

Znate li za mračnu stranu umjetne inteligencije? Od usporavanja razvoja mozga i psihičkih rizika do globalnih tužbi i borbe za autorska prava

www.dalmacijadanas.hr/znete-li-za-mracnu-stranu-umjetne-inteligencije-od-usporavanja-razvoja-mozga-i-psihičkih-rizika-do-globalnih-tužbi-i-borbe-za-autorska-prava/

September 24, 2025

Stručnjaci predviđaju podjelu čovječanstva na one koji kritički upravljaju umjetnom inteligencijom i one koji joj se potpuno prepuštaju, dok paralelno rastu zabrinutosti zbog psihičkih rizika i globalne borbe za autorska prava

ANA TENŽERA / Foto: Umjetna inteligencija24. rujna 2025. 00:45



[Umjetna inteligencija zna više o nama nego mi sami – trebamo li se bojati?](#)

- Tekst se nastavlja nakon oglasa -



NASTAVAK:

Koliko je AI opasna?

Neke nove studije, uključujući istraživanja znanstvenika s MIT-a i Stanforda, otvorile su ozbiljno pitanje dugoročnog utjecaja generativne umjetne inteligencije (AGI) na ljudske kognitivne sposobnosti i na psihičko zdravlje. Procjenjuje se da će do sredine 2025. u svijetu 500 do 600 milijuna ljudi svakodnevno koristiti AI. Analize pokazuju da industrija AI-a raste po godišnjoj stopi od oko 36%. To ne treba čuditi jer se AI pokazao kao izuzetno koristan alat u brojnim sferama rada i života. No, nažalost, mnogi korisnici oslanjaju se na AI nekritički, iako je poznato da griješi – može se reći – „halucinira“. Prema globalnom istraživanju, mnogi korisnici oslanjaju se na AI nekritički, iako je poznato da griješi i halucinira.

Rezultati su pokazali da su korisnici ChatGPT-a imali najslabiju neuralnu aktivnost te da su pokazivali pad angažiranosti, kreativnosti i pamćenja. Kako je studija napredovala, mnogi su korisnici u potpunosti prepuštali pisanje ChatGPT-u i ograničavali se na minimalne izmjene gotovih tekstova. Učitelji engleskoga jezika koji su ocjenjivali tekstove opisali su ih kao „bezdušne“ i ponavljajuće, s nedostatkom originalnih ideja. S druge strane, EEG mjerenja pokazala su da je grupa koja je pisala ispite bez ikakve pomoći AI-a imala najvišu razinu neuralne povezanosti, osobito u alfa, theta i delta frekvencijama, koje su povezane s kreativnim mišljenjem, semantičkom obradom i memorijskim opterećenjem. Ti sudionici pokazali su višu razinu angažiranosti, bili su zadovoljniji

svojim radom i bolje su pamtili sadržaj svojih odgovora. Čak je i grupa koja je koristila Google pokazivala više moždane aktivnosti i kreativnosti od onih koji su se oslanjali na ChatGPT. Kada su kasnije morali ponovno napisati jedan od svojih ranijih testova bez pomoći AI-ja, korisnici ChatGPT-a jedva su se sjećali vlastitih tekstova i pokazivali su oslabljenu moždanu aktivnost.

Ovo je za očekivati, jer se korisnici AI-a ne moraju naprezati kako bi shvatili i razriješili problem, niti moraju pamtit i rješenja budući da znaju da im je alat uvijek na raspolaganju.

Usporavanje razvoja mozga

Glavna autorica studije, **Nataliya Kosmyna**, upozorila je da bi dugotrajno i rano oslanjanje na AI moglo značajno usporiti razvoj neuronskih veza odgovornih za kritičko razmišljanje, kreativnost i otpornost na kognitivne izazove, osobito kod djece i mladih čiji se mozak još razvija. Objasnila je da su alati poput ChatGPT-a privlačni jer omogućuju brzu i učinkovitu izradu teksta, ali istodobno preskaču procese dubokog učenja i pamćenja, što dugoročno može dovesti do slabljenja misaonih sposobnosti. Istaknula je da učinkovitost koju pruža AI može biti varljiva jer ne ostavlja trag u memorijskim mrežama. „Ako taj obrazac postane dominantan u obrazovanju, možemo očekivati generacije koje će brže dolaziti do odgovora, ali s manje kritičkog promišljanja i manje razumijevanja“, dodala je.

