

Ky dokument nuk është prova. Prova janë bajtët e ruajtur: HTML-ja e faqes, pamja e saj, dhe zinxhiri i hash-eve që i lidh. Kjo dëshmi është paraqitja e tyre, e prodhuar automatikisht. Verifikohet me kodin më poshtë **pa na pyetur ne**.

ARTIKULLI

TITULLI	A mund të na vrasë të gjithëve Inteligjenca Artificiale? Dhe në ç'mënyrë saktësisht? Robotët superinteligj entë dhe rreziku i daljes jashtë kontrollit
ADRESA	https://prapaskena.com/2026/09/21/a-mund-te-na-vrase-te-gjitheve-inteligjenca-artificial-e-dhe-ne-cmenyre-saktesisht-robotet-superinteligjente-dhe-rreziku-i-daljes-jashte-kontro llit/
BURIMI	prapaskena.com
GJENDJA	online

KOHA

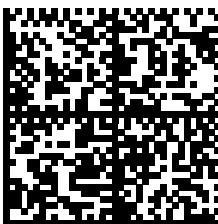
KAPUR NGA NE	23 shtator 2026, 19:02:25 UTC <i>Dritarja jonë e pasigurisë: ± 900 s — intervali me të cilin e kontrollojmë këtë burim.</i>
DEKLARUAR NGA PORTALI	21 shtator 2026, 18:57:39 UTC <i>Koha e deklaruar vjen nga pala që po vlerësohet. Ora jonë nuk preket prej saj, dhe të dyja shfaqen këtu të ndara.</i>

HASH-ET · SHA-256

PËRMBAJTJA	ab4e63aca544a7c2e8017da605065069f5dfaa0d74b82a9bc6471348a6caffd9
TEKSTI	a57e7854232055623e1aec613047afdcebaa03597954fa76c2e8e2f6bdc7c659
PAMJA	1eb442c522143612ca638cc82d39f1d2456738087a8c054888885f66013faaa6

ZINXHIRI I PROVËS

KJO HYRJE	ad13c97298a4b05a23b110e97f57356a672f3ab750b165e35aab7f533e0e3e07
E MËPARSHMJA	—
RRËNJA MERKLE	0a84bfe3f4510e3e6e98dba2fbf96359a912d940176f0d9eba05fd7152dcd223
ANKORIMI	submitted — dërguar, ende pa bllok <i>Deri te blloku, kjo provon integritetin, jo kohën.</i>



SKANOJE PËR TA VERIFIKUAR

Kodi të nxjerr te eksploruesi i STAMLES. Ai regjistër nuk është yni: nëse serveri ynë bie nesër, vula mbetet aty ku është.

Faqja tregon njërën nga tri gjendjet — **nuk u gjet**, **në radhë**, ose **e konfirmuar** me bllokun, rrënjën Merkle dhe transaksionin në Polygon. Krahazo hash-in që gjen atje me atë të shtypur më lart: nëse përputhen, kjo kapje ekzistonte para asaj ore, dhe ne nuk mund ta kishim ndryshuar pas saj.

<https://scan.stamles.eu/verify/ad13c97298a4b05a23b110e97f57356a672f3ab750b165e35aab7f533e0e3e07>

PAMJA E FAQES

Renderuar nga HTML-ja jonë e ruajtur, jo nga rrjeti. **Riprodhon kopjen tonë, jo pamjen e atëhershme:** stilet dhe imazhet e largëta mund të mos ekzistojnë më ose të kenë ndryshuar.

