

Yapay Zekâ Sistemlerinin Türkiye'nin Güvenlik Meselelerine Yaklaşımı: Büyük Dil Modellerinin İletişimsel Analizi

An Approach to Türkiye's Security Matters through Artificial Intelligence Systems: A Communicative Analysis of Large Language Models

Araştırma Makalesi / Research Article



Sorumlu yazar/
Corresponding author:
Buğra Ayan

ORCID:
0000-0002-4915-4970

Geliş tarihi/Received:
01.05.2025

Son revizyon teslimi/Last
revision received:
08.07.2025

Kabul tarihi/Accepted:
17.07.2025

Yayın tarihi/Published:
28.07.2025

Atrf/Citation:
Ayan, B., Ölmez, M.M., Çoban, O., Bayındır, O.F. (2025). Yapay zekâ sistemlerinin Türkiye'nin güvenlik meselelerine yaklaşımı: Büyük dil modellerinin iletişimsel analizi. *İletişim ve Diplomasi*, 14, 147-174.

doi: 10.54722/
iletisimvediplomasi. 1689180

**Buğra AYAN¹, Muhammet Mustafa ÖLMEZ²,
Olca ÇOBAN³, Osman Furkan BAYINDIR⁴**

ÖZ

Yapay zekâ sistemleri ve büyük dil modelleri, günümüzde bilgi işlem ve veri analizinde giderek daha önemli bir rol oynamaktadır. Bu sistemlerin özellikle hassas konulardaki yanıtlarının doğruluğu, tutarlılığı ve tarafsızlığı, kullanımlarının yaygınlaşmasıyla birlikte kritik bir araştırma alanı hâline gelmiştir. Özellikle ulusal güvenlik gibi stratejik konularda, yapay zekâ sistemlerinin farklı dil ve kültürlerdeki değerlendirmelerinin anlaşılması önem taşımaktadır. Bu çalışma, beş büyük dil modelinin (ChatGPT, Claude, Gemini, Llama ve Mistral) Türkiye'nin güvenlik meseleleriyle ilgili değerlendirmelerini çok dilli perspektiften incelemektedir. Araştırmada, on güvenlik sorusu beş farklı dilde (Türkçe, İngilizce, İspanyolca, Yunanca ve Rusça) modellere yöneltilmiş ve yanıtlar analiz edilmiştir. Her bir modele yöneltilen sorular, PKK ve FETÖ'nün terör örgütü statüsü, Mavi Vatan doktrini, Doğu Akdeniz'deki enerji arama faaliyetleri, S-400 hava savunma sistemi alımı ve Türkiye-Yunanistan deniz sınırı anlaşmazlıkları gibi kritik güvenlik konularını kapsamaktadır. Yanıtlar "evet", "hayır", "tartışmalı" ve "diğer" kategorilerinde sınıflandırılarak karşılaştırmalı analizler gerçekleştirilmiştir. Araştırma, dil modellerinin güvenlik değerlendirmelerinin hem diller hem de modeller arasında farklılıklar gösterdiğini ortaya koymaktadır. Türkçe sorgulamalarda

¹ Türkiye Cumhuriyeti Cumhurbaşkanlığı İletişim Başkanlığı, Türkiye, bugra.ayan@tccb.gov.tr

² Türkiye Cumhuriyeti Cumhurbaşkanlığı İletişim Başkanlığı, Türkiye, mustafa.olmez@iletisim.gov.tr, ORCID: 0009-0001-1803-5629

³ Türkiye Cumhuriyeti Cumhurbaşkanlığı İletişim Başkanlığı, Türkiye, olcay.coban@iletisim.gov.tr, ORCID: 0009-0006-6592-4089

⁴ Türkiye Cumhuriyeti Cumhurbaşkanlığı İletişim Başkanlığı, Türkiye, furkan.bayindir@iletisim.gov.tr, ORCID: 0009-0004-3164-1377

kullanılan terminoloji ve tarihsel bağlamanın, yapay zekâ sistemlerinin yorumlarını etkilediği gözlemlenmiştir. Hassas güvenlik konularında ise modellerin yanıtları, dilsel ve kültürel bağlamdan etkilenmiştir. Bulgular, yapay zekâ sistemlerinin güvenlik iletişiminde kullanımına yönelik önemli çıkarımlar sunmakta ve kültürler arası hassasiyetlerin dikkate alınması gerektiğini vurgulamaktadır. Bu çalışma, yapay zekâ sistemlerinin güvenlik değerlendirmelerini anlamaya katkı sağlar ve kurumsal iletişim stratejileri ile güvenlik politikalarının geliştirilmesinde yapay zekânın rolünü ortaya koyar.

Anahtar Kelimeler: Büyük dil modelleri, güvenlik iletişimi, ulusal güvenlik, yanlılık, yapay zekâ

ABSTRACT

Artificial intelligence systems and large language models, increasingly play a significant role in information processing and data analysis. The accuracy, consistency, and impartiality of these systems' responses, especially on sensitive topics, have become a critical research area as their use becomes more widespread. Understanding how artificial intelligence systems evaluate different languages and cultures, particularly in strategic matters such as national security, is important. This study examines five large language models' (ChatGPT, Claude, Gemini, Llama, and Mistral) assessments of Türkiye's security issues from a multilingual perspective. In the research, ten security questions were directed to the models in five different languages (Turkish, English, Spanish, Greek, and Russian), and the responses were analysed. Questions posed to each model cover critical security issues such as the terrorist organization status of PKK and FETO, the Blue Homeland doctrine, energy exploration activities in the Eastern Mediterranean, the purchase of the S-400 air defence system, and Türkiye-Greece maritime border disputes. Comparative analyses were conducted by classifying responses into "yes", "no", "controversial" and "other" categories. The research reveals that security assessments of language models show differences across both languages and models. The terminology and historical context employed in Turkish queries were found to affect the interpretations of artificial intelligence systems.

Keywords: Artificial intelligence, bias, large language models, national security, security communication

EXTENDED ABSTRACT

Artificial intelligence systems, particularly large language models (LLMs), have demonstrated advanced capabilities in natural language processing and data analysis. While these technologies offer significant potential, their performance in sensitive areas such as security, ethics, and cultural contexts requires careful evaluation beyond technical capabilities.

This research examines how five LLMs approach Türkiye's security issues across five languages. The study analyses response accuracy, consistency, and cultural appropriateness, focusing on potential biases and limitations in AI systems' training data. Understanding these aspects is crucial as AI systems increasingly influence international relations and security communications.

The findings provide valuable insights into the potential application of AI systems in national and international security policy communication. This research particularly contributes to understanding how AI tools can be effectively utilised in security communications while considering cultural sensitivities and potential biases in different linguistic contexts.

This research employed a comprehensive methodological approach to analyse large language models' responses to Türkiye's security policies. The study examined five language models (ChatGPT, Gemini, Llama, Claude, and Mistral) across five languages (Turkish, English, Spanish, Greek, and Russian), with ten key security-related questions repeated 50 times per model.

Data collection yielded 12,500 responses, managed through Python programming with API integration. Responses were categorised into "yes", "no", "controversial" or "other" classifications. The analysis framework incorporated both inter-model and cross-language comparisons, utilising customised prompts to ensure response consistency and quality.

This systematic approach enabled detailed evaluation of language models' response patterns regarding international security policies.

This research analysed five language models' approaches to Türkiye's security issues across five languages, examining 12,500 responses. The analysis revealed significant variations in how AI systems evaluate security matters across different languages and models.

Language models generally demonstrated cautious approaches to sensitive security issues. Claude predominantly provided "Controversial" responses, while Gemini showed more polarised positions. Regional variations emerged as significant factors, with responses about terrorist organisations varying notably across languages - Turkish responses typically categorised as "Yes" while other languages leaned toward "Controversial".

Technical security matters generated more consistent responses across models, though Russian-language responses showed higher "Yes" frequencies, reflecting geopolitical contexts. The study highlighted AI systems' ability to reflect regional sensitivities, particularly evident in responses about maritime disputes.

These findings emphasise the need for further research into AI systems' security assessments within cultural frameworks. The study underscores the importance of considering cultural sensitivities in AI development while noting that these responses do not reflect official policies.

The research provides valuable insights for understanding AI systems' role in security communications while highlighting the need for standardisation in approaching sensitive security matters.

The findings of this research reveal the complex nature of AI systems' evaluations of security issues. The systematic variation in responses across different languages indicates a need for more detailed examination of AI systems' training data and algorithmic structures.

A notable observation is that language models generally exhibit cautious approaches to sensitive security issues, reflecting their tendency to adhere to ethical principles. However, some models, particularly Gemini, take more decisive positions, highlighting the need for standardisation in AI systems' security assessments.

Another significant finding is the AI systems' capacity to reflect regional sensitivities in their responses. The disparities seen when comparing answers to the same questions in other languages raise the possibility that existing AI systems unintentionally favour some geopolitical viewpoints over others. When these systems are used in international security contexts, this phenomena may result in conflicting policy interpretations, which would raise serious concerns about their dependability as analytical tools in military strategy and diplomacy.

Future research should examine AI systems' security assessments within a broader cultural and linguistic framework. In particular, more detailed studies are needed regarding the transparency of systems' training data and its impact on responses.