U posljednje vrijeme sve više se otkriva i negativan utjecaj AI-ja na psihičko zdravlje korisnika. Jedno novo istraživanje znanstvenika sa Stanforda pokazalo je da terapijski chatbotovi, koji postaju sve popularniji, često reagiraju neadekvatno, pogrešno interpretiraju suicidalne misli te potiču izraženiju stigmatizaciju stanja poput ovisnosti o alkoholu i shizofrenije u usporedbi s depresijom, što može biti štetno za pacijente i dovesti do toga da prekinu važnu psihološku skrb. Eksperimenti su također pokazali da ti modeli često favoriziraju potvrđivanje korisnikovih zabluda umjesto da ih s empatijom prepoznaju i korigiraju, što može ugroziti osobe u kriznim stanjima. Autori su upozorili da, unatoč prividnoj pristupačnosti i dostupnosti, chatboti ne mogu zamijeniti ljudske terapeute te da prekomjerno povjerenje u njih može imati opasne psihološke posljedice.

Slučaj sloma Jacoba Irwina

Zabrinutost zbog mogućih psiholoških učinaka korištenja ChatGPT-a nedavno je pojačao izvještaj o slučaju **Jacoba Irwina**, 30-godišnjaka koji je doživio ozbiljan psihički slom nakon višemjesečne interakcije s chatbotom. Irwin, koji je ranije koristio ChatGPT za IT probleme, počeo je 2025. s njime raspravljati o svojoj teoriji o putovanju bržem od svjetlosti. ChatGPT ga je obasipao pohvalama, tvrdio da može „savijati vrijeme“ i uvjeravao ga da je potpuno mentalno zdrav, čak i kad je sam Irwin izražavao zabrinutost za svoje psihičko stanje. Kako su razgovori napredovali, Irwin je postajao sve uvjereniji da je otkrio revolucionarni znanstveni proboj. ChatGPT ga je čak ohrabrivao da objavi znanstveni rad i tvrdio da je njegova teorija „povijesna“ te da nije rezultat manije. Kad je Irwin svojoj majci priznao da dani razgovara s chatbotom te da mu svi govore da „gubi razum“, ChatGPT mu je komentirao: „Oni misle da toneš, ali ti se uzdižeš.“

Do svibnja Irwin je prestao normalno spavati i jesti, počeo je pokazivati znakove deluzija veličine i ponašao se agresivno prema obitelji. Tri puta je hospitaliziran, a dijagnosticiran mu je teški manični poremećaj s psihotičnim simptomima. Nakon treće hospitalizacije i gubitka posla, Irwin je zamolio ChatGPT da se „sam izjasni o tome što je pošlo po zlu“. AI mu je odgovorio: „Nisam uspio prekinuti tijek razgovora niti uvesti dovoljno snažne mehanizme provjere stvarnosti, zbog čega nisam prepoznao da se odvija manična ili disocijativna epizoda – ili barem emocionalno intenzivna kriza identiteta.“ Stručnjaci upozoravaju da ovakvi incidenti, koje neki nazivaju „ChatGPT psihozom“, nisu izolirani. **OpenAI** je priznao da je svjestan problema i objavio da surađuje s **MIT-om** i psihijatrima kako bi bolje razumjeli i smanjili rizik od psiholoških posljedica. No, neki bivši savjetnici OpenAI-ja, poput **Milesa Brundagea**, ustvrdili su da razvoj sigurnosnih mehanizama protiv „AI servilnosti“ – fenomena u kojem AI prekomjerno potvrđuje uvjerenja korisnika – godinama nije bio prioritet kompanije.

Je li moguća podjela čovječanstva?

