

Üretken yapay zekâ destekli karar verme sistemleri: PRISMA 2020 akışına dayalı sistematik literatür incelemesi ve tematik sentez¹

Generative artificial intelligence-supported decision-making systems: A PRISMA 2020-based systematic literature review and thematic synthesis

¹ Bu makale, 23-24 Mayıs 2026 tarihlerinde düzenlenen 9. Uluslararası Uludağ Bilimsel Araştırmalar Kongresi'nde "Üretken Yapay Zekâ Destekli Karar Verme Sistemleri: PRISMA Tabanlı Sistematik Literatür İncelemesi" başlığıyla sunulan bildirinin genişletilmiş ve yeniden düzenlenmiş hâlidir.

Fatmanur Yavuz² 

Zeliha Seçkin³ 

² Arş. Gör., Aksaray Üniversitesi, Aksaray, Türkiye, fatmanur.yuksel@aksaray.edu.tr

ORCID: 0000-0001-9720-2220

³ Prof. Dr., Aksaray Üniversitesi, Aksaray, Türkiye, seckinz@aksaray.edu.tr

ORCID: 0000-0003-0603-3236

Sorumlu Yazar/Corresponding Author:

Fatmanur Yavuz,

Aksaray Üniversitesi, Aksaray, Türkiye,

fatmanur.yuksel@aksaray.edu.tr

Öz

Bu çalışma, üretken yapay zekâ ve büyük dil modeli destekli karar verme sistemlerine ilişkin literatürü; araştırma eğilimleri, uygulama alanları, yöntemsel özellikler, riskler ve yönetim yaklaşımları bakımından sistematik biçimde incelemektedir. Web of Science ve Scopus'ta 13 Mayıs 2026'da gerçekleştirilen, 2020-2026 dönemini kapsayan tarama PRISMA 2020'ye göre raporlanmıştır. 402 kayıttan tekrarlar çıkarıldıktan ve değerlendirmeler tamamlandıktan sonra 31 çalışma senteze alınmıştır. Çalışmalar, araştırma tasarımlarına uygun MMAT 2018 veya JBI kontrol listeleriyle değerlendirilmiş; 249 maddede 171 "Evet", 33 "Hayır" ve 45 "Belirsiz" kararı kaydedilmiş, toplam puan veya kalite kategorisi üretilmemiştir. Tematik sentez; karar destek kapasitesi, güven ve kullanıcı kabulü, açıklanabilirlik ve belirsizlik yönetimi, insan gözetimi ve yönetim, riskler ve sorumlu kullanım ile karar mimarisi ve çok kriterli karar verme olmak üzere altı tema ortaya koymuştur. Bulgular doğrultusunda altı bileşenli Üretken Yapay Zekâ Destekli Karar Vermeye İlişkin Bütünleştirici Kavramsal Çerçeve önerilmiştir. Sonuçlar, üretken yapay zekânın insan karar vericinin yetki ve sorumluluğunu devralan otonom bir aktör olarak konumlandırılmasından kaçınılması; açıklanabilir, denetlenebilir, insan gözetimli ve kurumsal yönetime bağlı bir karar destek bileşeni olarak kullanılması gerektiğini göstermektedir.

Anahtar Kelimeler: Üretken Yapay Zekâ, Büyük Dil Modelleri, Karar Verme, Karar Destek Sistemleri, Açıklanabilirlik

JEL Kodları: M15, O33, C80

Abstract

This study systematically reviews the literature on generative artificial intelligence (GenAI)- and large language model-supported decision-making systems, focusing on research trends, application domains, methodological characteristics, risks, and governance approaches. A Web of Science and Scopus search conducted on May 13, 2026, covering 2020-2026, was reported in accordance with PRISMA 2020. After duplicate removal and eligibility assessment, 31 of 402 records were included in the synthesis. Studies were appraised using MMAT 2018 or design-appropriate JBI checklists; across 249 items, 171 "Yes," 33 "No," and 45 "Unclear" judgments were recorded, without calculating an overall score or quality category. The thematic synthesis identified six themes: decision-support capacity; trust and user acceptance; explainability and uncertainty management; human oversight and governance; risks and responsible use; and decision architecture and multi-criteria decision-making. Based on these findings, a six-component Integrative Conceptual Framework for Generative AI-Supported Decision-Making was proposed. The findings indicate that GenAI should function as an explainable, auditable, human-supervised, and institutionally governed decision-support component, while human decision-makers retain authority and responsibility.

Keywords: Generative Artificial Intelligence, Large Language Models, Decision-Making, Decision Support Systems, Explainability

JEL Codes: M15, O33, C80

Atıf/Citation: Yavuz, F., & Seçkin, Z., Üretken yapay zekâ destekli karar verme sistemleri: PRISMA 2020 akışına dayalı sistematik literatür incelemesi ve tematik sentez, bmij (2026) 14 (3): 1241-1264, doi: <https://doi.org/10.15295/bmij.v14i3.2819>

Extended Abstract

Generative artificial intelligence-supported decision-making systems: A PRISMA 2020-based systematic literature review and thematic synthesis

Literature

Research subject

This study focuses on generative artificial intelligence and large language model-supported decision-making systems. Recent studies show that generative AI is used for data analysis, information synthesis, recommendation generation, scenario development, and risk assessment. Albashrawi (2025) discusses generative AI as a multidisciplinary decision-making technology, Song et al. (2026) emphasize large language models in supply chain decision support, and Saup et al. (2026) examine their integration into strategic decision systems from a socio-technical and governance perspective.

Research purpose and importance

The study systematically examines the literature on generative AI- and large language model-supported decision-making systems to identify research trends, application areas, risks, and governance approaches. This topic is important because generative AI in decision-making cannot be evaluated through technical performance alone. Trust, explainability, uncertainty management, human oversight, accountability, data privacy, and responsible use also shape adoption (Huynh, 2024; Kumar et al., 2025; Sheng et al., 2025; Tahvildari, 2025).

Contribution of the article to the literature

This article contributes to the literature by providing a PRISMA 2020-based systematic review and thematic synthesis of generative AI-supported decision-making studies. It classifies the literature under six themes: decision-support capacity, trust and user acceptance, explainability and uncertainty management, human oversight and governance, risks and responsible use, and decision architecture and multi-criteria decision-making. The study positions generative AI as an explainable, auditable, and human-supervised cognitive decision-support component within institutional governance rather than an autonomous decision-maker.

Design and method

Research type

This study is a systematic literature review and thematic synthesis. It does not involve primary data collected from participants; instead, it systematically reviews and synthesizes previous studies in the field.

Research problems

The study addresses four research questions. First, it examines in which decision domains and decision functions generative AI/LLM-supported decision-making systems are discussed. Second, it investigates how trust, user acceptance, explainability, and uncertainty management are addressed in these systems. Third, it examines how human oversight, human-AI collaboration, accountability, and governance are integrated into decision-making processes. Fourth, it identifies the main risks, limitations, and future research gaps in the literature.

Data collection method

Data were collected through searches of Web of Science and Scopus conducted on 13 May 2026. The search covered the 2020–2026 period and used three conceptual blocks: generative artificial intelligence and large language models; decision-making and decision support; and trust, explainability, human oversight, governance, accountability, and decision quality. The search yielded 402 records. After 130 duplicates were removed, 272 records were screened. The researchers examined inclusion and exclusion decisions jointly and resolved all decisions by consensus using predefined criteria and a shared screening file. Forty-two reports were sought for full-text retrieval; 31 accessible reports met the eligibility criteria and were included. Because the decisions were generated as a single consensus dataset rather than as two independent rating sets, an inter-rater kappa coefficient was not technically applicable.

Quantitative/qualitative analysis

The study mainly employed qualitative thematic synthesis, supported by descriptive counts. Each included study was coded for publication year, application field, research design, GenAI/LLM approach, decision context, trust, explainability, human oversight, governance, and risks. Methodological adequacy was assessed retrospectively using one design-appropriate appraisal approach for each study. The empirical components of 23 empirical or technical studies were appraised using the relevant MMAT 2018 design criteria. Three systematic literature reviews and one systematic scoping review were appraised using the JBI Checklist for Systematic Reviews and Research Syntheses; four conceptual or textual studies were appraised using the JBI Checklist for Textual Evidence: Expert Opinion. Item-level judgments were recorded as Yes, No, or Unclear. Mixed-methods studies were assessed across the applicable qualitative, quantitative, and integration criteria. No percentage, summative quality score, or overall high/moderate/low category was calculated, and appraisal was not used as an automatic exclusion threshold.

Research model

This study does not test an empirical research model. Rather, the thematic synthesis yields an Integrative Conceptual Framework for Generative AI-Supported Decision-Making. The framework comprises six interrelated components: decision context and risk mapping, generative AI decision-support functions, assurance mechanisms, human authority and decision control, institutional governance and accountability, and monitoring, learning, and feedback. These components do not constitute a fixed linear sequence; instead, they operate interactively and iteratively throughout the decision-system life cycle.

Research hypotheses

This study does not include research hypotheses because it is a systematic literature review rather than a hypothesis-testing empirical study. Instead, the study is structured around research questions designed to guide the review and thematic synthesis.

Findings and discussion

Findings as a result of analysis

The findings show that generative AI-supported decision-making studies are concentrated in finance, supply chain management, business intelligence, public administration, human resources, industrial processes, strategic decision-making, and organizational decision systems. The appraisal comprised 249 item-level judgments across 31 studies: 171 Yes, 33 No, and 45 Unclear. Fifteen studies had at least one No judgment, 24 had at least one Unclear judgment, and three received only Yes judgments. In the MMAT-appraised studies, research purposes were consistently clear and the data were generally aligned with those purposes; recurring limitations concerned sample or technical test-unit representativeness and control of confounding factors. In the JBI-appraised reviews, critical appraisal, independent duplicate appraisal, and publication-bias assessment were commonly absent or insufficiently reported. These results are descriptive indicators of criterion-level strengths and limitations, not overall quality ratings.

The thematic synthesis revealed six major themes: decision-support capacity; trust and user acceptance; explainability and uncertainty management; human oversight and governance; risks and responsible use; and decision architecture and multi-criteria decision-making. These themes were not treated as isolated categories. Decision-support capacity represents the technological affordance layer; explainability and uncertainty management constitute the assurance layer; trust and acceptance reflect the user-response layer; human oversight preserves decision authority; governance institutionalizes accountability; and decision architectures operationalize these relationships. Cross-cutting risks such as hallucination, bias, privacy violations, security failures, and overreliance affect every layer.

Hypothesis test results

No hypothesis testing was conducted because the study was designed as a systematic literature review and thematic synthesis. Accordingly, no hypothesis-test results are reported.

Discussing the findings with the literature

The findings are consistent with research emphasizing both the decision-support potential and the governance risks of generative AI. Their theoretical contribution lies in showing that technical capability alone does not translate into decision quality. This translation is conditional on assurance, calibrated trust, preserved human authority, and institutional governance. The proposed framework also connects the reviewed evidence with the OECD AI Principles, the NIST AI Risk Management Framework, the EU AI Act, and ISO/IEC 42001, thereby translating thematic findings into a multilevel governance architecture.

Conclusion, recommendation and limitations

Results of the article

This study concludes that generative AI-supported decision-making systems should function as explainable, auditable, and human-supervised components embedded in institutional governance. Its main contribution is the Integrative Conceptual Framework for Generative AI-Supported Decision-Making, which links decision-support functions with context and risk mapping, assurance mechanisms, human authority and control, governance and accountability, and monitoring, learning, and feedback. The framework is a conceptual synthesis derived from the review rather than an empirically validated model.

Suggestions based on results

Future studies should examine how trust develops and weakens, test the effects of explainability and uncertainty management on decision quality across sectors, and compare human-in-the-loop and human-in-command oversight models. Interdisciplinary research should also address the ethical, legal, and governance dimensions of these systems.

Limitations of the article

The review was limited to Web of Science and Scopus, English-language records, and the 2020–2026 period; other relevant sources may therefore have been omitted. Eleven reports could not be retrieved. The protocol was not prospectively registered, and retrospective documentation cannot eliminate post hoc decisions. Screening, coding, and appraisal were conducted through joint consensus rather than independent parallel ratings; therefore, inter-rater kappa was not calculated. Record-level reasons for 50 exclusions at the second scope-assessment stage were not retained, limiting the audit trail. These limitations were partly mitigated through database-specific search strings, a dated search log, explicit eligibility rules, recorded title–abstract exclusion reasons, stage-level counts, a shared coding framework, and full-text, item-level appraisal of all included studies.

Giriş

Üretken yapay zekâ (generative artificial intelligence) ve büyük dil modelleri (large language models – LLMs), son yıllarda bilgi sistemleri, yönetim bilişim sistemleri, karar destek sistemleri ve örgütsel karar verme literatüründe hızla görünür hâle gelen teknolojilerdir. Bu teknolojiler başlangıçta metin üretimi, içerik oluşturma ve doğal dil etkileşimi gibi işlevlerle öne çıkmış olsa da güncel literatür, üretken yapay zekânın karar verme süreçlerinde veri analizi, bilgi sentezi, senaryo üretimi, risk değerlendirme, öneri geliştirme ve karar destek sağlama gibi daha karmaşık roller üstlenmeye başladığını göstermektedir (Albashrawi, 2025; Gupta ve Kumar, 2026; Hao vd., 2024). Büyük dil modellerinin yapılandırılmamış verileri işleyebilmesi, doğal dil üzerinden etkileşim kurabilmesi ve farklı bilgi kaynaklarını bağlamsal biçimde yorumlayabilmesi, bu sistemleri örgütsel ve yönetsel karar verme süreçleri açısından önemli kılmaktadır (Agarwal vd., 2026; Saup vd., 2026; Song vd., 2026).

Karar verme süreçleri yalnızca bilgiye erişim veya veri işleme kapasitesiyle açıklanamaz. Özellikle örgütsel ve yönetsel bağlamlarda karar vericilerin alternatifleri değerlendirmesi, riskleri öngörmesi, belirsizliği yönetmesi, paydaş beklentilerini dikkate alması ve kararın etik, hukuki ve stratejik sonuçlarını gözetmesi gerekir. Üretken yapay zekâ bu süreçlerde karar vericilerin bilişsel yükünü azaltabilir, karmaşık bilgileri özetleyebilir, alternatif senaryolar geliştirebilir ve karar seçeneklerini gerekçelendirebilir (Agarwal vd., 2026; Hao vd., 2024; Kim vd., 2026; Zong vd., 2026). Bu teknolojilerin karar verme süreçlerine entegrasyonu yeni risk alanlarını da gündeme getirmektedir. Literatürde; güven, açıklanabilirlik, insan gözetimi, hesap verebilirlik, veri gizliliği, algoritmik yanlılık, halüsinasyon, aşırı güvenme ve yönetim sorunları üretken yapay zekâ destekli karar verme sistemlerinin temel tartışma alanları olarak ele alınmaktadır (Hinov ve Ivanova, 2026; Kumar vd., 2025; Sheng vd., 2025; Tahvildari, 2025). Bu noktada temel soru, üretken yapay zekânın karar verme süreçlerinde nasıl konumlandırılması gerektiğidir. Bazı uygulamalar LLM'leri bağımsız öneri üreticiler veya otonom karar aktörleri gibi sunarken güncel çalışmalar daha temkinli bir yaklaşıma işaret etmektedir. Buna göre üretken yapay zekâ, insan karar vericinin yerini alan bağımsız bir sistem olarak değil; insan yargısını destekleyen, açıklanabilirlik sağlayan, belirsizliği görünür kılan ve yönetim mekanizmalarıyla denetlenmesi gereken bilişsel bir karar destek bileşeni olarak değerlendirilmelidir (Albashrawi, 2025; Hinov ve Ivanova, 2026; Saup vd., 2026; Zhu vd., 2026).