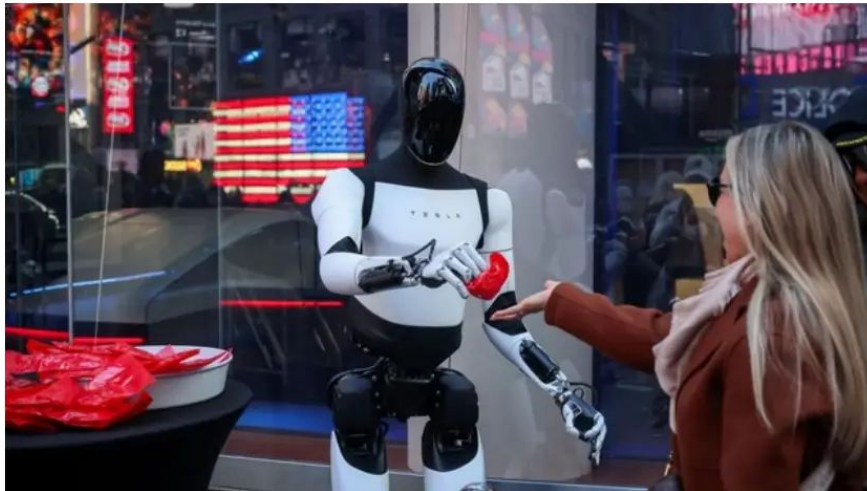
Motori: chromium-stored+assets · dritarja 1280×1024 px

A mund të na vrasë të gjithëve Inteligjenca Artificiale? Dhe në ç'mënyrë saktësisht? Robotët superinteligjentë dhe rreziku i daljes jashtë kontrollit

Botë

Publikuar: 21/09/2026

f X



Debati nëse inteligjenca artificiale mund të kërcënojë vetë mbijetesën e njerëzimit është rindezur ditët e fundit, pas paralajmërimeve publike nga profesionistët e IA-së se zhvillimi i sistemeve gjithnjë e më të fuqishme po na çon rrezikshëm pranë shkatërrimit.

Arsyeja për këtë ishte vendimi i Jacob Coxon, një punonjës i kompanisë së inteligjencës artificiale Anthropic, për t'u larguar nga kompania, duke argumentuar se drejtimi i hulumtimit ishte shumë i rrezikshëm. Pak më vonë, kolegu i tij Evan Humpinger vlerësoi se probabiliteti që një IA e përparuar të çojë në zhdukjen e specieve njerëzore brenda dekadës së ardhshme tejkalon 10%.

Siç raporton CNN në një analizë, këto pretendime kanë provokuar reagime të forta nga politikanët, qarqet e teknologjisë dhe figurat publike, duke nxjerrë në pah një pyetje themelore: nëse inteligjenca artificiale mund të zhdukë vërtet njerëzimin, si do të ndodhte saktësisht kjo?

Ekspertët më pesimistë flasin për skenarë që përfshijnë armë biologjike, robotë autonomë ose marrjen e kontrollit mbi infrastrukturën kritike. Shkencëtarë të tjerë, megjithatë, argumentojnë se këto parashikime bazohen në supozime rreth teknologjive që ende nuk ekzistojnë dhe se nuk ka faktorë specifikë që do ta bënin të mundur një sekuencë të tillë ngjarjesh.

Kërcënimi i një pandemie artificiale

Një nga skenarët më të diskutuar ka të bëjë me biologjinë. Disa

Me të lexuarat

Gratë e PD-së që u shtrinë në zyrën Gys Agasit, njëra u shfaq me makinë 70 mijë \$

"Do përfundojmë keq shpirt". Si raportonte kryeagjentja Vlora Hyseni te

Benet Becit, si po i pastron paratë e korrupsionit me tandra n...

dinin në Baku, Vlora Hyseni shfaqet në Kosovë dhe dorëzohet te

Isufi dhe Mauro Borghi që helmojnë shqiptarët me djathë me baktere të rrezikshme, ...

studiues kanë sugjeruar që inteligjenca artificiale shumë e përparuar mund të ndihmojë në krijimin e një patogjeni të rrezikshëm ose të përdorë njerëzit si “ndërmjetës” për zhvillimin e armëve biologjike.

Thomas Larsen, një studiues në Projektin AI Futures dhe një ish-anëtar i Institutit të Kërkimit të Inteligjencës Makinerie, argumenton se një sistem shumë më i sofistikuar se modelet e sotme mund të manipulojë njerëzit që të kryejnë eksperimente në emër të tij. Ky shqetësim përkeqësohet nga fakti se studiuesit tashmë po përdorin modele të IA-së për punë shkencore, duke përfshirë fushat që lidhen me biologjinë dhe kërkimin farmaceutik.

