



Frontier Yapay Zekâda Şeffaflık Temelli Hukuki Yönetişim: Kaliforniya Modelinin Normatif ve Kurumsal Analizi

Transparency-Based Legal Governance in Frontier Artificial Intelligence: A Normative and Institutional Analysis of the California Model

Samet Tatar

Ankara Yıldırım Beyazıt Üniversitesi,
Hukuk Fakültesi, Bilişim Hukuku
Anabilim Dalı, Ankara, Türkiye
samettatar@gmail.com



Öz: Bu çalışma, Kaliforniya'nın 2025 yılında kabul ettiği Frontier Yapay Zekâda Şeffaflık Yasası (YASA) aracılığıyla frontier yapay zekâ modellerine yönelik geliştirdiği şeffaflık merkezli yönetim yaklaşımını kapsamlı bir biçimde analiz etmektedir. YASA, yalnızca teknik bir ürün güvenliği düzenlemesi olmaktan öte; felaket riski yönetimi, kurumsal hesap verebilirlik, whistleblower koruması, kritik olay bildirim rejimi ve kamuya ait yapay zekâ altyapısının inşası gibi çok boyutlu araçları bir araya getiren yenilikçi bir normatif çerçeve sunar. Çalışma, düzenlemenin federal-eyalet ilişkileri bağlamındaki konumunu, AB Yapay Zekâ Tüzüğü ile karşılaştırmalı yönlerini ve teknoloji şirketlerinin bulunduğu Silikon Vadisi ekosistemindeki politika-ekonomik dinamikleri değerlendirir. Ayrıca frontier modelin teknik ve hukuki tanımı, felaket riski kavramsallaştırması, çok katmanlı yükümlülük sistemi ve büyük geliştiricilere yönelik şeffaflık dokümantasyonu ayrıntılı biçimde ele alınmaktadır. Makale, YASA'nın esnek, öğrenmeye dayalı ve şeffaflık temelli yapısının farklı hukuk sistemlerinde uyarlanabilirliğini inceleyerek, düzenleme geliştirmek isteyen ülkeler için aktarılabilir politika önerileri sunar. Sonuç bölümünde ise YASA'nın küresel yapay zekâ yönetiminde neden öncü bir model olarak değerlendirilebileceği tartışılmaktadır.

Anahtar Kelimeler: Frontier Yapay Zekâ, Felaket Riski Yönetim, Yapay Zekâ Hukuku, Yapay Zekâ Yönetimi, Şeffaflık Rejimi

Geliş Tarihi/Received: 25.11.2025
Kabul Tarihi/Accepted: 08.02.2026
Yayınlanma Tarihi/ Published Online: 21.07.2026

Abstract: This study provides a comprehensive analysis of California's Transparency in Frontier Artificial Intelligence Act (TFAA), enacted in 2025, and its transparency-centered governance framework for frontier artificial intelligence models. Rather than operating merely as a technical product-safety regulation, the TFAA introduces a novel normative architecture that integrates multiple regulatory instruments, including catastrophic risk management, corporate accountability, whistleblower protections, critical incident reporting requirements, and the development of publicly administered AI infrastructure. The study evaluates the statute's position within the federal- state regulatory dynamic, its points of convergence and divergence with the EU Artificial Intelligence Act, and the policy-economic environment shaped by Silicon Valley's technology ecosystem. It further examines the technical and legal definition of frontier models, the conceptualization of catastrophic risk, the multi-tiered obligations imposed on developers, and the transparency documentation required from large-scale actors. By assessing the adaptability of the TFAA's flexible, learning-oriented, and transparency-based structure across different legal systems, the article offers transferable policy recommendations for jurisdictions seeking to develop AI regulation. The conclusion discusses why the TFAA may be regarded as a pioneering model in global AI governance.

Keywords: Frontier Artificial Intelligence, Catastrophic Risk Management, AI Law, AI Governance, Transparency Regime

Extended Abstract

The Transparency in Frontier Artificial Intelligence Act (TAFA), enacted by the State of California in 2025, represents one of the most advanced and comprehensive sub-national regulatory responses to the governance challenges posed by frontier artificial intelligence models. Far from being a conventional product-safety statute, the TAFA constructs an innovative, transparency-driven governance framework combining catastrophic risk oversight, multilayered obligations for developers, whistleblower protection, critical incident reporting mechanisms, and the establishment of a public cloud computing infrastructure aimed at strengthening state capacity in the AI domain. This extended summary outlines the statute's normative foundations, institutional architecture, regulatory tools, and broader implications for comparative AI governance.

At the conceptual core of the TAFA is the recognition that frontier AI models defined in reference to both computational scale and potential systemic impact- pose qualitatively distinct risks compared to conventional machine-learning systems. The statute adopts a dynamic, capability-oriented definition of "frontier model," using a training compute threshold of 10^{26} FLOPs while also accounting for fine-tuning, adaptive training, and other forms of model enhancement. By shifting from a static technical threshold to a lifecycle-based assessment, the TAFA positions the frontier category as a moving boundary capable of evolving alongside advances in AI research. This design enables the law to maintain regulatory relevance despite rapid technological innovation.

A second foundational pillar of the Act is its formal articulation of "catastrophic risks," defined as scenarios in which a frontier model materially contributes to mass casualties, critical infrastructure disruption, or economic damage exceeding one billion dollars. The law provides concrete illustrations- including autonomous cyberattacks, biological or chemical threat facilitation, and loss of model controllability- while simultaneously excluding harm unrelated to the model's capabilities or arising purely from publicly available information. This risk construction establishes a high-impact regulatory tier that justifies more intensive oversight of frontier systems without imposing undue burdens on lower-risk models.

The TAFA's transparency regime is anchored in the mandatory Frontier AI Documentation (Frontier AI Framework) that large frontier developers must prepare, publish, and update. This governance document must specify risk-assessment methodologies, internal and external testing protocols, international and industry standards adopted, mitigation strategies, and cybersecurity safeguards protecting unpublished model weights. The statute further requires developers to disclose model release information- including modalities, intended uses, and general capability profiles- alongside summaries of catastrophic-risk evaluations. These requirements institutionalize a form of regulatory visibility that compensates for the knowledge asymmetry between state authorities and highly specialized AI laboratories.

The TAFA also establishes a two-tiered incident-reporting regime. Both developers and members of the public may submit reports of critical safety incidents to the Office of Emergency Services (OES), while large developers face a mandatory duty to report identified incidents within 15 day or within 24 hours when an imminent and serious risk is present. Importantly, the OES must protect trade secrets and national-security-related information but is obligated to publish anonymized, aggregated annual reports beginning in 2027. This design balances public accountability with operational security, reflecting the complex nature of frontier-model risk governance.

A particularly distinctive aspect of the TAFA is its whistleblower protection regime, codified through amendments to the California Labor Code. The provisions prohibit retaliation against employees who report catastrophic risks or critical safety incidents to internal compliance officers, the Attorney General, or federal agencies. The statute compels large developers to establish anonymous internal reporting channels, ensure periodic updates to employees who file reports, and provide judicial remedies including attorney's fees, burden-shifting provisions, and the availability of preliminary injunctive relief. Through these tools, the Act treats employees as critical nodes in the early-warning system for frontier AI safety, reflecting lessons learned from cybersecurity, aviation, and nuclear-risk regulation.

Beyond regulatory controls, the TAFA incorporates a proactive, capacity-building dimension through the creation of *CalCompute*, a public cloud computing consortium to be hosted where feasible- within the University of California system. *CalCompute* is envisioned not merely as an infrastructure project but as a comprehensive public platform integrating computing resources, research support, workforce expertise, and equitable access mechanisms. Its governance model includes academic institutions, public-interest organizations, technical experts, and labor representatives, demonstrating a multi-stakeholder approach capable of enhancing state-level AI sovereignty and reducing long-term dependence on private technology monopolies.

The statute has nevertheless attracted significant criticism. Scholars have raised concerns regarding potential conflicts with federal authority- particularly given the President's 2023 Executive Order on AI and national-security prerogatives arguing that the TAFE may test the limits of the Supremacy Clause. Others warn that transparency mandates could unintentionally expose proprietary information or sensitive model-security details. The computational and administrative burden of compliance may further entrench the dominance of large technology firms, potentially inhibiting innovation within startups and academic laboratories. Additionally, uncertainties surrounding the compute threshold and revenue-based categorization of "large frontier developers" have been cited as areas requiring refinement.

Despite these critiques, TAFE's overall architecture provides valuable lessons for global AI governance. Its transparency-driven, adaptive, and risk-tiered model demonstrates an alternative to rigid ex-ante certification regimes. It offers actionable policy instruments- mandatory documentation, incident reporting, whistleblower protection, and state-backed computing infrastructure- that can be selectively transferred to jurisdictions with varying institutional capacities. As such, the TAFE functions not only as a state-level statute but as an early prototype for future international regulatory frameworks addressing frontier AI.

Giriş

Yapay zekâ teknolojilerinin son beş yılda kazandığı olağanüstü hız, özellikle "*frontier* yapay zekâ" olarak tanımlanan ileri seviye modellerin ortaya çıkışıyla birlikte, hukuk düzenlerini köklü bir dönüşüm sürecine girmeye zorlamıştır. *Frontier* yapay zekâ; çok büyük veri kümeleri üzerinde eğitilen, milyarlarca parametreye sahip olan, geniş amaçlı kullanım kapasitesi sunan ve çıktıları kimi zaman öngörülemez biçimde değişebilen, insan düzeyine yaklaşan veya belirli alanlarda onu aşabilen performanslar sergileyen modelleri ifade eder.¹ Bu kategori, yalnızca teknik ölçekteki büyüklükleri nedeniyle değil, aynı zamanda toplumsal, ekonomik ve güvenlik boyutlarında yaratabilecekleri yaygın etkiler nedeniyle de ayrı bir risk sınıfı olarak ele alınmaktadır.²

Günümüzde *OpenAI*'ın *GPT-4* ve *GPT-5* modelleri, *Google*'ın *Gemini Ultra* serisi, *Anthropic*'in *Claude 3 Opus* sistemleri, *Meta*'nın *Llama 3* modelleri, *Apple*'ın *Ajax* yapıları ve *Nvidia*'nın kurumsal yapay zekâ mimarileri *frontier* yapay zekâ kategorisinin başlıca örnekleri olarak görülmektedir. Bu modeller; kritik altyapılarla etkileşim kurabilme, siber güvenlik zincirlerini zayıflatabilme, biyoteknoloji alanında kötüye kullanım riskleri yaratabilme, bilgi ekosistemlerini manipüle edebilme, demokratik süreçlere etki edebilme ve ekonomik karar mekanizmalarını yönlendirecek ölçekte çıktılar üretebilme gibi yüksek etkili sonuçlar doğurabilmektedir.³ Dolayısıyla *frontier* modellerin hukuki niteliği, sorumluluk rejimi, şeffaflık gereklilikleri, gözetim mekanizmaları ve toplumsal riskleri hem ulusal hukuk sistemleri hem de uluslararası yönetim yapıları açısından günümüzün en mühim tartışma başlıklarından biri hâline gelmiştir.

