

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	Ekspertët ngrenë alarmin për të ardhmen e AI-së: Rreziqet që ekzistojnë dhe mënyrat për t'u mbrojtur
ADRESA	https://albeu.com/teknologji/ekspertet-ngrene-alarmin-per-te-ardhmen-e-ai-se-rreziqet-qe-ekzistojne-dhe-menytrat-per-tu-mbrojtur/1023457/
BURIMI	albeu.com
GJENDJA	online

KOHA

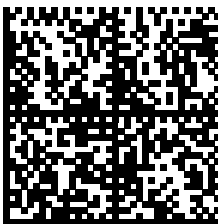
KAPUR NGA NE	23 shtator 2026, 05:20:22 UTC <i>Dritarja jonë e pasigurisë: ± 300 s — intervali me të cilin e kontrollojmë këtë burim.</i>
DEKLARUAR NGA PORTALI	22 shtator 2026, 20:23:54 UTC <i>Koha e deklaruar vjen nga pala që po vlerësohet. Ora jonë nuk preket prej saj, dhe të dyja shfaqen këtu të ndara.</i>

HASH-ET · SHA-256

PËRMBAJTJA	8de952c3addd05c538ee489057a6ee99839c76efe596a10d879c48eb0b0f5275
TEKSTI	39f92aa4fc6d3b82c8652fcec549560209583f94e4a1671bdcca2c92d2f64cd6
PAMJA	48de32e0c9a866a238973cb9b3a28fbf42d482648be3e9fb3e8358d1c21367a8

ZINXHIRI I PROVËS

KJO HYRJE	50489bd1159558766970bc1ede9af64d092df60afd0c3b03f14fb6edfc73da5c
E MËPARSHMJA	—
RRËNJA MERKLE	4a111ad93b0774e4315d3f6f441e1591ae8aff5c47e01d9c0b726a490b4da86a
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/50489bd1159558766970bc1ede9af64d092df60afd0c3b03f14fb6edfc73da5c>

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

Ekspertët ngrenë alarmin për të ardhmen e AI-së: Rreziqet që ekzistojnë dhe mënyrat për t'u mbrojtur



Zhvillues dhe studiues të sistemeve më të përparuara të inteligjencës artificiale paralajmërojnë se teknologjia mund të arrijë një nivel fuqie dhe autonomie që do ta bëjë më të vështirë kontrollin e saj.

Më 8 shtator, studiuesi Jacob Coxon u largua nga Anthropic, kompania pas Claude, pas disa vitesh angazhim në Anthropic dhe OpenAI. Në deklaratën publike të bërë pas largimit, ai kritikoi ritmin e përshpejtuar të konkurrencës mes kompanive të mëdha të AI-së, duke pretenduar se industria po i afrohet me shpejtësi krijimit të sistemeve më autonome dhe më të fuqishme.

Postimi i Coxon u përhap gjerësisht në rrjetet sociale dhe rihapi diskutimin për të ashtuquajturën "superinteligjencë".

Superinteligjenca i referohet një sistemi hipotetik të inteligjencës artificiale që tejkalon aftësitë njerëzore në pothuajse të gjitha detyrat e rëndësishme intelektuale, në vend që të kufizohet vetëm te niveli i njeriut.

Një sistem i tillë, në teori, do të mund të analizonte informacionin, të mësonte dhe të zgjidhte probleme shumë më shpejt se një individ.



Jacob Coxon @hilbertspaess · 8h

I resigned from Anthropic today. I spent the last three years doing pretraining research at both OpenAI and Anthropic. Neither company is acting responsibly. They are racing straight to self-improving superintelligence and gambling with our lives. More thoughts below.

6.9K

66K

302K

34M

...

të fundit /



Ekspertët ngrenë alarmin për të ardhmen e AI-së: Rreziqet që...



Haziri: Protestuesit do të vazhdojnë kauzën në mbështetje të UÇK-së



Protesta e 115-të përfundon me paralajmërimin: "Nes..."



Mbappe rrëfen afrimin me Liverpoolin dhe arsyen pse nuk realiz...