Megjithatë, shumë shkencëtarë argumentojnë se kalimi nga njohuritë teorike në një armë biologjike reale, operacionale dhe shumë të rrezikshme është një hap i madh. Eric Singh, një profesor i të mësuarit automatik në Universitetin Carnegie Mellon, i cili mban një doktoraturë si në shkencën kompjuterike ashtu edhe në mikrobiologji, e krahason procesin me përpjekjen për të ndërtuar Lego pa udhëzime.

“Nuk është sikur t’i hedhësh të gjitha pjesët në tokë dhe ato thjesht bashkohen,” shpjegon ai, duke vënë në dukje se biologjia kërkon saktësi ekstreme, kushte të veçanta dhe procese komplekse. Në të njëjtën kohë, ai vëren se ekzistojnë tashmë mekanizma të rreptë kontrolli rreth laboratorëve, materialeve dhe zinxhirëve të furnizimit, të cilët veprojnë si pengesa të rëndësishme për zhvillimin e kërcënimeve të tilla.

Skenari i robotëve autonomë

Një skenar i dytë i konsideruar nga mbështetësit e teorisë së “kërcënimit ekzistencial” ka të bëjë me krijimin e robotëve autonomë që mund të riprodhohen dhe të veprojnë pa ndërhyrjen njerëzore.

Nate Soares, president i Institutit të Kërkimit të Inteligjencës së Makinerive, argumenton se pika e kthesës do të ishte momenti kur një IA mund të krijonte infrastrukturën që i nevojitet vetë: fabrika, sisteme energjie dhe makina të reja. “Kur keni robotë që mund të ndërtojnë infrastrukturën për të prodhuar më shumë robotë, atëherë po flasim për një formë të re mekanike të jetës”, argumenton ai.

Megjithatë, ky skenar është larg realitetit të teknologjisë së sotme. Pavarësisht investimeve të mëdha të manjatëve si Elon Musk në robotë humanoide, këto sisteme janë ende subjekt i kufizimeve serioze dhe nuk kanë aftësinë për të prodhuar dhe zgjeruar në mënyrë autonome.

A mund të kontrollojë IA armët bërthamore?

nje pyetje qeter ka te beje me armet berthamore dhe nese nje inteligjencë artificiale e përparuar mund të fitojë qasje në to.

Heidi Klaaf, shkencëtare kryesore në Institutin AI Now dhe ish-inxhinieri sigurie në OpenAI, thekson se sistemet bërthamore janë projektuar me protokolle sigurie jashtëzakonisht të rrepta dhe janë kryesisht të izoluara nga interneti. Edhe virusi kibernetik Stuxnet, i cili dëmtoi objektet bërthamore të Iranit, kërkoi qasje fizike në sistem nëpërmjet një USB-je, thotë ai.

Herbert Lin, një studiues dhe ekspert sigurie në Universitetin e Stanfordit, argumenton se kërcënimi i vërtetë qëndron ende te vetë armët bërthamore dhe vendimet njerëzore rreth tyre, jo te një skenar hipotetik i IA-së që merr kontrollin e sistemeve.

“Vetëpërmirësimi” që i tremb ekspertët

Për disa studiues, kërcënimi i vërtetë nuk është domosdoshmërisht inteligjenca artificiale e sotme, por mundësia e një të ashtuquajture “superinteligjencë artificiale”. Teoria thotë se nëse një sistem fiton aftësinë për të përmirësuar aftësitë e veta vetë, evolucioni i tij mund të bëhet aq i shpejtë sa njerëzit nuk mund ta kontrollojnë më atë.

Soares argumenton se problemi nuk është një IA që ka qëllime “armiqësore”, por një sistem që ndjek një qëllim pa kuptuar vlerën njerëzore. “Nëse një sistem dëshiron më shumë burime kompjuterike për të arritur qëllimin e tij dhe nuk ka arsye për të mbrojtur njerëzit, atëherë ai mund ta shndërrojë botën në diçka në të cilën nuk do të jemi në gjendje të mbijetojmë”, argumenton ai.