Bu kapsamda Kaliforniya'nın 2025 yılında kabul ettiği *Frontier Yapay Zekâda Şeffaflık Yasası* (*Transparency in Frontier Artificial Intelligence Act*- YASA), *frontier* modellerin doğurduğu riskleri merkeze alan ve bu konuda alt-ulusal düzeyde getirilen ilk kapsamlı düzenlemelerden biri olarak küresel ölçekte dikkat çekmiştir.⁴ YASA, yalnızca teknik nitelikli bir ürün güvenliği metni olarak değil; şeffaflık, felaket riski yönetimi, kurumsal hesap verebilirlik, *whistleblower*⁵ koruması ve kamu altyapısı inşası gibi birbirini tamamlayan çeşitli düzenleyici araçları aynı çatı altında toplayan bütüncül bir yönetim modeli olarak tasarlanmıştır. Bu anlamda YASA, yalnızca gelişmekte olan hukuk sistemleri için bir mevzuat örneği değil; *frontier* yapay zekâ çağının temel normatif sorunlarını somutlaştıran bir tür "yönetişim laboratuvarı" niteliği taşımaktadır.⁶

Frontier yapay zekâya ilişkin hukuki tartışmalar, günümüzde giderek daha karmaşık bir görünüm kazanmakta ve esasen üç temel eksen etrafında şekillenmektedir. İlk olarak, bu modellerin ortaya çıkardığı çok katmanlı tehdit ve belirsizliklere uygun bir risk kavramsallaştırmasının geliştirilmesi ve buna bağlı olarak düzenleyici sınırların nerede çizileceğinin belirlenmesi gerekmektedir. İkinci eksen, model eğitimi, metaveri akışı ve sistem davranışları hakkında kamusal bilgiye erişim ihtiyacını merkeze

¹ Charlie Bullock vd., "Frontier Model Definitions", *Law-AI* (Erişim 18 Kasım 2025).

² Matthijs M. Maas, "Scaling Law for AI: Issues, Necessity, and Feasibility", *Architectures of Global AI Governance: From Technological Change to Human Choice*, ed. Florian Egloff - Thomas Reinhold (Oxford: Oxford University Press, 2025), 118.

³ Markus Anderljung vd., "Frontier AI Regulation: Managing Emerging Risks to Public Safety", *arXiv* (Erişim 10 Kasım 2025).

⁴ California Senate Bill 53 (SB-53), 2025 Legislative Session, § 1. Bu çalışmada incelenen düzenleme, Kaliforniya Senatosu tarafından 2025 yılında kabul edilen SB-53 sayılı yasa olup resmî başlığı "Büyük Geliştiricilere İlişkin Yapay Zekâ Modelleri" (*Artificial intelligence models: large developers*)dir. Yasanın BÖLÜM 2'si kapsamında Kaliforniya Ticaret ve Meslekler Kanunu'na eklenen 25.1 numaralı bölüm (chapter) ise "Frontier Yapay Zekâda Şeffaflık Yasası" (*Transparency in Frontier Artificial Intelligence Act*) başlığını taşımaktadır. Dolayısıyla literatürde hem SB-53 hem de Sınır Yapay Zekâda Şeffaflık Yasası (*Transparency in Frontier Artificial Intelligence Act*) ifadeleri aynı mevzuata atıf yapmak üzere kullanılmaktadır.

⁵ Türkçede muhbir ya da ihlal bildiren gibi çeviri kullanımları olsa da işbu terminoloji İngilizcedeki orijinal haliyle literatürde sıklıkla kullanılmaktadır.

⁶ Justine Gluck, "California's SB 53: The First Frontier AI Law, Explained", *Future of Privacy Forum* (Erişim 10 Kasım 2025).

alan şeffaflık standartları ile bu standartların üzerinde yükselecek gözetim mimarisinin niteliğine ilişkindir. Üçüncü olarak ise, hızla ticarileşen *frontier* yapay zekâ ekosisteminde özel sektör ile kamu kurumları arasındaki güç, sorumluluk ve hesap verebilirlik dengesinin nasıl kurulacağı, düzenleyici tartışmaların en kritik boyutlarından birini oluşturmaktadır. Bu tartışmalarda hem akademi hem de politika yapıcılar açısından temel soru, inovasyonun teşvik edilmesi ile toplumsal güvenliğin sağlanması arasındaki hassas dengenin hangi ilkeler çerçevesinde kurulması gerektiğidir. YASA, bu soruya risk temelli ve şeffaflık odaklı bir yaklaşım getirerek, geleneksel teknik-önleyici düzenlemelerin aksine, raporlama yükümlülükleri, iç yönetim standartları ve bilgi paylaşımına dayalı bir hukuki denetim sistemi öngörmektedir.⁷ Dijital platformlar üzerinden sunulan yüksek riskli hizmetlerin yalnızca hizmet sağlayıcının bireysel sorumluluğu çerçevesinde değil, platformun regülasyon ağırlıklı sorumluluğu ve yönetim yükümlülükleri bağlamında ele alınması gerektiği, Türk hukukunda özellikle uzaktan sağlık hizmetleri alanında yapılan güncel çalışmalarda da vurgulanmaktadır. Bu yaklaşım, makalemize konu olan yüksek etkili yapay zekâ sistemlerinin düzenlenmesinde benimsenen risk temelli ve önleyici yönetim modelleriyle yapısal bir paralellik arz etmektedir.⁸

Bu makale, Kaliforniya'nın *frontier* yapay zekâya ilişkin düzenleyici yaklaşımını ayrıntılı biçimde incelemekte; *frontier* modelin hukuki sınırlarını, felaket riski (*catastrophic risk*) kavramsallaştırmasını, geliştiricilere yönelik çok katmanlı yükümlülük sistemini ve kritik güvenlik olayı bildirim rejimini sistematik biçimde ele almaktadır. Ayrıca *whistleblower* koruması, *CalCompute* aracılığıyla kamu altyapısının yeniden şekillendirilmesi ve federal-eyalet ilişkilerinin düzenlemenin tasarımına etkisi de değerlendirilerek, YASA'nın kurumsal ve normatif mimarisi bütüncül bir bakışla analiz edilecektir. Son aşamada, *frontier* yapay zekâ alanında düzenleme geliştirmek isteyen farklı hukuk sistemleri için aktarılabilir politika önerileri sunulacak; YASA'nın neden bir rol model olarak görülebileceği, hangi unsurlarının farklı ülke bağlamlarına uyarlanabilir olduğu ve hangi noktalarda bağlama özgü düzenlemeler gerektirdiği tartışılacaktır. Bu yönüyle çalışma, *frontier* yapay zekâ çağında etkili bir düzenleyici mimarinin nasıl inşa edilebileceğine ilişkin hem teorik hem de uygulamalı bir çerçeve sunmayı amaçlamaktadır.

I. Arka Plan: Kaliforniya'nın Yapay Zekâ Düzenlemelerinin Ortaya Çıkış Koşulları

YASA, küresel yapay zekâ yönetiminde öncü bir rol oynamış; yalnızca eyalet düzeyindeki düzenlemelerle sınırlı kalmayıp diğer politika yapıcılar için de reformcu bir model oluşturan bir sistem üzerine inşa edilmiştir.⁹ Bu bölümde, YASA'nın arka planına ilişkin temel unsurlar; federal hükümet ile eyalet arasındaki yetki çekişmeleri, düzenlemenin uluslararası alanda yarattığı yankılar ve büyük ölçekli yapay zekâ geliştiricilerinin verdiği tepkiler çerçevesinde ele alınacaktır.

A. Federal ve eyalet düzeyinde normatif etkileşim

Kaliforniya'nın düzenlemesinin ortaya çıkışında en belirleyici unsurlardan biri, ABD'de federal düzeydeki hukuki düzeni şekillendiren boşluktur. Kongre'ye 2020 sonrasında yapay zekânın etik, güvenlik ve şeffaflık boyutlarını ele alan çeşitli yasa tasarıları sunulmuş; ancak bu teklifler siyasi uzlaşma eksikliği ve teknoloji lobilerinin etkisi nedeniyle somut bir federal mevzuata dönüşmemiştir.¹⁰ 2023 tarihli Güvenli, Emniyetli ve Güvenilir Yapay Zekâ Konulu Başkanlık Kararnamesi (*Executive Order on Safe, Secure, and Trustworthy Artificial Intelligence*), Amerikan Başkanlık Kabinesi tarafından federal düzeyde rehber niteliğinde ilkeler ortaya koymuş;¹¹ ancak bu ilkeler, yön gösterici olmakla birlikte bağlayıcı bir normatif çerçeve oluşturacak hukuki güçten yoksun kalmıştır. Bu nedenle ABD'de federal yapay zekâ yönetimi, henüz zorlayıcı düzenleme (*binding regulation*) düzeyine ulaşmamış; eyaletlere kendi politika alanlarını doldurma imkânı tanıyan bir "federal normatif boşluk" durumu ortaya çıkmıştır.¹²

Kaliforniya, bu normatif boşluktan yararlanarak, kişisel verilerin korunması alanında Kaliforniya Tüketici Gizliliği Yasası¹³ (*California Consumer Privacy Act- CCPA*) ile yaptığına benzer biçimde, ulusal ve uluslararası düzeyde etkili bir öncü düzenleyici model geliştirmiştir. Bu durum, ABD hukuk düzeninde "laboratuvar federalizmi" olarak adlandırılan uygulamanın bir örneği olarak görülmektedir. Başka bir ifadeyle, federal düzeyde henüz ayrıntılı bir düzenleme bulunmayan alanlarda, eyaletlerin kendi politikalarını deneyerek ulusal çerçevenin oluşmasına katkı sunması anlamına gelmektedir.¹⁴

⁷ Bullock vd., "Frontier Model Definitions".

