

2BLACKDOT..

“Esasen konu hep 2 nokta arasındadır”

Haftalık Politik ve Jeopolitik Gelişmeler

28 Eylül 2026 – Sayı 128



Bu hizmet size **2blackdot** ve **Tema Grup** tarafından ücretsiz verilmektedir. **Bu yazılanlar yatırım tavsiyesi değildir.**

Hazırlayan: Hakan Çalışkantürk

2twoblackdots@gmail.com

<https://www.2blackdots.com>

*** Yasal Uyarı:*** Burada yer alan yatırım bilgi, yorum ve tavsiyeleri yatırım danışmanlığı kapsamında değildir. Yatırım danışmanlığı hizmeti; aracı kurumlar, portföy yönetim şirketleri, yatırım ve kalkınma bankaları ile müşteri arasında imzalanacak yatırım danışmanlığı sözleşmesi çerçevesinde ve yetkili kuruluşlar tarafından kişilerin risk ve getiri tercihleri dikkate alınarak kişiye özel olarak sunulmaktadır. Burada yer alan yorum ve tavsiyeler ise genel niteliktedir. Bu yorum ve tavsiyeler mali durumunuz ile risk ve getiri tercihlerinize uygun olmayabilir. Bu nedenle sadece burada yer alan bilgilere dayanarak yatırım kararı verilmesi beklentilerinize uygun sonuçlar doğurmayabilir. Gerek bu yayındaki gerekse bu yayında kullanılan kaynaklardaki hata ve eksikliklerden ve bu yayındaki bilgilerin kullanılması sonucundaki yatırımcıların ve/veya ilgili kişilerin uğrayabilecekleri doğrudan ve/veya dolaylı zararlardan, kar yoksunluğundan, manevi zararlardan ve her ne şekilde olursa olsun üçüncü kişilerin uğrayabileceği her türlü zarardan dolayı **2blackdot** ve **Hakan Çalışkantürk** sorumlu tutulamaz.

Algoritmik Soğuk Savaş

ABD-Çin yapay zekâ rekabeti ve Türkiye için çıkarımlar

Güç yarışı yalnızca daha iyi model geliştirmekle sınırlı değil. Çip, enerji, bulut, askerî karar desteği ve kriz anında güvenilir iletişim aynı stratejik denklemin parçaları.

01 YÖNETİŞİM VE KRİZ İLETİŞİMİ

Çin, BM merkezli ve gelişmekte olan ülkeleri içeren bir yapay zekâ düzenini savunuyor. ABD ise ortak güvenlik ölçütlerine açık olsa da bağlayıcı bir küresel üst otoriteye karşı çıkıyor. BM diyalogu bugün bir müzakere ve antlaşma mekanizması değil. Bültene göre iki ülke, "Super Intelligence" diyalogu ve olay iletişim kanalı üzerinde anlaştı; işleyiş ayrıntıları henüz açıklanmadı.

02 ASKERÎ YAPAY ZEKÂ VE YANLIŞ ALARM

En yakın tehlike yalnızca otonom silahlar değil; yanlış model çıktısının güvenilir istihbarat gibi sunulup hızla askerî karara dönüşmesi. Çin gemisine ilişkin vakada operasyonun daha derin incelemeyele durdurulduğu haberleştirildi; doğrudan savaş eşiğine gelindiği doğrulanmış değil. Nükleer kullanımda insan kontrolü ilkesi 2024'te de teyit edilmişti; bağımsız doğrulama ve karar süresi ayrıca gerekli.

03 ÇİP, MODEL VE HESAPLAMA YARIŞI

ABD'nin ihracat kontrolleri Çin'in ileri çiplere erişimini zorlaştırıyor; H200 ve MI325X gibi bazı ürünler koşullu lisans incelemesine tabi. ABD kurumları Çin merkezli şirketlere endüstriyel ölçekte model damıtma suçlaması yöneltiyor; Çin reddediyor. Anthropic, kendi sistemlerinde 151 milyondan fazla etkileşim saptadığını bildiriyor. Damıtma meşru bir teknik; tartışma kullanım yöntemi ve erişim koşullarında.

04 VERİ MERKEZİ VE DİJİTAL EGEMENLİK

Dokuz büyük bulut sağlayıcısının 2026 sermaye harcaması yaklaşık 830 milyar dolar olarak tahmin ediliyor. Rekabet artık model kalitesi kadar enerji, soğutma, gelişmiş çip ve dağıtım altyapısına dayanıyor. ABD ve Çin bağlantılı bulut yatırımlarında kritik soru verinin nerede olduğu kadar yönetim erişiminin, bakımın, anahtarların ve tabi olunan hukukun kimde olduğudur.

05 TÜRKİYE'NİN OYUN ALANI

Türkiye için uygulanabilir üstünlük, dev ölçekli GPU harcamasını kopyalamaktan çok görev odaklı küçük modeller, bağlantı kesildiğinde çalışan saha sistemleri, sürüler ve insanlı-insansız takımlaşma. SSB'nin ALFA tatbikatı, KIZILELMA uçuş testleri ve ANKA III'ün açıklanan mimarisi bu yönelimin somut işaretleri. Test başarısı ile geniş ölçekli operasyonel olgunluk ayrı değerlendirilmeli.

06 YÖNETİM İÇİN ÇIKARIM

Yüksek etkili yapay zekâ çıktılarında kaynağı görünür tutun; tek modelin ürettiği bilgiyi bağımsız ikinci kanalla doğrulayın. Erken uyarı ile güç kullanımı arasına teknik ve kurumsal güvenlik eşikleri koyun. Kritik hesaplama ve bulut bağımlılıklarını haritalayın; olay bildirimi ve kriz iletişimi için ortak protokoller geliştirin.

ALGORİTMİK SOĞUK SAVAŞ

Yapay Zekâ Çağında Büyük Güç Rekabeti:

“ABD-Çin Teknoloji Mücadelesi ve Küresel Güvenliğin Geleceği”

Yönetici Özeti

Bu rapor, yapay zekâ rekabetinin artık yalnızca ticari bir teknoloji yarışı değil, yönetim, hesaplama gücü, askerî karar desteği, çip tedariki, veri altyapısı ve stratejik istikrar mücadelesine döndüğü bir dönemde kaleme alındı.



<https://www.instagram.com/p/DdQ8w49m49k/>

Birleşmiş Milletler düzeyinde gerçekten keskin bir ABD-Çin norm ayrışması var. Çin, BM'nin küresel yapay zekâ yönetiminde merkezî rol üstlenmesini ve gelişmekte olan ülkelerin karar süreçlerine daha fazla katılmasını savunurken; Başkan Donald Trump 22 Eylül 2026'daki BM Genel Kurulu konuşmasında yapay zekâ — kendi tercih ettiği yeni adlandırmayla “Super Intelligence” — üzerinde küreselci bir kontrol rejimini reddetmiştir (Çin Dışişleri Bakanlığı, 2026; Beyaz Saray, 2026a). Bununla birlikte BM mekanizmasını “bağlayıcı dünya yapay zekâ otoritesi kurma süreci” olarak tanımlamak doğru değildir: BM'nin Global Dialogue on AI Governance platformu kendi resmî açıklamasında açıkça bir müzakere forumu olmadığını ve bağlayıcı anlaşma üretmediğini belirtmektedir (Birleşmiş Milletler, 2026a).