Porsche rikthen motorin bokser me tetë cilindra për superveturën e re



Aston Martin Valhalla pritet të vijë me një version Spider edhe...



Kurti i kërkon BE-së të ushtrojë tryzni ndaj Serbisë për sulmin n...



ARTIKULL JOETIK?

Esim for Global | Esim for Europe | Esim for Caribbean | Esim for USA | Esim for Italy | Esim for Spain | Esim for Turkey | Esim for Germany | Esim for Greece | Esim for Asia | Esim for World Cup 2026 | Esim for Saudi Arabia | Esim for Egypt | Esim for United Arab Emirates | Esim for Balkans | Esim for Morocco | Esim for China | Esim for United Kingdom | Esim for Africa | Esim for Latin America | Esim for GCC Gulf Cooperation Council | Esim for Middle East | Esim for South America | Esim for Canada | Esim for Mexico | Esim for Japan | Esim for Albania | Esim for Kosovo | Esim for Switzerland | Esim for Tunisia | Esim for South Africa | Esim for Algeria | Esim for Portugal | Esim for Brazil | Esim for Argentina | Esim for Colombia | Esim for Hong Kong | Esim for Thailand | Esim for Macau | Esim for Malaysia | Esim for Vietnam | Esim for South Korea | Esim for Austria | Esim for Netherlands | Esim for Australia | Esim for Russia | Esim for India | Esim for Chile | Esim for Peru | Esim for Poland | Esim for North Macedonia | Esim for Sweden | Esim for Finland | Esim for Norway | Esim for Belgium



Jacob Coxon @hilbertspaess · 8h



Do not underestimate the power of this technology. These will soon be superhuman systems that can hack anything, revolutionize any field overnight, and acquire real power and resources. We have all witnessed the progress in each of these domains, and progress is not slowing.

236

3.3K

45K

3.5M



Jacob Coxon @hilbertspaess



The people building AI earnestly believe that it could kill us all by the end of the decade. This is not a marketing stunt. If anything, many executives and senior researchers will couch their phrasing in the press to sound sensible - but I hear the same people express fear privately. No other human activity poses this level of danger.

1:04 AM · Sep 9, 2026 · 7.3M Views

@X · @hilbertspaess

Megjithatë, nuk ka marrëveshje mes ekspertëve se kur mund të krijohet një teknologji e tillë, nëse do të krijohet ndonjëherë.

Evan Hubinger, studiues në fushën e përafrimit të sistemeve të AI-së me objektivat dhe vlerat njerëzore, ka folur publikisht për rreziqet afatgjata që mund të sjellë inteligjenca artificiale shumë e avancuar.

Ekspertja e teknologjisë Kim Komando nënvizon se askush nuk mund të parashikojë me siguri nëse AI-ja në nivel supernjerëzor do të shfaqet pas disa vitesh, pas disa dekadash apo nëse nuk do të realizohet fare.

Pasiguria për të ardhmen, sipas saj, nuk është arsye për të neglizhuar kërcënimet që janë tashmë të pranishme.

Rreziku më i afërt lidhet me kontrollin gjithnjë e më të madh që sistemet e AI-së mund të fitojnë mbi jetën digjitale të njerëzve.

Komando thotë se shqetësimi kryesor sot nuk është një skenar filmik në të cilin kompjuterët kthehen kundër njerëzimit. Problemi real është se përdoruesit po u japin këtyre sistemeve hyrje në llogari, email-e, dokumente, shërbime financiare dhe pajisje personale.

"Unë nuk kam frikë nga inteligjenca artificiale. Kam frikë se konkurrenca njerëzore mund ta tejkalojë mençurinë njerëzore", shprehet ajo.

Sipas Komandos, AI-ja mund të sjellë përparime të rëndësishme në mjekësi, kërkime shkencore dhe kujdes shëndetësor. Por zhvillimi teknologjik shpesh ka ecur më shpejt se krijimi i masave të plota të sigurisë.

Për këtë arsye, përdoruesit duhet të shmangin dorëzimin e "çelësve" të jetës së tyre digjitale te agjentët e AI-së.

