... ve Noble, I. (2013). Sustainable Development Goals for People and Planet. *Nature*, 495(7441), 305-307.

- Grimmer, J. ve 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.
- Guidance, W. H. O. (2021). Ethics and Governance of Artificial Intelligence for Health. World Health Organization.
- Hatch, S. G., Goodman, Z. T., Vowels, L., Hatch, H. D., Brown, A. L., Guttman, S., ... ve Braithwaite, S. R. (2025). When ELIZA Meets Therapists: A Turing Test for The Heart and Mind. *PLOS Mental Health*, 2(2), e0000145.
- Hull, D., Wolstencroft, K., Stevens, R., Goble, C., Pocock, M. R., Li, P. ve Oinn, T. (2006). Taverna: A Tool for Building and Running Workflows of Services. *Nucleic Acids Research*, 34(suppl_2), W729-W732.
- Jagla, B., Wiswedel, B. ve Coppée, J. Y. (2011). Extending KNIME for Next-Generation Sequencing Data Analysis. *Bioinformatics*, 27(20), 2907-2909.
- Jeong, C. (2025). Beyond Text: Implementing Multimodal Large Language Model-Powered Multi-Agent Systems Using a No-Code Platform. arXiv preprint arXiv:2501.00750.
- Komtaş. (2025). Agentic Yapay zekâ (Agentic AI) Nedir?, Erişim Tarihi:28.10.2025, Erişim adresi: <https://www.komtas.com/glossary/agentic-yapay-zekâ-nedir>.
- Krippendorff, K. (2018). Content Analysis: An Introduction to Its Methodology. Sage Publications.
- Leal-Arcas, R. (2025). The Future of Global Economic Governance: Balancing Trade, Sustainability, and Social Justice. *Sustainability, and Social Justice*, (February 16, 2025).
- Lee, B. X., Kjaerulf, F., Turner, S., Cohen, L., Donnelly, P. D., Muggah, R., ... ve Gilligan, J. (2016). Transforming Our World: Implementing The 2030 Agenda Through Sustainable Development Goal Indicators. *Journal of Public Health Policy*, 37(Suppl 1), 13-31.
- Linke, B. (2012). Conveyor-a Workflow Engine for Bioinformatics Analyses.
- Ludäscher, B., Altintas, I., Berkley, C., Higgins, D., Jaeger, E., Jones, M., ... ve Zhao, Y. (2006). Scientific Workflow Management and The Kepler System. *Concurrency and Computation: Practice and Experience*, 18(10), 1039-1065.
- Mahela, O. P., Khosravy, M., Gupta, N., Khan, B., Alhelou, H. H., Mahla, R., ... ve Siano, P. (2020). Comprehensive Overview of Multi-Agent Systems for Controlling Smart Grids. *CSEE Journal of Power and Energy Systems*, 8(1), 115-131.
- Nakicenovic, N., Messner, D., Zimm, C., Clarke, G., Rockström, J., Aguiar, A. P., ... ve Yillia, P. (2019). TWI2050-The World in 2050 (2019). The Digital Revolution and Sustainable Development: Opportunities and Challenges. Report prepared by The World in 2050 initiative.
- Nelson, L. K. (2020). Computational Grounded Theory: A Methodological Framework. *Sociological Methods & Research*, 49(1), 3-42.
- Néron, B., Ménager, H., Maufrais, C., Joly, N., Maupetit, J., Letort, S., ... ve Letondal, C. (2009). Mobyle: A New Full Web Bioinformatics Framework. *Bioinformatics*, 25(22),

3005-3011.