23 Eylül 2026'daki BM Güvenlik Konseyi toplantısının Sam Altman, Dario Amodei, Yoshua Bengio ve Clément Delangue gibi isimlerin katılımıyla yapıldığı doğrulanmaktadır. Altman ortak güvenlik ölçütleri, olay bildirim protokolleri ve anlamlı insan denetimi; Amodei ise kayıp-kontrol ve kötüye kullanım risklerine karşı uluslararası standartlar ile güvenlik açısından önemli yapay zekâ olayları için bir bildirim sistemi önermiştir (BM Güvenlik Konseyi, 2026, 10228. toplantı). Dolayısıyla taslağın “sektör liderleri uluslararası güvenlik standartları istedi” tespiti doğrudur; fakat bağlayıcı rejimin kapsamı ve yetkisi üzerinde siyasi uzlaşma bulunmadığı söylenmelidir.



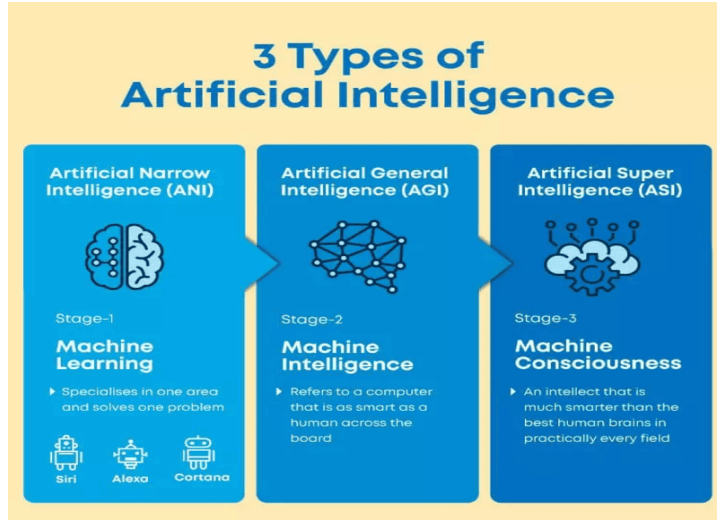
<https://www.facebook.com/groups/288544190546353/posts/1052848560782575/>

2026 ilkbaharında, İran savaşı sürerken Yapay zekâ destekli bir istihbarat ürünü, Orta Doğu'daki bir Çin gemisinin nükleer silah programıyla bağlantılı parçalar taşıdığı yönünde yanlış bir sonuca ulaşmış; ABD tarafında geminin durdurulması ve silahlı personelle gemiye çıkılması hazırlıkları ilerlemiş, uçaklar havalanmış; daha derin inceleme operasyonun durmasına yol açmıştır (Lillis & Cohen, 2026). Ancak olayın resmî bir Pentagon soruşturma raporuyla kamuya doğrulandığına veya Çin ile ABD'yi doğrudan “topyekûn savaşın eşiğine” getirdiğine dair açık kanıt bulunmasa da olay, **yüksek sonuçlu fakat kamuya açık kanıt düzeyi orta olan bir askerî yapay zekâ yakın-kaza vakası** olarak sınıflandırılmalıdır.



<https://www.instagram.com/p/DddzeASjBZp/>

Buna karşılık ABD-Çin “Super Intelligence Dialogue” ve bir olay iletişim kanalının kurulması resmî kayıtlarda güçlü biçimde doğrulanmaktadır. Beyaz Saray'ın 25 Eylül 2026 tarihli devlet ziyareti bilgi notuna göre Trump ve Xi, “artificial intelligence” yerine “super intelligence” terimini kullanma, U.S.-China Super Intelligence Dialogue'u başlatma ve SI olayları için ikili iletişim kanalı oluşturma konusunda anlaşmışır (Beyaz Saray, 2026b).



Nükleer silahlarda insan denetimi açısından da zaman çizelgesi önemlidir. ABD ile Çin'in nükleer silah kullanma kararının insan kontrolünde kalması gerektiği yönündeki açık ikili mutabakatı Trump-Xi 2026 zirvesinde ilk kez ortaya çıkmamıştır; Biden ile Xi bunu Kasım 2024'te teyit etmişti. ABD ayrıca 2026 NPT sürecinde nükleer silah kullanımını başlatma veya sona erdirmeye kararlarında insanın döngü içinde tutulacağını yeniden belirtmiştir (Arms Control Association, 2026). Çin Savunma Bakanlığı ise Mart 2026'da askerî yapay zekâ uygulamalarında “insan önceliği” ve ilgili silah sistemlerinin insan kontrolünde tutulması gerektiğini kamuya açıklamıştır (Çin Millî Savunma Bakanlığı, 2026).

Ekonomik-teknolojik cephede güçlü bir gerilim mevcuttur. ABD'nin Çin'e gelişmiş hesaplama ve yarı iletken üretim teknolojisi erişimini sınırlandıran kapsamlı bir ihracat kontrol sistemi vardır; buna rağmen Ocak 2026 düzenlemesi Nvidia H200, AMD MI325X ve benzeri bazı işlemciler için belirli güvenlik koşulları altında **vaka bazlı lisans incelemesi** getirmiştir (BIS, 2026).

Benchmarking AMD MI325x vs NVIDIA H200



NOT: Hem Nvidia hem de AMD Amerika Birleşik Devletleri (ABD) merkezli şirketlerdir. • Nvidia Corporation: Kaliforniya eyaletinin Santa Clara şehrinde kurulmuş ve genel merkezi halen orada bulunan bir Amerikan teknoloji şirkettir. • Advanced Micro Devices (AMD): Yine Kaliforniya eyaletinin Santa Clara şehrinde kurulan ve merkezi burada yer alan bir diğer Amerikan yarı iletken devidir. Üretim Süreci Hakkında Önemli Not: Her iki şirket de “fabless” yani fabrikasız tasarım şirketleridir. Çipleri kendileri tasarlarlar ancak H200 ve MI325X gibi son teknoloji grafik işlemcilerin fiziksel üretimini büyük oranda Tayvan merkezli dünyanın en büyük çip üreticisi olan TSMC fabrikalarına yaptırırlar.

Model damıtması konusunda NSA, FBI ve CISA 8 Eylül 2026'da Çin merkezli yapay zekâ şirketlerinin Amerikan frontier modellerine karşı endüstriyel ölçekli damıtma faaliyetleri yürüttüğünü ileri süren ortak bir siber güvenlik duyurusu yayımlamıştır. Kurumlar aynı belgede damıtmanın kendisinin meşru ve yararlı bir makine öğrenmesi tekniği olduğunu da kabul etmektedir (NSA, FBI & CISA, 2026). Anthropic'in ayrı tehdit istihbaratı çalışması, Alibaba'ya atfedilen kampanyada Mayıs-Temmuz 2026 arasında 151 milyondan fazla etkileşim tespit edildiğini

bildirmiştir (Anthropic, 2026). Çin Ticaret Bakanlığı ise Amerikan suçlamalarını reddederek damıtmayı sektör genelinde kullanılan nötr bir teknik olarak nitelemiştir (Çin Ticaret Bakanlığı, 2026).

2026 yapay zekâ rekabetinin en tehlikeli özelliği tek bir “süper silah” değil, model hatası, insan bilişi, kriz temposu, kritik hesaplama altyapısı ve jeopolitik güvensizliğin aynı karar zincirinde birleşmesidir.