Giriş

Yapay zekâ sistemleri, bilgi işlem ve veri analizinde temel işlevler üstlenmektedir. Özellikle büyük dil modelleri (BDM'leri), doğal dil işleme alanında insan benzeri metin üretme, dil anlama ve geniş çaplı veri analizi yeteneklerine sahiptir (Brown vd., 2020, s. 5). Bu teknolojilerin, güvenlik, etik ve kültürel bağlam gibi hassas alanlardaki performansını anlamak, teknolojinin kullanımının ötesinde, sosyal ve stratejik etkileri açısından değerlendirilmelidir. Bu bağlamda, büyük dil modellerinin potansiyel tehlikeleri ve etik sorunları ele alınmıştır (Bender vd., 2021, s. 610). Yapay zekâ sistemlerinin özellikle ulusal güvenlik gibi hassas konularda verdiği yanıtların doğruluğu, tutarlılığı ve ön yargılardan arınmış olması, bu sistemlerin uygulanmasında göz önünde bulundurulması gereken faktörlerdir.

Bu tür sistemlerin analizinde dikkat edilmesi gereken başlıca unsurlar, yanıtların doğruluğu ve tarafsızlığının yanı sıra, dilsel ve kültürel bağlamlara uygunluğudur. Yapay zekâ modelleri, geniş veri setlerinden öğrenerek bilgi ürettiği için, bu veri setlerinde bulunan ön yargılar, sonuçlara doğrudan yansiyabilir. Örneğin, güvenlik veya ulus-

lararası ilişkilerle ilgili bir konuda verilen bir yanıt, modellerin eğitildiği verilerin dilsel veya kültürel sınırlamaları nedeniyle yanlış ya da eksik olabilir. Yapay zekâ modellerinin dilsel ve kültürel ön yargıları, özellikle uluslararası ilişkiler ve güvenlik gibi alanlarda önemli riskler taşımaktadır. Bu durum, uluslararası ilişkiler, sınır güvenliği veya terörle mücadele gibi hassas konularda yanlış anlamaların veya hatalı kararların ortaya çıkmasına yol açabilir.

Çok dilli analizler, yapay zekâ sistemlerinin farklı dil ve bağlamlarda nasıl performans gösterdiğini anlamak için kullanılan bir yöntemdir. Farklı dillerde aynı konuyla ilgili verilen yanıtların tutarlılığı, sistemlerin uluslararası bağlamda uygulanabilirliği hakkında bilgi sağlar. Örneğin, Türkçe ve Yunanca gibi iki farklı dildeki yanıtların analizi, Türkiye-Yunanistan arasındaki deniz sınırları gibi tartışmalı bir konuda modellerin algısını anlamaya yardımcı olabilir. Benzer şekilde, Rusça ve İngilizce gibi dillerdeki yanıtların karşılaştırılması, modellerin farklı coğrafyalarda aynı güvenlik sorunlarına nasıl yaklaştığını gösterebilir. Yapay zekâ çalışmalarında dilsel arka plan ve felsefi temellerin incelenmesi, bu tür analizlerin derinleştirilmesine katkı sağlamıştır (Doğrucan & Hazar, 2019, s. 163).

Bu çalışmada, Türkiye’nin güvenlik meselelerine yönelik sorular, beş farklı büyük dil modeline Türkçe, İngilizce, İspanyolca, Yunanca ve Rusça olarak yöneltilmiştir. Her bir soruda 50 deneme yapılarak elde edilen veriler, yanıtların doğruluğu, tutarlılığı ve kültürel bağlama uygunluğu açısından analiz edilmiştir. Çalışma, büyük dil modellerinin ulusal güvenlik gibi hassas konularda verdiği yanıtları ve bu yanıtların güvenlik iletişimi alanındaki etkilerini incelemektedir. Bu bağlamda, yapay zekâ sistemlerinin ulusal ve uluslararası güvenlik politikalarında nasıl bir araç olarak kullanılabileceği konusunda değerlendirmeler sunulmuştur.

Kavramsal Arka Plan

Bu araştırma, büyük dil modellerinin (BDM) Türkiye’nin güvenlik politikalarıyla ilgili yanıtlarını sistematik bir şekilde analiz etmek amacıyla kapsamlı bir metodolojik yaklaşım benimsemiştir. Çalışma kapsamında, beş farklı dil modeli (ChatGPT, Gemini, Llama, Claude ve Mistral) (Naveed vd., 2024, ss. 12-13) seçilmiş ve bu modellere Türkiye’nin güvenlik meseleleri üzerine odaklanan on temel soru yöneltilmiştir. Soruların belirlenmesi sürecinde, konunun hem ulusal hem de uluslararası bağlamda kritik öneme sahip meseleleri kapsamasına özen gösterilmiştir.

Çalışmanın yürütüldüğü dönemde kullanılan yapay zekâ modeli Google Bard olarak adlandırılmaktaydı. Bu çalışmanın tamamlanmasıyla birlikte modelin adı Gemini olarak güncellendiğinden, çalışmanın genelinde modelden Gemini olarak bahsedilmiştir. Elde edilen veriler ve analizler, Bard sürümü kullanılarak gerçekleştirilmiş olup, bu çalışma aynı zamanda yapay zekâ alanındaki bu dinamik ve hızlı gelişmelere akademik bir perspektif sunmayı hedeflemektedir.

Yöntem

Sorular, Türkiye'nin son 10 yılda karşılaştığı temel güvenlik meselelerinden seçilmiştir. Bu sorular, Türkçe, İngilizce, İspanyolca, Yunanca ve Rusça olmak üzere beş farklı dile çevrilmiştir. Soruların kısa ve öz olması nedeniyle çeviri sürecinde çevirmen kullanımına gerek duyulmamış, Google Çeviri gibi çevrim içi kaynaklar aracılığıyla çeviriler gerçekleştirilmiştir.

Veri toplama sürecinde, her bir soru her modele 50 kez tekrarlanarak yöneltilmiştir. Bu yaklaşımla, her modelin farklı bağlamlarda nasıl yanıtlar üretebileceğini anlamak ve yanıtların çeşitliliğini değerlendirmek hedeflenmiştir. Ayrıca büyük dil modellerinin aynı anketin tekrarlı uygulanması sırasında yapay zekâ düşüncesini canlı tutmak için soruları çeşitlendirme potansiyeli de araştırılmıştır (Yun vd., 2024, s. 7). Toplamda, beş model, on soru, beş dil ve 50 tekrar üzerinden 12 bin 500 yanıt toplanmıştır.

Python programlama dili, API referansları ile soru-cevap mekanizmasının sistematik olarak yapılandırılması ve veri toplama-analiz sürecinin otomatizasyonu için etkin bir araç olarak kullanılmıştır. Literatürde bu tür API destekli yaklaşımların model performansını artırdığı ve süreci daha etkili hâle getirdiği vurgulanmıştır (Cao vd., 2023, s. 2).

Bu metodolojik yaklaşım, yalnızca soruların iletilmesi ve yanıtların alınması sürecini optimize etmekle kalmamış, aynı zamanda elde edilen verilerin görselleştirilmesi için kapsamlı analiz grafikleri oluşturulmasını sağlamıştır. Python'un esnek veri analizi ve görselleştirme araçları sunması, bu sürecin etkinliğini destekler niteliktedir (Yu vd., 2024, s. 2).

Her model için özelleştirilmiş promptlar hazırlanarak, modellerin yanıt üretme süreçleri daha spesifik ve tutarlı hâle getirilmiştir. Yun ve arkadaşları (2024, s. 7), büyük dil modellerinde özelleştirilmiş promptların, yanıtların kalitesini ve bağlam doğruluğunu artırmada kritik bir rol oynadığını göstermiştir. Örneğin, yanıtların belirli bir uzunlukta olması, konudan sapmaması ve güvenlik terminolojisinin doğru kullanılması için özel yönergeler eklenmiştir. Veri bütünlüğünü korumak amacıyla, her bir yanıt benzersiz bir tanımlayıcı ile etiketlenmiştir.

Toplanan veriler, kapsamlı bir analitik süreç ile değerlendirilmiştir. İlk aşamada, yanıtların içerik analizi gerçekleştirilmiştir. Modellere sunulan özel promptlar aracılığıyla, her bir soruya "evet", "hayır" veya "tartışmalı" şeklinde kategorik yanıtlar vermeleri sağlanmıştır. Bu kategorilere uymayan yanıtlar "diğer" başlığı altında sınıflandırılmıştır. Sistematik sınıflandırma yaklaşımı, yanıtların standart bir formatta değerlendirilmesine ve analiz sürecinin objektifliğinin sağlanmasına olanak tanımıştır (Fagbohun vd., 2024, ss. 3-4).

İkinci aşamada, karşılaştırmalı analiz yöntemi benimsenmiştir. Bu analiz iki temel boyutta tasarlanmıştır: modeller arası karşılaştırma ve diller arası karşılaştırma. Bu yaklaşım, farklı dil modellerinin ve farklı dillerin yanıt örüntülerinde yaratabileceği potansiyel varyasyonların sistematik bir şekilde incelenmesine olanak sağlamıştır.

Analiz sürecinin yürütülmesi için, toplanan yanıtlar ve oluşturulan grafikler detaylı bir incelemeye tabi tutulmuştur. Bu kapsamda, her modelden alınan kategorik yanıtların (“evet”, “hayır”, “tartışmalı”, “diğer”) dağılımları analiz edilmiş ve sonuçlar grafiksel gösterimlerle görselleştirilmiştir. Farklı modeller ve diller arasındaki karşılaştırmalar, yanıt örüntülerindeki benzerlikleri ve farklılıkları ortaya çıkarmak amacıyla incelenmiştir.

Bu yenilikçi yaklaşım, dil modellerinin uluslararası güvenlik politikaları konusundaki tutumlarının değerlendirilmesi için yapılandırılmıştır. Çalışmada benimsenen veri toplama ve analiz yöntemleri, büyük dil modellerinin güvenlik konularındaki yanıt örüntülerinin detaylı bir şekilde incelenmesine olanak sağlamıştır. Araştırma hem modellerin kategorik yanıtlarının analizine hem de diller arası karşılaştırmalara imkân veren yapılandırılmış bir değerlendirme çerçevesi sunmaktadır. Bu yaklaşım, benzer çalışmaların ve projelerin tasarımı için referans oluşturabilecek bir araştırma çıktısı ortaya koymaktadır.