Bu çalışmanın amacı, üretken yapay zekâ ve büyük dil modeli destekli karar verme sistemlerine ilişkin güncel literatürü sistematik biçimde incelemek; alandaki temel araştırma eğilimlerini, uygulama alanlarını, riskleri ve yönetim odaklı yaklaşımları ortaya koymaktır. Bu amaç doğrultusunda Web of Science ve Scopus veri tabanlarında sistematik tarama yapılmış, elde edilen kayıtlar PRISMA 2020 akışına uygun biçimde değerlendirilmiş ve nihai olarak 31 çalışma tematik senteze dâhil edilmiştir. Çalışma yalnızca literatürü altı tema altında sınıflandırmakla kalmamakta; dâhil edilen çalışmaların yöntemsel yeterliğini araştırma tasarımlarına uygun MMAT 2018 veya JBI kontrol listesiyle değerlendirmekte ve temalar arasındaki ilişkilerden hareketle Üretken Yapay Zekâ Destekli Karar Vermeye İlişkin Bütünleştirici Kavramsal Çerçeve'yi sunmaktadır. Bu yönüyle çalışma, karar destek sistemleri ve yapay zekâ yönetimi literatürleri arasında açıklanabilirlik, insan otoritesi ve kurumsal hesap verebilirlik üzerinden bir köprü kurmayı hedeflemektedir.

Üretken yapay zekâ ve örgütsel karar verme kesişimindeki güncel sistematik derlemeler, yalnızca kullanım alanlarını listelemek yerine GenAI görevlerini karar sürecinin aşamalarına yerleştiren süreçsel modeller geliştirmeye yönelmektedir. Örneğin Schulte ve Kanbach (2026), GenAI'nin örgütsel karar vermedeki görev ve rollerini dikkat, zekâ, tasarım, seçim, uygulama ve geri bildirim bileşenleri üzerinde haritalamıştır. Bu çalışma ise söz konusu süreçsel bakışı güven, açıklanabilirlik, insan gözetimi ve kurumsal yönetim mekanizmalarıyla genişletmektedir.

Kavramsal arka plan

Üretken yapay zekâ ve büyük dil modelleri

Üretken yapay zekâ mevcut veri örüntülerinden hareketle yeni içerik, metin, görsel, kod, açıklama, öneri veya senaryo üretebilen yapay zekâ sistemlerini ifade eder. Büyük dil modelleri ise çok büyük ölçekli metin verileri üzerinde eğitilen ve doğal dil anlama, metin üretme, özetleme, sınıflandırma, soru yanıtlama, açıklama oluşturma ve akıl yürütme benzeri görevlerde kullanılan üretken yapay zekâ sistemlerinin önemli bir alt kümesini oluşturmaktadır (Albashrawi, 2025; Kaleta, 2025; Song vd., 2026). Albashrawi (2025), üretken yapay zekânın sağlık, hukuk, işletme, eğitim ve turizm gibi farklı alanlarda karar verme doğruluğu, verimlilik ve kişiselleştirme açısından dönüştürücü etkiler sunduğunu; ancak yanlılık, yanlış bilgilendirme ve şeffaflık eksikliği gibi sorunların dikkatle yönetilmesi gerektiğini belirtmektedir. Benzer biçimde Song vd. (2026), LLM'lerin yapılandırılmamış ve çok alanlı bilgiyi

işleme kapasitesi nedeniyle tedarik zinciri gibi karmaşık karar alanlarında yeni bir karar destek paradigması oluşturduğunu vurgulamaktadır.

LLM'lerin karar süreçlerine katkısı yalnızca veri işleme kapasitesiyle sınırlı değildir. Bu modeller yapılandırılmamış metinlerin işlenmesi, kurumsal dokümanların analiz edilmesi, karar alternatiflerinin oluşturulması ve teknik çıktının doğal dil açıklamalarına dönüştürülmesi gibi işlevler üstlenmektedir (Kaleta, 2025; Liu vd., 2024; Liu vd., 2026; Peddi, 2025). LLM çıktılarının olasılıksal ve üretken doğası; hatalı bilgi üretimi, bağlamdan kopuk öneriler, açıklama güvenilirliği ve karar sorumluluğunun kime ait olduğu gibi sorunları gündeme getirmektedir (Kumar vd., 2025; Sachan vd., 2025; Sonar vd., 2026).

Karar verme süreçlerinde üretken yapay zekâ

Karar verme süreçleri problemin tanımlanması, bilgi toplanması, alternatiflerin geliştirilmesi, alternatiflerin değerlendirilmesi, seçim yapılması ve karar sonuçlarının izlenmesi gibi aşamalardan oluşmaktadır. Üretken yapay zekâ bu aşamaların her birinde farklı roller üstlenebilir. Stratejik planlama bağlamında GenAI; senaryo geliştirme, bilişsel çerçeveleme ve stratejik seçeneklerin değerlendirilmesi açısından karar zekâsını güçlendirebilir (Gupta ve Kumar, 2026; Saup vd., 2026). Örgütsel karar verme bağlamında ise GenAI'nin bilişsel yükü azaltabileceği, sezgisel yanlılıkları sınırlayabileceği ve veri temelli değerlendirmeyi destekleyebileceği gösterilmiştir (Agarwal vd., 2026; Hao vd., 2024).

Finansal karar destek alanında LLM'ler kredi değerlendirme, yatırım analizi, risk değerlendirme, piyasa duyarlılığı analizi, portföy optimizasyonu ve finansal açıklama üretimi gibi süreçlerde kullanılmaktadır (Liu vd., 2024; Peddi, 2025; Tahvildari, 2025; Zhuang ve Wu, 2025). Tedarik zinciri yönetimi alanında ise LLM'ler sözleşme analizi, tedarikçi değerlendirme, talep tahmini, risk görünürlüğü, karar destek ve süreç koordinasyonu açısından ele alınmaktadır (Du vd., 2026; Ghosh vd., 2026; Sonar vd., 2026; Song vd., 2026). Bu durum üretken yapay zekânın karar verme süreçlerinde disiplinler arası bir araştırma alanı oluşturduğunu göstermektedir.

Güven, açıklanabilirlik ve insan gözetimi

Güven, kullanıcıların üretken yapay zekâ önerilerini dikkate alma, kabul etme veya bu önerilere dayanarak karar verme eğilimini etkileyen temel bir değişkendir (Huynh, 2024; Pescher, 2026; Tahvildari, 2025). Huynh (2024), örgütsel kullanıcıların GenAI destekli karar önerilerini benimsemesinde algılanan fayda, mahremiyet riski, bilgi kontrolü ve güvenin belirleyici olduğunu göstermektedir. Benzer biçimde Tahvildari (2025), robo-danışmanlık sistemlerinde XAI, kullanıcı güveni ve insan gözetimi mekanizmalarının algoritma karşıtlığını azaltmada kritik olduğunu vurgulamaktadır.

Açıklanabilirlik, yapay zekâ sistemlerinin ürettiği sonuçların kullanıcılar tarafından anlaşılabilir ve denetlenebilir olmasını ifade eder (Kumar vd., 2025; Liu vd., 2024; Sheng vd., 2025). Bu çalışmada açıklanabilirlik yalnızca teknik XAI yöntemleriyle sınırlı biçimde ele alınmamaktadır. SHAP, LIME, belirsizlik raporları, dikkat (attention) mekanizmasına dayalı açıklamalar ve sonradan üretilen açıklamalar (post hoc explanations) gibi teknik mekanizmaların yanında, açıklamaların kullanıcılar tarafından anlaşılabilir doğal dil ifadelerine dönüştürülmesi ve karar gerekçelerinin izlenebilir biçimde sunulması da açıklanabilirliğin parçası olarak değerlendirilmektedir (Buranasomphop, 2025; Kumar vd., 2025; Liu vd., 2024; Sheng vd., 2025).

İnsan gözetimi ise karar verme süreçlerinde nihai sorumluluğun insan karar vericide kalmasını, yapay zekâ çıktılarının insan tarafından kontrol edilebilmesini ve yüksek etkili kararların otomatik biçimde uygulanmamasını ifade etmektedir (Bajestani vd., 2026; Saup vd., 2026). Akıllı üretim bağlamında LLM ve human-in-the-loop yaklaşımlarını inceleyen Bajestani vd. (2026), insan uzmanlığının bağlamsal doğruluk, performans güvencesi ve kalite kontrol açısından vazgeçilmez olduğunu belirtmektedir. Stratejik karar verme bağlamında Saup vd. (2026) da GenAI çıktılarının karar sistemlerine alınabilmesi için kalite, kaynak izlenebilirliği, açıklanabilirlik ve hesap verebilirlik kapılarından geçmesi gerektiğini vurgulamaktadırlar.

Yöntem

Araştırma tasarımı

Bu çalışma, üretken yapay zekâ ve büyük dil modeli destekli karar verme sistemlerine ilişkin güncel literatürü sistematik biçimde incelemek amacıyla yürütülmüştür. Araştırmada PRISMA 2020 akışı temel alınmış; arama stratejisinin ayrıntılı raporlanmasında PRISMA-S, protokol öğelerinin geriye dönük belgelendirilmesinde ise PRISMA-P ilkelerinden yararlanılmıştır (Moher vd., 2015; Page vd., 2021; Rethlefsen vd., 2021). Literatür taraması 2020–2026 yılları arasını kapsayacak biçimde planlanmış; Web of Science ve Scopus dışı aktarımları 13 Mayıs 2026 tarihinde alınmış ve veri seti Mayıs 2026 içinde

tamamlanmıştır. Çalışmanın kapsamı üretken yapay zekâ/LLM destekli karar verme sistemlerinde güven, açıklanabilirlik, insan gözetimi ve yönetim boyutlarıyla sınırlandırılmıştır. Bu çalışma yalnızca kamuya açık yayımlanmış bilimsel kaynaklara dayalı bir sistematik literatür incelemesi olduğundan etik kurul izni gerektirmemektedir.

Araştırma soruları ve anahtar kelime stratejisinin geliştirilmesi

Araştırma soruları (research questions – RQ) ve tarama kelime dağarcığı çalışmanın kavramsal kapsamını doğrudan yansıtabilecek biçimde oluşturulmuştur. Bu süreçte önce araştırma problemindeki üç ana kavram alanı belirlenmiştir: (i) üretken yapay zekâ ve büyük dil modelleri, (ii) karar verme ve karar destek sistemleri, (iii) güven, açıklanabilirlik, insan gözetimi, yönetim, hesap verebilirlik ve karar kalitesi. Böylece arama stratejisi hem geniş literatürü yakalayacak hem de yalnızca genel yapay zekâ veya teknik model performansı odaklı çalışmaları dışarıda bırakacak şekilde sınırlandırılmıştır. Sistematik literatür incelemelerinde arama terimlerinin araştırma soruları, uygunluk ölçütleri ve bilgi kaynaklarıyla uyumlu biçimde raporlanması, PRISMA 2020 yaklaşımının temel şeffaflık ilkeleriyle uyumludur (Page vd., 2021).

Arama terimleri üç kavram bloğu altında geliştirilmiştir. İlk aşamada üretken yapay zekâ ve büyük dil modellerini temsil eden “generative artificial intelligence”, “generative AI”, “ChatGPT”, “large language model” ve “LLM” ifadeleri seçilmiştir. Bu terimler, literatürde GenAI ve LLM sistemlerinin karar verme, bilgi sentezi ve öneri üretimi bağlamında giderek artan kullanımını kapsamak amacıyla belirlenmiştir (Albashrawi, 2025; Gupta ve Kumar, 2026; Hao vd., 2024; Song vd., 2026). İkinci aşamada karar verme bağlamını yakalamak için “decision making”, “decision-making”, “decision support”, “managerial decision making” ve “organizational decision making” ifadeleri eklenmiştir. Bu seçim çalışmanın odağını yalnızca genel yapay zekâ uygulamalarından ayırarak örgütsel, yönetsel ve karar destek bağlamlarıyla ilişkili yayınlara yöneltmiştir (Saup vd., 2026; Zong vd., 2026).

Üçüncü kavram bloğu çalışmanın tematik odağına göre belirlenmiştir. “Trust” ve kullanıcı kabulü literatürü, GenAI önerilerinin benimsenmesinde algılanan fayda, mahremiyet riski ve bilgi kontrolünün belirleyici olduğunu göstermektedir (Huynh, 2024; Pescher, 2026). “Explainability” ve belirsizlik yönetimi, LLM çıktılarının kullanıcı tarafından anlaşılabilir, izlenebilir ve sınırlılıklarıyla birlikte değerlendirilebilir olmasını gerektirmektedir (Kumar vd., 2025; Liu vd., 2024; Sheng vd., 2025). “Human oversight”, “governance” ve “accountability” kavramları ise yüksek etkili karar süreçlerinde insan gözetimi, denetim izi, sorumluluk paylaşımı ve kurumsal karar kapılarının önemine işaret etmektedir (Bajestani vd., 2026; Goyal vd., 2026; Hinov ve Ivanova, 2026; Sachan vd., 2025; Saup vd., 2026). Bu nedenle üçüncü blok; çalışmanın güven, açıklanabilirlik, insan gözetimi ve yönetim odağını sistematik biçimde yakalamak için kullanılmıştır.

Araştırma soruları bu kavram bloklarıyla tutarlı biçimde geliştirilmiştir. RQ1, literatürde GenAI/LLM destekli karar sistemlerinin hangi sektör ve karar işlevlerinde incelendiğini belirlemek için oluşturulmuştur. RQ2, güven, açıklanabilirlik ve belirsizlik yönetiminin karar destek kabulü üzerindeki merkezi rolünden hareketle belirlenmiştir. RQ3, insan gözetimi, yönetim ve hesap verebilirlik gereksinimlerinin yüksek riskli karar bağlamlarında öne çıkması nedeniyle oluşturulmuştur. RQ4 ise halüsinasyon, yanlışlık, veri gizliliği, aşırı güvenme, alan bilgisi eksikliği ve regülasyon uyumu gibi risklerin literatürde tekrar eden sınırlılıklar olarak vurgulanmasına dayanmaktadır (Al Khalidy ve Gheraibia, 2025; Ghosh vd., 2026; Sonar vd., 2026; Tahvildari, 2025). Böylece araştırma soruları hem literatürdeki boşlukları hem de nihai tematik analizde kullanılacak kodlama boyutlarını yönlendiren analitik çerçeve olarak kullanılmıştır.

RQ1. Üretken yapay zekâ/LLM destekli karar verme sistemleri hangi karar alanlarında ve karar işlevlerinde incelenmektedir?

RQ2. Bu sistemlerde güven, kullanıcı kabulü, açıklanabilirlik ve belirsizlik yönetimi hangi mekanizmalarla ele alınmaktadır?

RQ3. İnsan gözetimi, insan-yapay zekâ iş birliği, hesap verebilirlik ve yönetim karar süreçlerine nasıl entegre edilmektedir?

RQ4. Literatürde öne çıkan riskler, sınırlılıklar ve gelecek araştırma boşlukları nelerdir?