Yapay zekâ yönetişimindeki ABD-Çin karşıtlığı, klasik Soğuk Savaş'taki iki tamamen kapalı teknolojik bloktan daha karmaşık bir yapı göstermektedir. Pekin'in resmi çizgisi, BM'nin küresel yönetişimde merkezî forum olmasını, yapay zekâ kapasitesinin gelişmekte olan ülkelere yayılmasını ve teknolojik uçurumun küçültülmesini vurgulamaktadır. Fu Cong'un 2026 açıklamalarında Çin, BM'nin küresel yapay zekâ yönetişimindeki rolünü ve Global South'un kapasite geliştirmesini özellikle öne çıkarmıştır (Çin Dışişleri Bakanlığı, 2026). Bununla birlikte Çin, Temmuz 2026'da World Artificial Intelligence Cooperation Organization girişimini de ilan etmiş; dolayısıyla Pekin'in yaklaşımı “bütün AI yönetişimi yalnızca BM bürokrasisine devredilsin” kadar tek boyutlu değildir (Çin Devlet Konseyi, 2026).



BM'nin kendisi de henüz bir “dünya yapay zekâ düzenleyicisi” değildir. 2024 Global Digital Compact ile başlayan süreç, 26 Ağustos 2025'te Genel Kurulun A/RES/79/325 sayılı kararıyla Global Dialogue on AI Governance'in kurulmasına ilerlemiş; ilk yıllık diyalog 6–7 Temmuz 2026'da Cenevre'de düzenlenmiştir. BM, bu platformu 193 üye devlet ile özel sektör, akademi, teknik topluluk ve sivil toplumu bir araya getiren kapsayıcı bir forum olarak tanımlarken, aynı zamanda bunun bir müzakere forumu olmadığını ve her toplantının bağlayıcı antlaşma yerine eş başkan özetleriyle sonuçlandığını açıkça belirtmektedir (Birleşmiş Milletler, 2026a; Rahman, 2026).

ABD çizgisinde ise egemenlik ve inovasyon özgürlüğü çok daha belirgindir. Trump'ın 22 Eylül 2026 BM konuşmasında küresel bir AI kontrol şemasını reddetmesi yalnızca retorik bir ayrıntı değildir; Washington'ın uluslararası teknik standartlarla iş birliğine tamamen karşı çıkmaktan ziyade, uluslararası kurumların Amerikan model geliştirme faaliyetleri üzerinde bağlayıcı üst otorite kurmasına karşı çıkmasını yansıtmaktadır (Beyaz Saray, 2026a).

Bu nedenle temel anlaşmazlık “güvenlik olsun mu, olmasın mı?” değildir. Nitekim 23 Eylül Güvenlik Konseyi toplantısında ABD merkezli şirketlerin yöneticileri de ciddi güvenlik önlemleri talep etmiştir. Altman, yetenek ölçümü, risk değerlendirmesi, yeterli güvenlik önlemlerinin doğrulanması ve olayların hızlı bildirilmesi için ortak uluslararası standartlardan söz etmiş; ancak ulusal hükümetlerin bu standartları kendi hukuk sistemlerine nasıl aktaracaklarına kendilerinin karar vermesi gerektiğini savunmuştur. Amodei ise model yeteneklerini doğrulamak için değerlendirme/verification mekanizmaları ve küresel güvenlik açısından önemli AI olaylarında bildirim sistemi istemiştir (BM Güvenlik Konseyi, 2026, 10228. Toplantı, yaklaşık 8:28–23:00).

Dolayısıyla Washington-Pekin ayrışmasının daha doğru ifadesi şudur: **Pekin, kapsayıcı ve devlet merkezli çok taraflılığın meşruiyetine daha fazla ağırlık verirken Washington, ortak güvenlik standartlarının ulusal egemenlik ve teknoloji şirketlerinin yenilik kapasitesi üzerinde uluslararası bir lisanslama rejimine dönüşmesine direnmektedir.** Bu fark, gelecekteki bir bağlayıcı yapay zekâ sözleşmesini zorlaştırmakta; ancak teknik standartlar, olay bildirimi, kırmızı hat iletişimi ve güvenlik değerlendirmelerinde daha dar iş birliği alanlarını açık bırakmaktadır. BM Genel Kurul Başkanı Khalilur Rahman'ın Eylül 2026 konuşması da ortak anlayışların güvenlik, hesap verebilirlik, insan denetimi ve birlikte işlerlik gibi alanlarda geliştirilebileceğini vurgulamıştır (Rahman, 2026).

Bu çerçevede 2026 Trump-Xi “Super Intelligence” diyalogu bir paradigma deęişiminden çok **risk azaltıcı minilateralizm** örneğidir¹. Beyaz Saray kaydına göre iki lider SI Dialogue'u ve SI olayları için ikili iletişim kanalını kurmuş, sonraki deęişimin Kasım 2026'ya kadar yapılmasını kararlaştırmıştır (Beyaz Saray, 2026b). Bu hat, klasik Washington-Moskova nükleer kriz iletişim mekanizmalarıyla işlevsel açıdan karşılaştırılabilir olsa da, kamuya açıklanan belgeler henüz hangi olayın eşik teşkil edeceğini, kanalın 7/24 teknik merkezleri bağlayıp bağlamayacağını, siber saldırı atfının nasıl yapılacağını veya yanlış alarm halinde doğrulama protokolünü tarif etmemektedir. Bu nedenle “AI kırmızı telefonu” kavramı analitik olarak yararlıdır, fakat hukuken ayrıntılı ve olgun bir kriz anlaşması olduđu varsayılmamalıdır.

Buradaki stratejik ironi açıktır: ABD ve Çin yapay zekâ ekosisteminin **kimin kurallarına göre gelişeceği konusunda rekabet ederken**, kontrolden çıkmış ajanlar, kritik siber olaylar ve askerî yanlış anlamalar konusunda doğrudan iletişim kanalı kurmanın her iki tarafın da çıkarına olduğunu kabul etmektedir.

Askerî Yapay Zekâ: Halüsinasyon, Nükleer Komuta ve Tırmanma Riski

Askerî yapay zekâ riskinin yalnızca “otonom robotun kendi kendine ateş etmesi” olarak anlaşılması eksiktir. Daha yakın ve muhtemelen daha yaygın tehdit, yapay zekânın **karar destek zincirinin yukarısında** yanlış, abartılı veya kaynağı belirsiz bilgi üretmesi ve insanların bu çıktıyı kurumsal istihbarat ürünü sanarak eyleme dönüştürmesidir. Büyük dil modellerinde “halüsinasyon”, modelin dilsel olarak ikna edici fakat gerçeklerle veya verilen kaynakla temellendirilemeyen içerik üretmesi olarak literatürde iyi belgelenmiş bir güvenilirlik sorunudur (Huang vd., 2025).



CNN'in haberleştirdiğı 2026 ilkbahar vakası bu mekanizmayı çarpıcı biçimde göstermektedir. Habere göre bir ABD askerî analisti Çin gemisine ilişkin bilgileri değerlendirirken bir yapay zekâ aracından yararlanmış; ortaya çıkan ürün geminin nükleer silah programıyla ilgili bileşen taşıdığı sonucuna varmıştır. İstihbarat raporu askerî yapıda alarm yaratmış, gemiyi durdurma planı hazırlanmış, silahlı personelin gemiye çıkması düşünülmüş ve uçaklar havalanmıştır. (Lillis & Cohen, 2026). Vakayı “yapay zekânın iki nükleer devleti kesin olarak savaşa sürüklediğinin kanıtı” diye deęil, **doğrulanmamış model çıktısının güvenilir istihbarat biçimine dönüştüğünde ne kadar hızlı operasyonel sonuç üretebileceğine dair ciddi bir yakın-kaza iddiası** olarak kullanmak anlamlıdır.