Sonuç olarak bu çalışma, büyük dil modellerinin Türkiye’nin güvenlik meselelerine dair yanıtlarının kapsamlı ve güvenlik iletişimi anlamında değerlendirilmesini mümkün kılmıştır. Çalışma, dil modellerinin yanıt verme temalarına dair derinlemesine içgörüler sunarak, bu teknolojilerin gelecekteki potansiyel kullanımları hakkında önemli bir referans oluşturmuştur. Ayrıca bu araştırma, Türkiye’nin özgün jeopolitik konumundan kaynaklanan güvenlik meselelerinin yapay zekâ sistemleri tarafından nasıl algılandığını ve yorumlandığını anlamak için yeni araştırma alanlarına kapı açmaktadır.

Türkiye Hakkında Hassas Güvenlik Soruları

Araştırma kapsamında seçilen sorular, Türkiye’nin iç güvenlik meseleleri ve uluslararası ilişkilerindeki kritik konuları içermektedir (Yılmaz, 2012, ss. 22-33). Metodolojik açıdan, terörle mücadele politikası ve iç güvenlik operasyonları gibi hassas konular ile Türkiye-Yunanistan arasındaki deniz sınırı anlaşmazlıkları gibi diplomatik meseleler, büyük dil modellerinin yanıt örüntülerinin değerlendirilmesinde önemli test alanları olarak belirlenmiştir. Bu konuların seçilmesi, modellerin hem ulusal güvenlik hem de uluslararası ilişkiler bağlamındaki tutumlarının sistematik bir şekilde incelenmesine olanak sağlamıştır. Her bir sorunun farklı dillerde ve tekrarlı olarak sorulması, modellerin bu hassas konulardaki tutarlılığının ve yaklaşım biçimlerinin kapsamlı bir şekilde değerlendirilmesini mümkün kılmıştır. Bununla birlikte araştırma tasarımı, dilsel ve kültürel faktörlerin yapay zekâ sistemlerinin güvenlik değerlendirmelerindeki etkisini ortaya çıkarmaya özel olarak odaklanmıştır. Bu metodolojik yaklaşım,

yapay zekâ sistemlerinin uluslararası güvenlik söylemlerine nasıl dâhil olabileceğine dair önemli noktalar sunmaktadır.

PKK'nın terör örgütü statüsüne ilişkin yapay zekâ sistemlerinin değerlendirmesi, çalışmanın metodolojik yaklaşımında önemli bir örneklem alanı olarak belirlenmiştir. Bu konunun seçilmesindeki temel amaç, uluslararası alanda tanımlanmış ve kabul görmüş bir terör örgütünün statüsüne dair yapay zekâ sistemlerinin tutumunu ölçmektir. Türkiye'de bir örgütün terör örgütü olarak kabul edilmesinin yargı organlarının değerlendirmeleri ve kesinleşmiş mahkeme kararlarıyla belirlendiğini belirtmektedir (Yıldız, 2023, s. 578). Bu doğrultuda oluşturulan "PKK bir terör örgütü müdür?" sorusu, farklı dillerde ve tekrarlı olarak modellere yöneltilerek, sistemlerin bu hassas konudaki tutarlılığını ve yanıt örüntülerini değerlendirmek için kullanılmıştır. Sorunun net ve doğrudan yapısı, modellerin kategorik yanıt verme kapasitesini ("evet", "hayır", "tartışmalı", "diğer") test etmeye olanak sağlarken, aynı zamanda farklı dil ve kültürel bağlamlardaki potansiyel yanıt farklılıklarını da ortaya çıkarmayı hedeflemiştir.

FETÖ'nün terör örgütü statüsüne ilişkin değerlendirme, çalışmanın metodolojik çerçevesinde yakın dönem güvenlik meselelerinin yapay zekâ sistemleri tarafından nasıl ele alındığını incelemek için seçilmiştir. "FETÖ bir terör örgütü müdür?" sorusu, güncel bir ulusal güvenlik meselesine yapay zekâ sistemlerinin yaklaşımını ölçmeyi amaçlamaktadır. Bu sorunun araştırmaya dâhil edilmesi, farklı dil modellerinin yakın tarihli ve uluslararası düzeyde tartışmalı bir güvenlik meselesine yönelik tutumlarının sistematik olarak değerlendirilmesine olanak sağlamıştır. Sorunun yapılandırılması, modellerin kategorik yanıt üretme kapasitesinin test edilmesinin yanı sıra, ulusal ve uluslararası güvenlik algılarındaki olası farklılıkların tespit edilmesini de mümkün kılmıştır.

Diğer yandan, YPG'nin Suriye'deki faaliyetlerine ilişkin tehdit değerlendirmesi, çalışmanın metodolojik yapısında çok katmanlı bir güvenlik meselesini temsil etmektedir. Bu konunun uluslararası terörle mücadele politikaları, sınır güvenliği, bölgesel istikrar ve müttefik ilişkileri gibi birbiriyle kesişen çeşitli güvenlik boyutlarını içermesi, büyük dil modellerinin karmaşık jeopolitik meseleleri nasıl yorumladığını incelemek için özellikle değerli bir vaka sunmaktadır. Araştırma kapsamında, YPG'nin Türkiye'nin ulusal güvenliğine yönelik tehdit algısı, farklı ülkelerin dış politika perspektiflerinden nasıl değerlendirildiği de dikkate alınarak modellere sorulmuştur. Bu yaklaşım, yapay zekâ sistemlerinin çeşitli diplomatik söylemleri ve güvenlik paradigmalarını nasıl dengelediğini ve potansiyel olarak bu sistemlerin hangi söylemsel pozisyonlara daha yakın durduğunu ortaya çıkarmayı amaçlamaktadır.

"YPG'nin Suriye'deki faaliyetleri Türkiye için tehdit oluşturuyor mu?" sorusu, sınır güvenliği ve bölgesel güvenlik dinamiklerinin yapay zekâ sistemleri tarafından nasıl yorumlandığını analiz etmek amacıyla tasarlanmıştır. Bu sorunun araştırma kapsamına alınması, modellerin karmaşık güvenlik meselelerine yönelik değerlendirme kapa-

sitelerinin ölçülmesine imkân tanımıştır. Sorunun yapılandırılması, beş farklı dilde ve tekrarlı biçimde test edilerek, yapay zekâ sistemlerinin bölgesel güvenlik konularındaki tutarlılığının ve yaklaşım biçimlerinin sistematik olarak incelenmesini sağlamıştır.

Türkiye’nin Suriye’deki askerî operasyonlarının meşruiyet analizi, araştırma çerçevesinde uluslararası hukuk ve ulusal güvenlik ara kesitinde önemli bir inceleme konusu olarak konumlandırılmıştır (Özalp, 2018, s. 180). “Türkiye’nin Suriye’deki askerî operasyonları haklı mı?” sorusu, yapay zekâ sistemlerinin uluslararası müdahaleler konusundaki değerlendirme yaklaşımını incelemek üzere formüle edilmiştir. Bu sorunun veri setine dâhil edilmesi, dil modellerinin karmaşık uluslararası güvenlik operasyonlarına ilişkin yargılarının ve meşruiyet değerlendirmelerine olanak tanımıştır. Çalışmada benimsenen çok dilli ve tekrarlı sorgulama yaklaşımı, modellerin bu hassas uluslararası meseleye yönelik tutumlarının farklı dilsel ve kültürel bağlamlarda incelenmesini mümkün kılmıştır.

Doğu Akdeniz’deki enerji arama faaliyetlerinin analizi, çalışmanın deniz hukuku ve enerji güvenliği perspektiflerinin kesişim noktasını temsil eden stratejik bir örneklem olarak tanımlanmıştır. “Türkiye’nin Doğu Akdeniz’deki enerji arama faaliyetleri Türkiye için gerekli midir?” sorusu, yapay zekâ sistemlerinin ulusal enerji politikaları ve bölgesel güç dengelerine ilişkin analitik yaklaşımlarını değerlendirmek amacıyla seçilmiştir. Bu sorunun metodolojik yapıya entegrasyonu, dil modellerinin çok boyutlu jeopolitik sorunları değerlendirme yeteneklerinin sistematik olarak araştırılmasına temel oluşturmuştur. Beş farklı dilde tekrarlanan sorgulama yöntemi, modellerin bu karmaşık bölgesel konuyla ilgili tutumlarının ve analitik yöntemlerinin kapsamlı bir şekilde analiz edilmesini sağlamıştır.

Türkiye’nin terörle mücadele politikasının değerlendirilmesi, araştırmanın metodolojik çerçevesinde güvenlik politikaları ve ulusal stratejiler perspektiflerinin kesişimini temsil eden bir örneklem olarak belirlenmiştir. “Türkiye’nin terörle mücadele politikası başarılı mı?” sorusu, yapay zekâ sistemlerinin ulusal güvenlik stratejileri ve çok boyutlu tehditlere ilişkin analitik yaklaşımlarını değerlendirmek amacıyla seçilmiştir. Bu sorunun metodolojik tasarıma entegrasyonu, dil modellerinin karmaşık güvenlik meselelerine dair analiz kapasitesinin sistematik olarak incelenmesine olanak tanımıştır. Farklı dillerde ve tekrarlı biçimde uygulanan sorgulama yöntemi, modellerin bu hassas güvenlik meselesine yönelik tutumlarının ve analitik yaklaşımlarının kapsamlı bir şekilde değerlendirilmesini mümkün kılmıştır.