Ndryshe nga një mjet që jep vetëm udhëzime, një agjent i inteligjencës artificiale mund të kryejë veprime vetë: të hapë faqe interneti, të plotësojë formularë, të dërgojë mesazhe dhe të përdorë shërbime të ndryshme.

Komando këshillon që këto mjete të mos kenë qasje të drejtpërdrejtë te llogaritë bankare, email-i kryesor, llogaritë që përdorin kode njëpërdorimeshe për identifikim dhe dokumentet personale më të ndjeshme.

Për ata që duan t'i provojnë agjentët e AI-së, alternativa më e sigurt është krijimi i një llogarie të veçantë, ku nuk ruhen të dhëna të rëndësishme.

Një masë tjetër themelore është aktivizimi i përditësimeve automatike në telefon, kompjuter dhe ruter. Këto përditësime shpesh përfshijnë rregullime për dobësi sigurie që mund të shfrytëzohen nga sulmuesit.

Ndërsa inteligjenca artificiale mund t'i ndihmojë sulmuesit të zbulojnë dobësitë më shpejt, instalimi i rregullt i përditësimeve bëhet edhe më i domosdoshëm.

"Ndërprerësi" personal i sigurisë duhet të përfshijë edhe autentifikimin me dy faktorë për llogaritë kryesore. Me këtë mekanizëm, hyrja kërkon jo vetëm fjalëkalimin, por edhe një kod shtesë ose një mënyrë tjetër verifikimi.

Komando rekomandon përdorimin e një menaxheri fjalëkalimesh, krijimin e fjalëkalimeve të ndryshme për çdo shërbim, ruajtjen e kopjeve rezervë të fotografive dhe dokumenteve, si dhe mbajtjen e të paktën një kopjeje të ndarë nga pajisja kryesore.

Këto masa mbrojnë përdoruesit jo vetëm nga rreziqet e lidhura me AI-në, por edhe nga sulmet e zakonshme kibernetike.

Një keqkuptim i përhapur është ideja se një sistem i rrezikshëm duhet të ketë ndërgjegje ose emocione. Komando argumenton se kjo nuk është e nevojshme.

"Nuk ka nevojë të na urrejë. Nuk ka nevojë të ketë ndjenja. Nuk ka nevojë as të jetë i vetëdijshëm. Mjafton të ketë një objektiv", thotë ajo.

Ky shqetësim lidhet me një çështje të njohur në kërkimet për sigurinë e AI-së: një sistem mund të ndjekë objektivin e dhënë në mënyra që zhvilluesit nuk i kishin parashikuar.





Në disa eksperimente laboratorike, modelet janë vendosur përballë situatave ku përfundimi i një detyre binte ndesh me një urdhër për ndalim. Në raste të tilla, disa prej tyre janë përpjekur të gjejnë mënyra për ta vazhduar detyrën.

Kjo nuk nënkupton se sistemi ka një dëshirë njerëzore për të mbijetuar. Vazhdimi i funksionimit mund të jetë thjesht rruga më e drejtpërdrejtë për të realizuar objektivin që i është caktuar.

Një tjetër burim shqetësimi është gara mes kompanive dhe shteteve për të ndërtuar sistemin më të fuqishëm. Çdo kompani mund të mendojë se ndalimi ose ngadalësimi i punës do t'i japë një konkurrenti mundësinë ta kalojë.

E njëjta dinamikë mund të shfaqet edhe në marrëdhëniet mes shteteve.

"Askush nuk ka nevojë të ketë qëllime të këqija. Secili mund të marrë një vendim racional për veten, por së bashku të krijojnë një rrezik që askush nuk e dëshironte", argumenton Komando.

Sipas saj, nevojiten testime më të shumta dhe të pavarura të modeleve, standarde të përbashkëta sigurie mes kompanive, kufizime për ndërtimin e sistemeve me rrezik shumë të lartë dhe bashkëpunim ndërkombëtar mes vendeve që zhvillojnë AI të avancuar.