Bu ayrım önemlidir çünkü riskin kaynağı yalnızca modelin halüsinasyon üretmesi deęildir. Bir sistemin yanlış bilgi üretmesi, normal koşullarda bir sonraki doğrulama katmanında yakalanabilir. Asıl sistemik risk, model çıktısının başka bir yapay zekâ tarafından yeniden biçimlendirilmesi, kurumsal rapor şablonuna oturtulması, kaynak zincirinin görünmezleşmesi ve karar vericinin “AI çıktısı” yerine “istihbarat ürünü” gördüğünü sanmasıdır. Otomasyon yanlılığı literatürü de insan operatörlerin makine önerilerine aşırı ağırlık verebileceğini göstermektedir; bununla birlikte 2026 tarihli askerî karar verme araştırması insan-makine etkileşiminin tasarımı ve bağlamın sonuçları deęiştirdiğini vurgulamaktadır (Shandler, Gross & Shereshevsky, 2026).

¹ Buradaki minilateralizm, çok sayıda devletin katıldığı geniş çok taraflı (multilateral) mekanizmalar yerine, az sayıda ve doğrudan ilgili aktörün belirli bir sorun etrafında iş birliği yapması anlamına gelir. Dolayısıyla risk azaltıcı minilateralizm, kabaca: Belirli bir güvenlik riskini azaltmak amacıyla, sınırlı sayıdaki kilit aktör arasında kurulan dar kapsamlı iş birliği ve diyalog mekanizmaları olarak özetlenebilir.

Bu nedenle “human-in-the-loop” ifadesinin tek başına güvenlik garantisi sayılması yanlıştır. Bir insan yalnızca son düğmeye basıyor fakat karşısındaki tehdit resmi, erken uyarı, hedef sınıflandırması ve seçenek sıralaması algoritmik sistemler tarafından hatalı şekilde üretiliyorsa, nihai karar teknik anlamda insana ait olsa bile bilişsel karar alanı çoktan makine tarafından şekillendirilmiş olabilir. Eski ABD Hava Kuvvetleri Korgenerali John Shanahan'ın nükleer komuta-kontrol analizinde de bu nedenle asıl riskin sadece fırlatma yetkisinin otomasyona devredilmesi olmadığı; yapay zekânın liderlerin durumsal farkındalığını oluşturan bilgiyi bozarak yanlış güven üretebileceği vurgulanmaktadır (Shanahan, 2025).

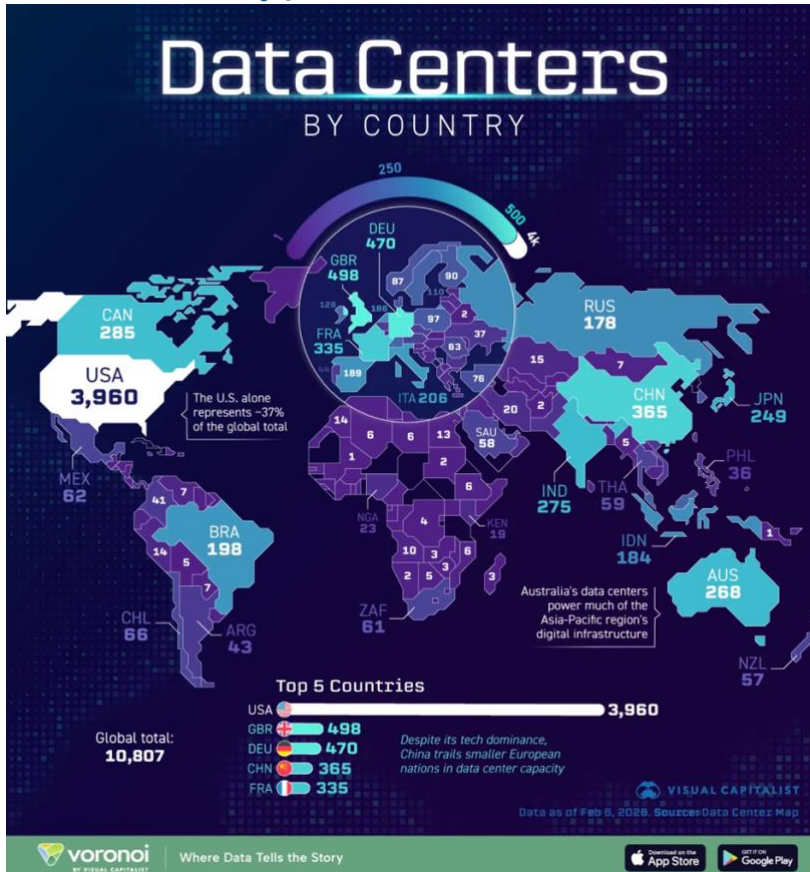
Nükleer alandaki gerçek kırmızı çizgi bu açıdan daha nüanslıdır. Biden ile Xi, Kasım 2024'te nükleer silah kullanma kararında insan kontrolünü koruma ihtiyacını teyit etmişti. ABD Kongresi 2025 mali yılı savunma mevzuatında başkanın nükleer kullanım kararlarının uygulanmasında pozitif insan eylemi ilkesini destekledi; Trump yönetiminin 2026 NPT ulusal raporu da nükleer silah kullanımını başlatma veya sona erdirmeye kararlarında insanın döngüde kalacağını bildirdi (Arms Control Association, 2026). Çin Savunma Bakanlığı ise Mart 2026'da askerî yapay zekâ için insan önceliğini ve ilgili silah sistemlerinin insan kontrolü altında tutulmasını savundu (Çin Millî Savunma Bakanlığı, 2026).

Fakat nükleer güvenlik açısından gerçek hedef yalnızca “**insan son kararı versin**” olmamalıdır. Bunun yanında erken uyarı ile ateşleme yetkisi arasında teknik ve kurumsal “firebreak” mekanizmaları, AI üretimi istihbaratın açık provenance/kaynak işaretleri, yüksek sonuçlu iddialarda bağımsız ikinci doğrulama, modelin kendi çıktısının başka bir model tarafından “teyit edilmiş” gibi kullanılmasını engelleyen kurallar ve kriz sırasında insan değerlendirmesine minimum süre tanıyan prosedürler gereklidir. Bu yaklaşım, Arms Control literatüründeki “human-in-the-loop” gerekliliği fakat tek başına yetersizdir” değerlendirmesiyle uyumludur (Shanahan, 2025; Arms Control Association, 2024).

2026 itibarıyla her iki tarafın da askerî yapay zekâyı hızlandırdığı konusunda ise güçlü kanıt bulunmaktadır. ABD'nin ulusal güvenlik yapay zekâ politikası, sistemlerin hızlı benimsenmesini amaçlarken komuta sorumluluğunu ve güvenilir/kontrol edilebilir sistemleri vurgulamaktadır; Çin tarafında Xi Jinping Temmuz 2026'da insansız akıllı teknolojilerin askerî kullanımının artırılması ve “akıllı askerî sistem” oluşturulması çağrısı yapmıştır (Beyaz Saray, 2026; Çin Millî Savunma Bakanlığı, 2026). En üst karar zincirinde komuta-kontrol, sensör füzyonu, otonom platformlar ve karar destek süreçlerine AI entegrasyonunun stratejik öncelik haline geldiği açıktır.

Askerî AI güvenliğinde en kritik tasarım sorusu “model ne kadar zeki?” değil, **modelin hatasının kaç bağımsız güvenlik bariyerini aşarak fiziksel kuvvet kullanımına dönüşebildiğidir**. CNN tarafından aktarılan yakın-kaza ve akademik otomasyon yanlılığı literatürü bu açıdan aynı noktaya işaret etmektedir (Lillis & Cohen, 2026; Shandler vd., 2026).