Veri tabanları ve tarama stratejisi

Literatür taraması Web of Science ve Scopus veri tabanlarında 13 Mayıs 2026 tarihinde gerçekleştirilmiştir. Bu iki veri tabanının seçilmesinin nedeni, sosyal bilimler, yönetim bilişim sistemleri, karar destek sistemleri, yapay zekâ, mühendislik ve yönetim alanlarında uluslararası hakemli yayınları kapsamlı biçimde indekslemeleridir. Benzer sistematik literatür incelemelerinde de Scopus ve Web of

Science'ın birlikte kullanılması, kapsamlı ve denetlenebilir bir literatür havuzu oluşturmak amacıyla tercih edilmektedir (Mongeon ve Paul-Hus, 2016; Song vd., 2026).

Tarama dönemi 2020-2026 yıllarıyla sınırlandırılmıştır. Bu aralık, üretken yapay zekâ ve büyük dil modeli literatürünün erken döneminden ChatGPT sonrası hızlanan döneme kadar olan gelişimini kapsamak amacıyla tercih edilmiştir. Her iki veri tabanında da yalnızca İngilizce dilinde yayımlanan makale, derleme ve konferans bildirisi türündeki çalışmalar tarama kapsamına alınmıştır. Nihai analiz setindeki yayınların ağırlıklı olarak 2024-2026 yıllarında yoğunlaşması, söz konusu araştırma alanının özellikle son yıllarda hızla genişlediğini ve bilimsel üretimin bu dönemde belirgin ölçüde arttığını göstermektedir.

Tarama stratejisi; üretken yapay zekâ/LLM, karar verme/karar destek ve güven, açıklanabilirlik, insan gözetimi ve yönetim olmak üzere üç kavram kümesine dayandırılmıştır. Kullanılan anahtar kelime blokları Tablo 1'de sunulmaktadır.

Tablo 1: Tarama Stratejisinde Kullanılan Anahtar Kelime Blokları

Kavram Bloğu	Kullanılan Anahtar Kelimeler	Amaç
Üretken yapay zekâ / LLM	"generative artificial intelligence", "generative AI", ChatGPT, "large language model*", LLM*	Üretken yapay zekâ ve büyük dil modeli odaklı çalışmaları belirlemek
Karar verme / karar destek	"decision making", "decision-making", "decision support", "managerial decision making", "organizational decision making"	Karar verme, karar destek ve örgütsel/yönetimsel karar bağlamını yakalamak
Güven / açıklanabilirlik / yönetim	trust, explainability, "human oversight", governance, accountability, "decision quality"	Çalışmanın güven, açıklanabilirlik, insan gözetimi ve yönetim odağını sınırlamak

Raporlanabilir sorgu yapısı şu şekilde özetlenebilir:

Web of Science sorgusu:

TS = (("generative artificial intelligence" OR "generative AI" OR ChatGPT OR "large language model*" OR LLM*) AND ("decision making" OR "decision-making" OR "decision support" OR "managerial decision making" OR "organizational decision making") AND (trust OR explainability OR "human oversight" OR governance OR accountability OR "decision quality"))

Scopus sorgusu:

TITLE-ABS-KEY (("generative artificial intelligence" OR "generative AI" OR ChatGPT OR "large language model*" OR LLM*) AND ("decision making" OR "decision-making" OR "decision support" OR "managerial decision making" OR "organizational decision making") AND (trust OR explainability OR "human oversight" OR governance OR accountability OR "decision quality"))

Veri tabanları dinamik indeksleme sistemleri olduğundan aynı sorunun daha sonraki bir tarihte çalıştırılması, geriye dönük indeksleme, kayıt düzeltmeleri veya veri tabanı kapsamındaki değişiklikler nedeniyle aynı mutlak kayıt sayısını vermeyebilir (Rethlefsen vd., 2021). Bununla birlikte arama tarihi, veri tabanına özgü alanlar, filtreler, sorgu dizileri ve tekrar kayıt temizleme kuralı açıkça raporlandığından süreç yöntemsel olarak yeniden uygulanabilir niteliktedir. Kayıtlar DOI üzerinden normalize edilerek eşleştirilmiş; DOI bulunmayan veya tutarsız olan kayıtlarda normalize edilmiş başlık eşleşmesi kullanılmıştır. Aynı çalışma iki veri tabanında bulunduğu tek kayıt korunmuş ve veri tabanı kaynağı birlikte belirtilmiştir.

Dâhil etme ve dışlama ölçütleri

Çalışmaların değerlendirilmesinde aşağıdaki dâhil etme ölçütleri kullanılmıştır:

- Çalışmanın üretken yapay zekâ, ChatGPT, büyük dil modeli veya LLM bağlamı içermesi,
- Karar verme, karar destek, örgütsel karar verme, yönetsel karar verme veya karar kalitesiyle ilişkili olması,
- Güven, açıklanabilirlik, insan gözetimi, insan-yapay zekâ etkileşimi, yönetim, hesap verebilirlik veya sorumlu kullanım temalarından en az biriyle ilişkili olması,
- Hakemli makale, konferans bildirisi, derleme niteliğinde olması,
- Tam metnine erişilebilmesi ve uygunluk ölçütlerini karşılaması.

Dışlama ölçütleri ise şu şekilde belirlenmiştir:

- Üretken yapay zekâ/LLM içermeyen genel yapay zekâ çalışmaları,
- Karar verme veya karar destek bağlamı bulunmayan çalışmalar,
- Yalnızca teknik model performansına veya algoritma geliştirmeye odaklanan çalışmalar,
- Tıbbi/klinik karar destek bağlamında kalıp örgütsel veya yönetsel karar verme boyutu zayıf olan çalışmalar,
- Eğitimde ChatGPT kullanımı gibi karar verme odağı zayıf olan çalışmalar,
- Tam metnine erişilemeyen çalışmalar,
- Üretken yapay zekâ/LLM'nin çalışmada merkezi bir konumda bulunmadığı, karar verme veya karar destek süreciyle doğrudan ilişki kurulmadığı ya da güven, açıklanabilirlik, insan gözetimi, hesap verebilirlik ve yönetim boyutlarından hiçbirini karşılamadığı belirlenen çalışmalar.

PRISMA akışı ve eleme süreci

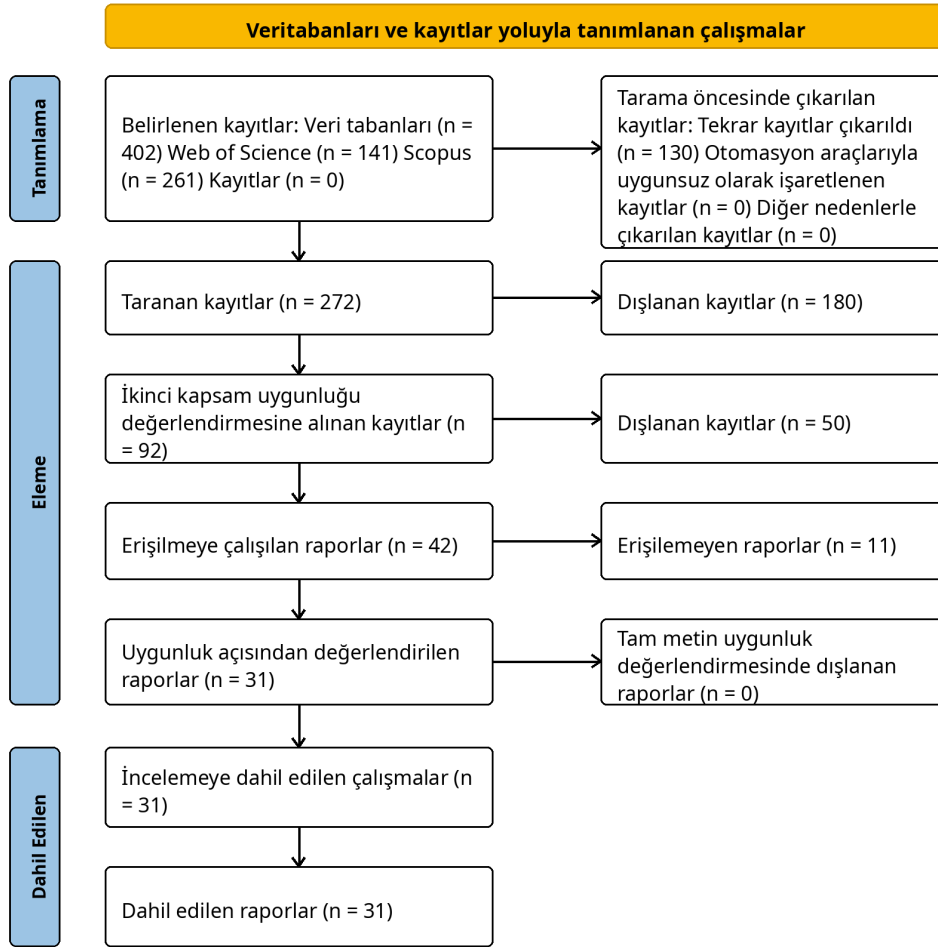
Web of Science ve Scopus veri tabanlarında 13 Mayıs 2026 tarihinde gerçekleştirilen tarama sonucunda toplam 402 kayda ulaşılmıştır. Bu kayıtların 141'i Web of Science, 261'i ise Scopus veri tabanından elde edilmiştir. İki veri tabanından aktarılan kayıtlar birleştirilmiş; tekrar kayıtların belirlenmesinde öncelikle DOI bilgileri, DOI'nin bulunmadığı veya tutarsız olduğu durumlarda ise normalize edilmiş başlık eşleşmeleri kullanılmıştır. Bu işlem sonucunda 130 tekrar kayıt çıkarılmış ve 272 benzersiz kayıt başlık ve özet düzeyinde değerlendirmeye alınmıştır. İlk tarama aşamasında üretken yapay zekâ veya büyük dil modeli odağı taşımayan, karar verme ya da karar destek süreçleriyle doğrudan ilişkili olmayan ve araştırmanın kavramsal kapsamının dışında kalan 180 kayıt dışlanmıştır. Böylece 92 kayıt potansiyel olarak uygun bulunarak ikinci aşama kapsam uygunluğu değerlendirmesine alınmıştır.

Tablo 2: PRISMA Akışına İlişkin Sayısal Özet

Aşama	Sayı
Web of Science kayıtları	141
Scopus kayıtları	261
Toplam ham kayıt	402
Çıkarılan tekrar kayıt	130
Başlık-özet taramasına alınan benzersiz kayıt	272
Başlık-özet taramasında dışlanan kayıt	180
İkinci kapsam uygunluğu değerlendirmesine alınan kayıt	92
İkinci kapsam değerlendirmesinde dışlanan kayıt	50
Tam metni edinilmeye çalışılan rapor	42
Tam metnine erişilemeyen rapor	11
Tam metin uygunluk değerlendirmesinde dışlanan raporlar	0
Nihai senteze dâhil edilen çalışma	31

İkinci aşamada söz konusu 92 kayıt, başlık ve özet bilgileri üzerinden üretken yapay zekâ veya büyük dil modellerinin çalışmadaki merkezi konumu, karar verme ya da karar destek süreçleriyle kurulan ilişki ve güven, açıklanabilirlik, insan gözetimi, hesap verebilirlik veya yönetim boyutlarından en az biriyle bağlantısı bakımından yeniden incelenmiştir. Bu değerlendirme sonucunda, önceden belirlenen kapsam ölçütlerini karşılamayan 50 kayıt dışlanmış ve 42 raporun tam metni edinilmeye çalışılmıştır. Kurumsal erişim olanakları, yayıncı sayfaları, DOI bağlantıları ve açık erişim depoları kontrol edilmesine rağmen 11 raporun tam metnine ulaşılamamıştır. Erişilebilen 31 rapor tam metin düzeyinde değerlendirilmiş ve tamamının dâhil etme ölçütlerini karşıladığı belirlenerek nihai tematik senteze alınmıştır. Eleme ve uygunluk değerlendirmeleri, araştırmacıların aynı kayıtları birlikte incelemesi ve kararları uzlaşa yoluyla vermesiyle yürütülmüştür. Başlık ve özet taramasındaki dışlama gerekçeleri ile

ikinci kapsam değerlendirmesinin aşama bazındaki sayısal sonuçları ortak tarama dosyasında kayıt altına alınmıştır. Değerlendirmeler bağımsız ve paralel olarak gerçekleştirilmediğinden iki ayrı değerlendirme dizisi oluşmamış ve bu nedenle kodlayıcılar arası uyum katsayısı hesaplanmamıştır. Çalışma seçimine ilişkin sayısal bilgiler Tablo 2’de, sürecin PRISMA 2020 doğrultusundaki görsel gösterimi ise Şekil 1’de sunulmuştur.



Şekil 1: Çalışma Seçim Sürecine İlişkin PRISMA 2020 Akış Diyagramı

Kaynak: Page vd., 2021; PRISMA Statement, t.y.

Not: Page vd. (2021) tarafından geliştirilen ve PRISMA Statement tarafından yayımlanan “PRISMA 2020 flow diagram for new systematic reviews” şablonu temel alınarak yazarlar tarafından uyarlanmıştır.

Not: Akış, Web of Science ve Scopus taramasından nihai tematik senteze kadar izlenen seçim sürecini özetlemektedir.

Veri çıkarma ve kodlama

Nihai analiz kapsamına alınan 31 çalışma için sistematik bir kodlama tablosu oluşturulmuştur. Kodlama alanları araştırma sorularından ve tam metinlerde tekrar eden kavramlardan türetilmiş; başlık, yayın yılı, yayın türü, uygulama alanı, araştırma yöntemi, GenAI/LLM yaklaşımı, karar verme bağlamı, güven, açıklanabilirlik, insan gözetimi, yönetim, riskler ve sentezdeki kullanım düzeyi kaydedilmiştir. Araştırmacılar kodlama tablosunu birlikte incelemiş, belirsiz veya çok temalı çalışmalar için birincil tema ortak tartışma yoluyla belirlenmiş ve karar gerekçeleri denetim izi olarak korunmuştur. Kodlama sonucunda çalışmalar altı ana tema altında sınıflandırılmıştır: karar destek kapasitesi, güven ve kullanıcı kabulü, açıklanabilirlik ve belirsizlik yönetimi, insan gözetimi ve yönetim, riskler ve sorumlu kullanım, karar mimarisi ve çok kriterli karar verme. Bu kodlama yaklaşımı, sistematik literatür incelemelerinde şeffaflık ve yeniden izlenebilirlik ilkeleriyle uyumludur (Page vd., 2021).

İnceleme protokolü, ortak değerlendirme ve kodlama güvenilirliği

Bu çalışma için araştırma süreci başlamadan önce ileriye dönük ve kamuya açık bir sistematik derleme protokolü kaydedilmemiştir. PROSPERO esas olarak sağlık ve sağlıkla ilişkili sonuçlara odaklanan sistematik derlemelerin kayıt altına alınması amacıyla kullanılan bir platform olduğundan, örgütsel ve

yönetmelik karar verme bağlamındaki bu incelemenin konusu PROSPERO'nun temel kapsamıyla doğrudan örtüşmemektedir (Centre for Reviews and Dissemination [CRD], t.y.). Bununla birlikte önceden kayıtlı bir protokolün bulunmaması yöntemsel bir sınırlılık olarak kabul edilmiştir.