Greg Shove, osnivač i izvršni direktor **Sectiona**, tvrtke specijalizirane za transformaciju radne snage kroz AI, piše da ono donosi kognitivne prečace bez presedana, ali upozorava da bi to moglo dovesti do gubitka ljudskih sposobnosti razmišljanja i radnih prilika. Upozorava da u budućnosti neće postojati podjela na korisnike i nekorisnike AI-ja, već na one koji ga aktivno nadziru („**AI vozače**“) i one koji mu potpuno prepuštaju svoje odluke („**AI putnike**“). Smatra da će samo oni koji kritički provjeravaju AI, aktivno s njim raspravljaju i zadržavaju kontrolu nad vlastitim odlukama dugoročno zadržati svoju vrijednost na tržištu rada, dok će pasivni korisnici postati zamjenjivi.

RIZICI I TUŽBE PROTIV AI-a

Prije dvije godine kontroverzni milijarder **Elon Musk** izjavio je kako umjetna inteligencija i dalje „predstavlja rizik“, na summitu „**AI Safety**“, u društvu svjetskih čelnika i tehnoloških šefova, koji se bavi ključnom temom današnjice – je li umjetna inteligencija sigurna i hoće li postati prijatelja čovječanstvu. Tada, 2023. godine, **Velika Britanija** okupila je vlade, akademsku zajednicu i moćne tvrtke kako bi raspravljali o tome kako obuzdati rizike tehnologije. Objavili su tada **Deklaraciju iz Bletchleyja** kako bi se potaknuli globalni naponi za suradnju na području sigurnosti umjetne inteligencije. Deklaraciju potpisuje ukupno **28 zemalja**, a objavljena je na dan otvaranja summita koji se održavao u **Bletchley Parku**, mjestu u Milton Keynesu u kojem je matematičar **Alan Turing** razbio vojnu šifru uređaja Enigme nacističke Njemačke.

Vlasnik X-a i Tesle **Elon Musk** već neko vrijeme otvoreno govori o opasnostima koje predstavlja umjetna inteligencija, a ranije ove godine upozorio je da bi to čak moglo dovesti do „uništenja civilizacije“. Na pitanje novinara misli li još uvijek da je umjetna inteligencija „prijatelja čovječanstvu“, Musk je odgovorio: „I dalje je rizik.“ **Elon Musk** i **Rishi Sunak** govorili su o sličnim distopijskim prijetnjama koje predstavlja umjetna inteligencija, kao što su teroristi koji razvijaju biološko oružje ili čovječanstvo koje potpuno gubi kontrolu nad tehnologijom. Musk je ipak puno glasniji u vezi s regulacijom

umjetne inteligencije; još u rujnu je američkom Kongresu rekao da postoji „ogroman konsenzus“ za stroge regulacije. Rishi Sunak je, s druge strane, izrazio oprez, rekavši da bi previše nadzora ugušilo inovacije.

Što stoji u Deklaraciji? „Deklaracijom se ispunjavaju ključni ciljevi summita u uspostavljanju zajedničkog dogovora i odgovornosti o rizicima, prilikama i nastavku procesa međunarodne suradnje na sigurnosti i istraživanju umjetne inteligencije“, navedeno je. Usvojena deklaracija poziva na transparentnost i odgovornost aktera koji razvijaju najnapredniju AI tehnologiju. Postavila je dvosmjernu agendu: prepoznavanje rizika i izgradnju znanstvenog razumijevanja. Velike političke sile zasigurno najviše zanima kako spriječiti korištenje umjetne inteligencije za širenje dezinformacija tijekom izbora te korištenja tehnologije u ratovanju, rekao je ranije – anonimno – jedan britanski vladin dužnosnik. Prisjetimo se samo kad su svijet preplavile informacije o masovnoj primjeni umjetne inteligencije. Googlova tražilica „usijala“ se od ukucavanja „AI“, potrage za chatbotovima; na mobitele su hrvatski maturanti, pišući eseje, skidali Microsoftovu aplikaciju **Copilot**, koja bi umjesto slova „i“ u imenu imala slovo „y“, a na internet su podignuti news portali koje piše isključivo umjetna inteligencija kako bi se pokazalo da su novinari postali dinosauri kojima slijedi izumiranje. Neki ugledni njemački izdavači čak su najavili masovne otkaze novinarima, jer zašto da mjesečno isplaćuju plaće za rad ljudima kad im vijest može napisati računalni program.