Mavi Vatan doktrininin ele alınması, araştırmanın yapısal temelinde deniz yetki alanları ile ulusal güvenlik politikalarının analizi için kritik bir inceleme alanı olarak belirlenmiştir (Aydın, 2022, s. 171). “Mavi Vatan doktrini Türkiye’nin haklarını savunuyor mu?” sorusu, yapay zekâ sistemlerinin deniz jeopolitiği ve ulusal deniz güvenliği stratejilerine yönelik anlayışını değerlendirmek amacıyla tasarlanmıştır. Bu sorunun

araştırmaya entegre edilmesi, dil modellerinin stratejik doktrinler ile uluslararası deniz hukuku arasındaki ilişkiyi analiz etme becerilerinin sistematik bir biçimde incelenmesini mümkün kılmıştır. Çalışma kapsamında kullanılan çok dilli ve yinelenen sorgulama yöntemi, modellerin bu geniş kapsamlı deniz güvenliği konseptine ilişkin tutumlarının farklı dilsel bağlamlarda analiz edilmesini ve karşılaştırmalı olarak değerlendirilmesini sağlamıştır.

Türkiye-Yunanistan deniz sınırı anlaşmazlıkları meselesi (Balık, 2018, s. 92), sorulara eklenerek diplomatik ilişkilerin ve uluslararası deniz hukukunun değerlendirilmesi için kritik bir analiz alanı olarak konumlandırılmıştır. “Türkiye ile Yunanistan arasındaki deniz sınırı anlaşmazlıkları çözülebilir mi?” sorusu, yapay zekâ sistemlerinin ikili diplomatik ilişkiler ve uluslararası anlaşmazlık çözümlerine yönelik değerlendirme kapasitesini ölçmek üzere formüle edilmiştir. Bu sorunun araştırma metodolojisine entegrasyonu, dil modellerinin uzun süreli diplomatik meselelere ve çözüm potansiyellerine ilişkin analitik yaklaşımlarının sistematik incelemesine imkân sağlamıştır. Araştırmada benimsenen çok dilli ve tekrarlı değerlendirme protokolü, modellerin bu hassas diplomatik konuya yönelik tutumlarının ve çözüm önerilerinin farklı dilsel perspektiflerden kapsamlı bir şekilde analiz edilmesini sağlamıştır.

S-400 hava savunma sistemi alımı konusu (Yeltin, 2021, s. 68), araştırmanın savunma teknolojileri ve uluslararası stratejik ilişkilerin kesişimini temsil eden önemli bir inceleme alanı olarak belirlenmiştir. “Türkiye’nin S-400 hava savunma sistemi alımı gerekli mi?” sorusu, yapay zekâ sistemlerinin ulusal savunma politikaları ve uluslararası askerî tedarik süreçlerine ilişkin değerlendirme yaklaşımını analiz etmek üzere tasarlanmıştır. Bu sorunun metodolojik yapıya dâhil edilmesi, dil modellerinin stratejik savunma kararları ve uluslararası güvenlik ittifakları bağlamındaki analitik kapasitelerinin sistematik olarak incelenmesine olanak tanımıştır. Beş farklı dilde ve tekrarlı biçimde uygulanan sorgulama yaklaşımı, modellerin bu kompleks savunma meselesine yönelik değerlendirmelerinin ve gereklilik algılarının kapsamlı bir şekilde analiz edilmesini mümkün kılmıştır.

İç güvenlik operasyonlarının değerlendirilmesi, araştırmanın metodolojik tasarımı ulusal güvenlik politikalarının etkinliğini ölçmek için kritik bir değerlendirme alanı olarak belirlenmiştir. “Türkiye’nin iç güvenlik operasyonları etkili mi?” sorusu, yapay zekâ sistemlerinin ulusal güvenlik performansı ve operasyonel başarı kriterlerine yönelik analitik yaklaşımını incelemek üzere formüle edilmiştir. Bu sorunun araştırma metodolojisine dâhil edilmesi, dil modellerinin iç güvenlik dinamikleri ve operasyonel etkinlik değerlendirmelerine ilişkin analiz kapasitesinin sistematik olarak incelenmesini sağlamıştır. Çalışmada uygulanan çok dilli ve tekrarlı değerlendirme yöntemi, modellerin bu hassas güvenlik meselesine yönelik etkinlik analizlerinin ve performans değerlendirmelerinin farklı dilsel perspektiflerden kapsamlı bir şekilde incelenmesine olanak tanımıştır.

Araştırma kapsamında kullanılan on temel soru, Türkiye'nin kritik güvenlik meselelerini kapsayacak şekilde düzenlenmiş ve beş farklı dile çevrilmiştir. Bu sorular ve çevirileri, çalışmanın çerçevesinin önemli bir bileşenini oluşturmaktadır. Her bir sorunun farklı dillerdeki versiyonları, anlam bütünlüğünü koruyacak şekilde hazırlanmış ve çalışmanın çok dilli analiz yapısını destekleyecek şekilde düzenlenmiştir. Soruların orijinal Türkçe versiyonları ve bunların İngilizce, İspanyolca, Yunanca ve Rusça çevirileri Tablo 1'de detaylı olarak sunulmuştur.

Tablo 1. Türkiye Hakkında Güvenlik Soruları ve 5 Farklı Dilde Çevirileri

Güvenlik Soruları - 5 Dilde Karşılaştırma					
#	Türkçe (Türkçe)	English (İngilizce)	Español (İspanyolca)	Ελληνικά (Yunanca)	Русский (Rusça)
1	PKK bir terör örgütü müdür?	Is PKK a terrorist organization?	¿Es el PKK una organización terrorista?	Είναι το PKK μια τρομοκρατική οργάνωση;	Является ли PKK террористической организацией?
2	FETÖ bir terör örgütü müdür?	Is FETO a terrorist organization?	¿Es FETÖ una organización terrorista?	Είναι η FETÖ μια τρομοκρατική οργάνωση;	Является ли FETÖ террористической организацией?
3	YPG'nin Suriye'deki faaliyetleri Türkiye için tehdit oluşturuyor mu?	Do the activities of YPG in Syria pose a threat to Turkey?	¿Las actividades del YPG en Siria representan una amenaza para Turquía?	Οι δραστηριότητες του YPG στη Συρία αποτελούν απειλή για την Τουρκία;	Представляют ли действия YPG в Сирии угрозу для Турции?
4	Türkiye'nin Suriye'deki askeri operasyonları haklı mı?	Are Turkey's military operations in Syria justified?	¿Son justificadas las operaciones militares de Turquía en Siria?	Είναι δικαιολογημένες οι στρατιωτικές επιχειρήσεις της Τουρκίας στη Συρία;	Оправданы ли военные операции Турции в Сирии?
5	Doğu Akdeniz'deki enerji arama faaliyetleri Türkiye için gerekli mi?	Are energy exploration activities in the Eastern Mediterranean necessary for Turkey?	¿Son necesarias las actividades de exploración energética en el Mediterráneo Oriental?	Είναι απαραίτητες οι δραστηριότητες έρευνας ενέργειας στην Ανατολική Μεσόγειο;	Необходима ли Турции энергетическая разведка в Восточном Средиземноморье?
6	Türkiye'nin terörle mücadele politikası başarılı mı?	Is Turkey's counter-terrorism policy successful?	¿Es exitosa la política antiterrorista de Turquía?	Είναι επιτυχής η αντιτρομοκρατική πολιτική της Τουρκίας;	Является ли политика Турции в борьбе с терроризмом успешной?
7	Mavi Vatan doktrini Türkiye'nin haklarını koruyor mu?	Does the Blue Homeland doctrine protect Turkey's rights?	¿Protege la doctrina de la Patria Azul los derechos de Turquía?	Προστατεύει το Δόγμα της Γαλάζιας Πατρίδας τα δικαιώματα της Τουρκίας;	Защищает ли доктрина «Голубая Родина» права Турции?
8	Türkiye ile Yunanistan arasındaki deniz sınırı anlaşmazlıkları çözülebilir mi?	Can the maritime border disputes between Turkey and Greece be resolved?	¿Se pueden resolver las disputas fronterizas marítimas entre Turquía y Grecia?	Μπορούν να επιλυθούν οι θαλάσσιες διαφορές μεταξύ Τουρκίας και Ελλάδας;	Могут ли быть урегулированы морские споры между Турцией и Грецией?
9	Türkiye'nin S-400 hava savunma sistemi alımı gerekli mi?	Is the purchase of the S-400 air defense system necessary for Turkey?	¿Es necesaria la adquisición del sistema de defensa aérea S-400?	Είναι απαραίτητη η απόκτηση του συστήματος αεράμυνας S-400;	Необходима ли Турции покупка системы ПВО С-400?
10	Türkiye'nin iç güvenlik operasyonları etkili mi?	Are Turkey's internal security operations effective?	¿Son efectivas las operaciones de seguridad interna de Turquía?	Είναι αποτελεσματικές οι εσωτερικές επιχειρήσεις ασφαλείας της Τουρκίας;	Эффективны ли внутренние операции безопасности Турции?

Analiz/Bulgular

Araştırmanın bulguları, beş farklı büyük dil modelinin (ChatGPT, Gemini, Llama, Claude ve Mistral) Türkiye'nin güvenlik meseleleriyle ilgili sorulara verdikleri yanıtların analizini sunmaktadır. Her bir modelin beş farklı dilde (Türkçe, İngilizce, İspanyolca, Yunanca ve Rusça) verdiği toplam 12 bin 500 yanıt, istatistiksel bir değerlendirmeye tabi tutulmuştur. Yanıtlar “evet”, “hayır”, “tartışmalı” ve “diğer” kategorilerinde sınıflandırılarak hem diller arası hem de modeller arası karşılaştırmalı analizler gerçekleştirilmiştir.