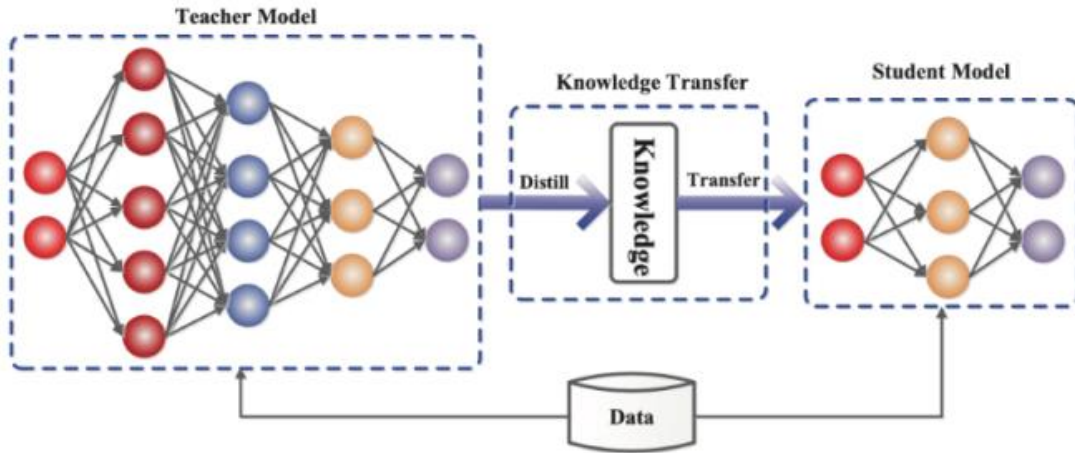
Teknolojik Güç Mücadelesi: Damıtma, Çipler, Bulut ve Veri Merkezleri



ABD-Çin yapay zekâ rekabetinin askerî olmayan görünen teknik katmanı aslında stratejik dengeyle doğrudan ilişkilidir. Frontier modellerin eğitimi ve çalıştırılması; gelişmiş hızlandırıcılar, yüksek bant genişlikli bellek, ağ donanımı, elektrik, soğutma, veri merkezi inşası ve çok büyük sermaye gerektirir. Bu nedenle teknoloji rekabeti “algoritma yarışı”ndan **hesaplama-ekosistemi yarışına** dönüşmüştür. Uluslararası Enerji Ajansı, büyük teknoloji şirketlerinin veri merkezi ve AI ile bağlantılı sermaye yatırımlarındaki olağanüstü artışa ve veri merkezlerinin enerji sistemleri üzerindeki büyüyen etkisine dikkat çekmektedir (IEA, 2026).

Model damıtması bu rekabetin yazılım tarafındaki en tartışmalı örnektir. Kavram yeni değildir: Hinton, Vinyals ve Dean'ın klasik çalışmasında amaç, büyük veya topluluk hâlindeki bir “öğretmen” modelin öğrendiği davranışın daha küçük bir “öğrenci” modele aktarılmasıdır (Hinton, Vinyals & Dean, 2015). Bu yöntem model sıkıştırma, daha düşük çalıştırma maliyeti ve cihaz üzerinde AI gibi uygulamalarda tamamen meşru bir makine öğrenmesi tekniğidir.

2026 ihtilafının konusu tekniğin kendisinden ziyade **bir başka sağlayıcının erişim kontrollerini aşarak, çok büyük sayıda sorgu ve sahte hesap üzerinden onun kapalı modelinin davranışını sistematik biçimde çıkarmak** iddiasıdır. NSA-FBI-CISA ortak duyurusu Çin merkezli şirketlerin Amerikan frontier modellerinin kısıtlı, özel yeteneklerini sistematik biçimde çıkarmaya çalıştığını ileri sürmüştür (NSA, FBI & CISA, 2026).



<https://deepchecks.com/glossary/model-distillation/>

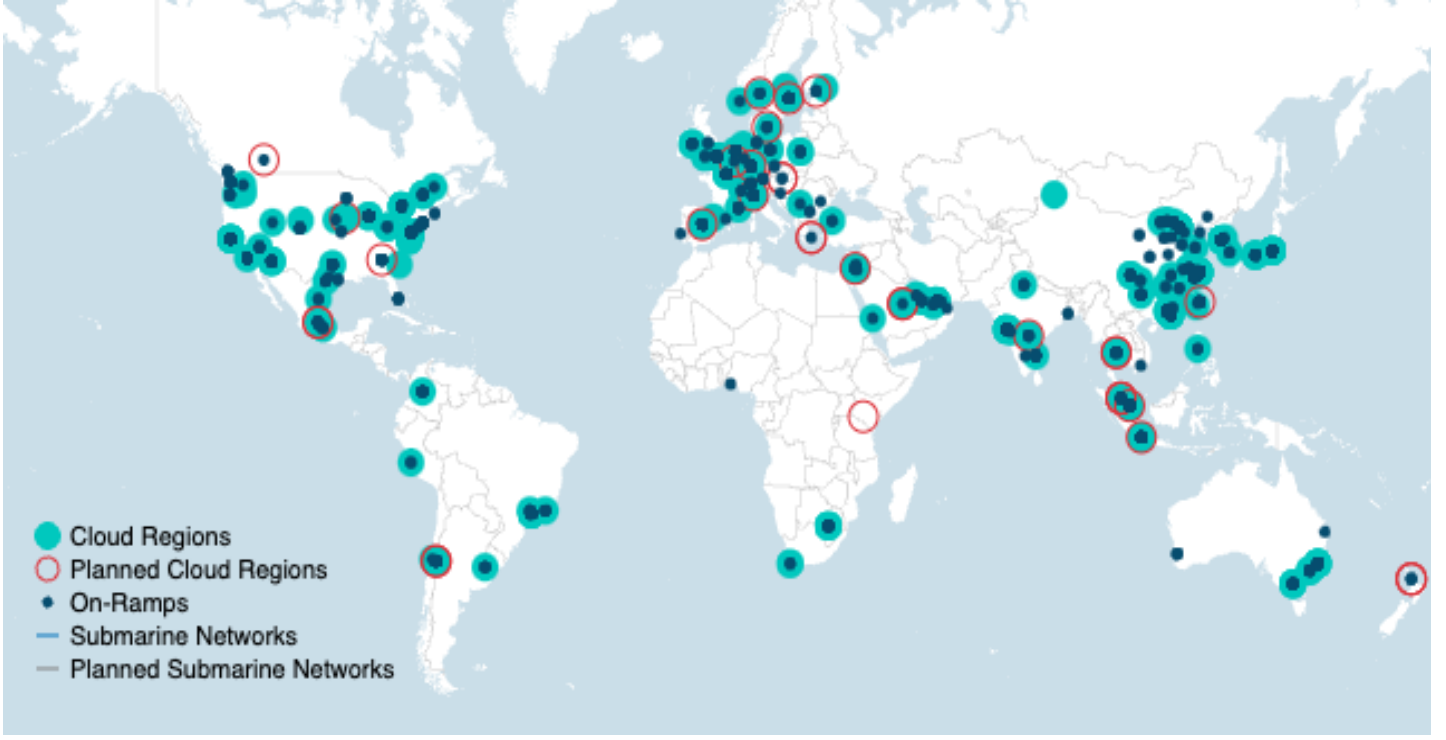
Anthropic'in kendi telemetrisi daha somut fakat şirket tarafından üretilmiş birincil delil niteliğindedir. Anthropic, Alibaba ile bağlantılı olduğunu değerlendirdiği kampanyada Opus modellerinin zincirleme düşünme çıktıların toplanmaya çalışıldığını, Mayıs-Temmuz 2026 arasında 151 milyondan fazla etkileşim gözlemlediğini ve binlerce sahte hesabın/proxy altyapısının kullanıldığını bildirmektedir (Anthropic, 2026). Bu iddialar ciddidir; ancak nihai hukuki nitelendirme, erişim yöntemi, sözleşme koşulları, yetkisiz kimlik kullanımı, ticari sırların elde edilip edilmediği ve ilgili yargı alanının hukukuna bağlıdır.