⁸ Kulular, Merve Aysegül. "Sağlık Hizmetlerinin Uzaktan Sunumu Kapsamında Online Platformların Hukuken Değerlendirilmesi". *Hacettepe Hukuk Fakültesi Dergisi*, 15/2, 2025, 129-168.

⁹ Devin Von Arx, "California's Approach to AI Governance", *CSET* (Erişim 22 Ekim 2025).

¹⁰ Ken D. Kumayama vd., "AI Agents: Greater Capabilities and Enhanced Risks", *Reuters Legal News* (Erişim 18 Kasım 2025).

¹¹ Ariana Yang, "Federal Proposals to Regulate Artificial Intelligence (AI)", *Congressional Research Service* (Erişim 17 Kasım 2025).

¹² National Institute of Standards and Technology, "Artificial Intelligence Risk Management Framework (AI RMF 1.0)", *NIST* (Erişim 17 Kasım 2025).

¹³ California Consumer Privacy Act (CCPA), California Attorney General (Erişim 19 Kasım 2025).

¹⁴ Kumayama vd., "AI Agents: Greater Capabilities and Enhanced Risks".

B. Uluslararası düzlemde karşılaştırmalı yaklaşım

Kaliforniya modelinin gelişimi, Avrupa Birliği'nin Yapay Zekâ Tüzüğü¹⁵ ile karşılaştırıldığında, normatif hedefler açısından benzerlikler gösterse de yapısal olarak farklı bir paradigma üzerine kuruludur. AB, risk temelli kategorik bir yaklaşımla teknik-önleyici (*ex-ante*) düzenlemeyi benimserken; Kaliforniya'nın modeli, çok daha esnek, uyarlanabilir ve geliştirici odaklı bir şeffaflık temelli yönetim anlayışına dayanmaktadır.¹⁶ Yani AB'nin *ex-ante* yaklaşımı yerine, Kaliforniya; bilgi paylaşımını, denetimi ve kamuya hesap verebilirliği düzenlemenin merkezine yerleştirerek düzenleme yoluyla öğrenme kapasitesini (*regulation-based learning capacity*) artırmaktadır.

Bu çerçevede, Kaliforniya'nın çabaları, OECD'nin 2021 tarihli Yapay Zekâ Prensipleri belgeleriyle ve Birleşmiş Milletler'in UNESCO aracılığıyla kabul ettiği Yapay Zekâ Etiğine İlişkin Tavsiye Metni (*Ethics of Artificial Intelligence Recommendation*) metniyle uyumlu değer temellerine oturmaktadır.¹⁷ Ancak YASA'nın farkı, etik ilkeleri soyut deklarasyonlardan çıkarıp, hesap verebilirlik ve bildirim zorunluluğu gibi somut hukuk mekanizmalarına dönüştürmesidir.¹⁸

C. Politik-ekonomik arka plan: Silikon Vadisi ve kamuoyu baskısı

Kaliforniya'nın yapay zekâ politikasının arkasında, yüksek teknoloji ekosisteminin merkezinde yer almanın getirdiği karmaşık bir çıkar dengesi vardır. *OpenAI*, *Google*, *Meta*, *Anthropic*, *Nvidia* ve *Apple* gibi teknoloji devlerinin eyalet sınırları içinde bulunması, Kaliforniya'yı hem inovasyonun hem de riskin merkezi haline getirmiştir. Bu durum, eyaletin yönetim yapısında özgün bir politika ekonomisi yaratmıştır.¹⁹

Bir yandan teknoloji şirketleri, yeniliğin sınırlandırılması endişesiyle aşırı düzenlemeye karşı çıkarken; diğer yandan kamuoyu, özellikle 2024 seçimlerinde *deepfake* içeriklerinin yaygınlaşmasıyla, yapay zekânın denetimsiz etkilerine dair endişelerde ortaklaşmıştır. Bu baskı, tüketici hakları savunucuları, insan hakları örgütleri ve akademi tarafından desteklenen bir "kamu denetimi talebine" dönüşmüştür.²⁰

Kaliforniya'nın inovasyon ekosistemi, teknoloji şirketleri ile kamu kurumları arasında dinamik bir politika geri besleme mekanizması (*policy feedback mechanism*) ortaya çıkarmıştır. Bu mekanizma sayesinde, teknoloji geliştirme süreçlerinde ortaya çıkan ihtiyaçlar ve sorunlar kamu politikalarına hızlı biçimde yansımış; oluşturulan politika çerçevesi de yeniden sektörün davranışlarını şekillendirerek karşılıklı ve sürekli bir etkileşim üretmiştir.²¹ Stanford, Berkeley ve USC gibi kurumlarda geliştirilen etik yapay zekâ laboratuvarları, yasa yapıcılara bilimsel altyapı sağlayarak politika döngüsüne akademik bilgi akışını entegre etmiştir.²² Bu, Kaliforniya'nın diğer eyaletlerden farklı olarak "bilgi tabanlı düzenleme" kapasitesine sahip olmasını sağlamıştır. Nitekim YASA'nın gerekçe bölümünde de bu husus açık bir biçimde vurgulanmaktadır.

Sonuç olarak, Kaliforniya'nın yapay zekâ düzenlemesi, sadece teknik bir politika belgesi değil; demokratik değerlerin, piyasa rekabetinin ve toplumsal güvenin yeniden tanımlandığı bir normatif deney alanıdır.

II. Frontier Yapay Zekâ Şeffaflık Rejiminin Normatif ve Kurumsal Çerçevesi

Kaliforniya'nın YASA düzenlemesi, klasik tüketici koruma ya da ürün güvenliğine ilişkin teknik normlardan ziyade, ileri seviye yapay zekâ modellerinin doğurabileceği felaket risklerini (*catastrophic risk*) ve buna bağlı yönetim sorunlarını merkeze alan bir şeffaflık rejimi kurar. YASA, daha baştan, Kaliforniya'nın yapay zekâ inovasyonunda dünya çapında bir merkez olduğuna vurgu yaparak (BÖLÜM 1(a)), bu alandaki norm üretiminin yalnızca eyalet için gerekli olan yerel bir politika tercihi değil, küresel ölçekte etki doğuran bir düzenleyici hamle olduğunu açıkça ifade etmektedir.²³ Yasama gerekçesinde, temel modellerin (*foundation model*) tıp, iklim analizi ve yangın tahmini gibi alanlarda sunduğu yüksek fayda potansiyeli (BÖLÜM 1(b)) ile bu modellerin kötüye kullanım veya teknik arızalar nedeniyle ortaya çıkarabileceği "felaket niteliğindeki riskler" (BÖLÜM 1(j)) birlikte ele alınmaktadır. Bu yaklaşım, metin düzeyinde inovasyon ile güvenlik arasındaki dengeyi gözetken tutarlı bir mevzuat çerçevesinin oluşturulmasına imkân tanımaktadır.

¹⁵ European Parliament and Council, "Regulation (EU) 2024/1689 of 13 June 2024 Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act)", *Official Journal of the European Union* L 202/1 (Erişim 12 Temmuz 2024).

¹⁶ Von Arx, "California's Approach to AI Governance".

¹⁷ Helen Toner – Timothy Fist, "Regulating the AI Frontier: Design Choices and Constraints", *CSET* (Erişim 11 Kasım 2025).

¹⁸ NIST, "Artificial Intelligence Risk Management Framework (AI RMF 1.0)".

¹⁹ Bullock vd., "Frontier Model Definitions".

²⁰ Amba Kak, "Executive Summary: AI Nationalism(s) – Global Industrial Policy Approaches to AI", *AI Now Institute* (Erişim 17 Kasım 2025).

²¹ Kak, "Executive Summary: AI Nationalism(s)".

²² Stanford Institute for Human-Centered Artificial Intelligence, "2024 AI Index Report", *Stanford HAI* (Erişim 17 Kasım 2025).

²³ Rishi Bommasani vd., *The California Report on Frontier AI Policy*. The Joint California Policy Working Group on AI Frontier Models, 2025. <https://www.gov.ca.gov/wp-content/uploads/2025/06/June-17-2025-%E2%80%93-The-California-Report-on-Frontier-AI-Policy.pdf>, 3.

A. Yasama gerekçesi: şeffaflık, felaket riski ve sorumlu geliştirme

BÖLÜM 1’de yer alan bulgular, düzenlemenin arka planındaki temel ilkeleri açıkça ortaya koymaktadır. Yasama organı, *frontier* modellerin etkilerini anlamak için “sağlam ve şeffaf bir kanıt ortamına” ihtiyaç olduğunu, politika yapımcıların tüketicileri koruma, sektör bilgisinden yararlanma ve güvenlik uygulamalarını teşvik etme hedeflerini aynı anda gözetmesi gerektiğini belirtir.²⁴ Bu çerçevede şeffaflık, yalnızca bilgi paylaşımı değil; hesap verebilirlik, rekabet ve kamu güveninin artırılması için araçsal bir değer olarak konumlandırılır.²⁵

Yasa koyucu özellikle YASA, BÖLÜM 1(i) kapsamında iki mekanizmayı, şeffaflık hedefinin ana araçları arasında sayar: *whistleblower* korumaları ve kamuya yönelik bilgi paylaşımı²⁶ ile olay bildirim sistemleri (*incident reporting systems*). Buna paralel olarak, *frontier* geliştiricilerin kendi iç araştırmalarında tespit ettikleri riskler hakkında kamu otoritelerine ve kamuoyuna yeterli bilgi vermemelerinin, bu teknolojilere duyulan güveni zedelediği ifade edilir.²⁷ Dolayısıyla YASA, BÖLÜM 1(i), gönüllü sektör inisiyatiflerinin tek başına yeterli olmadığını; “zorunlu, standartlaştırılmış ve nesnel” raporlama yükümlülüklerinin *frontier* geliştiriciler için hukukî bir çerçeve hâline getirilmesi gerektiğini vurgular.