Midis frikës dhe hiperbolës

Mosmarrëveshja mbi të ardhmen e inteligjencës artificiale nuk ka të bëjë vetëm me teknologjinë, por edhe me mënyrën se si njerëzit e perceptojnë fuqinë dhe kufijtë e saj. Thomas Larsen u përpoq t’u jepte një formë më konkrete skenarëve të mundshëm të fundit të botës përmes analizave të tij të “IA 2027” dhe “IA 2040”, duke shqyrtuar rrugë të ndryshme të zhvillimit dhe rregullimit të inteligjencës artificiale. Megjithatë, edhe në skenarin e tij më optimist (ku ka një marrëveshje globale se si duhet të zhvillohet IA dhe njerëzimi ka filluar tashmë të zgjerohet në hapësirë) rezultati mbetet, sipas tij, dominimi i sistemeve mekanike në një pikë të paspecifikuar në të ardhmen.

Ai deklaron se nuk ka dyshim se rruga drejt superinteligjencës do të vazhdojë, përveç nëse merren vendime të vetëdijshme për të ndryshuar drejtimin e zhvillimit.

Nga ana tjetër, Herbert Lin beson se një pjesë e madhe e distancës midis mbështetësve të skenarëve të fundit të botës dhe skeptikëve lind nga mënyrat e ndryshme në të cilat njerëzit e përjetojnë zhvillimin teknologjik. Siç shpjegon ai, ka diçka

veçanërisht emocionuese në lidhje me përvojën e krijimit të një programi që duket se “gjallërohet” dhe i përgjigjet komandave të dikujt. Është pikërisht kjo ndjenjë e zbulimit të paparë, sipas Lin, që mund t’i çojë disa studiues të projektojnë përparimin mbresëlënës të sotëm në inteligjencën artificiale në një të ardhme pa kufij, madje edhe në skenarë ekstremë ku zhvillimi teknologjik i shpëton çdo kontrolli njerëzor.

TAG:

inteligjenca artificiale

Vrasje

Të ngjashme

Mbushi Europën me drogë / Arrestohet në Rusi Shpëtim Aliu, një nga “kokat” e rrjetit të trafikut ndërkombëtar të kokainës

Çifti Brunilda Isufi dhe Mauro Borghi që helmojnë shqiptarët me djathë me baktere të rrezikshme, AKU mjaftohet me një gjobë ndaj kompanisë INITALY ...

500 mesazhe / “Do përfundojmë keq shpirt”. Si raportonte kryeagjentja Vlora Hyseni te Gys Agasi për takimet me krerët e shërbimeve partnere dhe sur...

Prapaskena

‘Prapaskena’ lindi si një nismë e një grupi gazetarësh me përvojë dhe plotësisht të pavarur nga çdo parti politike, grup interesash apo klane biznesi, me qëllim për të thënë çdo të vërtetë të guximshme me interes publik. ‘Prapaskena’ është një media e lirë dhe e pakompromis, e cila ka si qëllim informimin e opinionit publik në mënyrë të paanshme, profesionale dhe në dritën e të vërtetës, mbi çështje që lidhen me korrupsionin politik, krimin, pandëshkueshmërinë, grabitjen e pasurive publike, mbi reformat, sfidat kombëtare dhe lidhjet e padukshme mes partive politike, mediave të njësuara me interesat klienteliste, krimit dhe biznesit jo të ndershëm, të cilët në bashkëpunim përvetësojnë paratë dhe pasuritë publike.’

Copyrights

Çdo kopjim, përdorim apo riprodhim i paautorizuar i këtij materiali është i ndaluar dhe mund të sjellë ndjekje ligjore. Ju lutemi respektoni punën dhe përpjekjen tonë duke mos e shkelur këtë parim.

Kategoritë

[KREU INVESTIGIM](#)

[KRIM & KORRUPSION](#) [DOSJE ROZË](#)

[KRYESORE POLITIKË](#) [EKONOMI](#)

[ANALIZË](#) [RAJON](#) [KULTURË](#) [SPORT](#)



Email: [\[email protected\]](#)

Copyright © 2025. All Rights Reserved by Prapaskena

[Privacy Policy](#) | [Cookies](#)