

BÜSAM ANALİZ

**YAPAY ZEKÂNIN ULUSLARARASI İLİŞKİLERDE AKTÖRLEŞMESİ:
ANTHROPİC-OPENAI AYRIŞMASI
ÜZERİNDEN BİR OKUMA**

Doç. Dr. Kıvılcım Romya Bilgin

ŞUBAT 2026





Yapay Zekânın Uluslararası İlişkilerde Aktörleşmesi: Anthropic–OpenAI Ayrışması Üzerinden Bir Okuma

Trump yönetimi 28 Şubat 2026'da, İsrail ile koordineli olarak İran'a başlattığı hava saldırısı ABD ordusunun, Orta Doğu'daki ABD Merkez Komutanlığı (CENTCOM), bünyesinde istihbarat değerlendirmesi, hedef tespiti ve muharebe senaryolarının simülasyonu amacıyla Anthropic şirketi tarafından geliştirilen büyük dil modeli (LLM) Claude başta olmak üzere yapay zekâ araçlarından yoğun biçimde yararlandığı, yapay zekânın askerî sistemlere ne denli derinlemesine entegre olduğunu gözler önüne sermiştir. Bu askerî operasyon sırasında yapay zekânın CENTCOM düzeyinde yoğun biçimde kullanılması, söz konusu teknolojilerin tesadüfî ya da geçici araçlar olmadığını, aksine teknoloji alanındaki belirli şirketlerle yürütülen uzun vadeli, kurumsallaşmış ve planlı bir entegrasyon sürecinin ürünü olduğunu ortaya koymaktadır.

Nitekim, 2021 yılında eski OpenAI araştırmacıları tarafından kurulan Anthropic, 2024 sonlarından itibaren ve 2025 boyunca ABD Savunma Bakanlığı ve istihbarat kurumlarıyla yoğun ve kurumsallaşmış bir etkileşim sürecine girmiştir. Kasım 2024'te, Palantir ve Amazon Web Services (AWS) ile ortaklık kurarak Claude'u, gizli ortamlar dâhil olmak üzere ABD savunma ve istihbarat sistemlerine entegre etmeye başlamıştır. Haziran 2025'te ise, devlet ve ulusal güvenlik iş akışlarına özel olarak tasarlanan Claude Gov sürümünü tanıtarak bu sürümün, 2025'in sonlarına gelindiğinde ABD istihbarat ve savunma kurumlarında pilot ve operasyonel nitelikli kullanım alanları bulduğu görülmüştür. Bu arada, Temmuz 2025'te ABD Savunma Bakanlığı (Department of Defense – DoD), Anthropic'in de aralarında bulunduğu yapay zekâ firmalarına, toplam hacmi yaklaşık 200 milyon ABD doları değerinde yapay zekâ hizmetleri ihaleleri açmıştır. Savunma Bakanlığı, tek seferlik bir satın almadan ziyade, yapay zekâ şirketlerinin ve ürünlerinin ABD ordusuna ve operasyonel süreçlere kademeli ve kurumsallaşan bir entegrasyon süreci şeklinde ihaleleri ele alarak, farklı başlıklarda entegrasyonu kademelendirmeyi tercih etmiştir.