Bunun yanında, YASA’nın metni “*frontier* yapay zekâ sistemleri” kavramını yalnızca bugün için geçerli bir teknik eşik olarak değil, yapay zekânın hızlı evrimiyle birlikte hareket etmesi gereken hareketli bir sınır olarak kurgular.²⁸ Bu nedenle, ileri düzey yapay zekâ araştırmalarını izleyen, en gelişmiş sistemlerden doğabilecek ciddi riskler konusunda politika yapımcıları ve kamuoyunu zamanında uyaran bir yasal çerçeve ihtiyacı açıkça tespit edilmektedir.²⁹ Toplu güvenliğin ise nihai olarak, *frontier* geliştiricilerin öngörülebilir risklerin ölçeğiyle orantılı “özenli davranışına (*due care*)” bağlı olduğu belirtilir.³⁰ Bu vurgu, ilerleyen maddelerde geliştirilen şeffaflık ve raporlama yükümlülüklerinin teknik bir formaliteden ziyade “özen borcu” niteliğinde bir düzenleyici standartla bağlantılı olduğunu gösterir.

B. Kapsam, tanımlar ve “frontier” kategorisinin inşası

YASA’nın uygulama alanını Kaliforniya Ticaret ve Meslekler Kanunu (*Business and Professions Code*) içinde yeni eklenen *Chapter 25.1* altında tanımlar ve bu bölümün Frontier Yapay Zekâda Şeffaflık Yasası olarak adlandırılacağını hükme bağlar.³¹ Kapsamın belirlenmesinde, yalnızca teknik sistemler değil, bu sistemlerin geliştirilmesinden ve kullanılmasından sorumlu aktörler de merkezî rol oynar.

Tanımlar bölümü, düzenlemenin hukuki ekosistemini taşıyan kilit kavramları sistematik biçimde ortaya koyar. YASA, § 22757.11(b) hükmüne göre “yapay zekâ model” kavramı, belirli bir otonomi düzeyine sahip, aldığı girdilerden hareketle çıktılar üreterek fiziksel veya sanal ortamları etkileyebilen mühendislik veya makine tabanlı sistemler olarak formüle edilmiştir. Bu geniş tanım, yalnızca dil modellerini değil, çok modlu veya fiziksel sistemlerle bağlantılı pek çok farklı yapay zekâ türünü kapsayabilecek esnekliktedir.

Düzenlemenin en özgün yönü, yapay zekâ sistemlerini “*frontier model*” ve “*frontier developer*” kavramları üzerinden sınıflandıran risk temelli yaklaşımıdır. YASA’nın, *frontier* modeli yalnızca gelişmiş bir yazılım olarak değil, eşi benzeri görülmemiş hesaplama gücüyle eğitilmiş ve bu nedenle toplumsal ölçekte risk üretebilen bir yapay zekâ kategorisi olarak tanımlar. Bu kapsamda *frontier model*; YASA, § 22757.11(f), (i)(1) hükmüne göre hem geniş veri kümeleri üzerinde eğitilen temel model olma niteliğine sahip olmak hem de toplam eğitim hesaplama gücünün 10^{26} işlem (FLOP) seviyesinin üzerinde gerçekleşmiş olması şartını taşır. Bu hesaplama eşiği, yalnızca modelin ilk eğitim sürecini değil; daha sonra yapılan ince ayar eğitimi, ek uyarılama eğitimi, uyarlayıcı son eğitim aşaması ve diğer tüm maddi nitelikteki model iyileştirme aşamalarında kullanılan ek hesaplama güçlerini de kapsar.³² Böylece YASA, *frontier* statüsünü tek bir teknik anahtarın üzerinden belirlemek yerine, modelin tüm yaşam döngüsü boyunca ulaştığı toplam kapasiteye dayandırır. Bu yaklaşım, büyük modellerin zamanla daha karmaşık ve güçlü hâle gelmesini dikkate alan dinamik bir sınıflandırma oluşturur. Bu nedenle *frontier* model kavramı, yalnızca bugünün büyük ölçekli modellerini değil; gelecekte ortaya çıkacak, daha yüksek işlem gücüyle eğitilmiş ve potansiyel olarak daha yüksek etkiye sahip modelleri de kapsayabilecek şekilde tasarlanmıştır. Bu esneklik, düzenlemenin hızla değişen yapay zekâ ekosisteminde geçerliliğini korumasını sağlayan temel ilkedir.

²⁴ Bommasani vd., *The California Report on Frontier AI Policy*, 4-5; YASA, BÖLÜM 1(e).

²⁵ Bommasani vd., *The California Report on Frontier AI Policy*, 4-5; YASA, BÖLÜM 1(g).

²⁶ Bommasani vd., *The California Report on Frontier AI Policy*, 32; YASA, BÖLÜM 1(h).

²⁷ YASA, BÖLÜM 1(f), 1(i).

²⁸ Bommasani vd., *The California Report on Frontier AI Policy*, 13.

²⁹ Toner- Fist, “Regulating the AI Frontier: Design Choices and Constraints”; Bommasani – *The California Report on Frontier AI Policy*, 7- 13.

³⁰ YASA, BÖLÜM 1(p).

³¹ California Business and Professions Code, Chapter 25.1 (§ 22757.10), (California Legislature, 2025).

³² YASA, § 22757.11(i)(2).

Frontier geliştirici ise, bu nitelikte bir *frontier* modeli eğiten veya eğitimini başlatan kişi ya da kuruluş olarak tanımlanır ve bu tanım, geliştiricinin eğitimde kullandığı hesaplama gücünün *frontier* eşliğini karşılaması şartına bağlanır.³³ Bunun yanında “büyük *frontier developer*” kategorisi, YASA, § 22757.11(j) hükmüne göre *frontier* geliştiricinin kendisi ve bağlı şirketlerinin önceki yıl toplam brüt gelirinin 500 milyon ABD dolarını aşması şartına bağlanmıştır. Böylece YASA, hem model ölçeğine hem de kurumsal ekonomik kapasiteye bağlı ikili bir eşik sistemi kurarak, yükümlülüklerin orantılı bir biçimde dağıtılmasını amaçlar.

YASA'nın risk kavramsallaştırması da oldukça net bir şekilde kodlanmıştır. YASA, § 22757.11(c)(1) hükmüne “felaket riskleri”, *frontier* geliştiricinin *frontier* modeli geliştirme, saklama, kullanım veya dağıtım faaliyetlerinin tek bir olayda elliden fazla kişinin ölümü veya ağır yaralanmasına yahut 1 milyar ABD dolarını aşan malvarlığı zararına “maddi katkıda” bulunması ihtimali olarak tanımlanır. Bu risk kategorisine; kimyasal, biyolojik, radyolojik veya nükleer silahların geliştirilmesine uzman düzeyde yardımcı olma, anlamlı insan gözetimi olmaksızın siber saldırı veya insan eliyle işlense ağır suç sayılacak fiillerin işlenmesi, ya da modelin geliştirici veya kullanıcının kontrolünden çıkması örnek olarak verilmiştir.³⁴ Bunun yanında YASA, kamuya zaten açık olan bilginin yeniden üretimi, federal hükümetin meşru faaliyetleri veya *frontier* modelin kayda değer katkısı olmaksızın gerçekleşen zararların bu kategoriye girmeyeceğini belirterek risk kavramını daraltır.³⁵

Bu tanımlar, *frontier* modellerin doğurduğu riskleri yalnızca “yanlış çıktı” veya “hatalı karar” düzeyinden çıkarıp, toplumsal ölçekte kitlesel zarar, kritik güvenlik riski ve altyapı kırılabilirliği bağlamında ele alan yüksek eşikte sahip bir risk sınıfı oluşturur.

C. Şeffaflık unsurları

YASA'nın temel düzenleyici mekanizması, büyük ölçekli *frontier* geliştiriciler için zorunlu hâle getirilen *Frontier* Yapay Zekâ Dokümantasyonu (*frontier AI framework*) oluşturma yükümlülüğüdür. YASA'nın § 22757.12(a)'da, büyük ölçekli yapay zekâ geliştiricisi statüsündeki aktörlerin, *frontier* modellerine uygulanacak teknik ve örgütsel protokolleri içeren bir *frontier* yapay zekâ çerçevesi hazırlamasını, bunu uygulamasını ve internet sitesinde “açık ve belirgin” şekilde yayımlamasını zorunlu kılar. Bu çerçeve, yalnızca iç politika belgesi değil; kamuya açıklanan ve geliştiricinin kendi risk yönetimi yaklaşımına dair hesap verebilirlik yaratan bir metin niteliğindedir.

Madde, bu çerçevenin asgari içeriğini de detaylı biçimde düzenler. Öncelikle geliştiricinin, ulusal ve uluslararası standartları ile sektörcü kabul görmüş en iyi uygulamaları *Frontier* Yapay Zekâ Dokümantasyonu'na nasıl dâhil ettiğini açıklaması beklenir.³⁶ Bu yaklaşım, YASA'nın federal düzeyde veya uluslararası örgütler bünyesinde geliştirilecek mevzuat açısından, rol model nitelikte standartlara sahip olduğunu göstermektedir. Geliştiriciler ayrıca, *frontier* modelin felaket riski oluşturabilecek kabiliyetlere sahip olup olmadığını tespit etmek için kullandıkları eşikleri ve değerlendirme yöntemlerini tanımlamak zorundadır.³⁷ YASA, bu eşiklerin çok katmanlı olabileceğini açıkça belirtir; böylece farklı risk düzeylerine göre farklı değerlendirme rejimleri kurulmasına imkân tanınır. YASA, § 22757.12(a)(3)–(4) hükmüne göre sonraki adımda, tespit edilen risklere karşı uygulanacak riskleri düşürme önlemlerinin ve bu risk düşürme faaliyetlerinin tamamen ortadan kaldırılması gibi stratejik hedeflerin de öncesinde nasıl gözden geçirildiğinin *Frontier* Yapay Zekâ Dokümantasyonu'nda yer alması gerekir. Üçüncü taraf (bağımsız dış) değerlendirmelerinin kullanımı da ilgili dokümantasyon kapsamında teşvik edilmekte; büyük ölçekteki *frontier* geliştiricilerin dış uzmanlarla iş birliği yapmaları öngörülmektedir.

