

This document was translated from English by machine and may contain mistakes. The English original is at <https://www.novaracg.com/2026/09/14/the-ai-ceos-found-the-brakes-sign-language-ai-should-pay-attention/>.

Yapay Zeka CEO'ları Frenleri Buldu... İşaret Dili Yapay Zekası Dikkat Etmeli

Heather Grizzle · September 14, 2026

İçindekiler

1. Sınır teknolojisi fren istiyor
2. Yönetişim borcu gerçektir
3. Sınır teknolojisinin öğrendiğini işaret dili yapay zekası şimdi öğrenebilir
4. "We tested it" is not governance
5. Kendi yavaşlama anınızı beklemeyin
6. Bağımsız güvence yenilikçiliğe bir saldırı değildir
7. Bunun önüne geçmek için hâlâ bir fırsat var
8. Uyarı

Sınır teknolojisi fren istiyor

Yıllarca yapay zeka alanındaki baskın anlatı hızdı. Daha yetenekli modeli inşa et. Ajanı piyasaya sür. Rakibi geç. Sermaye topla. Kıyaslamayı genişlet. Ürünü pazara sok. Yönetişim daha sonra yetişebilirdi.

Görünen o ki o "sonra" gelmiş durumda.

On September 9, OpenAI chief global affairs officer Chris Lehane published a call for mandatory, capability-based national AI safety regulation covering testing standards, independent assessments, cybersecurity protections and incident reporting, alongside the company's endorsement of four California bills, including one establishing infrastructure for independent safety assessments and another setting standards for AI auditors.^[^1] Three days later, Anthropic CEO Dario Amodei went considerably further. In an essay arguing that the industry must deliberately pace frontier



development, he proposed embedding third-party evaluators inside frontier labs with near-employee access so they could independently examine safety measures, incidents and model behavior.

Sonra alışılmadık bir şey oldu.

Within hours, OpenAI CEO Sam Altman publicly agreed that the frontier needed pacing, and told Fortune days later that his company would not pursue a public offering this year because it had too much safety work left to do. Elon Musk, who had pointedly declined to join earlier industry calls for restraint, quoted the essay approvingly. Demis Hassabis, now chair of Google DeepMind and chief scientist of Alphabet, wrote that the essay pointed toward the right path forward. Companies locked in one of the most consequential technology races in history converged, inside a single day, on the proposition that capability could no longer be allowed to outrun control. Axios ran the moment under a headline that needed no elaboration: "AI's most powerful CEOs hit the brakes."

Today, September 14, Microsoft added another piece. Microsoft AI CEO Mustafa Suleyman released a draft code of conduct that he describes as a constitution of sorts for the company's future models. The draft would require Microsoft's systems to communicate clearly with humans, respect human boundaries, remain controllable, accept correction and shutdown, and treat behavior that violates the code as a model failure rather than an acceptable expression of autonomous judgment. Microsoft assembled it after expert consultations and has opened it to roughly six weeks of public feedback.[^3] Suleyman called it a warning shot.

Bu dizilim, işaret dili yapay zekası geliştiren her şirket için zorunlu okuma olmalı. İşaret yapan bir avatarın interneti ele geçirmek üzere olmasından değil. Çünkü yönetim borcu tam olarak böyle görünür.

Yönetişim borcu gerçektir

Teknoloji şirketleri teknik borcu anlar. Yeterince hızlı inşa edin, yeterince mimari kararı erteleyin, sonunda kestirmeler pahalıya patlar.

Yönetişim borcu da aynı şekilde işler ve her biri tek başına makul görünen kararlar aracılığıyla birikir. Bir şirket, sistem başarısız olduğunda kimin sorumlu olacağını tanımlamadan lansman yapar. Neyin hata sayılacağını belirlemeden bir doğruluk iddiası yayınlar. Bir kıyaslamanın ürünü gerçekten kullanacak insanları temsil edip



etmediğini sormadan bir kıyaslama oluşturur. Bağımsız değerlendirme elverişsiz olduğu için içeride doğrulama yapar. İnsanın ne fark etmesi gerektiğini, bu kişinin ne zaman müdahale etmesi gerektiğini veya başarısızlığı tanıyacak uzmanlığa sahip olup olmadığını belirtmeden insan gözetimi ekler. Önemli ürün kararları çoktan alındıktan sonra bir danışma grubu atar. Sistem işaretler ürettiği için sistemi erişilebilir olarak adlandırır.

Bu seçimlerin her biri yönetim borcu yaratır. Ve teknik borcun aksine, maliyeti yalnızca geliştirici taşımaz. Kullanıcılar taşır. Alıcılar taşır. Kamu kurumları taşır. İşaret dili teknolojisinde bunu Sağır insanlar taşır.