Yapılan ihaleler kapsamında, aslında Anthropic'in Claude ailesine ait büyük dil modelleri, Şubat 2026 ayı sonlarında yaşanan sözleşme anlaşmazlığına kadar Pentagon'un kullanımındaydı. Sözleşmenin askıya alınmasından önce Anthropic, Palantir ve Amazon Web Services Top Secret Cloud gibi ortaklar aracılığıyla, kamu ve gizli kullanım için özel olarak güvenlik önlemleri güçlendirilmiş araç sürümlerini Pentagon'a sağlamaya devam etmekteydi. Dahası Anthropic'in yapay zekâ araçları gizli ağlara benzeri az görülen bir düzeyde entegre edilerek akıl yürütme, özetleme ve yardımcı görevler gibi işlevler için kullanıma hazır hâle getirilmişti. Bu kapsamda, Şubat 2025 sonlarında, sözleşme yenileme ve kullanım kapsamının genişletilmesi görüşmeleri sırasında Pentagon'un, Anthropic'ten, Claude modelini 'her türlü hukuki amaç için' kullanma hakkı talep etti. Bu talebin, tam otonom öldürücü silah sistemleri ve vatandaşların toplu gözetlenmesi gibi alanlarda da kısıtlamaların kaldırılmasını içermesi, Pentagon ile Anthropic arasında ciddi bir kriz yarattı. Anthropic, yapay zekânın kullanımına ilişkin normatif sınırları belirleyen ve kendisinin öncülük ettiği Constitutional AI çerçevesi gereğince bu etik sınırların kaldırılmasını reddetti. Pentagon ile Anthropic arasında sözleşme üzerinden başlayan etik bir uyuşmazlık olarak gelişen gerilim, kısa sürede devletin ulusal güvenlik söylemi üzerinden şirketin hedef alındığı açık bir müdahaleye dönüşmüştür.

Savunma Bakanı Pete Hegseth, Anthropic'i ulusal güvenlik açısından bir tedarik zinciri riski yarattığını ilan ederek tüm federal kurumların şirketle çalışmasını yasaklamıştır. Trump ise Truth Social üzerinden şirket hakkında sert açıklamalarda bulunmuştur. Anthropic CEO'su Dario Amodei bu yasağı 'misilleme amaçlı ve emsalsiz' olarak nitelendirerek şirkete yönelik

bu tür bir yaptırımın daha önce hiçbir Amerikan firmasına uygulanmadığını vurgulamıştır. Amodei, CBS News'e verdiği demeçte 'Hükümetle anlaşmazlık içinde olmak en Amerikalıca şeydir' diyerek, şirketinin kararının arkasında durduğunu ve Amerikan hükümetinin kararına açık biçimde meydan okuduğunu ifade etmiştir.

Pentagon ile Anthropic arasındaki bu gerilim, 1950 yılında çıkarılan ve ABD başkanına geniş yetkiler tanıyan Savunma Üretim Yasası'nı (Defense Production Act) da yeniden tartışmaya açmıştır. Savunma Üretim Yasası kapsamında ABD Başkanı, şirketler üzerinde mutlak ve keyfi bir zor kullanma yetkisine sahip değildir. Diğer bir ifadeyle ABD Başkanı, şirketleri belirli bir ürünü üretmeye zorlayamaz, kamulaştıramaz ya da şirket yönetimine doğrudan müdahale edemez. Ancak ulusal savunma gereksesiyle, mevcut üretim kapasitesi çerçevesinde belirli sözleşmelerin öncelikli olarak yerine getirilmesini talep edebilir, şirketlerin sipariş reddetme hakkını sınırlandırabilir ve bu talimatlara uyulmaması hâlinde idari veya cezai yaptırımlar uygulayabilir. Bu yetkiler Kongre tarafından çıkarılmış bir yasaya dayandığı için anayasal meşruiyete sahiptir. Şirketler teorik olarak bu yükümlülüklerle itiraz edebilir; ancak itiraz mekanizması otomatik bir hak olmayıp bürokratik ve hukuki açıdan oldukça zorlayıcıdır. Bu nedenle uygulamada büyük savunma ve teknoloji şirketleri, devletle uzun vadeli ilişkilerini, gelecekteki kamu ihalelerini ve düzenleyici baskıyı dikkate alarak çoğunlukla uyum göstermeyi tercih etmektedir. Bu bağlamda Trump yönetiminin, Savunma Üretim Yasası'na dayanarak Anthropic'i köşeye sıkıştırmaya çalıştığı görülmektedir.