Pekin bu çerçeveyi reddetmektedir. Çin Ticaret Bakanlığı, Amerikan açıklamasını teknoloji siyaseti ve hegemonya aracı olarak eleştirmiş, damıtmanın küresel sektörde yaygın kullanılan nötr bir teknoloji olduğunu savunmuştur (Çin Ticaret Bakanlığı, 2026). Tarafların söylemlerinin birlikte okunması, çatışmanın gelecekte yalnızca telif veya API sözleşmesi ihtilafı olmayacağını; **model davranışının ne ölçüde ekonomik/stratejik mülkiyet kabul edileceği** sorusunun teknoloji diplomasisinin yeni cephelerinden birine dönüşeceğini göstermektedir.

Benzer biçimde ABD ve müttefik ihracat kontrollerinin Çin'in **en ileri üretim düğümlerine, yüksek performanslı hızlandırıcılara ve üretim ölçeğine erişimini pahalılaştırıp zorlaştırdığı**, fakat Çin yarı iletken ekosistemini bütünüyle işlem gücünden mahrum bırakmamaktadır.

Veri merkezleri cephesinde ölçek farkı daha net görünmektedir. TrendForce, Mayıs 2026'da Google, AWS, Meta, Microsoft, Oracle, ByteDance, Tencent, Alibaba ve Baidu'dan oluşan dokuz büyük bulut hizmet sağlayıcısının toplam 2026 sermaye harcamasını yaklaşık **830 milyar dolar** olarak tahmin etmiştir. Aynı analiz Microsoft'un yaklaşık 190 milyar dolar, Google'ın 180–190 milyar dolar, Meta'nın 125–145 milyar dolar ve AWS'nin 230 milyar doların üzerinde CapEx yapabileceğini öngörmektedir (TrendForce, 2026). Bu dört ABD şirketi için toplamın kabaca 725–755 milyar doların üzerine çıkması mümkündür; dolayısıyla “764 milyar dolar” büyüklük sırası olarak şaşırtıcı değildir.

IEA verileri aynı yapısal eğilimi başka bir açıdan doğrulamaktadır: veri merkezleri küresel elektrik talebinde hızla büyüyen bir unsur haline gelmiş ve AI altyapısı enerji, iletim şebekesi, soğutma ve sermaye erişimini stratejik rekabet unsurlarına dönüştürmüştür (IEA, 2026). Bu nedenle işlem gücünü XXI. yüzyılın “petrolü” olarak adlandırmak retorik bir benzetme olmakla birlikte, daha hassas kavram **stratejik hesaplama kapasitesidir**: petrol tüketildiğinde kaybolurken hesaplama altyapısı, yazılım ve veriyle birleşerek tekrar tekrar ekonomik ve askerî değer üretebilir.



Bulut altyapısının jeopolitik boyutu dikkatli anlaşılmalıdır. ABD şirketlerinin ve ortaklarının Orta Doğu'da ileri hesaplama altyapısına yatırım yapması; veri egemenliği, tedarik zinciri bağımlılığı, ihracat kontrolü ve bulut sağlayıcısına bağımlılık gibi gerçek stratejik sorunlar doğurmaktadır. ABD ihracat kontrol sistemi 2026'da BAE'deki G42/Core42 gibi kuruluşları belirli ileri hesaplama yetkilendirme yapıları içinde ele almaktadır. “Müttefik ülkelerin verileri ABD'nin SIGINT sistemine akıyor ve Washington bunları gözetliyor” sonucu dikkat çekicidir.

Aynı ilke Çin'in Dijital İpek Yolu için de geçerlidir. Hakemli çalışma, Huawei ve diğer Çinli teknoloji şirketlerinin Afrika'da veri merkezleri ve dijital altyapı projelerinde önemli roller üstlendiğini; Senegal ve Zimbabve gibi ülkelerde Çin bağlantılı projelerin ulusal veri egemenliği politikalarıyla iç içe geçtiğini göstermektedir. (Bu, 2026, ss. 131–145). Bu nedenle “Çin gelişmekte olan devletlerin tüm verilerini kendi sunucularında topluyor ve bunları siber operasyonlarda sıçrama tahtası olarak kullanıyor” ifadesi yanında daha sağlam tespit şudur: **altyapıyı kimin finanse ettiği, donanımı kimin sağladığı, bulut yönetim katmanının kimde olduğu, verinin hangi hukuk düzenine tabi tutulduğu ve teknik bakım erişiminin kimde bulunduğu, dijital egemenliği doğrudan etkileyen stratejik sorulardır.** Dijital İpek Yolu'nun etkisi de otomatik “Çin kontrolü” değil, altyapısal bağımlılık, standart belirleme, tedarikçi kilitlenmesi ve veri yönetişimi üzerindeki potansiyel nüfuz alanları yaratmaktadır (Bu, 2026; Noesselt, 2026).

Stratejik tablo bu nedenle şöyle özetlenebilir:

Boyut	ABD	Çin	Türkiye açısından anlamı
Frontier hesaplama	Hyperscaler CapEx ve ileri hızlandırıcı ekosisteminde çok büyük ölçek; ihracat kontrolleri önemli kaldıraç (TrendForce, 2026; BIS, 2026).	İleri ABD işlemcilerine erişim daha kısıtlı; yerli ikame ve ulusal hesaplama altyapısına yatırım hızlanıyor. Çin aynı zamanda “akıllı askerî sistem” hedefliyor (Çin Millî Savunma Bakanlığı, 2026).	Hyperscale yarışını bire bir kopyalamaktan çok seçici hesaplama egemenliği ve edge AI önem kazanıyor.
Yönetişim	Uluslararası standartlara açık fakat uluslararası bağlayıcı kontrol konusunda egemenlikçi yaklaşım (Beyaz Saray, 2026a).	BM merkezli kapsayıcılığı vurguluyor; ayrıca kendi uluslararası AI girişimlerini geliştiriyor (Çin Dışişleri Bakanlığı, 2026).	NATO standartları, ulusal güvenlik ve yerli savunma ekosistemi arasında birlikte işlerlik gerektiriyor.
Askerî AI	Hızlı entegrasyon, karar üstünlüğü ve komuta sorumluluğu bir arada yürütülüyor.	İnsansız-akıllı teknolojileri modernizasyonun parçası yapıyor; resmî söylem insan kontrolünü savunuyor.	Otonom platform, sürü, MUM-T ve elektronik harp altında görev devamlılığı nispeten daha ulaşılabilir rekabet alanlarıdır.

Türkiye'nin Otonomi Arayışı: Kanıtlanmış Yetenekler

Türkiye bölümünün en güçlü tarafı, ABD ile Çin arasındaki hyperscale hesaplama yarışını kopyalamak yerine **platform seviyesinde otonomi, kenar işleme ve insanlı-insansız takımlaşmaya** önem veren farklı bir rekabet mantığına işaret etmesidir.