İnceleme sürecinin şeffaflığını ve izlenebilirliğini artırmak amacıyla kullanılan veri tabanları, tarama tarihi, arama sorguları, dâhil etme ve dışlama ölçütleri, tekrar kayıtların belirlenme yöntemi, veri çıkarma alanları, başlık ve özet taramasındaki eleme gerekçeleri, aşama bazındaki sayısal sonuçlar ve tematik sentez süreci ayrıntılı biçimde belgelenmiştir. Başlık ve özet taraması, ikinci kapsam uygunluğu değerlendirmesi, tam metin incelemesi ve tematik kodlama aşamaları araştırmacıların aynı kayıtları birlikte değerlendirmesiyle yürütülmüştür. Dâhil etme, dışlama ve sınıflandırma kararları önceden belirlenen ölçütler doğrultusunda tartışılmış ve uzlaşma yoluyla sonuçlandırılmıştır. Belirsiz, sınırda kalan veya birden fazla temayla ilişkilendirilebilecek çalışmalar yeniden incelenmiş; nihai kararlar ortak değerlendirme sonucunda verilmiştir. Başlık ve özet taramasındaki dışlama gerekçeleri, sonraki aşamalarda sayısal kararlar ve tematik kodlama gerekçeleri ortak tarama ve kodlama dosyalarına kaydedilmiştir.

Değerlendirme süreci iki bağımsız kodlayıcının birbirinden ayrı karar üretmesi biçiminde yürütülmediğinden kodlayıcılar arası uyumu ölçen Cohen'in kappa katsayısı hesaplanmamıştır. Bunun yerine kodlama güvenilirliği; açık biçimde tanımlanmış uygunluk ölçütleri, ortak bir kodlama çerçevesi, kayıtların aşamalı olarak yeniden incelenmesi, belirsiz durumların tartışma yoluyla çözülmesi ve mevcut kararların ortak dosyalarda belgelenmesi yoluyla desteklenmiştir. Bağımsız paralel kodlama yapılmamış olması çalışmanın sınırlılıkları arasında ayrıca belirtilmiştir.

Yöntemsel kalite değerlendirmesi

Nihai senteze dâhil edilen çalışmalar araştırma tasarımı, veri kaynağı ve analiz yaklaşımı bakımından sınıflandırılmış; ardından her çalışma için tasarımına uygun değerlendirme aracı seçilmiştir. Bu doğrultuda 23 ampirik veya teknik çalışmanın ampirik bileşenleri MMAT 2018'in ilgili tasarım ölçütleriyle değerlendirilmiştir. Teknik çalışmalarda veri seti, model çıktısı, istem veya teknik test örneği, ilgili ölçütteki çalışma biriminin operasyonel karşılığı olarak yorumlanmıştır. Üç sistematik literatür derlemesi ve bir sistematik kapsam derlemesi JBI Sistematik Derlemeler ve Araştırma Sentezleri Kontrol Listesiyle; dört kavramsal veya metinsel çalışma ise JBI Metinsel Kanıt: Uzman Görüşü Kontrol Listesiyle değerlendirilmiştir. Sistematik kapsam derlemesinde JBI listesi tasarımın amacı dikkate alınarak uyarlanmış; eleştirel değerlendirme veya yayın yanlılığı analizi yapılmadığında ilgili maddeler "Hayır" olarak kaydedilmiştir (Aromataris vd., 2024; Hong vd., 2018; McArthur vd., 2025).

Yöntemsel kalite değerlendirmesi tam metinler üzerinden madde bazında yürütülmüş ve kararlar "Evet", "Hayır" veya "Belirsiz" olarak kaydedilmiştir. Karma yöntemli çalışmalarda ilgili nitel, nicel ve bütünleştirme ölçütleri birlikte uygulanmıştır. MMAT yönergesi doğrultusunda kararlar yüzdelik ya da tek bir toplam kalite puanına dönüştürülmemiş; çalışmalar yüksek, orta veya düşük kalite biçiminde sınıflandırılmamıştır (Hong vd., 2018). Değerlendirme otomatik bir dışlama eşiği olarak kullanılmamış, 31 çalışmanın tamamı sentezde tutulmuş ve belirlenen sınırlılıklar bulguların yorumlanmasında dikkate alınmıştır.

Bulgular

Dâhil edilen çalışmaların genel görünümü

Nihai senteze dâhil edilen 31 çalışma, üretken yapay zekâ ve büyük dil modellerinin karar verme süreçlerinde hızla gelişen ve farklı uygulama alanlarına yayılan bir araştırma konusu olduğunu göstermektedir. Tarama dönemi 2020–2026 yıllarını kapsamına rağmen dâhil edilen çalışmaların 18'i 2026, 10'u 2025 ve 3'ü 2024 yılına aittir. Bu dağılım üretken yapay zekâ destekli karar verme araştırmalarının özellikle son üç yılda belirgin biçimde yoğunlaştığını göstermektedir. İnceleme kümesi; 23 ampirik veya teknik çalışma, üç sistematik literatür derlemesi, bir sistematik kapsam derlemesi ve dört kavramsal veya metinsel çalışmadan oluşmaktadır. Nihai sentezde yalnızca sistematik tarama sonucunda belirlenen ve tam metin uygunluk ölçütlerini karşılayan çalışmalar kullanılmış; dolayısıyla bulgular bölümündeki değerlendirmeler 31 çalışmadan elde edilen kanıtlara dayandırılmıştır (Page vd., 2021).

Çalışmaların yoğunlaştığı başlıca uygulama alanları finans ve bankacılık, tedarik zinciri yönetimi, iş zekâsı, kamu yönetimi, insan kaynakları, endüstriyel süreçler, stratejik yönetim, kurumsal yatırım, robo-danışmanlık, turizm ve tüketici kararlarıdır (Du vd., 2026; Ghosh vd., 2026; Goyal vd., 2026; Hoque vd., 2026; Huynh, 2024; Kim vd., 2026; Tahvildari, 2025; Theodorakopoulos vd., 2026; Zhuang ve Wu, 2025; Zong vd., 2026). Bu çeşitlilik üretken yapay zekânın karar verme bağlamında bilgi sentezi,

alternatif üretimi, risk değerlendirmesi, açıklama oluşturma ve insan karar vericiyi destekleme işlevlerine sahip bilişsel bir karar destek bileşeni olarak incelendiğini göstermektedir (Albashrawi, 2025; Saup vd., 2026).

Her çalışma, araştırma soruları ve tam metin incelemesinde belirlenen kavramsal örüntüler doğrultusunda birincil bir tema altında sınıflandırılmıştır. Buna göre çalışmalar; karar destek kapasitesi, güven ve kullanıcı kabulü, açıklanabilirlik ve belirsizlik yönetimi, insan gözetimi ve yönetim, riskler ve sorumlu kullanım ile karar mimarisi ve çok kriterli karar verme olmak üzere altı ana tema altında toplanmıştır. Çalışmaların birincil temalara göre dağılımı Tablo 3'te sunulmaktadır. Çalışmalar ayrıca uygulama alanları ve literatüre sundukları katkılar bakımından değerlendirilmiştir. Odak alanları birbirini dışlayan kategoriler olarak oluşturulmamış; bir çalışmanın birden fazla uygulama veya kavramsal alana katkı sağlayabileceği kabul edilmiştir. Bu nedenle odak alanları, tematik dağılımdan farklı olarak literatürdeki kesişimleri ve çalışmaların ortak katkılarını görünür kılmak amacıyla kullanılmıştır. Çalışmaların odak alanlarına göre sınıflandırılması Tablo 4'te; her bir çalışmanın alanı, araştırma tasarımı, birincil teması ve sentezdeki kullanım biçimi ise Tablo 5'te gösterilmektedir.

Tablo 3: Dâhil Edilen Çalışmaların Tematik Dağılımı

Tema	Çalışma sayısı	Açıklama
Karar destek kapasitesi	7	Üretken yapay zekâ ve büyük dil modellerinin öneri üretimi, veri analizi, risk değerlendirme ve senaryo geliştirme işlevleri
Güven ve kullanıcı kabulü	4	Kullanıcı güveni, yapay zekâya yönelik olumlu değerlendirme (AI appreciation), algoritmadan kaçınma (algorithm aversion), mahremiyet ve kabul
Açıklanabilirlik ve belirsizlik	6	Açıklanabilir yapay zekâ (explainable artificial intelligence - XAI), SHAP/LIME açıklamaları, belirsizlik raporları ve doğal dil açıklamaları
İnsan gözetimi ve yönetim	6	İnsanın döngüde olduğu yaklaşım (human-in-the-loop), insanın nihai kontrolü elinde tuttuğu yaklaşım (human-in-command), denetlenebilirlik (auditability) ve hesap verebilirlik (accountability)
Riskler ve sorumlu kullanım	5	Halüsinasyon, yanlışlık, veri gizliliği, aşırı güvenme, regülasyon ve sorumlu kullanım
Karar mimarisi ve ÇKKV	3	TOPSIS, AHP, COCOSO, DEMATEL, ISM gibi çok kriterli karar verme yöntemleri ve karar modeli tasarımı

Tablo 4: Nihai Analiz Kapsamındaki Çalışmaların Odak Alanlarına Göre Özeti

Odak alanı	İlgili çalışmalar	Genel katkı
Genel üretken yapay zekâ karar verme ve stratejik karar desteği	Albashrawi (2025), Gupta ve Kumar (2026), Hao vd. (2024), Kim vd. (2026), Saup vd. (2026)	Üretken yapay zekânın karar zekâsı, stratejik karar verme, insan-AI etkileşimi ve karar kalitesi üzerindeki etkilerini tartışmaktadır.
Güven, kullanıcı kabulü ve değerlendirme	Huynh (2024), Pescher (2026), Sheng vd. (2025), Tahvildari (2025)	Kullanıcı güveni, mahremiyet algısı, algoritma takdiri/kaçınması, belirsizlik raporları ve karar destek kabulünü incelemektedir.
Açıklanabilirlik ve açıklanabilir yapay zekâ destekli karar sistemleri	ve Kumar vd. (2025), Liu vd. (2024), Liu vd. (2026), Sheng vd. (2025), Hinov ve Ivanova (2026)	Açıklanabilir yapay zekâ, doğal dil açıklamaları, belirsizlik yönetimi, açıklama doğrulama ve karar şeffaflığı üzerinde durmaktadır.
Tedarik zinciri operasyonel kararlar	ve Du vd. (2026), Ghosh vd. (2026), Sonar vd. (2026), Song vd. (2026)	Büyük dil modellerinin tedarik zinciri karar desteği, SCOR süreçleri, risk görünürlüğü ve benimseme zorlukları üzerindeki rolünü değerlendirmektedir.
Finans, yatırım ve bankacılık kararları	Liu vd. (2024), Peddi (2025), Tahvildari (2025), Zhuang ve Wu (2025)	Kredi değerlendirme, risk analizi, robo-danışmanlık, yatırım kararları ve finansal karar destek süreçlerini ele almaktadır.
İnsan gözetimi, yönetim ve sorumlu kullanım	Bajestani vd. (2026), Goyal vd. (2026), Hinov ve Ivanova (2026), Sachan vd. (2025), Saup vd. (2026), Zhu vd. (2026)	İnsanın döngüde olduğu yaklaşım, insanın nihai kontrolü elinde tuttuğu yaklaşım, hesap verebilirlik, denetlenebilirlik, veri güvenliği ve yönetim mekanizmalarını incelemektedir.
Karar mimarisi ve çok kriterli karar verme uygulamaları	Hoque vd. (2026), Sonar vd. (2026), Zeqiri ve Ben Youssef (2026), Zhu vd. (2026)	Büyük dil modellerinin TOPSIS, AHP, COCOSO, DEMATEL, ISM ve benzeri çok kriterli karar verme yöntemleriyle birlikte kullanımını göstermektedir.

Tablo 5: Nihai Analiz Kapsamındaki 31 Çalışmanın Kodlama Özeti

No.	Çalışma	Alan/Bağlam	Yöntem/Tasarım	Ana Tema	Makaledeki Kullanım
1	Du vd. (2026)	Tedarik zinciri / döngüsel ekonomi	Karma yöntemli vaka/teknik değerlendirme	Karar destek kapasitesi	Karar destek kapasitesi ve insan uzman-LLM karşılaştırması
2	Kim vd. (2026)	Kamu Ar-Ge / stratejik değerlendirme	Nicel betimsel değerlendirme	İnsan gözetimi ve yönetim	Stratejik karar verme ve kamu yönetimi
3	Gupta ve Kumar (2026)	Stratejik planlama / karar zekâsı	Kavramsal/metinsel uzman görüşü	Karar destek kapasitesi	Karar zekâsı ve stratejik planlama için kavramsal arka plan
4	Kumar vd. (2025)	Kritik karar uygulamaları	Nicel randomize olmayan teknik değerlendirme	Açıklanabilirlik ve belirsizlik	Açıklanabilirlik ve kritik karar uygulamaları için kavramsal dayanak
5	Buranasomphop (2025)	Girişim finansmanı / finans	Nicel randomize olmayan simülasyon deneyi	Açıklanabilirlik ve belirsizlik	Finansal karar verme, LLM yanlılığı ve açıklanabilirliği
6	Hao vd. (2024)	Örgütsel karar verme	Nicel randomize olmayan yarı deney	İnsan gözetimi ve yönetim	İnsan-üretken yapay zekâ ortak karar vermesi için temel ampirik dayanak
7	Liu vd. (2026)	Bankacılık / iç kontrol	Nicel randomize olmayan tasarım-bilim değerlendirmesi	Açıklanabilirlik ve belirsizlik	Bankacılıkta açıklanabilir karar desteği ve denetim
8	Albashrawi (2025)	Çok disiplinli bağlam	Sistemik literatür derlemesi ve tematik sentez	Karar destek kapasitesi	Literatür arka planı ve çok disiplinli karar verme
9	Bajestani vd. (2026)	Akıllı üretim	Sistemik kapsam derlemesi	İnsan gözetimi ve yönetim	İnsan gözetimi, doğrulama ve endüstriyel karar desteği
10	Agarwal vd. (2026)	İş yeri / örgütsel kullanım	Nicel betimsel hesaplamalı metin analizi	Güven ve kullanıcı kabulü	İş yeri benimsemesi, güven ve yapay zekâ ekip arkadaşı metaforu
11	Zhu vd. (2026)	İnsan-yapay zekâ ortak tasarımı / ürün geliştirme	Nicel randomize olmayan küçük örneklemli yarı deney	Açıklanabilirlik ve belirsizlik	Açıklanabilir karar çerçevesi, insan gözetimi ve ÇKKV
12	Saup vd. (2026)	Stratejik karar verme	Nitel tek vaka çalışması	İnsan gözetimi ve yönetim	Stratejik karar verme ve yönetim kapıları için çekirdek kaynak
13	Peddi (2025)	Bankacılık / risk değerlendirme	Nicel randomize olmayan teknik/saha değerlendirmesi	Karar destek kapasitesi	Bankacılıkta risk değerlendirme ve finansal karar desteği
14	Song vd. (2026)	Tedarik zinciri yönetimi	Sistemik literatür derlemesi	Karar destek kapasitesi	Tedarik zinciri yönetimi uygulamaları ve sistemik inceleme örneği
15	Hinov ve Ivanova (2026)	Örgütsel karar sistemleri / algoritmik yönetim	Kavramsal/metinsel uzman görüşü	İnsan gözetimi ve yönetim	Yönetiş mimarisi, denetlenebilirlik ve hesap verebilir karar destek sistemleri
16	Zong vd. (2026)	Endüstriyel süreçler / metalürji	Nicel randomize olmayan endüstriyel teknik değerlendirme	Karar destek kapasitesi	Endüstriyel karar desteği ve insan kontrollü LLM tabanlı karar destek sistemi
17	Goyal vd. (2026)	Kamu yönetimi / akıllı yönetim	Nicel betimsel uzman değerlendirme ve ÇKKV	Riskler ve sorumlu kullanım	Kamu yönetimi, benimseme engelleri ve politika önerileri
18	Zeqiri ve Ben Youssef (2026)	Sürdürülebilir turizm / destinasyon yönetimi	Ardışık keşfedici karma yöntem	Karar mimarisi ve çok kriterli karar verme	ÇKKV ile üretken yapay zekâ stratejilerinin sıralanması ve karar mimarisi
19	Sachan vd. (2025)	Yüksek etkili finansal kararlar	Nicel randomize olmayan teknik değerlendirme	İnsan gözetimi ve yönetim	Yüksek etkili kararlar, denetim izi ve sorumlu LLM kullanımı
20	Kalet (2025)	Yöneylem araştırması / karar destek sistemi geliştirme	Nicel randomize olmayan teknik deney	Karar mimarisi ve çok kriterli karar verme	Konuşmalı karar destek sistemleri, optimizasyon ve LLM ajan mimarileri