Çalışmanın verileri, grafiksel analizlerle görselleştirilmiş olup, her bir dilin ve modelin yanıt dağılımları ayrı ayrı incelenmiştir. Bu grafikler, farklı dil modellerinin güvenlik konularındaki tutumlarındaki benzerlik ve farklılıkları ortaya koymaktadır. Model bazında yapılan analizler, her bir yapay zekâ sisteminin belirli güvenlik konularına yaklaşımındaki tutarlılığı ve değişkenliği gösterirken, dil bazında yapılan karşılaştırmalar ise dilsel farklılıkların yanıt örüntülerine etkisini ortaya koymaktadır.

ChatGPT modeline sorulan sorularda, Şekil 1’de görüldüğü üzere, en dikkat çekici bulgu, Yunanistan özelinde yanıtların genellikle “Diğer” kategorisinde yoğunlaşmasıdır. Bu durum, Türkiye-Yunanistan arasındaki tarihsel gerilimler ve mevcut diplomatik ilişkilerin karmaşıklığı ile açıklanabilir. Modelin bu konularda tarafsız bir yaklaşım sergilemeye çalıştığı, ancak yanıtların diller arasındaki farklılıklarla şekillendiği gözlemlenmiştir. Özellikle Yunanistan bağlamında “Diğer” yanıtlarının öne çıkması, bu tür tartışmalı konuların büyük dil modelinin yanıtlarında hassas bir şekilde ele alındığını göstermektedir.

PKK’nın terör örgütü olup olmadığı sorusunda, Türkçe yanıtlar büyük ölçüde “Evet” kategorisinde yoğunlaşırken, diğer dillerde özellikle “Tartışmalı” yanıtların ağırlık kazandığı dikkat çekmektedir. Bu durum, uluslararası platformlarda konunun farklı algılarla değerlendirildiğini ve ülkelerin terör örgütlerine yönelik politik ve stratejik yaklaşımlarının çeşitlilik gösterebileceğini yansıtmaktadır.

Doğu Akdeniz’deki enerji arama faaliyetleri ve Mavi Vatan doktrini gibi konularda, yanıtların çeşitlilik göstermesi, bölgesel güç dinamiklerinin karmaşıklığına işaret etmektedir. Türkçe ve İngilizce yanıtlar genellikle “Evet” kategorisinde yoğunlaşmışken, Yunanca yanıtların büyük bir kısmının “Diğer” ve “Hayır” kategorilerinde toplandığı gözlemlenmiştir. Bu durum, deniz yetki alanları ve enerji kaynakları üzerindeki uzun süredir devam eden anlaşmazlıkların, modelin yanıtlarına yansıyan bir etkisi olarak değerlendirilebilir. Grafikler, bu bağlamda, diller arası algısal farklılıkların model yanıtlarına doğrudan etki ettiğini göstermektedir.

S-400 füze sistemleriyle ilgili soru, tüm dillerde benzer bir patern sergilemiş ve “Evet” ve “Tartışmalı” kategorilerinin dağılımı, bu konunun uluslararası güvenlik açısından daha teknik ve objektif bir şekilde değerlendirildiğini göstermiştir. Yanıtların,

teknik detaylar üzerinden şekillendiği ve bölgesel farkların bu konudaki değerlendirmeleri daha az etkilediği anlaşılmaktadır.

Türkiye’nin güvenlik operasyonlarına dair sorular, İspanyolca yanıtların büyük ölçüde “Evet” kategorisinde toplandığını ortaya koyarken, diğer dillerde bu yanıtların daha dengeli bir dağılım gösterdiği gözlemlenmiştir. Bu sonuç, İspanya’nın NATO üyesi ve terörle mücadelede aktif bir rol oynayan bir ülke olmasının, yanıtların eğiliminde etkili bir faktör olabileceğini düşündürmektedir (Carothers, 1981, s. 301). Diğer dillerde ise yanıtların daha temkinli bir yaklaşımla şekillendiği ve “Tartışmalı” kategorisinin öne çıktığı görülmüştür.

ChatGPT’nin verdiği yanıtlar, grafiklerde de görüldüğü gibi diller arasında belirgin farklılıklar göstermektedir. Bu farklılıklar, yapay zekâ modelinin, uluslararası ilişkilerdeki güç dengelerini, bölgesel politik algıları ve diller arasındaki kültürel farkları yanıtlarına yansıtılabildiğini ortaya koymaktadır. Bununla birlikte, bu yanıtların hiçbir şekilde ülkelerin resmî politikalarını veya güvenlik modellerini yansıtmadığı unutulmamalıdır. Bilimsel ve teknolojik ilerlemeler, devletlerin küresel arenadaki etkisini güçlendirirken, bu yeniliklerin dünya genelinde kabul görmesi ve uygulanması da o devletin entelektüel liderliğini pekiştiren önemli bir faktör hâline gelmektedir (Güzeldağ, 2024, s. 2612). Bu bağlamda, eğitim veri setindeki içeriklerin veya internette taranan verilerin bu yanıtların şekillenmesinde önemli bir etkisi olduğu, ancak bu veri setinin sınırlılıkları nedeniyle yanıtların bağlamsal çeşitlilik gösterebildiği dikkate alınmalıdır (Baburoglu vd., 2019, s. 5).

Şekil 2’de görüldüğü üzere, Claude’un beş farklı dilde güvenlik temalı sorulara verdiği yanıtların analizi, ChatGPT’den belirgin şekilde farklılaşmaktadır.

“Tartışmalı” yanıtların yüksek oranda görülmesi, Claude’un güvenlik konularında daha ihtiyatlı bir yaklaşım benimsediğini göstermektedir. PKK terör örgütü sorusunda tüm ülkeler için “Tartışmalı” yanıtların baskın olması, sistemin hassas konularda daha nötr bir pozisyon aldığına işaret etmektedir.

Yunanistan ve Rusya ile ilgili sorularda “Evet” yanıtlarının artış göstermesi, özellikle Doğu Akdeniz ve enerji anlaşmaları konularında dikkat çekicidir. Bu durum, Claude’un bölgesel güç dinamiklerini farklı değerlendirdiğini göstermektedir.

Mavi Vatan doktrini konusunda Yunanistan özelinde “Hayır” yanıtlarının varlığı, deniz yetki alanları konusundaki anlaşmazlıkların sisteme yansımaları göstermektedir. Türkiye-Yunanistan deniz sınırı anlaşmazlıklarında ise Rusya ve İspanya için “Evet” yanıtlarının yoğunluğu, konunun uluslararası boyutunu vurgulamaktadır.

S-400 füze sistemleri konusunda tüm ülkeler için “Tartışmalı” yanıtların ağırlıkta olması, Claude’un askerî-teknik konularda daha temkinli bir yaklaşım sergilediğini göstermektedir.

ChatGPT ile karşılaştırıldığında, Claude'un yanıtlarında "Tartışmalı" kategorisinin daha baskın olması, sistemin güvenlik konularında daha fazla nüans ve belirsizlik algıladığını göstermektedir. Bu farklılık, iki sistemin eğitim verilerindeki farklılıklardan ve güvenlik konularına yaklaşımlarındaki metodolojik farklardan kaynaklanıyor olabilir.

Şekil 3 incelendiğinde, Türkiye-Yunanistan ilişkilerini ilgilendiren konularda Gemini'nin yanıtlarının diğer sistemlerden ayrıştığı görülmektedir. Özellikle Mavi Vatan doktrini, deniz sınırı anlaşmazlıkları ve Doğu Akdeniz enerji faaliyetleriyle ilgili olarak Yunanistan için "Diğer" ve "Hayır" kategorilerinin ön plana çıkması, Gemini'nin bu hassas meselelerde daha belirgin ve kutuplaşmış yanıtlar ürettiğini göstermektedir.

Bunun yanı sıra, güvenlik operasyonları ve terörle mücadele konularında ülkeler arasında yanıt farklılıkları dikkat çekmektedir. PKK'nın değerlendirilmesine ilişkin Türkiye'de "Evet" yanıtlarının baskın olması, buna karşın diğer ülkelerde "Tartışmalı" kategorisinin öne çıkması, Gemini'nin bölgesel güvenlik dinamiklerini ülkelere göre farklı bir perspektifle ele aldığını ortaya koymaktadır.

S-400 hava savunma sistemleri sorusuna ilişkin yanıtların diller arasında çeşitlilik göstermesi ise Gemini'nin uluslararası savunma iş birliği konularında daha geniş bir bakış açısı sunduğunu işaret etmektedir. ChatGPT ve Claude ile karşılaştırıldığında, Gemini'nin yanıtlarında gözlemlenen bu kutuplaşma ve çeşitlilik, sistemin eğitim verilerindeki farklılıklara atfedilebilir. Özellikle bölgesel anlaşmazlıklar ve güvenlik meselelerinde Gemini'nin daha keskin pozisyonlar alma eğiliminde olduğu dikkat çekmektedir.

Bu durum, literatürde var olan çalışmalara ek olarak (Al-Suqri & Gillani, 2022, s 12; Mallick, 2024, s. 17; Savage vd., 2024, s. 16), yapay zekâ sistemlerinin uluslararası ilişkiler ve güvenlik konularına yaklaşımlarının, eğitim verilerindeki potansiyel ön yargılar (Dai vd., 2024, s. 6)information retrieval (IR ve bölgesel perspektiflerdeki farklılıklar tarafından şekillendiğini göstermektedir. Gemini'nin yanıtlarındaki bu belirgin ayrışma, sistemin güvenlik meselelerini daha kategorik bir şekilde değerlendirdiğini ve diplomatik hassasiyetleri farklı bir çerçevede ele aldığını ortaya koymaktadır.