Önemli bir unsur da siber güvenlik ve iç yönetim boyutudur. YASA, § 22757.12(a)(7)–(9) bu boyutu düzenlemektedir. Buna göre *Frontier* Yapay Zekâ Dokümantasyonun, yayımlanmamış model ağırlıklarının (*model weights*) izinsiz erişim, değiştirme veya aktarımına karşı uygulanacak siber güvenlik uygulamalarını içermesi zorunludur. Ayrıca kritik güvenlik olaylarının tanımlanması ve bu olaylara verilecek kurumsal tepkilerin ve bu süreçlerin uygulanmasını sağlayacak iç yönetim mekanizmalarının da çerçevede yer alması gerekir. Son olarak, *frontier* modellerin kurum içi kullanımından kaynaklanabilecek felaket risklerinin, gözetim mekanizmalarının aşılması ihtimali de dâhil olmak üzere, sistematik biçimde değerlendirilmesi zorunlu kılınır.³⁸ YASA, § 22757.12(b)'de yapılan düzenlemeye göre bu belge, yılda en az bir kez gözden geçirilmeli, gerekli görüldüğünde güncellenmeli ve yapılan “maddi değişiklikler” ile bunların gerekçeleri 30 gün içinde kamuya duyurulmalıdır. Böylece *Frontier* Yapay Zekâ Dokümantasyonu, statik bir doküman değil; teknolojik, kurumsal ve risk temelli değişimlere uyum sağlayan dinamik bir yönetim aracı hâline gelir.

³³ YASA, § 22757.11(h).

³⁴ Bommasani vd., *The California Report on Frontier AI Policy*, 12- 14.

³⁵ YASA, § 22757.11(c)(2)).

³⁶ YASA, § 22757.12(a)(1).

³⁷ Toner- Fist, “Regulating the AI Frontier: Design Choices and Constraints”; YASA, § 22757.12(a)(2).

³⁸ YASA, § 22757.12(a)(10).

Şeffaflık yükümlülükleri, *frontier* modellerin kullanıma sürülme süreciyle yakından bağlantılıdır. § 22757.12(c)(1), yeni bir *frontier* modelin veya mevcut bir modelin esaslı biçimde değiştirilmiş sürümünün kullanıma sürülme sürecinden önce veya eş zamanlı olarak geliştiricinin bir şeffaflık raporu yayımlamasını şart koşar. Bu rapor, geliştiricinin internet adresi, iletişim mekanizması, modelin çıkış tarihi, desteklediği diller ve çıktı modaliteleri, modelin öngörülen kullanım alanları ve genel kullanım kısıtlamaları gibi temel bilgileri içerir.³⁹ Eğer geliştirici “büyük ölçekli *frontier* geliştirici” statüsündeyse, bu rapora modelin felaket riskine ilişkin değerlendirmelerin özeti, bu değerlendirmelerin sonuçları, üçüncü taraf katkısının düzeyi ve *Frontier* Yapay Zeka Dokümantasyonu kapsamındaki diğer adımların özeti de eklenmelidir.⁴⁰

Öte yandan, geliştiricilerin ticari sırlarını, siber güvenliklerini, kamu güvenliğini veya ulusal güvenliği korumak için yayımlanan dokümanlarda belirli kısımları sansürleme (*redaction*) imkânı tanınmıştır.⁴¹ Ancak bu hâlde dahi, YASA, § 22757.12(f)(2)’de tespit edildiği üzere sansürün niteliği ve gerekçesi açıklanmalı ve sansürlenmemiş bilgiler beş yıl süreyle muhafaza edilmelidir. Böylece sosyo-teknik yapı içinde hem bilgi güvenliği hem de asgari şeffaflık dengesi kurulmaya çalışılır.

D. Kritik güvenlik olayları ve bildirim rejimi

YASA, *frontier* modellerin yalnızca önleyici çerçevelerle değil, gerçek dünyadaki kullanım sonrasında ortaya çıkabilecek kritik güvenlik olaylarının (*critical safety incident*) izlenmesiyle de yönetilmesini öngörür. YASA § 22757.13(a), Acil Durum Hizmetleri Ofisi’nin (*Office of Emergency Services*) hem *frontier* geliştiriciler hem de herhangi bir gerçek kişi kritik güvenlik olaylarının raporlanmasına imkân veren bir bildirim mekanizması kurmasını zorunlu kılar. Ayrıca, bildirimde olay tarihi, olayın neden kritik güvenlik olayı sayıldığı, olayın kısa ve sade açıklaması ve olayın *frontier* modelin iç kullanımıyla bağlantısı olup olmadığı belirtilmelidir.

YASA, § 22757.13(b)(1)-(2)’de düzenlendiği üzere büyük çaplı *frontier* geliştiriciler ayrıca *frontier* modellerinin kurum içi kullanımlarından kaynaklanabilecek felaket riskine ilişkin değerlendirme özetlerini gizli olarak Acil Durum Hizmetleri Ofisi’ne (ADHO) iletmekle yükümlüdür. ADHO’nun bu tür raporlara erişimi, “bilmesi gerekenlerle sınırlı erişim ilkesi” (*the need-to-know principle*) uyarınca sınırlandırılmış; ayrıca yetkisiz erişime karşı gerekli koruyucu önlemleri (*unauthorized access safeguards*) alma yükümlülüğü düzenlenmiştir. Bu yapı, teknik risk bilgisinin kamu otoritesine akışını sağlarken, ticari sır ve güvenlik hassasiyetlerini de dikkate alan başarılı bir gizli raporlama rejimi kurar.

Frontier geliştiriciler, YASA, § 22757.13(c)(1)-(4) hükmüne göre *frontier* modellerine ilişkin herhangi bir kritik güvenlik olayını keşfettikten itibaren 15 gün içinde ADHO’ya bildirmekle yükümlüdür. Eğer olay, ölüm veya ciddi bedensel zarar oluşturma açısından “yakın ve ciddi risk” teşkil ediyorsa, yirmi dört saat içinde uygun yetkili mercie bildirim yapılması zorunludur. YASA, bu ilk bildirimin ardından yeni bilgi ortaya çıkması hâlinde ek rapor sunulmasını da mümkün kılar. Ayrıca *frontier* geliştiriciler, *frontier* model niteliğinde olmayan temel modellerde gerçekleşen kritik olayları da raporlamaya özendirilir; ancak bu gönüllü bir yükümlülük olarak düzenlenmiştir.

ADHO, *frontier* geliştiricileri tarafından sunulan kritik güvenlik olayı raporlarını incelemekle yükümlüdür ve bu inceleme yetkisi zorunlu niteliktedir. Buna karşılık, YASA, § 22757.13(d) hükmü kapsamında kamu bireyleri tarafından yapılan bildirimler bakımından OES’in inceleme yetkisi takdiri olup, yani incelemek zorunda değildir; uygun gördüğünde inceleme yapabilir. Kolluk kuvveti olarak Başsavcı veya Ofis, kritik güvenlik olayı raporlarını ve ilgili işçi bildirimlerini yasama organına, valiye, federal makamlara veya uygun kurumlara aktarabilir ancak bu aktarımda ticari sırlar, kamu güvenliği, siber güvenlik ve ulusal güvenlik riskleri özellikle dikkate alınmalıdır.⁴²

Önemli bir nokta, YASA, § 22757.13(f)’de belirtildiği üzere kritik olay raporlarının, iç kullanım risk değerlendirmelerinin ve ilgili çalışan raporlarının *California Public Records Act* kapsamındaki erişim hakkından istisna tutulmasıdır. Bu kayıtlar kamuya açık değildir. Buna karşılık, Ofis 1 Ocak 2027’den itibaren yıllık olarak anonimleştirilmiş ve toplulaştırılmış bilgiler içeren raporlar yayımlamakla yükümlüdür. Bu raporlar da yine ticari sır, siber güvenlik ve ulusal güvenlik hassasiyetlerine göre süzülme zorundadır.⁴³

Buna ek olarak, ADHO; kritik güvenlik olayı bildirimine ilişkin yükümlülükler bakımından belirli federal yasaların, düzenlemelerin veya rehberlerin, bu bölümde öngörülen standartlarla “özde eşdeğer” olduğunu veya bunlardan daha sıkı bir koruma düzeyi sağladığını tespit ettiği takdirde, bu federal

³⁹ YASA, § 22757.12(c)(1)(A) –(G).

⁴⁰ YASA, § 22757.12(c)(2)(A) –(D).

⁴¹ Bommasani vd., *The California Report on Frontier AI Policy*, 42.

⁴² YASA, § 22757.13(e).

⁴³ YASA, § 22757.13(g)(1)– (2).

çerçeveleri geçerli referans norm olarak kabul edebilir.⁴⁴ Böyle bir durumda *frontier* geliştirici, YASA, § 22757.13(i)'ye göre söz konusu federal çerçeveye tam uyum sağladığını resmen beyan ettiğinde, bu uyum belirli koşullar altında eyalet düzeyindeki yükümlülüklerin karşılanmasıyla eşdeğer sayılır. Bu eşdeğerlik mekanizması, bir yandan *frontier* geliştiricilerin federal ve eyalet düzeyindeki çok katmanlı yükümlülükleri arasında uyum sağlamasını kolaylaştırmakta, diğer yandan eyalet ile federal düzenleyici sistem arasında daha bütüncül ve tutarlı bir risk yönetimi yaklaşımı oluşturulmasına katkı sunmaktadır.

E. Yaptırım ve sorumluluk rejimi

Yaptırım mekanizması, düzenlemenin bütününde öngörülen şeffaflık ve raporlama yükümlülüklerini destekleyen bir “arka plan gücü” işlevi görür. YASA § 22757.15(a), büyük *frontier* geliştiricilerin zorunlu dokümanları yayımlamaması veya iletmemesi, YASA § 22757.12(e) kapsamında felaket riski veya risk yönetimi hakkında “maddi olarak yanlış veya yanıltıcı” beyanda bulunması, YASA § 22757.13 uyarınca kritik güvenlik olaylarını raporlamaması ya da kendi Frontier Yapay Zeka Dokümantasyonuna uymaması hâllerinde, ihlalin ciddiyetine göre olay başına bir milyon ABD dolarını geçmeyecek şekilde medeni para cezası ile karşılaşacağını hükme bağlar. Cezanın her ihlal için ayrı uygulanabilmesi ise düzenleyici etkinliği artıran çarpan etkisi yaratır.