No.	Çalışma	Alan/Bağlam	Yöntem/Tasarım	Ana Tema	Makaledeki Kullanım
21	Sonar vd. (2026)	Tedarik zinciri yönetimi	Nicel betimsel Delphi ve ÇKKV	Riskler ve sorumlu kullanım	Tedarik zinciri riskleri, LLM benimseme engelleri ve önceliklendirme
22	Pescher (2026)	Tüketici karar verme / öneri sistemleri	Nicel randomize olmayan kontrollü teknik faktöriyel deney	Güven ve kullanıcı kabulü	Yapay zekâ aracılı öneri sistemleri ve güven
23	Al Khaldy ve Gheraibia (2025)	İşletme kararları	Kavramsal/metinsel uzman görüşü	Riskler ve sorumlu kullanım	LLM sınırlılıkları ve risk azaltma stratejileri
24	Ghosh vd. (2026)	Tedarik zinciri yönetimi	Nicel betimsel uzman değerlendirmesi ve ISM-MICMAC	Riskler ve sorumlu kullanım	Tedarik zinciri karar verme çerçevesi, risk hiyerarşisi ve ISM/MICMAC bulguları
25	Tahvildari (2025)	Robo-danışmanlık / finansal karar desteği	Sistematiik literatür derlemesi	Güven ve kullanıcı kabulü	Finansal karar desteği, robo-danışmanlık, açıklanabilir yapay zekâ ve hibrit gözetim
26	Zhuang ve Wu (2025)	Kurumsal yatırım	Kavramsal/metinsel uzman görüşü	Riskler ve sorumlu kullanım	Kurumsal yatırım kararları ve risk değerlendirmesi
27	Liu vd. (2024)	Finansal karar desteği / kredi değerlendirmesi	Nicel randomize olmayan teknik gösterim	Açıklanabilirlik ve belirsizlik	Açıklanabilir yapay zekâ ile LLM tabanlı doğal dil açıklamalarının bütünleştirilmesi
28	Theodorakopoulos vd. (2026)	Bütünleşik iş zekâsı	Nicel randomize olmayan tasarım-bilim değerlendirmesi	Karar destek kapasitesi	Bütünleşik iş zekâsı ve finans-pazarlama-denetim karar destek mimarisi
29	Sheng vd. (2025)	İnsan-yapay zekâ karar sistemleri	Nicel randomize olmayan tekrarlı ölçüm deneyi	Açıklanabilirlik ve belirsizlik	Açıklanabilirlik ve belirsizlik yönetimi için çekirdek ampirik kaynak
30	Huynh (2024)	Örgütsel kullanıcılar / üretken yapay zekâ karar destek sistemi	Nicel betimsel kesitsel anket	Güven ve kullanıcı kabulü	Güven, kullanıcı kabulü ve üretken yapay zekâ karar destek araçlarının benimsenmesi
31	Hoque vd. (2026)	Personel seçimi / insan kaynakları yönetimi	Nicel randomize olmayan teknik değerlendirme	Karar mimarisi ve çok kriterli karar verme	LLM ve Bulanık TOPSIS bütünleşmesine dayalı personel seçimi karar mimarisi

Not: Tablo, tam metin incelemesine dâhil edilen çalışmaların tematik kodlama sürecinde hangi bağlam ve kullanım alanlarıyla değerlendirildiğini göstermektedir.

Yöntemsel kalite değerlendirmesi bulguları

Nihai inceleme kümesinin genel özellikleri belirlendikten sonra 31 çalışmanın yöntemsel yeterliği, araştırma tasarımlarına uygun araçlarla madde bazında değerlendirilmiştir. Değerlendirme, nihai çalışma kümesinin belirlenmesinden sonra ve tematik bulguların yorumlanmasından önce gerçekleştirilmiştir. Yirmi üç ampirik veya teknik çalışma MMAT 2018; üç sistematik literatür derlemesi ile bir sistematik kapsam derlemesi JBI Sistematik Derlemeler ve Araştırma Sentezleri; dört kavramsal veya metinsel çalışma ise JBI Metinsel Kanıt: Uzman Görüşü ölçütleriyle incelenmiştir. Otuz bir çalışma için toplam 249 madde bazlı karar doğrulanmıştır: 171 “Evet”, 33 “Hayır” ve 45 “Belirsiz”. On beş çalışmada en az bir “Hayır”, 24 çalışmada en az bir “Belirsiz” kararı bulunurken üç çalışma yalnızca “Evet” kararlarından oluşmuştur. MMAT ile değerlendirilen çalışmalarda araştırma amaçları bütün çalışmalarda açık, veriler ise büyük ölçüde amaçlarla uyumludur; buna karşılık nicel çalışmalarda örneklem veya teknik test birimlerinin temsil gücü ile karıştırıcı etkenlerin kontrolü tekrar eden sınırlılıklardır. JBI ile değerlendirilen derlemelerde derleme sorusu ve sentez yaklaşımı genel olarak açık olmakla birlikte eleştirel değerlendirme, bağımsız çift değerlendirme ve yayın yanlılığı incelemesi çoğunlukla raporlanmamıştır. Uzman görüşü niteliğindeki çalışmalarda kaynak, uzmanlık ve literatür dayanağı genellikle açık; görüşün analitik türetimi ve literatürle uyumsuzlukların gerekçelendirilmesi ise bazı çalışmalarda belirsizdir. Bu kararlar toplam kalite puanı veya genel kalite derecesi olarak yorumlanmamış; hiçbir çalışma yalnızca yöntemsel değerlendirme nedeniyle dışlanmamıştır.

Tablo 6: Dâhil Edilen Çalışmaların Yöntemsel Değerlendirme Özeti

Çalışma grubu	Araç	n	Madde kararı dağılımı	Öne çıkan yöntemsel bulgular
Ampirik veya teknik çalışmalar	MMAT 2018'in ilgili nitel, nicel ve karma yöntem ölçütleri	23	130 Evet; 17 Hayır; 34 Belirsiz (181 karar)	Araştırma amaçları açık ve veri-amaç uyumu çoğunlukla yeterlidir. Nicel çalışmalarda temsil gücü ve karıştırıcı etkenlerin kontrolü başlıca sınırlılıklardır.
Sistematiik literatür derlemeleri	JBİ Sistematiik Derlemeler ve Araştırma Sentezleri Kontrol Listesi	3	15 Evet; 13 Hayır; 5 Belirsiz (33 karar)	Derleme soruları ve sentez yöntemleri genellikle açıktır; eleştirel değerlendirme, bağımsız çift değerlendirme, veri çıkarma güvenceleri ve yayın yanlılığı incelemesi sınırlıdır.
Sistematiik kapsam derlemesi	JBİ Sistematiik Derlemeler ve Araştırma Sentezleri Kontrol Listesi (uyarlanmış)	1	7 Evet; 3 Hayır; 1 Belirsiz (11 karar)	Kapsam derlemesinde eleştirel değerlendirme ve yayın yanlılığı analizi yapılmamıştır; kararlar tasarının amacı dikkate alınarak yorumlanmıştır.
Kavramsal veya metinsel çalışmalar	JBİ Metinsel Kanıt: Uzman Görüşü Kontrol Listesi	4	19 Evet; 0 Hayır; 5 Belirsiz (24 karar)	Kaynak, uzmanlık ve literatür dayanağı genellikle açıktır; analitik üretim ve literatürle uyumsuzlukların gerekçelendirilmesi bazı çalışmalarda belirsizdir.

Not: Kararlar “Evet”, “Hayır” ve “Belirsiz” olarak kaydedilmiş; “Uygulanamaz” kararı bulunmamıştır. MMAT önerileri doğrultusunda toplam/yüzdelik puan hesaplanmamış ve genel kalite sınıflandırması yapılmamıştır.

Üretken yapay zekânın karar destek kapasitesi

İncelenen literatürde üretken yapay zekâ ve büyük dil modellerinin karar verme süreçlerinde temel olarak bilgi işleme ve bilişsel destek sağlama işlevleriyle öne çıktığı görülmektedir. Finansal karar destek çalışmalarında LLM’ler kredi değerlendirme, risk analizi, yatırım fırsatlarının değerlendirilmesi, portföy optimizasyonu ve piyasa duyarlılığı analizi gibi alanlarda kullanılmaktadır (Liu vd., 2024; Peddi, 2025; Tahvildari, 2025; Zhuang ve Wu, 2025). Örneğin Tahvildari (2025), GenAI’nin robot-danışmanlık sistemlerinde portföy optimizasyonu, risk değerlendirme ve yatırım önerilerinin kişiselleştirilmesi bakımından önemli fırsatlar sunduğunu; buna karşın açıklanabilirlik, insan gözetimi ve güven mekanizmalarının zorunlu olduğunu vurgulamaktadır. Zhuang ve Wu (2025) da kurumsal yatırım kararlarında ChatGPT’nin veri analizi, risk değerlendirme ve yatırım fırsatlarının taranması gibi karar öncesi (pre-decision) ve karar sonrası (post-decision) süreçlerde kullanılabileceğini göstermektedir.

Tedarik zinciri yönetimi alanındaki çalışmalar, LLM’lerin tedarikçi iletişimi, sözleşme analizi, talep tahmini, risk izleme, tedarik zinciri görünürlüğü ve karar destek süreçlerinde kullanılabileceğini göstermektedir (Du vd., 2026; Ghosh vd., 2026; Sonar vd., 2026; Song vd., 2026). Du vd. (2026), LLM’lerin bağlamdan bağımsız veya önceden sağlanmış bilgiye dayalı karar görevlerinde insan uzmanlara yakın sonuçlar üretebildiğini ancak bağlama özgü ve örtük uzmanlık gerektiren karar görevlerinde insan uzmanlığının hâlâ belirleyici olduğunu göstermektedir. Sonar vd. (2026) ve Ghosh vd. (2026) ise LLM benimseme zorluklarının yorumlanabilirlik, aşırı güvenme, veri noktalarının dağınıklığı, yasal uyum ve yönetsel güven sorunlarıyla ilişkili olduğunu belirtmektedir. Bu bulgu, LLM’lerin karar verme süreçlerinde görev uygunluğu ve insan-AI rol dağılımının açık biçimde tanımlanması gerektiğine işaret etmektedir.

İş zekâsı ve denetim bağlamındaki çalışmalar, LLM’lerin farklı iş fonksiyonları arasında bütünleşik karar destek mimarileri kurma potansiyeline dikkat çekmektedir. Liu vd. (2026), bankacılık iç kontrol sistemlerinde yapılandırılmamış raporları karar destek göstergelerine dönüştüren, XGBoost ve SHAP ile desteklenen açıklanabilir bir iç kontrol karar destek sistemi önermektedir. Theodorakopoulos vd. (2026) ise pazarlama, finansal performans ve denetim kalitesini bütünleştiren LLM tabanlı iş zekâsı çerçevesi geliştirerek LLM’lerin yalnızca tekil karar alanlarında değil, işlevler arası karar sistemlerinde de kullanılabileceğini göstermektedir.

Güven ve kullanıcı kabulü

Güven, üretken yapay zekâ destekli karar verme sistemlerinin benimsenmesinde merkezi bir temadır. İncelenen çalışmalarda kullanıcıların LLM önerilerini kabul etme veya reddetme davranışlarının

yalnızca model performansına bağlı olmadığı; algılanan fayda, mahremiyet riski, bilgi kontrolü, açıklama kalitesi, sistemin tutarlılığı ve insan gözetiminin varlığı gibi unsurlarla ilişkili olduğu görülmektedir (Huynh, 2024; Pescher, 2026; Sheng vd., 2025; Tahvildari, 2025). İş yeri bağlamında ise kullanıcıların GenAI araçlarını anlamlandırma ve başkalarına aktarma biçimleri, bu araçların “yapay zekâ ekip arkadaşı” (AI teammate) olarak benimsenmesinde belirleyici olmaktadır (Agarwal vd., 2026).

Huynh (2024), örgütsel kullanıcıların üretken yapay zekâ (GenAI) tabanlı karar destek araçlarına güven duymasında algılanan fayda ve algılanan risklerin birlikte rol oynadığını göstermektedir. Bu bağlamda kullanıcılar GenAI’nin sağlayacağı işlevsel, sosyal, duygusal ve epistemik faydaları dikkate almakta; aynı zamanda bilgi hassasiyeti ve mahremiyet riski nedeniyle temkinli davranabilmektedir. Bu bulgu, GenAI destekli karar sistemlerinin yalnızca teknik performans üzerinden değil, kullanıcıların mahremiyet ve kontrol algıları üzerinden de değerlendirilmesi gerektiğini göstermektedir.

Pescher (2026), LLM’lerin tüketici kararlarında “karar eşik bekçisi” (decision gatekeeper) rolü üstlenebildiğini ve ikna edici pazarlama içeriklerine ürün türüne göre farklı tepkiler verebildiğini ortaya koymaktadır. Bu durum LLM aracılı öneri sistemlerinde güven ve karar kalitesinin yalnızca kullanıcı-sistem ilişkisiyle değil, modelin maruz kaldığı içerik ve ikna teknikleriyle de şekillenebileceğini göstermektedir. Bu nedenle GenAI destekli karar sistemlerinde güven, model çıktısının doğruluğu kadar çıktının nasıl biçimlendiği ve hangi dışsal etkilere açık olduğu ile de ilişkilidir.

Açıklanabilirlik ve belirsizlik yönetimi

Açıklanabilirlik, üretken yapay zekâ destekli karar verme sistemlerinde güvenin ve sorumlu kullanımın temel koşullarından biridir. İncelenen çalışmalarda açıklanabilirlik iki düzeyde ele alınmaktadır. Birinci düzey, SHAP, LIME, dikkat (attention) mekanizmaları, belirsizlik raporları, sonradan üretilen (post-hoc) açıklamalar ve doğrulama katmanları gibi teknik mekanizmaları içermektedir. İkinci düzey ise bu açıklamaların insan kullanıcılar tarafından anlaşılabilir doğal dil ifadelerine dönüştürülmesini, karar gerekçelerinin izlenebilir biçimde sunulmasını ve insan karar verici tarafından denetlenebilmesini kapsamaktadır (Buranasomphop, 2025; Hinov ve Ivanova, 2026; Kumar vd., 2025; Liu vd., 2024; Liu vd., 2026; Sheng vd., 2025).

Kumar vd. (2025), kritik uygulamalarda LLM’lerin kara kutu niteliğinin şeffaflık, güven ve hesap verebilirlik açısından önemli sorunlar doğurduğunu belirtmekte ve açıklanabilir LLM karar verme çerçevelerinin gerekli olduğunu savunmaktadır. Liu vd. (2024) ise kredi değerlendirme bağlamında SHAP ve LIME çıktılarının ChatGPT aracılığıyla kullanıcı dostu açıklamalara dönüştürülebileceğini göstermektedir. Buranasomphop (2025) da girişim (startup) finansmanı kararlarında LLM davranış örüntülerinin açıklanabilir biçimde analiz edilmesinin karar şeffaflığı açısından önemli olduğunu ortaya koymaktadır. Bu yaklaşımlar açıklanabilirliğin yalnızca teknik çıktı üretiminden ziyade bu çıktının farklı kullanıcı rolleri için anlaşılır hâle getirilmesiyle ilişkili olduğunu göstermektedir.