- Newman, D., Asuncion, A., Smyth, P. ve Welling, M. (2009). Distributed Algorithms for Topic Models. *Journal of Machine Learning Research*, 10(8).
- n8n.io. (2025). Flexible AI Workflow Automation for Technical Teams. <https://n8n.io/>, Erişim tarihi: 10.10.2025
- Okuy, E. K. (2020). Understanding The Role of The National Human Rights Institutions in Implementing The Sustainable Development Agenda: The Case of Europe (Master's thesis, Middle East Technical University (Turkey)).
- Parmaksız, H. ve Akarsu, O. (2025). Diagnosing Core Topics in Digital Transformation Studies via Topic Model Approach. *Selçuk Üniversitesi Sosyal Bilimler Enstitüsü Dergisi*, (57), 262-282.
- Roll, I. ve Wylie, R. (2016). Evolution and Revolution in Artificial Intelligence in Education. *International Journal of Artificial Intelligence in Education*, 26(2), 582-599.
- Rolnick, D., Donti, P. L., Kaack, L. H., Kochanski, K., Lacoste, A., Sankaran, K., ... ve Bengio, Y. (2022). Tackling Climate Change with Machine Learning. *ACM Computing Surveys (CSUR)*, 55(2), 1-96.
- Rotar, C. ve Zhang, Q. (2025). A Design Science Research Approach to Large Language Model-Based Agents for Requirements Specification (LLMBA4RS) in Low-Code Applications. *Requirements Engineering*, 1-24.
- Russell, S. J. ve Norvig, P. (2021). Artificial Intelligence: A Modern Approach, Global Edition 4e.
- Sachs, J. D. (2015). The age of sustainable development. Columbia University Press.
- Sachs, J. D., Schmidt-Traub, G., Mazzucato, M., Messner, D., Nakicenovic, N. ve Rockström, J. (2019). Six Transformations to Achieve The Sustainable Development Goals. *Nature Sustainability*, 2(9), 805-814.
- Schaeffer, J. ve Plaat, A. (1997). Kasparov versus deep blue: The rematch.
- Schneider, J. (2025). Generative to Agentic Ai: Survey, Conceptualization, and Challenges. arXiv preprint arXiv:2504.18875.
- Searle, J. (1999). The chinese room.
- Shinn, N., Cassano, F., Gopinath, A., Narasimhan, K. ve Yao, S. (2023). Reflexion: Language Agents with Verbal Reinforcement Learning. *Advances in Neural Information Processing Systems*, 36, 8634-8652.
- Shneiderman, B. (2022). Human-centered AI. Oxford University Press.
- Silah, E. ve Eğilmez, Ö. (2025). Eğitimde Sürdürülebilir Kalkınma Hedefleri Üzerine Bir Değerlendirme. *Sürdürülebilirlik, Yönetim & Ekonomi Dergisi*, 1(1), 58-90.
- Sim AI Agent, (2025). Sim: Build and Deploy AI Agent Workflows in Minutes. GitHub. <https://github.com/simstudioai/sim>, Erişim tarihi: 10.10.2025
- Turan, T., Turan, G. ve Kırbaş, İ. Yapay zekânın Çevresel Ayak İzi: Enerji, Su ve Elektronik Atık Üzerine Çok Yönlü Bir Analiz. *Uluslararası Mühendislik Tasarım ve Teknoloji Dergisi*,

- 7(2), 78-89.
- Turing, A. M. (1950). *Mind*. *Mind*, 59(236), 433-460.
- Van Niekerk, A. J. (2020). Inclusive Economic Sustainability: SDGs and Global Inequality. *Sustainability*, 12(13), 5427.
- Vinuesa, R., Azizpour, H., Leite, I., Balaam, M., Dignum, V., Domisch, S., ... ve Fuso Nerini, F. (2020). The Role of Artificial Intelligence in Achieving The Sustainable Development Goals. *Nature communications*, 11(1), 233.
- Viswanadhapalli, V. (2025). The Future of Intelligent Automation: How Low-Code/No-Code Platforms are Transforming AI Decisioning. *International Journal Of Engineering And Computer Science*, 14(1), 26803-26825.
- Wooldridge, M. (1999). Intelligent Agents. *Multiagent Systems: A Modern Approach to Distributed Artificial Intelligence*, 1, 27-73.
- Yao, L., Mimno, D. ve McCallum, A. (2009, June). Efficient Methods for Topic Model Inference on Streaming Document Collections. In *Proceedings of the 15th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 937-946).
- Yu, C., Cheng, Z., Cui, H., Gao, Y., Luo, Z., Wang, Y., ... ve Zhao, Y. (2025, May). A Survey on Agent Workflow–Status and Future. In *2025 8th International Conference on Artificial Intelligence and Big Data (ICAIBD)* (pp. 770-781). IEEE.