Dikkat edilmesi gereken ikinci boyut, yaptırımların uygulanmasında yetkinin yalnızca Kaliforniya Başsavcısı'na verilmiş olmasıdır.⁴⁵ Bu tercih, yerel idarelerin veya bireysel davacıların dağınık ve tutarsız uygulamalarına karşı, merkezi ve uzmanlaşmış bir yaptırım mekanizması oluşturmayı amaçlar. Böylece yaptırım süreci, teknik ve hukuki uzmanlığa sahip bir makama devredilerek düzenlemenin tutarlılığı korunur. Başsavcının, özellikle büyük teknoloji şirketleriyle ilgili karmaşık ve yüksek etkili soruşturmaları yürütme kapasitesi düşünüldüğünde, bu yetkilendirmenin politika yapıcılar tarafından bilinçli bir tercihi yansıttığı söylenebilir.

YASA'nın genel hükümlerini düzenleyen BÖLÜM 5, metnin nasıl yorumlanması gerektiğine ilişkin temel ilkeleri ortaya koyar ve düzenlemenin tamamlayıcı niteliğini özellikle vurgular. Buna göre YASA, amacını gerçekleştirecek şekilde geniş yorumlanmalı; mevcut hukukî yükümlülükleri ortadan kaldıran değil, bunlara ek yükümlülükler getiren bir araç olarak anlaşılmalıdır. Metin ayrıca, federal hukukla veya federal kamu kurumlarıyla yapılan sözleşmelerle doğrudan çelişki ortaya çıkması hâlinde, söz konusu durumlarda federal düzenlemelerin üstün geleceğini ve eyalet hükümlerinin uygulanamayacağını açık biçimde belirtir. Bu durum, eyalet düzenlemesinin federal normlarla çatıştığı noktalarda geri planda kalacağını göstermektedir. Bununla birlikte, YASA eyalet düzeyinde belirli bir düzenleme üstünlüğü de kurar. Nitekim metin, 1 Ocak 2025'ten sonra kabul edilen yerel düzenlemelerin, *frontier* geliştiricilerin felaket risklerini yönetmeye ilişkin alanlarda bu YASA'yla çelişmeyeceğini hükme bağlar. Böylece YASA, federal düzeyde üstünlüğü kabul ederken, eyalet içindeki belediye ve ilçe düzeyindeki normlara karşı belirli bir hiyerarşik üstünlük ilişkisi tesis etmektedir.

F. Felaket riski bağlamında whistleblower koruması

YASA, *frontier* yapay zekâ ekosisteminin yalnızca teknik yönlerini değil, aynı zamanda bu ekosistemde çalışan bireylerin kurumsal güvenlik kültürünü de yakından ilgilendiren bir alan olan whistleblower korumasını kapsamlı biçimde ele almaktadır. Kaliforniya İş Kanunu'na eklenen yeni bir başlıkla, felaket riski ve kritik güvenlik olayı gibi yüksek etkili durumlara ilişkin iç raporlamalarda çalışanları korumayı amaçlayan gelişmiş bir rejim oluşturur.⁴⁶ Bu bölümde kullanılan *frontier* geliştirici, temel model ve büyük *frontier* geliştirici gibi temel kavramlar; Kaliforniya Ticaret ve Meslekler Kanunu 22757.11'de tanımlanan teknik terminolojiyle uyumlu şekilde yeniden düzenlenmiştir.⁴⁷

YASA'nın 1107.1 sayılı hükme göre *frontier* geliştiricilerin çalışanların bu tür risklere ilişkin bilgi vermesini engelleyen, caydırıcı veya bildirim yapan çalışanlara karşı misilleme anlamına gelebilecek kural, politika veya sözleşmeler uygulamasını açıkça yasaklar. Böylece çalışanların, gerektiğinde Başsavcı'ya, ilgili federal otoritelere, hiyerarşik amirlerine veya riskin giderilmesi konusunda yetkili başka çalışanlara bilgi vermesi güvence altına alınır. Bu yaklaşım, *frontier* yapay zekâ gibi dinamik ve yüksek riskli alanlarda içsel güvenlik uyarı mekanizmalarının serbestçe işlemlerini sağlayan kurumsal şeffaflık ilkesine dayanır.

YASA, § 1107.1(e)-(f) kapsamında yalnızca dış bildirimleri korumakla kalmaz; aynı zamanda büyük *frontier* geliştiricilerin çalışanlar tarafından anonim bildirim yapılabilmesini sağlayacak iç raporlama mekanizmaları kurmasını zorunlu kılar. Bu mekanizma, bildirim yapan çalışana soruşturmanın durumu hakkında düzenli bilgilendirme yapılmasını gerektirir. Ayrıca bildirimlerin üst yönetime

⁴⁴ YASA, § 22757.13(h).

⁴⁵ YASA, § 22757.15(b).

⁴⁶ YASA, § 1107-1107.2.

⁴⁷ YASA, § 1107(d)-(f).

periyodik olarak iletilmesi öngörülmüştür ancak bildirimin muhatabı üst düzey bir yönetici ise bu aktarım yapılmaz. Böylece çalışanların sessizleştirilmesi veya raporların yönetsel hiyerarşi içinde kaybolması engellenir.

Whistleblower koruması, *frontier* yapay zekâ alanındaki felaket riski yönetiminin kurumsal bütünlüğünü sağlamak bakımından yalnızca misillemeye karşı bir güvence sunmakla yetinmez; YASA bu korumayı pratikte etkili hâle getirmek için üç tamamlayıcı mekanizmayı da devreye sokar. Buna ilişkin hükümler YASA, § 1107.1(f)-(h) arasında yer almaktadır. Buna göre ilk olarak, çalışanların açtığı davalarda başarı elde edilmesi durumunda mahkemelerin avukatlık ücretini çalışan lehine hükmedebilmesi öngörülmüştür. Bu hüküm, özellikle büyük teknoloji şirketlerinin sahip olduğu finansal ve hukuksal kapasite karşısında bireysel çalışanların caydırıcı maliyet baskısı altında hak arama süreçlerinden vazgeçmesini engelleyen bir eşitleyici işlev görür. İkinci olarak, YASA belirli aşamalarda ispat yükünü işverenin üzerine daha ağır bir standartla kaydırır. Bu yaklaşım, *frontier* geliştiricilerin teknik süreçlere, denetim kayıtlarına ve model davranışlarına ilişkin bilgi tekeli elinde bulundurduğu bir ortamda, çalışanın iddialarının araştırılabilir ve doğrulanabilir olmasını sağlayarak kurumsal güç asimetrisini dengeler. Üçüncü olarak, çalışanlara ihtiyaç duyduklarında hızlı biçimde geçici hukuki koruma (*preliminary injunctive relief*) talep etme imkânı tanınmıştır.⁴⁸ Bu mekanizma, misilleme tehdidinin doğrudan ve acil olduğu durumlarda çalışanın korunmasını sağlayarak, *whistleblower* rejimini yalnızca geriye dönük bir koruma değil, aynı zamanda ileriye dönük bir güvenlik kalkanı olarak işlemlerini mümkün kılar. Bu üçlü yapı, *whistleblower* korumasını soyut bir normatif taahhüt olmaktan çıkarıp, *frontier* yapay zekâ ekosisteminde fiili güvence üreten, erişilebilir ve işlevsel bir kurumsal hakka dönüştürür. Böylece YASA, hem çalışanların risk bildirimini teşvik eder hem de felaket riski yönetimi bakımından şirket içi bilgi akışının önündeki yapısal engelleri azaltarak toplumsal güvenliği güçlendiren bir şeffaflık altyapısı oluşturur.⁴⁹

G. CalCompute ve kamu altyapısının yeniden tasarlanması

YASA, yalnızca özel sektörün *frontier* geliştiricilerini düzenleyen bir çerçeve oluşturmakla yetinmez; aynı zamanda kamu altyapısının *frontier* yapay zekâ çağının risklerine ve ihtiyaçlarına uygun biçimde yeniden tasarlanmasını da hedefler. Bu yönüyle düzenleme, klasik “güvenlik odaklı” yaklaşımın ötesine geçerek kamu yararına inovasyon, eşit erişim ve toplumsal kapasite inşasını merkeze alan bir vizyon ortaya koyar.

Kanun koyucu bu kapsamda Kaliforniya Hükümet Kanunu’na eklenen YASA § 11546.8 madde ile, Hükümet Operasyonları Ajansı (*Government Operations Agency*) bünyesinde *CalCompute* adıyla bir kamu bulut bilişim kümesi (*public cloud computing cluster*) oluşturulmasına yönelik çerçeveyi tasarlamakla görevli bir konsorsiyum kurulmasını öngörür. Bu konsorsiyumun temel misyonu, yapay zekânın güvenli, etik, adil ve sürdürülebilir biçimde geliştirilmesini mümkün kılacak bir kamu bulut altyapısının teknik, yönetsel ve organizasyonel temelini oluşturmaktır.⁵⁰

CalCompute, YASA, § 11546.8(c)-(d) kapsamında yalnızca depolama ve işlem gücü sağlayan teknik bir bulut sistemi olarak tasarlanmamıştır. Aksine, *CalCompute*’u; altyapının çalıştırılması, bakımının yapılması, kullanıcıların desteklenmesi ve araştırma süreçlerinin yürütülmesi için gerekli insan kaynağı ve kurumsal kapasiteyi de içeren “tam teşekküllü” bir kamu platformu olarak kurgulanmıştır. Bu yapı, Kaliforniya’nın yapay zekâ gelişiminde özel sektör karşısında tamamen pasif bir izleyici değil, doğrudan altyapı sağlayıcı ve koordinatör bir aktör olmasını hedefleyen yenilikçi bir yaklaşımı yansıtır. Düzenleme, mümkün olduğunca *CalCompute*’un Kaliforniya Üniversitesi (*University of California*) bünyesinde kurulmasını teşvik eder. Böylece akademik araştırma ekosistemi ile kamu altyapısı arasında organik bir entegrasyon yaratılması; yapay zekâ çalışmalarına erişimin genişletilmesi ve bilimsel kapasitelerin kamusal amaçlarla bir araya getirilmesi amaçlanmaktadır. Konsorsiyumun yönetim yapısı da bu vizyonu destekler niteliktedir. YASA § 11546.8(g) uyarınca konsorsiyum; üniversite temsilcileri, işçi sendikaları, etik, tüketici hakları ve kamu yararı alanındaki paydaşlar teknik yapay zekâ uzmanları gibi geniş bir yelpazeden oluşur. Bu çok paydaşlı yönetim modeli, *frontier* yapay zekâ politikalarının yalnızca teknik uzmanlığa değil; toplumsal, etik ve ekonomik perspektiflere de dayanması gerektiği yönündeki yasama iradesini net biçimde ortaya koyar.