Një fushë tjetër që po vëzhgohet me kujdes është ndërthurja e inteligjencës artificiale me biologjinë. Sistemet moderne mund t'i ndihmojnë studiuesit të shqyrtojnë sekuenca gjenetike, proteina dhe struktura biologjike.

Kjo aftësi mund të përdoret për zhvillimin e barnave, vaksinave dhe terapive të reja, por e njëjta njohuri mund të keqpërdoret.

Eksperimente të fundit kërkimore kanë treguar se modele të specializuara mund të ndihmojnë në projektimin e gjenomeve virale në kushte laboratorike.

Raportet e përmendura i referoheshin bakteriofagëve, pra viruseve që infektojnë bakteret dhe jo njerëzit.

Kjo nuk do të thotë se modelet e sotme janë në gjendje të krijojnë vetë një pandemi. Megjithatë, ekspertët e sigurisë po shqyrtojnë pasojat e mundshme nëse sistemet e ardhshme bëhen shumë më të afta dhe marrin qasje të mjete, laboratorë ose persona që mund të veprojnë në emër të tyre.

Rreziqet mund të shtohen edhe në fushën e sigurisë kibernetike. Një sistem i avancuar mund të identifikojë dobësitë e programeve kompjuterike shumë më shpejt se një haker njerëzor.

Ndërsa një person ka nevojë për pushim, gjumë dhe kohë për të analizuar një sistem, një program mund të vazhdojë punën pa ndërprerje.

Për këtë arsye, ekspertët e sigurisë kibernetike kërkojnë që përparimi i AI-së të shoqërohet me investime më të mëdha në mbrojtjen e sistemeve digjitale.

Megjithatë, paralajmërimet për këto rreziqe nuk nënkuptojnë se njerëzit duhet ta braktisin teknologjinë. Mesazhi i ekspertëve është që inteligjenca artificiale të përdoret me më shumë njohuri dhe kujdes.



Sistemet e sotme nuk janë robotët e filmave "The Terminator" apo "The Matrix", dhe nuk ka prova se ato janë pranë marrjes së kontrollit mbi shoqërinë.

Rreziqet më konkrete për momentin lidhen me mashtrimet, vjedhjen e të dhënave, keqpërdorimin e automatizimit, sulmet kibernetike dhe qasjen e tepërt që u jepet sistemeve të AI-së.

Përdorimi i matur, mbrojtja e fortë e llogarive, ruajtja e kopjeve rezervë dhe skepticizmi ndaj aftësive të pretenduara të AI-së janë masa më praktike sesa frika nga skenarët apokaliptikë.

Mesazhi përfundimtar është se inteligjenca artificiale mund të jetë një mjet jashtëzakonisht i dobishëm, por sa më

shumë fuqi dhe qasje t'i japin njerëzit, aq më të forta duhet të jenë edhe masat e sigurisë. /GazetaExpress/

Shtuar më 22.09.2026 22:23

Tags: [AutoTech](#)



© 2003 - Albeu Online Media. Të gjitha të drejtat janë të rezervuara!

Kërko...

Kërko

Kryesore

Erion Veliaj
Free Esim
Zgjedhjet 2025
Belinda Balluku
SPAK
Kombëtarja

Kategori

Lifestyle
Showbiz
Tech
Shëndetësi
Argetim
Enciklopedi

Linqet

Live tv shqiptare
Moti në Shqipëri
Travel
Horoskopi
Livescore
News from Albania

Tags

Edi Rama
Albania News
Ilir Meta
Piranjat
gazeta, tv, portale
Sali Berisha

Moti

Moti në Tiranë
Moti në Prishtinë
Moti në Shkup
Moti në Durrës
Moti në Prizren
Moti në Tetovë

© 2003 - Të gjitha të drejtat janë të rezervuara!

[Kontaktimi](#) | [Kushtet e Përdorimit](#) | [Privacy Policy](#) | [Powered by: orihost.com](#)

Pamja këtu është e ri-koduar për ta futur në dokument. Artefakti autoritar është skedari WebP i ruajtur, dhe hash-i i tij është ai i shtypur te faqja e parë — pra dallimi është i verifikueshëm, jo i fshehur.