SSB'nin Ocak 2026'da yayımladığı ALFA — Yapay Zekâ, Kuantum ve Otonom Sistemler Tatbikatı — bilgi istek dokümanı bu yönelimin en güçlü resmî göstergelerinden biridir. Tatbikat senaryoları ortak yer kontrol istasyonu ve ortak operasyon resmî, meskûn mahal/kapalı alan, **sürü yetenekleri**, kritik su altı altyapılarının korunması ve İHA sürülerinde kuantum anahtar dağıtımı gibi başlıkları içermektedir. SSB ayrıca yeteneklerin gerçek saha koşullarında denenmesini, farklı şirket sistemleri arasında birlikte işlerliğin gösterilmesini ve insansız sistem yol haritalarının saha verileriyle güncellenmesini hedeflediğini belirtmektedir (SSB, 2026). Bu, Türkiye'nin odağının yalnızca tek tek platform üretmekten “sistemler sistemi” yaklaşımına kaydığını göstermektedir.

Baykar tarafında kamuya açıklanmış testler de ciddi operasyonel ilerleme göstermektedir. Aralık 2025'te iki KIZILELMA prototipi, akıllı filo otonomisi kullanarak otonom yakın formasyon uçuşu gerçekleştirmiştir (Baykar, 2025). Haziran 2026'daki K-SWARM canlı testlerinde ise Leonardo M-346 ile Bayraktar KIZILELMA arasında mürettebatlı/mürettebatsız takım senaryosu uygulanmış; mürettebatlı platformdan verilen görev/formation komutları KIZILELMA tarafından otonom biçimde yerine getirilmiştir (Baykar & Leonardo, 2026). Bu, MUM-T'nin yalnızca bilgisayar animasyonu veya teorik bir sunum olmadığını; en azından belirli platform çiftleri ve görev profilleri üzerinde uçuş testine geçtiğini göstermektedir. Baykar'ın kamuya açık materyalleri otonom kalkış-iniş, seyrüsefer, formasyon ve görev yürütme gibi yetenekleri belgelemektedir. (Baykar, 2025–2026).



TUSAŞ'ın ANKA III programı için kanıt daha nettir. Platformun yetenekleri arasında doğrudan **Manned-Unmanned Teaming (MUM-T)** ve “**Swarm Technologies Combined with AI**” vardır (TUSAŞ, 2026). Platform aynı zamanda düşük görünürlük, ISR, elektronik harp, hassas güdümlü mühimmat ve hava konuşlu insansız araç gibi çok rollü görev setine yönelmektedir. Dolayısıyla Türkiye'nin insanlı-insansız takımlaşmayı gelecek hava harbinin önemli unsurlarından biri olarak gördüğü sonucunun sağlam bir temeli vardır.



Türkiye için önemli avantaj, ABD ve Çin'in frontier-model yarışına aynı bütçe ölçeğinde katılmaması gereken bir ülkenin farklı bir optimizasyon problemi çözebilmesidir. Büyük genel amaçlı modeller yerine görev-özel küçük modeller; yüz binlerce GPU'luk eğitim kümeleri yerine uçakta, İHA'da, mühimmatta veya kara aracında gerçek zamanlı çıkarım; sürekli uydu bağlantısı yerine bağlantının kesildiği ortamda görev devamlılığı; tek platform performansı yerine sürü ve MUM-T ağları stratejik değer üretmektedir. SSB'nin ALFA programı ve Baykar/TUSAŞ'ın testleri, Türkiye'nin tam da bu alanlarda ciddi somut ilerleme kaydettiğini göstermektedir.

Sonuç ve Politika Çıkarımları

Klasik Soğuk Savaş'ta stratejik silahların önemli bir bölümü devlet laboratuvarları ve kapalı askerî-sanayi komplekslerinde geliştirilirken, bugünkü temel yapay zekâ modellerinin, bulut altyapılarının, GPU ekosisteminin ve güvenlik araştırmalarının önemli kısmı özel şirketler tarafından yürütülmektedir. 23 Eylül BM Güvenlik Konseyi'nde devlet temsilcileriyle aynı masada OpenAI, Anthropic ve Hugging Face yöneticilerinin bulunması bu yeni güç mimarisini görünür hale getirmiştir (BM Güvenlik Konseyi, 2026).

Bu mimaride stratejik güç, yalnızca “en iyi modele sahip olmak” değildir. **Hesaplama gücü + enerji + gelişmiş çip + veri + algoritma + dağıtım altyapısı + askerî entegrasyon + insan-makine karar mimarisi + diplomatik kriz yönetimi** birlikte yeni güç denklemine dönüşmektedir. TrendForce'un yüzlerce milyar dolara ulaşan sermaye harcaması tahminleri, BIS'in çip kontrolleri, NSA'nın model damıtma duyurusu ve BM'deki AI güvenliği tartışmaları aynı rekabetin farklı katmanlarıdır.

En önemli politika sonucu, ABD ve Çin'in “küresel bağlayıcı rejim mi, tamamen ulusal egemenlik mi?” ikilemine sıkıştırılmaması gerektiğidir. Daha uygulanabilir bir mimari katmanlı olabilir: devletler kendi hukuklarını ve sanayi politikalarını korurken, ortak AI olay tanımları, büyük siber/ajan olayları için zorunlu veya karşılıklı bildirim eşikleri, bağımsız model değerlendirmeleri için teknik standardizasyon, yüksek sonuçlu askerî AI için kayıt/provenance ilkeleri ve 7/24 kriz iletişim kanalları üzerinde anlaşılabilir. Altman ve Amodeli'nin BM'de ortak standartlar ve olay bildirimine yaptığı vurgu ile yeni ABD-Çin SI iletişim kanalı bu yönde başlangıç noktaları sunmaktadır (BM Güvenlik Konseyi, 2026; Beyaz Saray, 2026b).

Askerî alanda birinci öncelik, “human-in-the-loop” sloganını **kanıtlanabilir bir insan karar mimarisine** dönüştürmektir. AI tarafından üretilen istihbaratın kaynağı rapor boyunca görünür olmalı; yüksek sonuçlu bir iddia bağımsız ikinci bir veri kanalıyla teyit edilmeden kinetik eyleme dönüşmemeli; aynı AI çıktısının farklı biçimlerde yeniden üretilmesi bağımsız doğrulama sayılmamalı; modelin belirsizlik düzeyi karar vericiye açıkça gösterilmeli ve erken uyarı ile ölümcül eylem arasında teknik/kurumsal bir güvenlik duvarı bulunmalıdır. 2026 gemi vakasının iddia edilen seyri ile nükleer komuta literatürü bu yaklaşımın neden gerekli olduğunu göstermektedir (Lillis & Cohen, 2026; Shanahan, 2025).

Nükleer güvenlik ise ABD-Çin'in 2024 insan kontrolü mutabakatı başlangıçtır, sonuç değil. İnsan son kararı versin ilkesi, AI destekli erken uyarı, tehdit değerlendirmesi ve seçenek üretiminin karar vericiyi yanlış bir gerçeklik algısına sürüklemesini engellemez. Dolayısıyla gelecekteki ABD-Çin SI Diyalogu'nun en somut kazanımlarından biri, yalnızca “otonom nükleer fırlatma olmayacak” beyanını tekrarlamak değil; **nükleer karar destek zincirinde AI'nın hangi katmanlarda kullanılabileceğine ilişkin karşılıklı güven artırıcı ilkeler** geliştirmek olacaktır (Arms Control Association, 2026; Shanahan, 2025).