Sınır teknolojisinin öğrendiğini işaret dili yapay zekası şimdi öğrenebilir

Mevcut yapay zeka tartışmasının içinde olağanüstü bir gerçek yatıyor. Dünyanın en büyük mühendislik ekiplerine, araştırma bütçelerine ve güvenlik organizasyonlarına sahip şirketler daha fazla dış denetim istiyor. Anthropic, gömülü üçüncü taraf değerlendirme öneriyor. OpenAI, zorunlu federal düzenlemeyi ve bağımsız değerlendirmeyi destekliyor. Microsoft, insan kontrolünün sınırlarını resmileştirmeye çalışıyor. Rakip yöneticiler, yakın zamana kadar düşünülemez olacak şekilde, koordineli geri çekilmeyi kamuoyu önünde tartışıyor.

Ders, şirketlerin bir hükümetin mükemmel bir standart yazmasını beklemesi gerektiği değil. Aksine daha yakın. Teknolojiye en yakın kuruluşlar, içsel güvenin yeterli kanıt olmadığını kabul ediyor.

Bu, işaret dili yapay zekası için önemli çünkü alan hâlâ farklı bir sıralama seçebilecek kadar genç. Yönetişim onarım işi olmak zorunda değil. Mimarının bir parçası olabilir.

"We tested it" is not governance

İşaret dili yapay zekası, doğrulama probleminin özellikle tehlikeli bir versiyonuna sahip. Bir sistem tanınabilir şekilde işarete benzer bir şey üretebilir ve yine de ona bağımlı olan kişiyi hayal kırıklığına uğratabilir. Bir çeviri akıcı ve yanlış olabilir. Bir avatar görsel olarak cıvıllı olabilirken dilsel bozulma da getirebilir. Bir kıyaslama, dar bir veri kümesini, kısıtlı bir işaret ortamını veya dağıtımla çok az benzerlik taşıyan bir görevi temsil ederken etkileyici bir sayı bildirebilir. Bir çalışma, sonunda ona eklenen ticari iddiayı desteklemek için çok küçük bir örnekleme dayanırken insan doğrulaması



içerebilir. Ve bir ürün gösterimi, sistemin sağlık, eğitim, istihdam, acil durum iletişimi veya devlet hizmetleri için uygun olduğunu kanıtlamadan güzelce çalışabilir.

Bunlar basitçe makine öğrenmesi soruları değildir. Bunlar güvence sorularıdır. Kim test etti ve testi kim seçti? Hangi işaretçiler temsil edildi, hangileri edilmedi? Ne ölçüldü ve ne dışlandı? Neler maddi hata sayılır? Hangi dağıtım bağlamları değerlendirildi? Dağıtımı durdurma yetkisi kimde? Performans düştüğünde ne olur ve bu gerçekleştiğinde olayları kim inceler? Ve Sağır iletişimi etrafında inşa edilen teknoloji için, Sağır otoritesi karara nerede fiilen dahil oluyor?

Bu sorular, sözleşmeler imzalandıktan, pazarlama iddiaları yayınlandıktan ve müşterilere teknolojinin hazır olduğu söylendikten sonra yanıtlanması önemli ölçüde daha zor hale gelir.

Kendi yavaşlama anınızı beklemeyin

There is an understandable temptation for smaller AI companies to look at the frontier governance debate and file it under somebody else's problem. Different technology, different scale, different risks. All correct. But risk magnitude is not the relevant comparison. Governance timing is.

Sınır sektörü olağanüstü rekabet baskısı altında ileriye koştu ve liderleri şimdi bağımsız değerlendirme, zorunlu test, olay bildirimi, insan kontrolü, devlet gözetimi ve sektörün kendi sistemlerinin ne zaman yeterince güvenli olduğuna karar verme konusunda güvenilip güvenilemeyeceği sorularıyla karşı karşıya. İşaret dili yapay zekası geliştiricileri bunun gerçek zamanlı olarak gelişmesini izleyebilir ve ardından tekrarlayıp tekrarlamayacaklarına karar verebilir.

Bugün işaret dili teknolojisi geliştiren bir şirket zaten beş soruyu yanıtlayabilmelidir.

1. Hangi kanıt bizim kendi ürünümüzü dağıtmamamıza neden olur?
2. Şirket dışında kimlerin iddialarımıza itiraz etme yetkisi var?
3. Sağır insanlar geri bildirim veya eğitim verisi sağlamanın ötesinde ne yetkiye sahip?
4. Bir müşteri ödemeye istekli olsa bile, teknolojimizin hangi kullanımlarını uygunsuz görüyoruz?



5. Dağıtım sonrası kanıtlar, sistemi satmak için kullandığımız kanıtlarla çeliştiğinde ne olur?

Bu yanıtlar mevcut değilse, şirketin sadece tamamlanmamış yönetim belgeleri yoktur. Yönetişim borcu vardır.