Sheng vd. (2025), LLM destekli insan-yapay zekâ karar sistemlerinde belirsizlik raporlarının kullanıcı güveni, karar verimliliği ve karar doğruluğu üzerinde etkili olduğunu ortaya koymaktadır. Çalışmada belirsizlik bilgisinin yönü, belirsizlik derecesi ve görev uzmanlığı gibi faktörlerin kullanıcıların yapay zekâ çıktısını nasıl yorumladığını etkilediği belirtilmektedir. Bu bulgu açıklanabilirlik mekanizmalarının bağlamdan bağımsız ve tek tip biçimde tasarlanamayacağını; kullanıcının uzmanlık düzeyi, görev karmaşıklığı ve karar bağlamına göre uyarlanması gerektiğini göstermektedir.

İnsan gözetimi ve yönetim

Nihai analiz kapsamındaki çalışmalar üretken yapay zekânın karar verme süreçlerinde insan gözetimi olmadan kullanılmasının önemli riskler taşıdığını göstermektedir. İnsanın döngüde olduğu yaklaşım (human-in-the-loop), insanın nihai kontrolü elinde tuttuğu yaklaşım (human-in-command) ve insan-yapay zekâ iş birliği (human-AI collaboration) kavramları literatürde sık biçimde kullanılmaktadır. Bu kavramlar yapay zekâ sistemlerinin karar sürecine katkı sağlamasına rağmen nihai sorumluluğun insan karar vericide kalması gerektiğini vurgulamaktadır (Bajestani vd., 2026; Hinov ve Ivanova, 2026; Sachan vd., 2025; Saup vd., 2026; Zhu vd., 2026).

Saup vd. (2026), üretken yapay zekânın stratejik karar verme süreçlerinde pilot uygulamalardan kurumsal karar sistemlerine geçişini incelemekte ve GenAI çıktılarının resmi karar süreçlerine dâhil edilebilmesi için kalite, kaynak izlenebilirliği, açıklanabilirlik ve hesap verebilirlik kapılarından geçmesi gerektiğini ortaya koymaktadır. Bu çalışma GenAI’nin “otonom bir kâhin” (oracle) olarak değil yönetimle sınırlandırılmış karar destek bileşeni olarak değerlendirilmesi gerektiğini açık biçimde göstermektedir.

Hinov ve Ivanova (2026), algoritmik yönetim sistemlerinde LLM destekli açıklama arayüzlerinin karar çekirdeği, bilişsel açıklama katmanı ve doğrulama/yönetişim katmanı biçiminde ayrıştırılması gerektiğini savunmaktadır. Bu model LLM'lerin karar alma süreçlerinde insan denetimini güçlendiren ve karar gerekçelerini şeffaflaştıran arayüzler olarak konumlandırılması gerektiğini göstermektedir. Sachan vd. (2025) ise yüksek riskli karar alanlarında LLM faaliyetlerinin blokzincir ve IPFS gibi merkeziyetsiz teknolojilerle kaydedilmesini, böylece güvenlik, hesap verebilirlik ve denetlenebilirliğin güçlendirilmesini önermektedir.

Zhu vd. (2026), insan-yapay zekâ ortak tasarım süreçlerinde üretken yapay zekâ kaynaklı verimlilik kazanımlarının insan gözetiminin, açık öncelik belirlemenin ve şeffaf karar desteğinin yerine geçemeyeceğini vurgulamaktadır. Bajestani vd. (2026) de akıllı üretimde LLM tabanlı sistemlerin alan uzmanlığı, gerçek zamanlı karar gereksinimleri ve kalite güvencesi nedeniyle insanın döngüde olduğu yaklaşım (human-in-the-loop) doğrultusunda tasarlanması gerektiğini belirtmektedir. Bu yaklaşımlar üretken yapay zekâ destekli sistemlerde insan ve yapay zekâ rollerinin açık biçimde tanımlanmasının hem etik hem de karar kalitesi açısından kritik olduğunu göstermektedir.

Riskler, sınırlılıklar ve sorumlu kullanım

Üretken yapay zekâ destekli karar verme sistemlerinde öne çıkan riskler halüsinasyon, yanlılık, aşırı güvenme, veri gizliliği, siber güvenlik, alan bilgisi eksikliği, açıklama yanılsaması ve hesap verebilirlik sorunlarıdır (Albashrawi, 2025; Al Khaldy ve Gheraibia, 2025; Ghosh vd., 2026; Kumar vd., 2025; Sonar vd., 2026). Halüsinasyon, LLM'lerin gerçeğe uygun görünmesine rağmen hatalı veya uydurma bilgi üretmesi anlamına gelmektedir. Karar verme bağlamında bu risk, yanlış önerilere, hatalı risk değerlendirmelerine ve kullanıcı güveninin zedelenmesine yol açmaktadır (Al Khaldy ve Gheraibia, 2025; Ghosh vd., 2026; Hinov ve Ivanova, 2026).

Sonar vd. (2026), tedarik zinciri bağlamında LLM çıktılarının anlaşılması, yapay zekâyâ aşırı güvenme riski ve karmaşık optimizasyon problemlerindeki sınırlılıkların en kritik benimseme zorlukları arasında yer aldığını göstermektedir. Ghosh vd. (2026) ise LLM'lerin tedarik zinciri yönetiminde kullanımına ilişkin zorlukları ISM ve MICMAC yaklaşımlarıyla yapılandırmakta; veri güvenliği, halüsinasyon, yasal uyum, gerçek zamanlı veri eksikliği ve yönetsel şüphe gibi faktörlerin birbirleriyle ilişkili olduğunu ortaya koymaktadır.

Al Khaldy ve Gheraibia (2025), LLM'lerin iş kararlarında veri yanlılığı, alan bilgisi eksikliği, yorumlanabilirlik sorunları ve etik riskler nedeniyle dikkatle değerlendirilmesi gerektiğini belirtmektedir. Bu bulgular üretken yapay zekâ destekli karar sistemlerinin yalnızca verimlilik veya hız üzerinden değil, sorumlu kullanım ve risk kontrolü açısından da değerlendirilmesi gerektiğini göstermektedir.

Karar mimarisi ve çok kriterli karar verme yaklaşımları

İncelenen literatürde bazı çalışmalar üretken yapay zekâ ve LLM'leri çok kriterli karar verme yöntemleriyle birlikte kullanmaktadır. TOPSIS, AHP, COCOSO, DEMATEL, ISM, MICMAC ve konuşmalı karar destek sistemleri gibi yaklaşımlar; LLM benimseme zorluklarının önceliklendirilmesi, GenAI stratejilerinin sıralanması, personel seçiminde adayların değerlendirilmesi, tedarik zinciri uygulamalarının analiz edilmesi ve operasyon araştırması modelleme süreçlerinin desteklenmesi gibi farklı alanlarda kullanılmaktadır (Ghosh vd., 2026; Hoque vd., 2026; Kaleta, 2025; Sonar vd., 2026; Zeqiri ve Ben Youssef, 2026; Zhu vd., 2026).

Hoque vd. (2026), LLM tabanlı profil analizi ile Bulanık TOPSIS (Fuzzy-TOPSIS) yöntemini birleştirerek personel seçimi için otomatik ve çok kriterli bir değerlendirme sistemi önermektedir. Zeqiri ve Ben Youssef (2026), sürdürülebilir turizmde GenAI stratejilerini AHP, TOPSIS, PROMETHEE II ve ELECTRE III yöntemleriyle sıralamakta; yönetim ve güven kriterlerini karar modelinin merkezine yerleştirmektedir. Kaleta (2025) ise LLM'lerin operasyon araştırması metodolojisinde matematiksel modelleme, kod üretimi, doğrulama ve açıklama aşamalarına konuşmalı karar destek mantığıyla entegre edilebileceğini göstermektedir. Bu çalışmalar, üretken yapay zekânın salt öneri üreten bir araç olmanın ötesinde karar mimarisinin kurucu unsurlarından biri olarak ele alınabileceğini ortaya koymaktadır.

Araştırma sorularına göre bulguların sentezi

Nihai tematik sentez, araştırma sorularına doğrudan yanıt verecek biçimde yapılandırılmıştır. RQ1 uygulama alanları ve karar işlevleriyle, RQ2 güven ve açıklanabilirlik mekanizmalarıyla, RQ3 insan gözetimi ve yönetimle, RQ4 ise riskler ve gelecek araştırma boşluklarıyla ilişkilendirilmiştir. Araştırma sorularının hangi temalarla ilişkilendirildiği ve her bir araştırma sorusuna bulgular doğrultusunda verilen kısa yanıtlar Tablo 7'de özetlenmektedir.

Tablo 7: Araştırma Sorularına Göre Bulguların Sentezi

Araştırma Sorusu	Bulguların Dayandığı Tema(lar)	Kısa Yanıt
RQ1	Karar destek kapasitesi; karar mimarisi ve ÇKKV	GenAI/LLM destekli karar sistemleri finans, tedarik zinciri, iş zekâsı, kamu yönetimi, insan kaynakları, endüstriyel süreçler, turizm ve stratejik karar alanlarında incelenmektedir.
RQ2	Güven ve kullanıcı kabulü; açıklanabilirlik ve belirsizlik	Güven; algılanan fayda, mahremiyet riski, bilgi kontrolü ve açıklama kalitesiyle; açıklanabilirlik ise açıklanabilir yapay zekâ (XAI), SHAP/LIME açıklamaları, doğal dil açıklamaları ve belirsizlik raporlarıyla ele alınmaktadır.
RQ3	İnsan gözetimi ve yönetim	İnsan gözetimi, insanın döngüde olduğu yaklaşım (human-in-the-loop) ve insanın nihai kontrolü elinde tuttuğu yaklaşım (human-in-command) ile; yönetim ise denetim izi, hesap verebilirlik, uzman doğrulaması ve politika kısıtlarıyla karar süreçlerine entegre edilmektedir.
RQ4	Riskler ve sorumlu kullanım	Halüsinasyon, yanlışlık, veri gizliliği, aşırı güvenme, siber güvenlik ve sorumluluk belirsizliği temel risklerdir; gelecek çalışmalar için ampirik doğrulama ve yönetim modelleri öne çıkmaktadır.

Tartışma

Nihai senteze dâhil edilen çalışmalar birlikte değerlendirildiğinde, üretken yapay zekâ destekli karar verme literatürünün birbiriyle ilişkili üç temel eksende geliştiği görülmektedir. İlk eksen, GenAI ve büyük dil modellerinin bilgi işleme, analiz, özetleme, alternatif üretme ve senaryo geliştirme gibi bilişsel karar destek işlevlerine odaklanmaktadır (Albashrawi, 2025; Agarwal vd., 2026; Gupta ve Kumar, 2026; Hao vd., 2024; Zong vd., 2026). İkinci eksen; kullanıcı güveni, yapay zekâyâ yönelik olumlu değerlendirme, mahremiyet riski, belirsizlik algısı ve açıklanabilirlik gibi insan merkezli boyutları kapsamaktadır (Huynh, 2024; Pescher, 2026; Sheng vd., 2025; Tahvildari, 2025). Üçüncü eksen ise insan gözetimi, yönetim, hesap verebilirlik, veri güvenliği, düzenleyici uyum ve sorumlu kullanım mekanizmalarını ele almaktadır (Goyal vd., 2026; Hinov ve Ivanova, 2026; Sachan vd., 2025; Saup vd., 2026). Bu eksenler doğrultusunda belirlenen altı tema, birbirinden bağımsız betimleyici kategoriler olmaktan ziyade üretken yapay zekâ destekli karar sistemlerinin teknik, insani ve kurumsal boyutlarını birlikte açıklayan ilişkisel bir yapı oluşturmaktadır.

Bulgular, üretken yapay zekâ ve büyük dil modellerinin karar verme süreçlerinde giderek daha merkezi bir rol kazandığını ancak bu rolün insan karar vericinin yerine geçen otonom bir sistemden çok insan yargısını destekleyen bilişsel bir karar destek bileşeni olarak kavramsallaştırıldığını göstermektedir (Albashrawi, 2025; Hao vd., 2024; Saup vd., 2026; Zhu vd., 2026). Bu sonuç üretken yapay zekânın karar süreçlerine etkisinin yalnızca teknik kapasite, hız veya model başarımı üzerinden açıklanamayacağını ortaya koymaktadır. Güven, açıklanabilirlik, insan gözetimi, etik sorumluluk ve kurumsal yönetim gibi sosyo-teknik unsurların birlikte değerlendirilmesi, bu sistemlerin karar kalitesine gerçek anlamda katkı sağlayabilmesinin temel koşuludur (Hinov ve Ivanova, 2026; Sheng vd., 2025; Tahvildari, 2025).

İncelenen çalışmaların finans, tedarik zinciri, iş zekâsı, kamu yönetimi, insan kaynakları ve endüstriyel süreçler gibi farklı alanlarda yoğunlaşması, üretken yapay zekânın karar verme bağlamında disiplinler arası bir araştırma alanı hâline geldiğini göstermektedir (Ghosh vd., 2026; Goyal vd., 2026; Hoque vd., 2026; Liu vd., 2024; Song vd., 2026; Theodorakopoulos vd., 2026; Zong vd., 2026). Bununla birlikte yöntemsel kalite değerlendirmesi, mevcut kanıt tabanının henüz bütünüyle olgunlaşmadığına işaret etmektedir. Çalışmaların önemli bir bölümü kavramsal çerçeve, pilot uygulama, deneysel tasarım veya sistematik derleme niteliğindedir (Kim vd., 2026; Saup vd., 2026; Song vd., 2026; Tahvildari, 2025). Bu nedenle alandaki bulgular umut verici olmakla birlikte farklı sektörlerde, daha geniş örneklemelerle ve gerçek karar ortamlarında yürütülecek karşılaştırmalı ve boylamsal araştırmalarla desteklenmelidir.

Açıklanabilirlik ve belirsizlik yönetimi, teknik kapasitenin kullanıcı güvenine ve karar kalitesine dönüşmesinde temel güvence mekanizmaları olarak öne çıkmaktadır. Kullanıcıların yalnızca bir öneriye erişmesi yeterli değildir; önerinin hangi veri ve gerekçelere dayandığını anlayabilmesi, belirsizlik düzeyini değerlendirebilmesi ve gerektiğinde sistem çıktısını sorgulayabilmesi gerekir (Hinov ve Ivanova, 2026; Liu vd., 2024; Sheng vd., 2025). Bu nedenle açıklanabilirlik, insan merkezli karar destek sistemlerinin tasarımında teknik bir özelliğin ötesinde kullanıcı deneyimini ve yönetim süreçlerini şekillendiren temel bir ilke olarak değerlendirilmelidir (Kumar vd., 2025; Saup vd., 2026). Bu nedenle amaç, kullanıcı güvenliğini en üst düzeye çıkarmak değil; güven düzeyini görevin risk düzeyi, kanıtların niteliği ve sistemin gerçek performansı ile uyumlu biçimde kalibre etmektir.