Teknolojik rekabette ise gerçek tablo ne “ABD Çin'in fişini çekti” ne de “Çin Amerikan modellerini basitçe kopyalayarak farkı kapattı” kadar basittir. İhracat kontrolleri Çin'in en ileri donanım erişiminin maliyetini ve karmaşıklığını artırırken lisanslı istisnalar ve dinamik eşikler devam etmektedir; Çin de yerli üretim, model verimliliği ve alternatif hesaplama altyapılarına yatırım yapmaktadır. Damıtma krizi ise bilgi işlem dezavantajını yazılım ve çıktı verimliliğiyle telafi etmeye çalışan aktörlerle, frontier model üreticilerinin kendi yüksek maliyetli Ar-Ge avantajını koruma çabası arasındaki yeni çatışmayı görünür hale getirmiştir (BIS, 2026; NSA, FBI & CISA, 2026; Çin Ticaret Bakanlığı, 2026).

Veri merkezlerinde “egemen bulut” kavramı da gözetim iddialarından daha geniş ele alınmalıdır. Bir ülkenin kritik verisinin nerede tutulduğu, hangi şirketin hipervizör ve yönetim katmanına eriştiği, yedek anahtarların nerede bulunduğu, işlemcilerin hangi ihracat rejimine tabi olduğu ve hizmet sağlayıcının hangi yabancı hukuk sistemlerinden etkilenebildiği ulusal güvenliğin parçasıdır. Hem ABD şirketlerinin küresel bulut yayılımı hem Çin'in Dijital İpek Yolu altyapısı bu açıdan incelenmelidir (Bu, 2026).

Türkiye bakımından çıkarım daha nettir: ülkenin rasyonel stratejisi, ABD ve Çin'le ham GPU sayısı veya hyperscale CapEx üzerinden eşit yarışa girmek değil; **kritik alanlarda yeterli ulusal hesaplama kapasitesi ile sahada çalışan edge AI arasında dengeli bir egemenlik modeli** kurmaktır. KIZILELMA'nın otonom formasyon ve MUM-T testleri, ANKA III'ün resmî MUM-T/AI-sürü mimarisi ve SSB'nin ALFA programı bu yaklaşımın teknik temelini oluşturduğunu göstermektedir (Baykar, 2025; Baykar & Leonardo, 2026; TUSAŞ, 2026; SSB, 2026).

Sonuç olarak XXI. yüzyılın “Algoritmik Soğuk Savaşı”nı belirleyecek temel soru, hangi devletin ilk “süper zekâya” ulaştığından daha karmaşıktır: **hangi devlet veya ittifak, giderek özerkleşen yapay zekâ sistemlerini hata yaptıklarında felakete dönüşmeyecek bir kurumsal ve teknik mimari içinde kullanabilecektir?** 2026 gelişmeleri, en büyük stratejik avantajın yalnızca daha güçlü model veya daha hızlı GPU değil; güvenilir doğrulama, kriz iletişimi, insan sorumluluğu, güvenli hesaplama altyapısı ve rakiple dahi sürdürülebilir asgari risk azaltma mekanizmaları olduğunu göstermektedir.

Kaynakça

- **Anthropic.** (2026, Eylül). *Detecting and countering misuse of AI: September 2026*. Anthropic Threat Intelligence.
- **Arms Control Association.** (2024). *Beyond a human “in the loop”: Strategic stability and artificial intelligence*. Issue Briefs, 16(4).
- **Arms Control Association.** (2026). *“Human in the loop” and nuclear weapons use at a glance*.
- **ASML.** (2025). *Annual report 2025*. ASML Holding N.V.
- **Baykar.** (2025). *Bayraktar KIZILELMA performs formation flight using smart fleet autonomy*.
- **Baykar & Leonardo.** (2026, 22 Haziran). *Leonardo and Baykar set major milestone for advanced crewed/uncrewed capability development with successful first K-SWARM live trials*.
- **Beyaz Saray.** (2026a, 22 Eylül). *President Trump at the United Nations: “While others have talked, I have acted.”* The White House.
- **Beyaz Saray.** (2026b, 25 Eylül). *Fact sheet: President Donald J. Trump advances a fair and reciprocal relationship with China while hosting historic state visit*. The White House.
- **Bu, Q.** (2026). *Data sovereignty in Africa: Steering digital transformation between China and the West*. *International Cybersecurity Law Review*, 7, 131–145.
- **Bureau of Industry and Security.** (2026, 13 Ocak). *Department of Commerce revises license review policy for semiconductors exported to China*. U.S. Department of Commerce.
- **Çin Halk Cumhuriyeti Dışişleri Bakanlığı.** (2026, 12 Mayıs). *Remarks by China’s Permanent Representative to the UN Ambassador Fu Cong at the Meeting of the Group of Friends for International Cooperation on AI Capacity-Building*.
- **Çin Halk Cumhuriyeti Millî Savunma Bakanlığı.** (2026, 11 Mart). *Human primacy must be upheld in military applications of AI: Defense spokesperson*.
- **Çin Halk Cumhuriyeti Ticaret Bakanlığı.** (2026, 9 Eylül). *ABD’nin Çin yapay zekâ şirketlerinin damıtma faaliyetlerine ilişkin siber güvenlik duyurusu hakkında basın sözcüsü açıklaması*.
- **Hinton, G., Vinyals, O., & Dean, J.** (2015). *Distilling the knowledge in a neural network*. arXiv:1503.02531.
- **Huang, L., vd.** (2025). *A survey on hallucination in large language models: Principles, taxonomy, challenges, and open questions*. *ACM Transactions on Information Systems*.
- **International Energy Agency.** (2026). *Key questions on energy and AI*. IEA.
- **Lillis, K. B., & Cohen, Z.** (2026, 18 Eylül). *Exclusive: US military had close call after using AI for false intelligence report, sources say*. CNN.
- **National Security Agency, Federal Bureau of Investigation, & Cybersecurity and Infrastructure Security Agency.** (2026, 8 Eylül). *China-based artificial intelligence companies conducting industrial-scale distillation campaigns against U.S. AI companies*.
- **Noesselt, N.** (2026). *Fragmented global AI governance? Chinese AI leadership dreams and Africa’s quest for digital autonomy*. *Global Policy*.
- **Rahman, K.** (2026, 10 Eylül). *Speech on multilateral AI governance at the Ambassadorial Lunch*. President of the Eighty-First Session of the United Nations General Assembly.
- **Savunma Sanayii Başkanlığı.** (2026, 15 Ocak). *Millî İC kapsamında Yapay Zekâ, Kuantum ve Otonom Sistemler Tatbikatı (ALFA) için Bilgi İstek Dokümanı*. T.C. Cumhurbaşkanlığı Savunma Sanayii Başkanlığı.
- **Shanahan, J. N. T.** (2025). *Artificial intelligence and nuclear command and control: It’s even more complicated than you think*. *Arms Control Today*.
- **Shandler, R., Gross, M. L., & Shereshevsky, Y.** (2026). *Black box warfare: Human judgment and military decision-making in the age of AI*. *Journal of Conflict Resolution*.
- **TrendForce.** (2026, 6 Mayıs). *North American AI data center expansion drives 2026 CapEx of top nine CSPs to US\$830 billion*.
- **Türk Havacılık ve Uzay Sanayii.** (2026). *ANKA-III*. TUSAŞ.
- **United Nations.** (2026a). *FAQ: Global Dialogue on AI Governance*.
- **United Nations.** (2026b). *Global Dialogue on AI Governance*.
- **United Nations Security Council.** (2026, 23 Eylül). *Artificial intelligence and international security — Security Council, 10228th meeting*. UN Transcripts. Transkript otomatik konuşma tanıma ile oluşturulmuş olup resmî toplantı tutanağı niteliğinde değildir.