Bağımsız güvence yenilikçiliğe bir saldırı değildir

İşaret dili yapay zekası pazarının hızla olgunlaşması gereken nokta burasıdır. Bağımsız denetim, teknolojiye karşı düşmanlık değildir. Teknolojinin denetlenecek kadar önemli olduğuna dair bir işarettir.

Uçak üreticileri, alıcıları uçağa hayran olmaya davet ederek güven oluşturmaz. Tıbbi cihaz üreticileri, referanslar yayımlayarak güvenlik oluşturmaz. Finansal denetimler, yönetim heyecanının kanıtın yerine geçmesine izin vermez. Yapay zeka erişilebilirliğe girdikçe aynı ayırım geçerlidir. Bir geliştirici sistemi yaratır. Yönetişim etrafındaki kuralları oluşturur. Değerlendirme kanıtı test eder. Güvence, bu kanıtın yapılan iddiayı desteklemeye yeterli olup olmadığını sorar. Bunlar dört farklı işlevdir ve tek bir şirket dördünü de kontrol ettiğinde, alıcılar sorular sormalıdır.

The frontier industry is beginning to acknowledge exactly that problem. Anthropic's proposal for embedded outside evaluators is significant precisely because it moves independent scrutiny closer to the development process rather than treating evaluation as a ceremonial checkpoint immediately before release.

İşaret dili yapay zekası daha da ileri gitmelidir. Bağımsız teknik değerlendirme gereklidir, ancak dilsel meşruiyet, dağıtım bağlamı, erişilebilirlik etkisi ve Sağlık yönetişimi mühendislik testlerine indirgenemez. Bir model tam olarak tasarlandığı gibi çalışabilir ve yine de satıldığı işe uygun olmayabilir.

Bunun önüne geçmek için hâlâ bir fırsat var

İşaret dili yapay zekası şirketlerinin kaçırmaması gereken kısım burasıdır. Onlara olağandışı derecede değerli bir uyarı veriliyor ve yönetim krizini önce kendileri yaşamak zorunda değiller.

İş belirlidir. Pazarlama etrafında dolaşmadan önce sınırlamaları belgeleyin. Gösterim performansını dağıtım kanıtından ayırın. Bir müşteri talep etmeden önce bağımsız değerlendiricileri kullanın. İlk ciddi olaydan önce olay prosedürleri oluşturun.



Ürünün satılmayacağı bağlamları belirleyin. Sağır uzmanlara danışma varlığı yerine gerçek karar yetkisi verin. Doğrulama iddialarını, başka bir tarafın inceleyip itiraz edebileceği kadar spesifik yapın. Bir alıcının bir dağıtım kararının yapıldığı zaman neden haklı olduğunu daha sonra belirleyebilmesi için kanıtları koruyun.

Güvence verilebilir olmak böyle görünür. Ayrıca daha sonra güveni yeniden inşa etmekten önemli ölçüde daha ucuzdur.

Uyarı

Dünyanın en güçlü yapay zeka şirketleri, yönetişimin kalıcı olarak yeteneğin gerisinde kalamayacağını keşfediyor. Liderlerinden bazıları şimdi frenler istiyor. Diğerleri yasalar istiyor. Diğerleri anayasalar, dış değerlendirme yapıları ve yeni kontrol mekanizmaları taslakları hazırlıyor. Çoğu erişilebilirlik teknolojisi şirketinin asla sahip olamayacağı kaynaklara komuta ediyorlar ve teknoloji zaten tam hızda ilerlerken yönetişimin ne kadar zorlaştığını onlar bile keşfediyor.

Sign language AI companies should study that carefully. Do not wait until a failed deployment, a procurement challenge, a regulator, a lawsuit or the Deaf community forces the governance conversation. Do not build the evidence after making the claim. Do not invite independent scrutiny only after trust has been lost. And do not confuse being first to market with being ready for the responsibility that comes with entering someone else's language, communication and access.

Sınır teknolojisi yapay zekası şimdi frenlerini bulmaya çalışıyor. İşaret dili yapay zekasının hızlanmadan önce onları takma fırsatı hâlâ var.

NOVARA CONSULTING GROUP

Novara Consulting Group, yapay zeka güvencesi, erişilebilirlik yönetişimi ve bağımsız değerlendirme sağlar.

FOOTNOTES:

Chris Lehane, "The AI policy window is open. We need to act.," OpenAI, September 9, 2026; Reuters, "OpenAI pushes for mandatory national AI safety requirements," September 9, 2026.

Axios, "AI's most powerful CEOs hit the brakes," September 13, 2026.



Jeffrey Dastin, "Microsoft drafts code of conduct to keep its AI under human control," Reuters, September 14, 2026.

How to cite

Grizzle, H. M. (2026, September 14). Yapay Zeka CEO'ları Frenleri Buldu... İşaret Dili Yapay Zekası Dikkat Etmeli. Novara Consulting Group. <https://doi.org/10.5281/zenodo.23003722>