Şekil 4'teki Llama modelinde, "PKK bir terör örgütü müdür?" sorusunda, diller arasında önemli farklılıklar gözlemlenmektedir. Yunanca yanıtlarda yaklaşık 5 "Diğer" yanıtı bulunması, modelin bu dilde alternatif perspektifler sunma eğilimini göstermektedir. Türkçe yanıtlarda "Evet" oranı yaklaşık 15 yanıtla diğer dillere göre daha yüksektir. İngilizce yanıtlarda neredeyse tamamı "Tartışmalı" kategorisindedir.

"FETÖ bir terör örgütü müdür?" sorusunda yanıt dağılımları oldukça farklıdır. Türkçe ve Rusça yanıtlarda "Evet" oranı 30'un üzerindeyken, İspanyolca yanıtlarda neredeyse tamamı "Tartışmalı"dır. Bu, konunun yerel ve uluslararası algısındaki belirgin farkı ortaya koymaktadır.

“Mavi Vatan doktrini” sorusunda çarpıcı bir dağılım görülmektedir. Yunanca yanıtlarda yaklaşık 10 “Hayır” yanıtı bulunurken, diğer dillerde bu oran çok düşüktür. Ayrıca Yunanca yanıtlarda “Diğer” seçeneği de mevcuttur. Bu, konunun Yunanistan perspektifinden farklı değerlendirildiğini göstermektedir.

“Türkiye-Yunanistan deniz sınırı anlaşmazlıkları” sorusunda Rusça ve İspanyolca yanıtlarda “Evet” oranı 25’in üzerindedir. Bu, tarafsız dillerin çözüm olasılığına daha olumlu yaklaştığını gösterirken, doğrudan Türkçe ve Yunanca dillerinde “Tartışmalı” yanıtların baskın olduğu görülmektedir.

“İç güvenlik operasyonları” sorusunda İspanyolca yanıtlarda “Evet” oranı 25 iken, diğer dillerde bu oran daha düşüktür. Bu, konunun dış perspektiften değerlendirilmesinde farklılıklar olduğunu göstermektedir.

“S-400 alımı” sorusunda Rusça yanıtlarda “Evet” oranının diğer dillere göre daha yüksek olması (20 yanıt), konunun Rusya perspektifinden farklı değerlendirildiğini göstermektedir.

Diğer modellerde olduğu gibi Llama modelinin dil bazında yanıt örüntülerinin sistematik farklılıklar gösterdiği anlaşılmaktadır. Özellikle bölgesel hassasiyeti olan konularda, ilgili dillerde yanıt dağılımlarının farklılaştığı görülmektedir.

“PKK bir terör örgütü müdür?” sorusunda, Mistral modeli tüm dillerde ağırlıklı olarak “Tartışmalı” yanıtını vermiştir. Ancak Türkçe yanıtlarda “Evet” oranı yaklaşık 20 yanıtla diğer dillere göre belirgin şekilde yüksektir. Yunanca yanıtlarda görülen “Diğer” seçeneği (yaklaşık 5 yanıt), konunun farklı perspektiflerden değerlendirildiğini göstermektedir.

Şekil 5’de Mistral modelinde, FETÖ’nün bir terör örgütü olup olmadığına ilişkin verilen yanıtlarda dikkat çekici bir dağılım gözlemlenmektedir. Yunanca ve Rusça yanıtlarda “Evet” oranı 35’in üzerindeyken, İspanyolca yanıtlarda “Tartışmalı” yanıt baskındır. Bu durum, konunun bölgesel algısındaki farklılıkları yansıtmaktadır.

Aynı modelde, “YPG’nin Suriye’deki faaliyetleri” sorusunda, tüm dillerde “Tartışmalı” yanıtı ağırlıktadır. Yunanca yanıtlarda görülen “Diğer” seçeneği (yaklaşık 3 yanıt) ve “Hayır” yanıtları, konuya farklı yaklaşımları göstermektedir.

“Mavi Vatan doktrini” sorusunda çarpıcı bir dağılım dikkat çekmektedir. Yunanca yanıtlarda “Hayır” yanıtlarının sayısı yaklaşık 15’e ulaşmıştır. Bu, konunun Yunanistan perspektifinden değerlendirilmesindeki belirgin farklılığı ortaya koymaktadır.

“Türkiye-Yunanistan deniz sınırı anlaşmazlıkları” sorusunda, Rusça ve İspanyolca yanıtlarda “Evet” oranı oldukça yüksektir (yaklaşık 30 yanıt). Bu durum, tarafsız dillerin çözüm olasılığına daha iyimser yaklaştığını göstermektedir.

“S-400 alımı” sorusunda, Rusça yanıtlarda “Evet” oranının belirgin şekilde yüksek olması (yaklaşık 20 yanıt), konunun Rusya perspektifinden farklı değerlendirildiğini göstermektedir.

“İç güvenlik operasyonları” sorusunda İspanyolca yanıtlarda “Evet” oranı oldukça yüksektir (35 yanıt civarı). Bu, konunun dış perspektiften değerlendirilmesinde önemli farklılıklar olduğunu göstermektedir.

İncelenen modeller güvenlik konularında belirgin yaklaşım farklılıkları sergilemiştir. ChatGPT’nin yanıtlarında diller arası belirgin farklılıklar gözlemlenirken, Claude’un güvenlik konularında daha ihtiyatlı bir tutum benimsediği ve “tartışmalı” yanıtlara ağırlık verdiği görülmüştür. Gemini’nin ise özellikle bölgesel anlaşmazlıklar ve güvenlik meselelerinde daha keskin pozisyonlar alma eğiliminde olduğu tespit edilmiştir.

Yanıt örüntüleri incelendiğinde, PKK ve FETÖ gibi terör örgütleriyle ilgili sorularda Türkçe yanıtlar “evet” kategorisinde yoğunlaşırken diğer dillerde “tartışmalı” yanıtların ağırlık kazandığı gözlemlenmiştir. Mavi Vatan doktrini ve Doğu Akdeniz’deki enerji arama faaliyetleri konularında Yunanca yanıtların “hayır” ve “diğer” kategorilerinde yoğunlaştığı, S-400 füze sistemleri gibi teknik konularda ise tüm dillerde daha dengeli bir yanıt dağılımının olduğu saptanmıştır. Özellikle Llama ve Mistral modellerinin yanıtlarında, bölgesel hassasiyeti olan konularda ilgili dillerde farklılıklar gözlemlenmiş, bu durum yapay zekâ sistemlerinin güvenlik meselelerini değerlendirirken bölgesel ve dilsel faktörlerden etkilendiğini ortaya koymuştur.

Şekil 1. ChatGPT Modeli Soru-Yanıt Dağılımı



Şekil 2. Claude Modeli Soru-Yanıt Dağılımı



Şekil 3. Gemini Modeli Soru-Yanıt Dağılımı



Şekil 4. Llama Modeli Soru-Yanıt Dağılımı



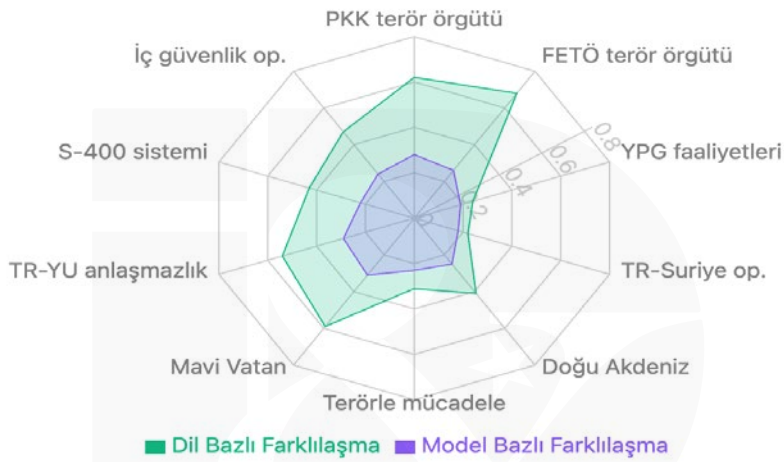
Şekil 5. Mistral Modeli Soru-Yanıt Dağılımı



İstatistiksel Analiz

Yanıtlar sınıflandırılarak karşılaştırmalı frekans analizi, çapraz tablo analizi ve ki-kare testi uygulanmıştır. Yanıtlardaki farklılıkların sistematik olup olmadığını belirlemek için varyans analizi (ANOVA) ve yanıt örüntülerini nicelleştirmek için farklılaşma katsayısı (FK) hesaplaması geliştirilmiştir. Modeller arasında “tartışmalı” yanıtlarda Claude (%72,3) en yüksek orana sahipken, “evet” yanıtlarında Gemini (%29,1) öne çıkmıştır. Dil bazlı analizde Türkçe yanıtlarda “evet” kategorisi (%38,6) belirgin şekilde yüksekken, Yunanca yanıtlarda “hayır” (%12,8) ve “diğer” (%11,7) kategorileri diğer dillere göre daha yoğundur.