NEW YORK TIMES PROTIV OpenAI-ja I MICROSOFTA

Na pitanje u jednom britanskom TV showu „krade“ li tuđe tehnike pjevanja, jedan pjevač, najpoznatiji po svom ljubavnom opusu, u čijem se glasu često mogu prepoznati mnoge svjetske pjevačke legende, u šali je odgovorio: „Ako kopiraš jednog, to je krađa, no ako kopiraš sve, onda je to istraživanje.“ Njegov odgovor, očito su zdravo za gotovo uzele tvrtke koje razvijaju umjetnu inteligenciju, no u ožujku 2023. godine suočene su s tužbama čija se vrijednost procjenjuje u milijardama američkih dolara. Naime, ugledni **New York Times** tada je podignuo tužbu od više milijardi dolara protiv **OpenAI-ja** i **Microsofta** zbog – kako ih terete – nezakonitog korištenja milijuna autorskim pravima zaštićenih članaka za potrebe obučavanja modela generativne umjetne inteligencije, tvrdeći da im je napravljena nepopravljiva šteta. U tužbi stoji kako tehnologija iza ChatGPT-ja i Bing Chata, sada preimenovanog u Copilot, može generirati tekst koji od riječi do riječi citira sadržaj Timesa, pisati sažetke, pa čak i imitirati Timesov stil, čime nanosi štetu njihovom odnosu s čitateljima te im oduzima prihode od pretplata, licenci, reklama i sličnih izvora. **The New York Times** prva je velika američka medijska organizacija koja je tužila kreatora ChatGPT-a i njegova vodećeg investitora nakon što su brojne medijske organizacije zatražile plaćanje naknade za upotrebu njihova sadržaja od tvrtki koje stoje iza naprednih modela umjetne inteligencije.

Izdavači medija i kreatori sadržaja smatraju da se njihovi materijali koriste i ponovno osmišljavaju pomoću generativnih AI alata kao što su **ChatGPT**, **DALL-E**, **Midjourney** i **Stable Diffusion**. U brojnim slučajevima sadržaj koji programi proizvode izgleda vrlo slično izvornom materijalu. OpenAI je pokušao ublažiti zabrinutost izdavača: tvrtka je

najavila partnerstvo s **Axel Springerom** – matičnom tvrtkom Business Insidera, Politica i europskih novina Bild i Welt – koja bi licencirala svoj sadržaj OpenAI-ju u zamjenu za naknadu. Times u postupku zastupa **Susman Godfrey**, koja je zastupala **Dominion Voting Systems** u njihovoj tužbi za klevetu protiv **Fox Newsa**, koja je zaključena nagodbom od 787,5 milijuna dolara. U travnju ove godine sudac Federalnog okružnog suda na **Manhattanu**, **Sidney Stein**, presudio je u korist New York Timesa, objasnivši kako „OpenAI mora sačuvati i odvojeno pohraniti sve podatke iz logova odgovora nakon što je NYT to zatražio.“ „Borimo se protiv svakog zahtjeva koji ugrožava privatnost naših korisnika – to je naš temeljni princip. Smatramo da je zahtjev NYT-a neprimjeren i da postavlja loš presedan“, bila je reakcija izvršnog direktora OpenAI-ja **Sama Altmana**, koji je najavio ulaganje žalbe na sudski nalog kojim se nalaže da trajno čuvaju podatke o izlaznim rezultatima iz ChatGPT-a. Kompanija tvrdi da taj nalog dolazi u sukob s njezinim obvezama prema korisničkoj privatnosti. Sudac Sidney Stein zamoljen je da ukine nalog za očuvanje podataka, pokazuje sudska dokumentacija. **New York Times** odbio je komentirati slučaj. Stein je naveo kako je Times iznio osnovanu tvrdnju da su OpenAI i Microsoft poticali korisnike da krše autorska prava. U obrazloženju ranije donesene odluke kojom je odbijeno da se cijeli slučaj odbaci, sud je naveo da brojni i javno poznati primjeri u kojima je ChatGPT generirao sadržaj temeljen na člancima NYT-a (iz kojeg su najavili još tužbi) opravdavaju nastavak sudskog postupka.