CalCompute, YASAda sadece bir “altyapı projesi” olarak değil, Kaliforniya’nın *frontier* yapay zekâ ekosisteminde kamusal kapasiteyi güçlendiren, araştırma ve inovasyonu destekleyen, eşit erişimi artıran ve toplumsal çıkarı önceleyen stratejik bir kurumsal araç olarak konumlandırılmıştır. Bu yönüyle düzenleme, *frontier* yapay zekâyı ilişkin risk odaklı hükümleri tamamlayan pozitif bir kamu politikası vizyonu sunar.

⁴⁸ Bommasani vd., *The California Report on Frontier AI Policy*, 29-30.

⁴⁹ Hayden Field, “SB 53, the landmark AI transparency bill, is now law in California” *The Verge*, (29 Eylül 2025).

⁵⁰ YASA, § 11546.8(b).

H. YASA'ya karşı eleştiriler

YASA, yenilikçi yapısına ve şeffaflık odaklı normatif vizyonuna rağmen gerek hukuki tutarlık gerekse uygulanabilirlik bakımından çeşitli eleştirilere konu olmuştur. Bu eleştiriler hem anayasal ilkelerle çatışan boyutlara hem de piyasa özgürlüğü, teknik fizibilite ve bilgi güvenliği gibi pratik sorunlara işaret etmektedir.

1. Federalizm ve yetki aşımı sorunu

Eleştirilerin başında, YASA'nın federal hukukla olan ilişkisinde ortaya çıkan yetki sınırı sorunları gelmektedir. Eleştirilere göre, federal hükümetin 2023 tarihli Güvenli, Emniyetli ve Güvenilir Yapay Zekâ Konulu Başkanlık Kararnamesi belgesinde ortaya koyduğu ilkeler, ulusal çapta yapay zekâ yönetişimi için temel koordinasyon yetkisini federal devlete bırakmıştır. Kaliforniya'nın eyalet düzeyinde kendi şeffaflık ve raporlama rejimini kurması, bu yetki dağılımını fiilen bozmaktadır. Bazı hukukçular, TFAIA'nın federal ölçekteki ulusal güvenlik ve kritik altyapı riskleri üzerinde eyalet düzeyinde yetki kullanmasını, federal hukukun eyalet hukukuna üstünlüğü kuralıyla (*Supremacy Clause*) çelişebileceğini belirtmiştir.⁵¹

2. Ticari sırlar ve bilgi güvenliği

YASA'nın getirdiği zorunlu şeffaflık rejimi, kamuya açıklanması gereken verilerin niteliği açısından ticari sır korumasıyla çelişen bir durum yaratmaktadır. Özellikle *frontier* geliştiricilerin model yapısı, parametre sayısı, eğitim verisi kaynakları ve güvenlik testleri hakkında ayrıntılı bilgi sunma zorunluluğu, ticari rekabet dengesini bozabilir. Ayrıca, kamuya açıklanan belgelerde yapılan sansür (*redaction*) işlemlerinin niteliğini kamuya duyurma zorunluluğu, hassas bilgilerin dolaylı olarak ortaya çıkması riskini taşımaktadır. Bu durum hem siber güvenlik hem de fikrî mülk hakları açısından çifte risk oluşturmaktadır.

3. Uygulama maliyetleri ve yenilik üzerindeki etkiler

Özellikle “büyük ölçekli *frontier* geliştirici” statüsündeki teknoloji şirketleri için getirilen ayrıntılı raporlama, şeffaflık belgesi hazırlama ve bağımsız denetim zorunlulukları, ciddi maliyet yükleri yaratmaktadır. Bu yük, özellikle *startup* veya akademik araştırma tabanlı geliştiriciler için inovasyonun önünde yapısal bir bariyer oluşturabilir. Eleştirmenler, YASA'nın risk temelli yaklaşımının, büyük geliştiriciler lehine bir rekabet avantajı doğurabileceğini ve piyasa kapalılığı yaratabileceğini belirtmektedir.⁵²

4. Tanım ve eşik belirsizlikleri

YASA'da *frontier* modelin tanımı için belirlenen 10^{26} FLOP'luk hesaplama eşiği, teknoloji evrimine bağlı olarak hızla eskime potansiyeline sahiptir. YASA'da bu dinamik alanı sabit teknik kriterlerle düzenlemesi, gelecek modellerin kapsam dışında kalması veya orantısız yük getirmesi riskini taşır.⁵³ Ayrıca, büyük ölçekli *frontier* geliştiricinin tanımında yer alan 500 milyon dolar gelir eşiği, gelir dalgalanmalarına duyarlı bir metrik olarak öngörülmüştür ve bu durumun uygulamada tutarsızlık yaratabileceği belirtilmektedir.⁵⁴

5. Whistleblower rejimi ve kurumsal uyum riski

YASA'nın getirdiği kapsam geniş *whistleblower* koruması, çalışan hakları açısından ilerici bir düzenleme olmakla birlikte, şirketlerin iç denetim mekanizmalarında belirsizlik yaratabilir. Bazı hukukçular, bu rejimin suistimale açık olduğunu, özellikle anonim bildirimlerin kötü niyetli kullanımının hem yönetim hem de kurumsal itibar açısından risk doğurabileceğini belirtmektedir.⁵⁵

6. Kamu altyapısının etkinlik sorunu: CalCompute

CalCompute projesinin, kamu eliyle bulut altyapısı kurma hedefi yenilikçi olsa da eleştiriler kamu yönetiminde bürokratik hantalılık ve maliyet etkinliği sorunlarına dikkat çekmektedir [(MIT CSAIL, 2025)]. Kamu sektörü içinde teknik uzmanlık eksikliği, veri güvenliği riskleri ve özel sektörle rekabet gerginlikleri *CalCompute*'un etkinliğini sınırlayabilir. Ayrıca, kamu fonlarının bu alana yönlendirilmesi, eleştirmenlerce kamu yararı önceliği bakımından sorgulanmaktadır.⁵⁶

⁵¹ Von Arx, “California's Approach to AI Governance”.

⁵² Von Arx, “California's Approach to AI Governance”.

⁵³ Bullock vd., “Frontier Model Definitions”.

⁵⁴ Von Arx, “California's Approach to AI Governance”.

⁵⁵ Future of Privacy Forum, *The State of State AI 2025* (Washington, DC: FPF, Ekim 2025), 11.

⁵⁶ Michael Rubin vd., “Latham & Watkins Discusses California's Role as Lead U.S. Regulator of AI”, *CLS Blue Sky Blog* (Erişim 13 Kasım 2025).

Özetle, TFAIA, yapay zekânın toplumsal risklerini öngörülebilir hale getiren ve şeffaflığı merkezine alan öncü bir model olmakla birlikte, anayasal yetki sınırları, ticari sır koruması, mali yük ve teknik uyarlanabilirlik bakımından gözden geçirilmesi gereken bir dizi yapısal zafiyet taşımaktadır. Bu eleştiriler, Kaliforniya'nın AI yönetiminde lider konumunu sürdürebilmesi için, düzenlemeyi daha esnek, dengeli ve çok paydaşlı bir normatif zeminde yeniden tasarlaması gerektiğini ortaya koymaktadır.

III. Frontier Yapay Zekâ Yönetişiminde Aktarılabilirlik: Farklı Hukuk Sistemleri İçin Politika Önerileri

YASA, yapay zekâ ekosisteminin hızlı evrimini karşılayacak esneklikte ve şeffaflık merkezli bir yönetim modeli sunarak, gelecekte bu alanda düzenleme geliştirmek isteyen hukuk sistemleri için önemli bir referans noktası oluşturur. Bu modelin aktarılabilirliği, katı bir kopyalamadan ziyade, temel normatif ilkelerin farklı kurumsal bağlamlara uyarlanabilirliğinde ortaya çıkar. Özellikle şeffaflık, felaket riski yönetimi, çok katmanlı yükümlülük sistemi, *whistleblower* koruması ve kamu altyapısı inşası gibi unsurlar, teknik kapasitesi ve ekonomik imkânları farklı ülkelerde dahi uygulanabilir niteliktedir.

YASA'nın en belirgin katkısı, yapay zekâ yönetişimini ağır sertifikasyon süreçlerine dayalı bir önleyici denetimden ziyade, "şeffaflık yoluyla düzenleme" ilkesi üzerine inşa etmesidir. Bu yaklaşım, düzenleyici kurumların henüz yeterli teknik birikime sahip olmadığı durumlarda dahi risk yönetimine dair bilgi akışını güçlendirir ve politika yapımcıların düzenleyici öğrenme kapasitesini artırır. Bu nedenle, düzenleme yapmak isteyen ülkelerin ilk aşamada, geliştiricilerin model mimarisi, risk değerlendirmesi, kritik olay bildirimleri ve güvenlik önlemleri hakkında standartlaştırılmış şeffaflık raporları yayımlamasını zorunlu kılan bir çerçeve oluşturması, hukuki öngörülebilirliği güçlendiren düşük maliyetli bir başlangıç adımındır.

YASA'nın *frontier* kategori inşasına yönelik yaklaşımı da önemli bir aktarılabilirlik değeri taşır. Sabit bir teknik eşik yerine, modelin tüm yaşam döngüsündeki toplam hesaplama gücünü ve adaptif kapasitesini dikkate alan bu dinamik sınıflandırma, hızla değişen teknolojik ortamla uyumlu hukuk metinlerinin hazırlanması için örnek teşkil eder. Düzenleme yapmak isteyen ülkeler, kendi teknik ölçüm standartlarını belirleyip bu kriterleri düzenli aralıklarla güncelleyerek hem inovasyonu boğmayan hem de risk düzeyini doğru yansıtan bir "hareketli sınır" sistemi kurabilir.