Agentic artificial intelligence: A multidimensional framework for autonomous systems, sustainable impact, and scholarly evolution

Extended Abstract

Introduction/Background

Agentic AI refers to the evolution of artificial intelligence from its role as a passive analysis tool to proactive, planning, and contextually aware agents. This transformation is not only critical for technological advancement but also for aligning with global sustainability and ethical imperatives. However, the current literature largely focuses on technical performance and sectoral applications; dimensions, such as sustainability, equity, and ethical governance are not adequately considered. This study addresses four key gaps in the literature by examining the theoretical foundations, operational structures, and relationship with the United Nations Sustainable Development Goals (SDGs) of Agent AI within a multi-dimensional framework: (1) the lack of a unified agent typology, (2) limited interest in low-code orchestration platforms, (3) weak theoretical and empirical connection with the SDGs, and (4) the absence of thematic mapping of the Agent AI literature.

Literature Review

The literature review highlights the key differences between Agent AI and Generative AI and traditional automation: Agent systems determine "how" to act rather than "what" to produce; they are contextually flexible rather than bound by fixed rules; and they offer a collaborative model that empowers rather than excludes humans. Low-code platforms (like Sim and n8n) provide critical infrastructure for the widespread adoption of these systems. However, the impacts of Agent AI on energy consumption, carbon footprint, algorithmic fairness, or digital inequality have not yet been systematically investigated. While some studies, such as Vinuesa et al. (2020), have mapped the potential contributions of AI to the SDGs, these analyses have not been adapted to the autonomous decision-making dynamics of agent systems.

Methodology

The study was conducted on a dataset of 218 academic publications indexed in the Scopus database between 2023 and 2025 using the keywords "agentic ai" or "agentic artificial intelligence." In the analysis, a hybrid content analysis approach was adopted by combining thematic analysis (Braun ve Clarke, 2019) with unsupervised topic modeling (Parallel LDA). Data preprocessing and analysis processes were performed using the KNIME Analytics Platform (v5.5.0). Abstract texts were subjected to a seven-stage preprocessing process including tokenization, lemmatization, stop word filtering, and lowercasing; then, thematic clusters were extracted by applying Parallel LDA for $k = 15$.

Key Findings

The analysis shows that Agent AI establishes strong connections with SDGs:

SDG 3 (Good Health and Well-being): Topic_8 and Topic_13 emphasize patient-centered care and clinical decision support.

SDG 4 (Quality Education): Topic_5 supports inclusivity in education through adaptive learning systems.

SDG 8 and 9 (Work, Innovation, Infrastructure): Topic_3, Topic_4, and Topic_11 highlight human-agent collaboration and smart infrastructure.

SDG 12 and 13 (Responsible Consumption, Climate Action): Terms like "drilling" and "semiconductor" in Topic_6 and Topic_1 refer to sustainable production and energy efficiency.

SDG 16 (Peace, Justice, and Strong Institutions): Topic_14 contributes to this goal through cybersecurity and autonomous threat response.

However, topics including fundamental sustainability and ethical concepts such as "carbon footprint," "algorithmic justice," and "accountability" have remained marginal in academic discourse.

Discussion/Analysis

While the findings support the theoretical alignment of Agent AI's technical sophistication with SDGs, they reveal a significant theory-practice gap in realizing this potential. The performance and efficiency paradigm is excessively prioritized over principles of justice and sustainability. This situation can be explained by Floridi's concept of "moral disruption": when technological progress advances in a way that is incompatible with social values, it can lead to harmful consequences. To mitigate these risks, it is essential to integrate ethical, environmental, and social dimensions into the system design.

Conclusion

The future of AI will be determined not only by how fast or accurately algorithms work, but also by how humanely, fairly, and respectfully they are designed within planetary boundaries. The study makes three key recommendations for realizing this vision: (1) Integrating ethics and sustainability from the design stage, (2) Implementing human-in-the-loop oversight mechanisms in high-risk areas, and (3) Fostering interdisciplinary research and co-creation involving computer engineering, social sciences, and environmental sciences. This approach will reveal the potential of Agent AI to build systems that are not only "smarter" but also "fairer and more sustainable."