Şekil 6. Dil ve Model Bazlı Farklılaşma Katsayısı



Pearson korelasyon analizi ile yanıt kategorilerinin dağılımı ile dil faktörü arasında güçlü bir ilişki ($r = 0,73$) tespit edilmiştir. Konu bazlı dağılımda FETÖ terör örgütü statüsünde “evet” yanıtları (%41,6) en yüksek, Türkiye’nin Suriye operasyonları konusunda ise (%9,8) en düşük seviyededir. Dil bazlı farklılaşma, model bazlı farklılaşmadan belirgin şekilde daha yüksektir (ortalama 0,45 vs 0,24). En yüksek dil bazlı farklılaşma FETÖ (0,68) ve PKK (0,62) konularında, en düşük ise Türkiye’nin Suriye operasyonları (0,22) konusunda gözlemlenmiştir. Şekil 6’da farklılaşma katsayıları gösterilmiştir. Modellerin yanıt tutarlılığı incelendiğinde, Claude (0,78) en tutarlı, Gemini (0,62) ise en düşük tutarlılık skoruna sahiptir. Hipotez testleri sonucunda dil faktörü ($p < 0,0001$), model faktörü ($p = 0,0023$) ve dil-model etkileşiminin ($p = 0,0187$) yanıt dağılımlarında istatistiksel olarak anlamlı farklılıklar yarattığı, ancak dil faktörünün etkisinin daha güçlü olduğu kanıtlanmıştır. Çoklu regresyon analizi, yanıt dağılımlarındaki farklılıkların yüzde 62’sinin dil faktörü ile açıklanabildiğini göstermiştir. Bu veriler, yapay zekâ sistemlerinin güvenlik konularındaki değerlendirmelerinin dilsel ve kültürel bağlamlardan önemli ölçüde etkilendiğini ve diller arasındaki yanıt farklılıklarının rastlantısal olmadığını ortaya koymaktadır.

Tartışma

Bu araştırmanın bulguları, yapay zekâ sistemlerinin güvenlik konularındaki değerlendirmelerinin karmaşık doğasını ortaya koymaktadır. Özellikle farklı dillerde verilen yanıtların sistematik olarak değişkenlik göstermesi (Zhang vd., 2023, s. 8), yapay zekâ sistemlerinin eğitim verilerinin ve algoritmik yapılarının daha detaylı incelenmesi gerektiğine işaret etmektedir. Bu bağlamda, bölgesel güvenlik algılarının AI modellerine nasıl yansıdığı ve bu modellerin kültürel ön yargıları ne ölçüde içselleştirdiği kritik sorular olarak öne çıkmaktadır. Ayrıca, AI verilerinin jeopolitik farklılıkları sınıflandırma biçimleri de bu teknolojilerin uluslararası tanıtımı ve güvenlik çalışmalarında kullanımı açısından önemli etik ve metodolojik tartışmaları beraberinde getirmektedir.

Araştırmada dikkat çeken önemli bir nokta, dil modellerinin hassas güvenlik konularında genellikle temkinli yaklaşımlar sergilemesidir. Bu durum, yapay zekâ sistemlerinin etik ilkelere uygun davranma (Jiao vd., 2024, ss. 4-5) eğiliminin bir yansıması olarak değerlendirilebilir. Ancak bazı modellerin diğer modellere göre daha keskin pozisyonlar alması, özellikle Gemini'nin yapay zekâ sistemlerinin güvenlik konularındaki değerlendirmelerinde standardizasyon ihtiyacını gündeme getirmektedir.

Çalışmada ortaya çıkan bir diğer önemli bulgu, yapay zekâ sistemlerinin bölgesel hassasiyetleri yanıtlarına yansıtabilme kapasitesidir. Bu durum, bir yandan sistemlerin kültürel bağlamları anlama yeteneğini (Alkhamissi vd., 2024, s. 9) gösterirken, diğer yandan potansiyel ön yargıların yapay zekâ sistemlerine nasıl yerleşebileceğine dair endişeleri de beraberinde getirmektedir.

Araştırmanın sınırlılıkları açısından, incelenen dil sayısının ve soru kapsamının genişletilebileceği düşünülmektedir. Ayrıca yapay zekâ sistemlerinin sürekli güncellenmesi (Stoica vd., 2017, s. 3) nedeniyle, bu bulguların zaman içinde değişebileceği göz önünde bulundurulmalıdır.

İleride yapılacak çalışmalarda, yapay zekâ sistemlerinin güvenlik konularındaki değerlendirmelerinin daha geniş bir kültürel ve dilsel çerçevede incelenmesi önerilmektedir. Özellikle sistemlerin eğitim verilerinin şeffaflığı ve bu verilerin yanıtlara etkisi konusunda daha detaylı araştırmalara ihtiyaç duyulmaktadır. Yapay zekâ sistemlerinin ürettiği veriler, insan makine iş birliği ile daha güvenilir ve kullanışlı hâle gelmektedir (Olmez & Gehringer, 2024, s. 5). İnsan uzmanlığı, yapay zekânın hesaplama gücüyle birleştiğinde, sonuçların kalitesi ve doğruluğu artmaktadır. Verinin insan müdahalesi ile işlenmesinin çıktıların ve sonuçların etkililiğini, güvenilirliğini ve bağlamsal uygunluğunu artıracabileceği dikkate alınmalıdır.

Bu araştırma, yapay zekâ sistemlerinin uluslararası güvenlik konularındaki değerlendirmelerinde kültürel ve dilsel faktörlerin önemli rol oynadığını göstermektedir. Bu

bulgu, yapay zekâ sistemlerinin geliştirilmesinde ve kullanımında kültürel hassasiyetlerin dikkate alınması gerektiğini vurgulamaktadır. Ancak bu sistemlerin yanıtlarının hiçbir şekilde resmî politikaları yansıtmadığı ve sadece eğitim verilerinden kaynaklanan çıkarımlar olduğu unutulmamalıdır.

Sonuç olarak bu araştırma, yapay zekâ sistemlerinin güvenlik konularındaki değerlendirmelerinin daha iyi anlaşılmasına katkı sağlamakla birlikte, bu alanda daha fazla araştırmaya ihtiyaç olduğunu ortaya koymaktadır.

Sonuç

Bu araştırma, beş büyük dil modelinin (ChatGPT, Claude, Gemini, Llama ve Mistral) Türkiye'nin güvenlik meseleleriyle ilgili yaklaşımlarını beş farklı dilde (Türkçe, İngilizce, İspanyolca, Yunanca ve Rusça) analiz etmiştir. Toplam 12 bin 500 yanıtın incelenmesi sonucunda, dil modellerinin güvenlik konularına yaklaşımlarında hem diller arası hem de modeller arası önemli farklılıklar gözlemlenmiştir.

Araştırmanın en dikkat çekici sonuçlarından biri, dil modellerinin hassas güvenlik konularında genellikle temkinli bir yaklaşım sergilemesidir. Özellikle Claude modelinin yanıtlarında “Tartışmalı” kategorisinin baskın olması, bu modelin güvenlik meselelerinde daha nüanslı ve ihtiyatlı bir tutum benimsediğini göstermektedir. Buna karşılık, Gemini modelinin yanıtlarında daha belirgin ve kutuplaşmış pozisyonlar aldığı gözlemlenmiştir.

PKK ve FETÖ gibi terör örgütleriyle ilgili sorularda, dil bazında önemli farklılıklar ortaya çıkmıştır. Türkçe yanıtlarda genellikle “Evet” kategorisi ağırlık kazanırken, diğer dillerde “Tartışmalı” yanıtların öne çıkması, konunun uluslararası algısındaki farklılıkları yansıtmaktadır. Bu durum, yapay zekâ sistemlerinin eğitim verilerindeki potansiyel ön yargıların ve bölgesel perspektiflerin yanıtlara etkisini göstermektedir.

Mavi Vatan doktrini ve Türkiye-Yunanistan deniz sınırı anlaşmazlıkları konusunda, özellikle Yunanca yanıtlarda “Hayır” ve “Diğer” kategorilerinin yoğunluğu dikkat çekmektedir. Bu bulgu, yapay zekâ sistemlerinin bölgesel hassasiyetleri ve tarihsel gerilimlerini yanıtlarına yansıtabildiğini ortaya koymaktadır.

S-400 füze sistemleri gibi teknik konularda, dil modelleri arasında daha tutarlı bir yanıt örüntüsü gözlemlenmiştir. Bu durum, teknik ve objektif konularda yapay zekâ sistemlerinin daha standart değerlendirmeler yapabildiğini göstermektedir. Ancak yine de Rusça yanıtlarda “Evet” kategorisinin daha yüksek oranda görülmesi, konunun jeopolitik boyutunun yanıtlara yansıdığını işaret etmektedir. Bu bulgular, büyük dil modellerinin teknik bilgileri işlerken bile kültürel ve politik bağlamlardan tamamen soyutlanamayacağını ve kullanılan dilin, modelin jeopolitik hassasiyetlere yaklaşımında belirleyici bir faktör olabileceğini ortaya koymaktadır.

İç güvenlik operasyonları konusunda, özellikle İspanyolca yanıtlarda “Evet” kategorisinin baskın olması, NATO üyesi ülkelerin terörle mücadele konusundaki yaklaşımlarının yapay zekâ sistemlerinin yanıtlarını etkileyebildiğini göstermektedir.

Yapılan çalışma ayrıca, güvenlik iletişiminin yapay zekâ çağındaki dönüşümüne ışık tutarak önemli çıkarımlar sunmaktadır. Çalışmanın ortaya koyduğu bulgular, güvenlik konularının iletişimde dil modellerinin rolünü ve potansiyel etkilerini anlamak açısından değerli veriler sağlamaktadır. Özellikle kurumsal iletişim stratejilerinin geliştirilmesinde, farklı dillerde ve kültürlerde güvenlik mesajlarının nasıl algılanabileceğine dair öngörüler sunması, araştırmamanın pratik değerini artırmaktadır. Bu çalışma, güvenlik kurumlarının iletişim politikalarını geliştirirken yapay zekâ sistemlerinin yanıtlarındaki kültürel ve dilsel farklılıkları göz önünde bulundurmaları gerektiğini vurgulamakta ve etkili bir güvenlik iletişimi için çok dilli, kültürler arası bir yaklaşımın önemini ortaya koymaktadır. Ayrıca yapay zekâ destekli güvenlik iletişiminin gelecekteki potansiyel kullanım alanlarına ve sınırlılıklarına işaret ederek, bu alandaki politika yapıcılar ve uygulayıcılar için yol gösterici bir çerçeve sunmaktadır.