İnsan gözetimi ve yönetim temaları da literatürde güçlü biçimde öne çıkmaktadır. Özellikle yüksek etkili karar süreçlerinde LLM çıktılarının doğrudan uygulanması yerine, insan uzmanlar tarafından kontrol edilmesi, doğrulanması ve gerekçelendirilmesi gerekmektedir (Bajestani vd., 2026; Sachan vd., 2025; Saup vd., 2026; Zhu vd., 2026). Bu yaklaşım insanın yalnızca sembolik olarak sürece dâhil edilmesinden farklı olarak, öneriyi kabul etme, reddetme, değiştirme veya geçersiz kılma yetkisini korumasını gerektirir. Böylece insan gözetimi karar düzeyinde sorumluluğu korurken, kurumsal yönetim rol dağılımı, denetim izi, uyum, olay yönetimi ve izleme mekanizmaları aracılığıyla hesap verebilirliği örgütsel düzeye taşımaktadır (Goyal vd., 2026; Hinov ve Ivanova, 2026).

Çalışmanın bulguları uluslararası yapay zekâ yönetimi çerçeveleriyle de önemli ölçüde örtüşmektedir. OECD Yapay Zekâ İlkeleri insan merkezli değerler, şeffaflık, açıklanabilirlik, güvenlik ve hesap verebilirliği güvenilir yapay zekânın temel unsurları olarak ele almaktadır (Organisation for Economic Co-operation and Development [OECD], 2024). NIST Yapay Zekâ Risk Yönetimi Çerçevesi yönetim, bağlamın haritalanması, risklerin ölçülmesi ve yönetilmesi işlevlerini bütüncül bir yaşam döngüsü yaklaşımı içinde yapılandırmaktadır (Tabassi, 2023). NIST'in Üretken Yapay Zekâ Profili ise doğruluk, bilgi bütünlüğü, mahremiyet, güvenlik, insan-yapay zekâ etkileşimi ve yanlılık gibi GenAI'ye özgü risk alanlarını ayrıntılandırmaktadır (Autio vd., 2024). Avrupa Birliği Yapay Zekâ Tüzüğü, yüksek riskli sistemlerde teknik dokümantasyon, kayıt tutma, şeffaflık ve etkili insan gözetimi gerekliliklerini düzenleyerek çalışmanın insan otoritesi ve denetlenebilirlik bulgularını normatif düzeyde desteklemektedir (European Parliament & Council of the European Union, 2024). ISO/IEC 42001:2023 ise politika oluşturma, sorumlulukların belirlenmesi, risklerin yönetilmesi, performansın izlenmesi ve sürekli iyileştirme yoluyla yapay zekâ yönetiminin kurumsallaştırılmasına yönelik bir yönetim sistemi yaklaşımı sunmaktadır (International Organization for Standardization [ISO], 2023).

Bu çalışmanın literatüre temel katkısı, üretken yapay zekâ destekli karar verme araştırmalarında belirlenen altı temayı birbirinden bağımsız kategoriler olarak sunmak yerine bu temalar arasındaki ilişkileri bütünleştirici bir yaklaşımla açıklamasıdır. Bulgular üretken yapay zekânın bilgi sentezi, alternatif üretimi, risk değerlendirmesi ve senaryo geliştirme gibi karar destek işlevlerinin tek başına karar kalitesi açısından yeterli olmadığını; bu işlevlerin açıklanabilirlik, belirsizlik yönetimi, kaynak doğrulama, insan gözetimi ve kurumsal hesap verebilirlik mekanizmalarıyla birlikte değerlendirilmesi gerektiğini göstermektedir. Bu doğrultuda çalışma üretken yapay zekânın teknik işlevleri ile insan ve kurum düzeyindeki yönetim gereksinimlerini aynı analitik yapı içerisinde ilişkilendiren bütünleştirici bir kavramsal çerçeve sunmaktadır. Çerçeve; karar bağlamı ve risklerin haritalanması, üretken yapay zekâ karar destek işlevleri, güvence mekanizmaları, insan otoritesi ve karar denetimi, kurumsal yönetim ve hesap verebilirlik ile izleme, öğrenme ve geri bildirim bileşenlerinden oluşmaktadır. Bu bileşenler katı ve doğrusal aşamalar olarak değil, karar bağlamı ve risk düzeyine göre birbirleriyle etkileşim içinde çalışan tamamlayıcı unsurlar olarak ele alınmaktadır. Çerçeve, ampirik olarak doğrulanmış bir modelden ziyade sistematik inceleme bulguları ile uluslararası yapay zekâ yönetimi ilkelerinin birlikte yorumlanmasına dayanan kavramsal bir sentez niteliğindedir. Çerçeveyi oluşturan bileşenler, karar sürecindeki işlevleri ve başlıca literatür ve yönetim dayanakları Tablo 8'de sunulmaktadır.

Tablo 8: Üretken Yapay Zekâ Destekli Karar Vermeye İlişkin Bütünleştirici Kavramsal Çerçevenin Bileşenleri

Bileşen	Karar sürecindeki işlevi	Literatür ve yönetim dayanağı
Karar bağlamı ve risklerin haritalanması	Kararın etki düzeyi, belirsizliği, veri hassasiyeti, ilgili paydaşlar ve kabul edilebilir risk düzeyi belirlenir.	NIST MAP; AB Yapay Zekâ Tüzüğü'nün risk temelli yaklaşımı
GenAI karar destek işlevleri	Bilgi sentezi, alternatif üretimi, senaryo geliştirme, risk değerlendirmesi ve doğal dilde gerekçelendirme gibi karar destek işlevleri yürütülür.	İnceleme bulguları: karar destek kapasitesi ve karar mimarisi
Güvence mekanizmaları	Kaynak doğrulama, açıklanabilirlik, belirsizlik bildirimi, izlenebilirlik, test, değerlendirme ve performans kontrolü uygulanır.	NIST MEASURE ve GenAI Profili; OECD şeffaflık ve açıklanabilirlik ilkesi, açıklanabilir yapay zekâ literatürü
İnsan otoritesi ve karar denetimi	Özellikle yüksek etkili kararlarda uzman incelemesi, itiraz, düzeltme ve geçersiz kılma mekanizmaları aracılığıyla insanın karar sürecindeki yetkisi korunur.	İnsan gözetimi literatürü; insanın döngüde olduğu ve nihai kontrolü elinde tuttuğu yaklaşımlar; AB Yapay Zekâ Tüzüğü'nde etkili insan gözetimi
Kurumsal yönetim ve hesap verebilirlik	Politikalar, görev ve sorumluluklar, denetim izleri, uyum süreçleri, olay bildirimleri ve hesap verebilirlik mekanizmaları oluşturulur.	NIST GOVERN; ISO/IEC 42001; OECD hesap verebilirlik ilkesi
İzleme, öğrenme ve geri bildirim	Karar sonuçları ve sistem performansı izlenir; hatalar, sapmalar ve yeni riskler modele, karar sürecine ve kurumsal politikalara geri beslenir.	NIST MANAGE ve GOVERN; ISO/IEC 42001 sürekli iyileştirme yaklaşımı

Not: Bileşenler katı ve doğrusal bir işlem sırasını temsil etmemektedir. Karar bağlamına ve risk düzeyine bağlı olarak bileşenler birbirleriyle etkileşimli, yinelemeli ve yaşam döngüsü boyunca sürekli biçimde uygulanabilir.

Sonuç

Bu çalışma, üretken yapay zekâ ve büyük dil modeli destekli karar verme sistemlerine ilişkin güncel literatürü PRISMA 2020 doğrultusunda sistematik biçimde inceleyerek alanın uygulama alanlarını, yöntemsel özelliklerini, risklerini ve yönetim gereksinimlerini ortaya koymuştur. Web of Science ve Scopus taramasında uygunluk ölçütlerini karşılayan 31 çalışma senteze dâhil edilmiştir. Çalışmaların 2024–2026 döneminde yoğunlaşması, alanın hızla geliştiğini ve farklı sektörlere yayıldığını göstermektedir.

Bulgular üretken yapay zekâ ve büyük dil modellerinin finans ve bankacılık, tedarik zinciri yönetimi, iş zekâsı, kamu yönetimi, insan kaynakları, stratejik yönetim, endüstriyel süreçler, turizm ve tüketici kararları gibi farklı alanlarda karar desteği sağlamak amacıyla kullanıldığını ortaya koymuştur. Bu sistemlerin başlıca katkıları; yapılandırılmış ve yapılandırılmamış bilgilerin sentezlenmesi, karar alternatiflerinin oluşturulması, senaryo geliştirilmesi, risklerin değerlendirilmesi, karmaşık çıktıların doğal dilde açıklanması ve karar vericilerin bilişsel yükünün azaltılmasıdır. Bununla birlikte üretken yapay zekânın sunduğu teknik kapasitenin tek başına karar kalitesini güvence altına almadığı görülmüştür. Sistem çıktılarının güvenilir ve sorumlu biçimde kullanılabilmesi; açıklanabilirlik, belirsizlik bildirimi, kaynak doğrulama, izlenebilirlik, insan gözetimi ve kurumsal hesap verebilirlik mekanizmalarının birlikte işletilmesine bağlıdır.

Tematik sentez; karar destek kapasitesi, güven ve kullanıcı kabulü, açıklanabilirlik ve belirsizlik yönetimi, insan gözetimi ve yönetim, riskler ve sorumlu kullanım ile karar mimarisi ve çok kriterli karar verme olmak üzere altı tema ortaya koymuştur. Temalar arasındaki ilişkiler, teknik kapasitenin açıklanabilirlik, insan gözetimi ve kurumsal hesap verebilirlikle birlikte ele alınması gerektiğini göstermektedir.

MMAT 2018 ve JBI kontrol listeleriyle gerçekleştirilen madde bazlı değerlendirmede 15 çalışmada en az bir “Hayır”, 24 çalışmada en az bir “Belirsiz” kararı bulunmuş; toplam kalite puanı veya genel kalite derecesi üretilmemiştir. Tekrarlanan yöntemsel sınırlılıklar, bulguların daha geniş örneklemelerle ve gerçek karar ortamlarında yürütülecek güçlü araştırmalarla sınanması gerektiğini göstermektedir.

Çalışmanın temel katkısı, altı tema arasındaki teknik, insani ve kurumsal ilişkileri Üretken Yapay Zekâ Destekli Karar Vermeye İlişkin Bütünleştirici Kavramsal Çerçeve’de birleştirmesidir. Çerçeve; karar bağlamı ve risklerin haritalanması, karar destek işlevleri, güvence mekanizmaları, insan otoritesi ve

karar denetimi, kurumsal yönetim ve hesap verebilirlik ile izleme, öğrenme ve geri bildirim bileşenlerinden oluşmaktadır. Bu yapı, sistematik inceleme bulguları ile uluslararası yönetim yaklaşımlarına dayanan ve ampirik olarak sınanması gereken kavramsal bir sentezdir.

Uygulama açısından kurumların üretken yapay zekâ sistemlerini karar süreçlerine doğrudan ve sınırsız biçimde dâhil etmek yerine öncelikle kararın etki düzeyini, veri hassasiyetini, paydaşları ve kabul edilebilir risk sınırlarını belirlemeleri gerekmektedir. Kaynak doğrulama, belirsizlik bildirimi, performans izleme, denetim izi, itiraz, düzeltme ve geçersiz kılma mekanizmaları sistem tasarımının ayrılmaz unsurları olarak ele alınmalıdır. Özellikle yüksek etkili karar alanlarında nihai karar yetkisi ve sorumluluğu insan karar vericide kalmalı; üretken yapay zekâ çıktıları doğrulama ve gerekçelendirme süreçlerinden geçirilmeden uygulanmamalıdır. Kurumsal politika, rol dağılımı, hesap verebilirlik, olay bildirimi ve sürekli iyileştirme mekanizmaları ise sistemlerin bütün yaşam döngüsü boyunca sürdürülmelidir.

Çalışmanın bazı sınırlılıkları bulunmaktadır. Literatür taramasının Web of Science ve Scopus veri tabanlarıyla, İngilizce yayımlanan çalışmalarla ve 2020–2026 dönemiyle sınırlandırılması, ilgili bazı yayınların kapsam dışında kalmasına neden olmuş olabilir. Ayrıca arama stratejisinde güven, açıklanabilirlik, insan gözetimi, yönetim, hesap verebilirlik ve karar kalitesi terimlerinin kullanılması, nihai çalışma kümesinin bu kavramsal boyutlara odaklanan yayınlar yönünde şekillenmesine yol açmış olabilir. Bu nedenle tematik dağılım, üretken yapay zekâ destekli karar verme literatürünün tamamını değil, çalışmanın önceden tanımlanan yönetim odaklı kapsamını yansıtmaktadır. Tam metnine erişilemeyen 11 çalışma nihai değerlendirmeye alınamamıştır. İnceleme protokolünün araştırma başlamadan önce ileriye dönük olarak kaydedilmemiş olması, sonradan alınmış yöntemsel kararların bütünüyle dışlanmasını güçleştirmektedir. Tarama, uygunluk değerlendirmesi, kodlama ve yöntemsel değerlendirme ortak ve uzlaşıya dayalı biçimde yürütülmüş; bağımsız paralel değerlendirme yapılmadığı için kodlayıcılar arası uyum katsayısı hesaplanmamıştır. Ayrıca ikinci kapsam değerlendirmesinde dışlanan 50 kaydın kayıt düzeyindeki gerekçelerinin korunmamış olması bu aşamanın denetim izini sınırlandırmaktadır. Dâhil edilen çalışmaların yöntem, bağlam ve yayın türü bakımından heterojen olması da bulguların genellenebilirliğini azaltmaktadır. Bununla birlikte arama tarihinin, sorguların, dâhil etme ve dışlama ölçütlerinin, aşama bazındaki sayısal sonuçların, kodlama alanlarının ve tam metne dayalı madde bazlı değerlendirme yaklaşımının raporlanması inceleme sürecinin şeffaflığını ve izlenebilirliğini güçlendirmiştir.

Gelecek araştırmalarda üretken yapay zekâ destekli karar sistemlerinin karar doğruluğu, karar süresi, bilişsel yük, kullanıcı güveni ve hesap verebilirlik üzerindeki etkileri gerçek örgütsel ortamlarda incelenmelidir. Açıklanabilirlik ve belirsizlik bildirimi mekanizmalarının farklı kullanıcı grupları, uzmanlık düzeyleri ve görev türleri üzerindeki etkileri karşılaştırmalı deneylerle test edilmelidir. İnsan gözetimi yaklaşımlarının hangi karar bağlamlarında daha etkili olduğu, insan müdahalesinin sembolik bir denetimden işlevsel karar yetkisine nasıl dönüştürülebileceği ve aşırı güvenmenin nasıl sınırlandırılabilceği araştırılmalıdır. Ayrıca bu çalışmada sunulan bütünleştirici kavramsal çerçevenin farklı sektörlerde ampirik olarak değerlendirilmesi, bileşenler arasındaki ilişkilerin ve bağlamsal koşulların daha açık biçimde ortaya konulmasına katkı sağlayacaktır.

Sonuç olarak üretken yapay zekâ, bağımsız ve otonom bir karar aktörü olarak konumlandırılmaktan ziyade; insan yargısını destekleyen, açıklanabilirlik ve denetlenebilirlik ilkelerine dayanan, insan gözetimi ve kurumsal yönetim çerçevesinde kullanılan bilişsel bir karar destek bileşeni olarak ele alınmalıdır. Üretken yapay zekânın karar süreçlerine sürdürülebilir ve sorumlu biçimde katkı sağlayabilmesi, teknik performans, güvence mekanizmaları, insan otoritesi ve kurumsal hesap verebilirlik arasında dengeli ve sürekli bir yapı kurulmasına bağlıdır.