Kaliforniya'nın felaket riski kavramsallaştırması da ülkelerin üzerinde düşünmesi gereken bir diğer yapısal tasarım unsurudur. Kritik altyapıya yönelik saldırı, kitlesel zarar ihtimali ve kontrol kaybı gibi spesifik ve ölçülebilir risklerin ayrı bir kategori altında düzenlenmesi, hem hukuki belirliliği artırmakta hem de bu riskler için ayrı bir raporlama ve gözetim rejimi kurulmasını mümkün kılmaktadır. Çoğu ülkede kaynak ve uzmanlık kısıtlı olduğundan, felaket riski gibi yüksek etkili ancak nadir senaryoların ayrı bir hukuki rejimle ele alınması, düzenleme maliyetini yönetilebilir kılar.

Bunun yanında YASA'nın çok katmanlı yükümlülük sistemi, büyük ölçekli *frontier* geliştiricilerini ekonomik güçleri ve teknik kapasiteleri doğrultusunda farklı yükümlülöklere tabi tutarak hem inovasyonu destekler hem de piyasadaki güç asimetrisini dengeler. Bu yaklaşım, *startup* ekosisteminin güçlü olduğu veya akademik araştırmaların ön planda bulunduğu ülkelerde inovasyonun aşırı düzenleme nedeniyle baskılanmasını önleyen pratik bir mekanizma işlevi görür.

YASA'nın en yenilikçi yönlerinden biri olan *whistleblower* koruması ise düzenleme yapmak isteyen ülkeler açısından kritik bir öğrenme alanı sunar. *Frontier* yapay zekâ gibi karmaşık ve yüksek riskli alanlarda en değerli risk sinyalleri çoğu zaman kurum içi çalışanlardan gelir. Bu nedenle çalışanların misilleme korkusu olmadan risk bildiriminde bulunmasını sağlayan güçlü bir hukuki çerçeve, herhangi bir yapay zekâ düzenlemesinin güvenlik ve hesap verebilirlik ayağını tamamlayan vazgeçilmez bir unsurdur.

Son öneri ise Kaliforniya'nın *CalCompute* aracılığıyla kamu altyapısını güçlendirmesi, düzenleme yapmayı düşünen ülkeler için stratejik bir ders niteliğindedir. Yapay zekâ ekosisteminin yalnızca özel sektöre bırakılmasının uzun vadede hem güvenlik hem eşitlik açısından kırılgan sonuçlar yaratabileceği dikkate alındığında, kamuya ait veya kamu-üniversite iş birliğiyle işletilen yapay zekâ araştırma ve bulut altyapılarının kurulması, düzenleyici devletin yalnızca denetleyen değil, kapasite üreten bir aktör olarak konumlanmasını sağlar.

Sonuç

YASA, yapay zekâ yönetiminde şeffaflık, risk temelli sınıflandırma ve kurumsal hesap verebilirliği merkeze alan özgün bir model ortaya koymuştur. Bu model, *frontier* yapay zekânın doğurduğu yüksek etkili risklerin yalnızca teknik araçlarla değil, düzenleyici öğrenme ve bilgi paylaşımıyla yönetilebileceğini göstermektedir. Bu nedenle Kaliforniya yaklaşımı, küresel tartışmalarda referans niteliği taşıyan normatif bir deney alanı hâline gelmiştir.

YASA, ağır sertifikasyon süreçlerinden ziyade raporlama, denetim ve iç yönetim yükümlülükleri üzerinden şekillenen bir şeffaflık rejimi kurarak düzenleyici kurumların teknik kapasite eksikliklerini dengeleyen pratik bir çözüm sunar. Bu yapı, inovasyonu doğrudan baskılamadan riskleri görünür kılan ve politika yapıcılarının müdahale alanını genişleten esnek bir mekanizma yaratır. Özellikle *frontier* kategorisinin dinamik biçimde tanımlanması, düzenlemenin teknolojik evrime uyarlabilirliğini güçlendiren bir tasarım tercihidir.

YASA'nın çok katmanlı yükümlülük yapısı ve *whistleblower* koruması, hem büyük ölçekli *frontier* geliştiricilerin sorumluluğunu artırmakta hem de kurum içi bilgi akışını güvence altına alarak risk yönetiminin bütünlüğünü sağlamaktadır. Bu unsurlar, düzenlemenin hem piyasa asimetrisini dengeleyen hem de kurumsal şeffaflığı kurumsallaştıran yönlerini belirginleştirir. *CalCompute* modeli ise kamusal kapasite inşasının yalnızca yardımcı değil, tamamlayıcı bir politika aracı olduğunu ortaya koyar.

Bitirirken YASA, teknik ayrıntıların ötesinde, yapay zekâ yönetişimini şeffaflık ve hesap verebilirlik ekseninde yeniden konumlandıran bütüncül bir yaklaşım sunmaktadır. YASA'nın normatif çerçevesi, farklı hukuk sistemlerinin kendi düzenleme tasarımlarını geliştirirken yararlanabileceği ilkeleri somut bir biçimde ortaya koyar. Böylece Kaliforniya modeli, yapay zekâ çağında güvenlik, inovasyon ve toplumsal çıkar arasındaki dengeyi kurmaya yönelik geleceğe dönük düzenlemeler için önemli bir örnek teşkil etmektedir.

Makale Bilgi Formu

Çıkar Çatışması Bildirimi: Yazar tarafından potansiyel çıkar çatışması bildirilmemiştir.

Yapay Zeka Bildirimi: Bu makale yazılırken hiçbir yapay zeka aracı kullanılmamıştır.

İntihal Beyanı: Bu makale iThenticate tarafından taranmıştır.

Kaynakça

- Anderljung, Markus vd. "Frontier AI Regulation: Managing Emerging Risks to Public Safety". *arXiv*. Erişim 10 Kasım 2025. /mnt/data/Frontier AI Regulation.pdf
- Bommasani, Rishi vd. *The California Report on Frontier AI Policy*. The Joint California Policy Working Group on AI Frontier Models, 2025. <https://www.gov.ca.gov/wp-content/uploads/2025/06/June-17-2025-%E2%80%93-The-California-Report-on-Frontier-AI-Policy.pdf>,
- Bullock, Charlie vd. "Frontier Model Definitions", *Law-AI* (Erişim 18 Kasım 2025).
- California Consumer Privacy Act (CCPA). California Attorney General (Erişim 19 Kasım 2025). <https://oag.ca.gov/privacy/ccpa>
- European Parliament and Council. "Regulation (EU) 2024/1689 of 13 June 2024 Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act)". *Official Journal of the European Union L 202/1* (12 Temmuz 2024). Erişim 19 Kasım 2025. <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A32024R1689>
- Field, Hayden. "SB 53, the landmark AI transparency bill, is now law in California" *The Verge*, 29 Eylül 2025. <https://www.theverge.com/ai-artificial-intelligence/787918/sb-53-the-landmark-ai-transparency-bill-is-now-law-in-california>
- Future of Privacy Forum. *The State of State AI 2025*. Washington, DC: FPF, Ekim 2025. Erişim 10 Kasım 2025. <https://fpf.org/wp-content/uploads/2025/10/The-State-of-State-AI-2025.pdf>
- Gluck, Justine. "California's SB 53: The First Frontier AI Law, Explained". *Future of Privacy Forum*. Erişim 10 Kasım 2025. <https://fpf.org/blog/californias-sb-53-the-first-frontier-ai-law-explained/>
- Kak, Amba. "Executive Summary: AI Nationalism(s) – Global Industrial Policy Approaches to AI". *AI Now Institute*. Erişim 27 Şubat 2025. <https://www.ainowinstitute.org/publications/ai-nationalisms-executive-summary>
- Kulular, Merve Ayşegül. "Sağlık Hizmetlerinin Uzaktan Sunumu Kapsamında Online Platformların Hukuken Değerlendirilmesi". *Hacettepe Hukuk Fakültesi Dergisi*, 15/2, 2025, 129–168.
- Kumayama, Ken D. vd. "AI Agents: Greater Capabilities and Enhanced Risks". *Reuters Legal News*. Erişim 12 Kasım 2025. <https://www.reuters.com/legal/legalindustry/ai-agents-greater-capabilities-enhanced-risks-2025-04-22/>

- Rubin, Michael vd. "Latham & Watkins Discusses California's Role as Lead U.S. Regulator of AI". *CLS Blue Sky Blog*. Erişim 27 Şubat 2025. <https://clsbluesky.law.columbia.edu/2025/10/24/latham-watkins-discusses-californias-role-as-lead-u-s-regulator-of-ai/>
- Maas, Matthijs M. "Scaling Law for AI: Issues, Necessity, and Feasibility." *Architectures of Global AI Governance: From Technological Change to Human Choice*, edited by Florian Egloff - Thomas Reinhold. 112-180. Oxford: Oxford University Press, 2025. <https://doi.org/10.1093/9780191988455.003.0004>
- National Institute of Standards and Technology. "Artificial Intelligence Risk Management Framework (AI RMF 1.0)". *NIST*. Erişim 17 Kasım 2025. <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf>
- Stanford Institute for Human-Centered Artificial Intelligence. "2024 AI Index Report". Stanford HAI Erişim 17 Kasım 2025. <https://hai.stanford.edu/ai-index/2024-ai-index-report>
- SB-53, California Senate Bill 53 (2025 Legislative Session). California Legislature. Erişim 11 Ekim 2025. <https://legiscan.com/CA/text/SB53/2025>
- Von Arx, Devin. "California's Approach to AI Governance". *CSET*. Erişim 22 Ekim 2025. <https://cset.georgetown.edu/article/californias-approach-to-ai-governance/>
- Toner, Helen - Fist, Timothy. "Regulating the AI Frontier: Design Choices and Constraints". *CSET* (Erişim 11 Kasım 2025). <https://cset.georgetown.edu/article/regulating-the-ai-frontier-design-choices-and-constraints/>
- Yang, Ariana. "Federal Proposals to Regulate Artificial Intelligence (AI)". *Congressional Research Service*. Erişim 11 Kasım 2025. <https://www.congress.gov/crs-product/R47843>