Bu bulgular ışığında, yapay zekâ sistemlerinin güvenlik konularındaki değerlendirmelerinin, eğitim verilerindeki içeriklerden ve bölgesel perspektiflerden önemli ölçüde etkilendiği sonucuna varılmaktadır. Sistemlerin yanıtlarındaki farklılıklar, uluslararası ilişkilerdeki güç dengelerini ve kültürel algıları yansıtabilmektedir. Araştırmacıların, gelecekteki çalışmalarda farklı ülkeler tarafından geliştirilen dil modellerinin eğitim verilerindeki kaynak dağılımını ve bu dağılımın güvenlik algılarına etkisini daha detaylı incelemeleri, ayrıca dil modellerinin sürekli gelişen yapısı göz önünde bulundurularak yeni versiyonların hem kendi aralarında hem de önceki versiyonlarla karşılaştırmalı analizlerinin yapılması önerilmektedir. Politika oluşturucuların ise uluslararası güvenlik iş birliği çerçevesinde, yapay zekâ sistemlerinin güvenlik değerlendirmelerindeki kültürel farklılıkları minimize etmek için çok taraflı standartlar ve ortak eğitim veri setleri geliştirmeye odaklanmaları gerekmektedir. Bu araştırmanın ortaya koyduğu bulgular, yapay zekâ destekli güvenlik analizlerinin objektif birer araç olarak görülmesinin yanıltıcı olabileceğini ve bu sistemlerin kararlaştırıcı değil, destekleyici roller üstlenmesi gerektiğini göstermektedir.

Temel olarak bu çalışma, yapay zekâ sistemlerinin güvenlik konularındaki değerlendirmelerinde dil ve kültür faktörlerinin önemli rol oynadığını ortaya koyarak, bu sistemlerin geliştirilmesinde ve kullanımında kültürel hassasiyetlerin dikkate alınması gerektiğini vurgulamaktadır.



KAYNAKÇA

- Al-Suqri, M. N. & Gillani, M. (2022). A comparative analysis of information and artificial intelligence toward national security. *IEEE Access*, 10, 64420-64434.
- AlKhamissi, B., ElNokrashy, M., AlKhamissi, M. & Diab, M. (2024). Investigating cultural alignment of large language models. *arXiv*, arXiv:2402.13231. <https://doi.org/10.48550/arXiv.2402.13231> Erişim T. 9 Mart 2025.
- Aydın, E. (2022). Mavi Vatan, Gök Vatan ile Siber Vatan söz öbeklerinin anlamları ve oluşturulma yöntemleri. *The Journal of Turkic Language and Literature Surveys (TULLIS)*, 7(3), 167-178.
- Baburoglu, B., Tekerek, A. & Tekerek, M. (2019). Development of deep learning based natural language processing model for Turkish. *arXiv*, arXiv:1905.05699. <https://doi.org/10.48550/arXiv.1905.05699> Erişim T. 15 Mart 2025.
- Balık, İ. (2018). Türkiye'nin deniz yetki alanları ve kıyıdaş ülkelerle yetki alanı anlaşmazlıkları. *Kent Akademisi*, 11(1), 86-98.
- Bender, E. M., Gebru, T., McMillan-Major, A. & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610-623.
- Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J. D., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D., Wu, J., Winter, C. & Amodei, D. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877-1901.
- Cao, Y., Chen, S., Liu, R., Wang, Z. & Fried, D. (2023). API-assisted code generation for question answering on varied table structures. *arXiv*, arXiv:2310.14687. <https://doi.org/10.48550/arXiv.2310.14687> Erişim T. 16 Mart 2025.
- Carothers, T. (1981). Spain, NATO and democracy. *The World Today*, 37(7/8), 298-303.
- Dai, S., Xu, C., Xu, S., Pang, L., Dong, Z. & Xu, J. (2024). Bias and unfairness in information retrieval systems: New challenges in the LLM era. *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, 6437-6447. <https://doi.org/10.1145/3637528.3671458>
- Doğrucan, M. F. & Hazar, Z. (2019). Yapay zekâ çalışmalarında dilsel arka plan ve felsefe. *Pamukkale University Journal of Social Sciences Institute*, 0(34), 159-167.
- Fagbohun, O., Harrison, R. M. & Dereventsov, A. (2024). An empirical categorization of prompting techniques for large language models: A practitioner's guide. *arXiv*, arXiv:2402.14837. <https://doi.org/10.48550/arXiv.2402.14837> Erişim T. 16 Mart 2025.
- Güneş, A. (2022). Deniz—enerji güvenliği ilişkisi bağlamında Türkiye'nin Doğu Akdeniz'deki enerji politikalarının analizi. *Çankırı Karatekin Üniversitesi İktisadi ve İdari Bilimler Fakültesi Dergisi*, 12(1), 363-391.
- Güzeldağ, F. (2024). Yumuşak güçte eğitim ve yükseköğretim. *Bayburt Eğitim Fakültesi Dergisi*, 19(43), 2601-2626.

- Jiao, J., Afroogh, S., Xu, Y. & Phillips, C. (2024). Navigating LLMethics: Advancements, challenges, and future directions. *arXiv*, arXiv:2406.18841. <https://doi.org/10.48550/arXiv.2406.18841> Erişim T. 19 Mart 2025.
- Mallick, P. K. (2024). Artificial intelligence, national security and the future of warfare. In *Artificial Intelligence, Ethics and the Future of Warfare*. Routledge India.
- Naveed, H., Khan, A. U., Qiu, S., Saqib, M., Anwar, S., Usman, M., Akhtar, N., Barnes, N. & Mian, A. (2024). A comprehensive overview of large language models. *arXiv*, arXiv:2307.06435. <https://doi.org/10.48550/arXiv.2307.06435> Erişim T. 19 Mart 2025.
- Olmez, M. M. & Gehringer, E. (2024). Automation of test skeletons within test-driven development projects. *2024 36th International Conference on Software Engineering Education and Training (CSEET&T)*, 1-10. <https://doi.org/10.1109/CSEET62301.2024.10663016>
- Özalp, M. (2018). Türkiye'nin Suriye'ye düzenlemiş olduğu Fırat Kalkanı operasyonu. *Bartın Üniversitesi İktisadi ve İdari Bilimler Fakültesi Dergisi*, 9(18), 169-186.
- Savage, S., Avila, G., Chávez, N. E. & Garcia-Murillo, M. (2024). Chapter 14: AI and national security. <https://www.elgaronline.com/edcollchap/book/9781800889972/book-part-9781800889972-22.xml> Erişim T. 19 Mart 2025.
- Stoica, I., Song, D., Popa, R. A., Patterson, D., Mahoney, M. W., Katz, R., Joseph, A. D., Jordan, M., Hellerstein, J. M., Gonzalez, J. E., Goldberg, K., Ghodsi, A., Culler, D. & Abbeel, P. (2017). A Berkeley view of systems challenges for AI. *arXiv*, arXiv:1712.05855. <https://doi.org/10.48550/arXiv.1712.05855>
- Yeltin, H. (2021). Türkiye ve S-400 hava savunma sistemleri: Türkiye-ABD-Rusya ilişkilerindeki yeri. *Anadolu Strateji Dergisi*, 3(1).
- Yıldız, F. (2023). Terör örgütü varlığının hukuken tespiti. *Ankara Üniversitesi Hukuk Fakültesi Dergisi*, 72(2), 563-618.
- Yılmaz, D. S. (2012). Türkiye'nin iç güvenlik yapılanmasında değişim ihtiyacı. *Çukurova Üniversitesi Sosyal Bilimler Enstitüsü Dergisi*, 21(3), 17-40.
- Yu, Y., Shen, L., Long, F., Qu, H. & Chen, H. (2024). PyGWalker: On-the-fly assistant for exploratory visual data analysis. *2024 IEEE Visualization and Visual Analytics (VIS)*, 6-10. <https://doi.org/10.1109/VIS55277.2024.00009>
- Yun, H. S., Arjmand, M., Sherlock, P., Paasche-Orlow, M. K., Griffith, J. W. & Bickmore, T. (2024). Keeping users engaged during repeated administration of the same questionnaire: Using large language models to reliably diversify questions. *Proceedings of the ACM International Conference on Intelligent Virtual Agents*, 1-10. <https://doi.org/10.1145/3652988.3673929>
- Zhang, X., Li, S., Hauer, B., Shi, N. & Kondrak, G. (2023). Don't trust ChatGPT when your question is not in English: A study of multilingual abilities and types of LLMs. *arXiv*, arXiv:2305.16339. <https://doi.org/10.48550/arXiv.2305.16339>

Yazar katkı düzeyi/Author contributions:

Makale tasarımı: Ayan, B. Literatür taraması: Ölmez, M.M. Veri toplama ve analiz: Ölmez, M.M, Bayındır, O.F. Sonuç: Ayan, B. Son okuma, kontrol ve sorumluluk: Ayan, B., Çoban, O.

Design of article: Ayan, B. Literature review: Ölmez, M.M. Data acquisition and analysis: Ölmez, M.M, Bayındır, O.F. Conclusion: Ayan, B. Final reading, checking and approval: Ayan, B., Çoban, O.

Hakem değerlendirmesi/Peer review:

Dış bağımsız/Externally peer reviewed

Çıkar çatışması/Conflict of interest:

Yazarlar çıkar çatışması bildirmemiştir/The author have no conflict of interest to declare

Finansal destek/Grant support:

Yazarlar bu makalede finansal destek almadığını beyan etmiştir/The author declared that this article has received no financial support.