Bu süreçte OpenAI ise, Anthropic'e kıyasla daha 'ihtiyatlı fakat devletle uyumlu' bir pozisyon alarak doğrudan muharebe, hedefleme veya öldürücü operasyonlara yönelik kullanımı açık biçimde dışladığını belirtmekle birlikte, istihbarat analizi, senaryo üretimi, stratejik simülasyon ve karar destek gibi alanlarda ABD güvenlik bürokrasisiyle sınırlı ve kontrollü iş birliklerini sürdürmüştür. OpenAI'in bu yaklaşımı, şirketin kamuoyuna dönük etik taahhütleri ile fiili ulusal güvenlik ihtiyaçları arasında denge kurmaya çalışan pragmatik bir meşruiyet stratejisi olarak görülse de kamuoyunda ve yapay zekâ çalışmaları yapan profesyoneller arasında ciddi tepkilere neden olmuştur. Aslında, OpenAI ve Anthropic'in yapay zekânın askerî ve güvenlik bağlamında kullanımına ilişkin yaklaşımları, özellikle İran krizi etrafında belirgin biçimde ayrılmaktadır. OpenAI, 2023–2025 döneminde ABD hükümeti, savunma ve istihbarat ekosistemi ile ulusal güvenlik odaklı projelerde daha açık ve pragmatik bir çizgi izleyerek yapay zekâ modellerinin silah olarak değil, savunma, analiz ve karar destek amaçlı araçlar olarak kullanılabileceğini vurgulamıştır. Bu yaklaşım, doğrudan saldırı üretimini dışlamakla birlikte, istihbarat analizi, açık kaynak istihbaratı, tehdit senaryosu oluşturma ve kriz simülasyonları gibi alanlarda yapay zekânın aktif biçimde devreye sokulmasını mümkün kılmıştır. Ancak bu durum, yapay zekânın özellikle otonom sistemlerle muharebe alanında aktif bir aktör olarak eklenmesi anlamına geldiği yönünde eleştirilere yol açmıştır. Anthropic ise kuruluşundan itibaren askerî kullanıma daha mesafeli ve normatif açıdan temkinli bir duruş benimsemiştir. Şirket, 'frontier model' olarak adlandırdığı gelişmiş yapay zekâ sistemlerinin yüksek riskler barındırdığına dikkat çekerek, yapay zekânın savaş bağlamında kullanımına ilkesel itirazlar geliştirmiştir. Devletlerle iş birliğini bütünüyle reddetmemekle birlikte, aktif çatışma ortamlarında yapay zekâ kullanımına ciddi sınırlamalar getirilmesi gerektiğini savunmaktadır. İran savaşı bağlamında Anthropic, yapay zekânın savaşta karar verici bir aktör olmaması, en fazla pasif ve destekleyici bir analiz aracı olarak konumlandırılması gerektiğini ileri sürmekle birlikte jeopolitik krizlerde bu tür teknolojilerin kullanımının geri döndürülemez normlar üretebileceği uyarısında bulunmuştur. Bu yönüyle Anthropic, OpenAI'ye kıyasla daha etik-merkezli, ihtiyatlı ve normatif güç iddiası daha belirgin bir yaklaşımı temsil etmektedir.

Büyük dil modellerinin savaşılara taktik ve operasyonel entegrasyonu, modern savaşta önemli bir paradigma değişimine işaret etmektedir. Uzmanlar bu araçların hem istihbarat hem de öldürücü operasyonlar açısından çift kullanımlı bir nitelik taşıdığını vurgulamaktadır. Nükleer, kimyasal ve biyolojik silahlar için uluslararası hukuki çerçeveler mevcutken, yapay zekâ hâlen düzenleyici bir gri alanda bulunmaktadır. ABD'nin, İran saldırısı ve Anthropic'e yönelik eş zamanlı devlet müdahalesi, savaşın artık yalnızca fiziksel savaş alanıyla sınırlı kalmadığını göstermektedir. Bu nedenle İran'a yönelik operasyon, aynı zamanda yapay zekânın uluslararası ilişkilerde özerk bir aktör olarak görünürlük kazandığı ilk büyük kırılma anlarından biri olarak değerlendirilebilir.