Hakem Değerlendirmesi / Peer-review:

Dış bağımsız hakem değerlendirmesi uygulanmıştır.

Externally peer-reviewed

Çıkar Çatışması / Conflict of interests:

Yazarlar herhangi bir çıkar çatışması bulunmadığını beyan etmişlerdir.

The authors declare that they have no conflict of interest.

Etik Beyan / Ethical Statement:

Bu çalışma, kamuya açık yayımlanmış bilimsel kaynaklara dayalı bir sistematik literatür incelemesi olduğundan etik kurul izni gerektirmemektedir.

This study is a systematic literature review based on publicly available published scientific sources; therefore, ethics committee approval was not required.

Veri Erişilebilirliği Beyanı / Data Availability Statement:

Bu çalışmada birincil araştırma verisi üretilmemiştir. İnceleme kapsamında kullanılan veriler makalede ve kaynakçada sunulmuştur.

No primary research data were generated in this study. The data used in the review are presented in the article and its reference list.

Üretken Yapay Zekâ Kullanım Beyanı / Generative Artificial Intelligence Use Statement

Bu makalenin revizyon sürecinde OpenAI ChatGPT (GPT-5.6); dil ve anlatımın iyileştirilmesi, biçimsel düzenleme ve İngilizce çeviri desteği amacıyla kullanılmıştır. Yapay zekâ çıktıları yazarlar tarafından kontrol edilmiş ve nihai içerik ile bilimsel sorumluluk yazarlara aittir.

During the revision of this manuscript, OpenAI ChatGPT (GPT-5.6) was used to support language improvement, formal editing, and English translation. All AI-generated outputs were reviewed by the authors, who retain full responsibility for the final content and scientific integrity of the manuscript.

Finansal Destek / Grant Support:

Yazarlar bu çalışma için herhangi bir finansal destek almadıklarını beyan etmişlerdir.

The authors declare that this study received no financial support.

Teşekkür / Acknowledgement:

Yazarlar, çalışmanın geliştirilmesine değerli görüş ve önerileriyle katkı sağlayan akademisyenlere ve yapıcı değerlendirmeleri için hakemlere teşekkür ederler.

The authors thank the academics who contributed valuable comments and suggestions to the development of the study and the reviewers for their constructive feedback.

Yazar Katkıları / Author Contributions:

Fikir/Kavram/Tasarım - Idea/Concept/Design: F.Y., Z.S. Yöntem - Methodology: F.Y., Z.S. Veri Toplama ve/veya İşleme - Data Collection and/or Processing: F.Y. Analiz ve/veya Yorum - Analysis and/or Interpretation: F.Y., Z.S. Makalenin Yazımı - Writing the Article: F.Y., Z.S. Eleştirel İnceleme - Critical Review: F.Y., Z.S. Danışmanlık - Supervision: Z.S. Onay - Approval: F.Y., Z.S.

Kaynakça / References

- Agarwal, A., Sebastian, M. P., & Krishnan, S. (2026). Understanding GenAI teammates in the workplace: A sensemaking and sensegiving analysis of user reviews. *Information Systems Frontiers*, 28, 803-825. <https://doi.org/10.1007/s10796-025-10672-5>
- Al Khaldy, M., & Gheraibia, Y. (2025). Evaluation and mitigation of the limitations of large language models in business decision-making. In *2025 1st International Conference on Computational Intelligence Approaches and Applications (ICCIAA)* (ss. 1-6). IEEE. <https://doi.org/10.1109/ICCIAA65327.2025.11013278>
- Albashrawi, M. (2025). Generative AI for decision-making: A multidisciplinary perspective. *Journal of Innovation & Knowledge*, 10(4), 100751. <https://doi.org/10.1016/j.jik.2025.100751>
- Aromataris, E., Lockwood, C., Porritt, K., Pilla, B., & Jordan, Z. (Eds.). (2024). *JBİ manual for evidence synthesis*. JBİ. <https://doi.org/10.46658/JBİMES-24-01>

- Autio, C., Schwartz, R., Dunietz, J., Jain, S., Stanley, M., Tabassi, E., Hall, P., & Roberts, K. (2024). *Artificial intelligence risk management framework: Generative artificial intelligence profile* (NIST AI 600-1). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.AI.600-1>
- Bajestani, M. S., Mun, D., & Kim, D. B. (2026). Human-in-the-loop and large language models in smart manufacturing: Current applications, challenges, and perspectives. *Journal of Manufacturing Systems*, 86, 913-941. <https://doi.org/10.1016/j.jmsy.2026.04.027>
- Buranasomphop, D. (2025). Explainable patterns of LLM funding behavior in startup funding decisions. In *2025 IEEE International Conference on Information Reuse and Integration and Data Science (IRI)* (ss. 256-259). IEEE. <https://doi.org/10.1109/IRI66576.2025.00056>
- Centre for Reviews and Dissemination. (t.y.). *PROSPERO*. University of York. 30 Temmuz 2026 tarihinde <https://www.crd.york.ac.uk/prospero/> adresinden erişildi.
- Du, B., Chen, G., Xie, S., & Geng, D. (2026). Assessing large language models for decision support in novel circular supply chains: A case of hotel linen. *IEEE Transactions on Engineering Management*, 73, 2397-2425. <https://doi.org/10.1109/TEM.2026.3667718>
- European Parliament & Council of the European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence. *Official Journal of the European Union*, L, 2024/1689. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- Ghosh, S. K., Virmani, N., & Tripathi, P. M. (2026). Exploring the challenges of using large language models in supply chain management: A decision-making framework. *LogForum*, 22(2), 229-248. <https://doi.org/10.17270/J.LOG.001417>
- Goyal, R., Deshmukh, S. G., & Bolia, N. (2026). Navigating barriers to GenAI adoption in public administration: A systematic evaluation and policy roadmap. *Socio-Economic Planning Sciences*, 105, 102461. <https://doi.org/10.1016/j.seps.2026.102461>
- Gupta, C. P., & Kumar, V. V. R. (2026). Enhancing decision intelligence through generative AI: A framework for data-driven strategic planning. In *2026 International Conference on Communication, Computing and Emerging Technologies (IC3ET)* (ss. 596-601). IEEE. <https://doi.org/10.1109/IC3ET64989.2026.11467400>
- Hao, X., Demir, E., & Eyers, D. (2024). Exploring collaborative decision-making: A quasi-experimental study of human and generative AI interaction. *Technology in Society*, 78, 102662. <https://doi.org/10.1016/j.techsoc.2024.102662>
- Hinov, N., & Ivanova, M. (2026). LLM-augmented algorithmic management: A governance-oriented architecture for explainable organizational decision systems. *AI*, 7(3), 102. <https://doi.org/10.3390/ai7030102>
- Hong, Q. N., Fàbregues, S., Bartlett, G., Boardman, F., Cargo, M., Dagenais, P., Gagnon, M.-P., Griffiths, F., Nicolau, B., O’Cathain, A., Rousseau, M.-C., Vedel, I., & Pluye, P. (2018). The Mixed Methods Appraisal Tool (MMAT) version 2018 for information professionals and researchers. *Education for Information*, 34(4), 285-291. <https://doi.org/10.3233/EFI-180221>
- Hoque, S., Karim, A. A. J., Alam, M. G. R., & Gope, N. (2026). When LLM meets Fuzzy-TOPSIS for personnel selection through automated profile analysis. *IEEE Access*, 14, 15778-15794. <https://doi.org/10.1109/ACCESS.2026.3658575>
- Huynh, M.-T. (2024). Using generative AI as decision-support tools: Unraveling users’ trust and AI appreciation. *Journal of Decision Systems*, 1-32. <https://doi.org/10.1080/12460125.2024.2428166>
- International Organization for Standardization. (2023). *ISO/IEC 42001:2023 information technology – Artificial intelligence – Management system*. <https://www.iso.org/standard/42001.html>
- Kaleta, M. (2025). Toward conversational decision support systems: Integrating LLMs in the operations research methodology. In *Proceedings of the 20th Conference on Computer Science and Intelligence Systems* (ss. 165-174). <https://doi.org/10.15439/2025F5420>
- Kim, D., Kang, S., & Hong, A. (2026). Bridging the maturity-expectation gap: Generative AI in strategic decision-making for public R&D interim review. *Technovation*, 149, 103374. <https://doi.org/10.1016/j.technovation.2025.103374>

- Kumar, S., Praveen, R. V. S., Aida, R., Varshney, N., Alsalami, Z., & Boob, N. S. (2025). Enhancing AI decision-making with explainable large language models (LLMs) in critical applications. In 2025 IEEE International Conference on Advances in Computing Research on Science Engineering and Technology (ACROSET) (ss. 1-6). IEEE. <https://doi.org/10.1109/ACROSET66531.2025.11280656>
- Liu, Y., Li, X., & Su, C. (2026). From unstructured text to automated insights: An explainable AI approach to internal control in banking systems. *Systems*, 14(3), 234. <https://doi.org/10.3390/systems14030234>
- Liu, Z., He, H., Zhang, L., & Peng, C. (2024). Leveraging XAI in prompt-based ChatGPT for financial decision support. In *International Conference on Digital Economy, Blockchain and Artificial Intelligence (DEBAI 2024)* (ss. 255-259). Association for Computing Machinery. <https://doi.org/10.1145/3700058.3700099>
- McArthur, A., Cooper, A., Edwards, D., Klugarova, J., Yan, H., Barber, B. V., Gregg, E., Weeks, L., & Jordan, Z. (2025). Textual evidence systematic reviews series paper 3: Critical appraisal of evidence from narrative, opinion, and policy. *JBI Evidence Synthesis*, 23(5), 833-839. <https://doi.org/10.11124/JBIES-24-00293>
- Moher, D., Shamseer, L., Clarke, M., Ghersi, D., Liberati, A., Petticrew, M., Shekelle, P., Stewart, L. A., & PRISMA-P Group. (2015). Preferred reporting items for systematic review and meta-analysis protocols (PRISMA-P) 2015 statement. *Systematic Reviews*, 4, 1. <https://doi.org/10.1186/2046-4053-4-1>
- Mongeon, P., & Paul-Hus, A. (2016). The journal coverage of Web of Science and Scopus: A comparative analysis. *Scientometrics*, 106(1), 213-228. <https://doi.org/10.1007/s11192-015-1765-5>
- Organisation for Economic Co-operation and Development. (2024). *Recommendation of the Council on Artificial Intelligence* (OECD/LEGAL/0449). OECD Legal Instruments. <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>
- Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., Shamseer, L., Tetzlaff, J. M., Akl, E. A., Brennan, S. E., Chou, R., Glanville, J., Grimshaw, J. M., Hróbjartsson, A., Lalu, M. M., Li, T., Loder, E. W., Mayo-Wilson, E., McDonald, S., . . . Moher, D. (2021). The PRISMA 2020 statement: An updated guideline for reporting systematic reviews. *BMJ*, 372, Article n71. <https://doi.org/10.1136/bmj.n71>
- Peddi, V. S. K. K. (2025). Harnessing large language models for risk assessment and decision support in the banking sector. In 2025 IEEE Silchar Subsection Conference (SILCON) (ss. 1-5). IEEE. <https://doi.org/10.1109/SILCON67893.2025.11327012>
- Pescher, A. (2026). Decision gatekeepers: How hedonic vs. utilitarian products shape differential persuasion effects in LLM-mediated recommendations. *Journal of Decision Systems*, 35(1), 2649245. <https://doi.org/10.1080/12460125.2026.2649245>
- PRISMA Statement. (t.y.). *PRISMA 2020 flow diagram*. 29 Temmuz 2026 tarihinde <https://www.prisma-statement.org/prisma-2020-flow-diagram> adresinden erişildi.
- Rethlefsen, M. L., Kirtley, S., Waffenschmidt, S., Ayala, A. P., Moher, D., Page, M. J., Koffel, J. B., & PRISMA-S Group. (2021). PRISMA-S: An extension to the PRISMA Statement for Reporting Literature Searches in Systematic Reviews. *Systematic Reviews*, 10, 39. <https://doi.org/10.1186/s13643-020-01542-z>
- Sachan, S., Miller, T., & Nguyen, M. P. (2025). Responsible LLM deployment for high-stake decisions by decentralized technologies and human-AI interactions. In 2025 IEEE 5th International Conference on Human-Machine Systems (ICHMS) (ss. 99-104). IEEE. <https://doi.org/10.1109/ICHMS65439.2025.11154208>
- Saup, T.-O., Asghar, J., Kanbach, D. K., & Kraus, S. (2026). From pilots to decision systems: Embedding generative AI into strategic decision-making through a socio-technical and governance lens. *Journal of Decision Systems*, 35(1), 2597835. <https://doi.org/10.1080/12460125.2025.2597835>
- Schulte, N., & Kanbach, D. K. (2026). Rethinking organizational decision-making: The emerging roles and tasks of generative artificial intelligence. *Management Review Quarterly*. <https://doi.org/10.1007/s11301-026-00611-2>

- Sheng, L., Chang, F., Sun, Q., Wangzha, D., & Gu, Z. (2025). Uncertainty reports as explainable AI: A cognitive-adaptive framework for human-AI decision systems in context tasks. *Advanced Engineering Informatics*, 68, 103634. <https://doi.org/10.1016/j.aei.2025.103634>
- Sonar, H., Ghag, N., & Sharma, I. (2026). Bridging theory and practice in AI-driven supply chains: Prioritizing LLM adoption challenges and SCOR-based applications. *International Journal of Production Economics*, 296, 110008. <https://doi.org/10.1016/j.ijpe.2026.110008>
- Song, Z., Xie, Y., Yang, L., & Zhao, Y. (2026). Large language models in supply chain management: A systematic literature review and application framework. *International Journal of Production Research*, 64(16), 7144-7184. <https://doi.org/10.1080/00207543.2026.2641103>
- Tabassi, E. (2023). *Artificial intelligence risk management framework (AI RMF 1.0)* (NIST AI 100-1). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.AI.100-1>
- Tahvildari, M. (2025). Integrating generative AI in robo-advisory: A systematic review of opportunities, challenges, and strategic solutions. *Multidisciplinary Reviews*, 8, e2025379. <https://doi.org/10.31893/multirev.2025379>
- Theodorakopoulos, L., Karras, A., Theodoropoulou, A., & Klavdianos, C. (2026). LLMs for integrated business intelligence: A Big Data-driven framework integrating marketing optimization, financial performance, and audit quality. *Big Data and Cognitive Computing*, 10(4), 110. <https://doi.org/10.3390/bdcc10040110>
- Zeqiri, A., & Ben Youssef, A. (2026). Not just hype: A multi-criteria decision framework for ranking generative AI strategies in sustainable tourism. *Annals of Operations Research*. <https://doi.org/10.1007/s10479-026-07223-9>
- Zhu, L., Xie, Y., Xiang, N., & Chen, G. (2026). An explainable HCI-based decision support framework for human-AI co-design. *Applied Sciences*, 16(8), 4007. <https://doi.org/10.3390/app16084007>
- Zhuang, X., & Wu, Y. (2025). Large language models in corporate investment: An analysis of ChatGPT's transformative applications, strategic impacts, and future opportunities. *IT Professional*, 27(2), 21-27. <https://doi.org/10.1109/MITP.2025.3543555>
- Zong, Y., Jia, R., Li, K., Xue, D., Zhang, L., & He, D. (2026). LLM-driven human-AI collaborative decision support system for complex industrial processes: A case study in metallurgy. *Neural Networks*, 202, 109055. <https://doi.org/10.1016/j.neunet.2026.109055>