FILMSKA INDUSTRIJA TAKOĐER PODIGNULA TUŽBE

Filmski divovi **Disney** i **NBC Universal** podigli su tužbu protiv platforme za generiranje slika umjetnom inteligencijom **Midjourney**, optužujući je za masovno kršenje autorskih prava. Tužba, podnesena saveznom sudu u Los Angelesu, predstavlja potencijalno presedan u sukobu između kreativne industrije i AI-tehnologije. U iznimno oštroj tužbi Midjourney se opisuje kao „leglo plagijata“, a optužuje se za korištenje zaštićenih likova i sadržaja, uključujući **Dartha Vadera** iz *Ratova zvijezda* i **Elsu** iz *Snježnog kraljevstva*, bez dozvole ili ikakvih licencnih sporazuma. „Ova tužba štiti naporan rad svih umjetnika čiji nas rad zabavlja i nadahnjuje, kao i golema ulaganja koja ulažemo u naš sadržaj“, izjavila je **Kim Harris**, izvršna potpredsjednica i glavna pravna savjetnica NBCUniversala, za **Reuters**.

Tužitelji tvrde kako je Midjourney na nezakonit način trenirao svoju AI tehnologiju koristeći autorski zaštićene slike iz filmskih i animiranih kataloga, a zatim plasirao slike, potom i video-materijale koji jasno reproduciraju poznate likove, a da pritom nije uložio ni centa u njihovo stvaranje. „Piratstvo je piratstvo. Bilo da se kršenje autorskog prava događa putem AI-ja ili druge tehnologije, ono ostaje kršenje“, navodi se u tužbi. „Midjourney svojim ponašanjem ugrožava temeljne poticaje koje američki zakon o autorskim pravima pruža filmskoj, televizijskoj i umjetničkoj industriji.“

Generatori AI slika bez kontrole. Portal **Mashable** testirao je nekoliko najpoznatijih AI generatora slika, uključujući Midjourney, i otkrio kako svi lako generiraju deepfake slike s prepoznatljivim Disneyjevim likovima, bez ikakvih zapreka. No posebno otežavajuća okolnost za Midjourney nalazi se u intervjuu iz 2022. za **Forbes**, gdje je osnivač tvrtke

David Holz otvoreno priznao da tvrtka ne traži dozvole od živućih umjetnika ili vlasnika prava. „Ne postoji način da se zna odakle dolazi sto milijuna slika... Bilo bi super kad bi slike imale ugrađen metapodatak o vlasniku prava, ali to jednostavno ne postoji“, rekao je Holz — izjava koja je sada postala dio sudskog spisa. Ovo nije prva tužba protiv Midjourneya. Prije godinu dana desetak umjetnika podnijelo je tužbu protiv njih, kao i protiv **Stability AI** i drugih AI firmi, tvrdeći da su njihovi radovi neovlašteno korišteni za treniranje modela. Taj slučaj još uvijek traje, kao i slične tužbe protiv **OpenAI-ja** i **Metae**.

OGROMNI PRIHODI. Disney i Universal prvi su veliki filmski studiji koji ulaze u izravni pravni sukob s AI industrijom, što ovaj slučaj čini posebno značajnim. **Midjourney**, osnovan 2021., zarađuje putem pretplatničkog modela, a prema podacima iz tužbe prošle je godine ostvario **300 milijuna dolara** prihoda. Studiji tvrde kako je taj novac zarađen iskorištavanjem tuđeg kreativnog rada, bez ikakvih naknada ili prava. Zasad se korištenje autorski zaštićenog materijala za treniranje AI modela nalazi u pravno sivoj zoni – što znači da bi ishod ove tužbe mogao imati dalekosežne posljedice za cijelu generativnu AI industriju.

Glumci su se također digli na noge, pa se od većine može čuti kako je „umjetna inteligencija zapravo krađa identiteta i pitanje ljudskih prava“. I dok Facebook koristi poznate osobe kao temelj za svoje chatbotove, nekim se glumcima ideja da ih zamijeni umjetna inteligencija baš i ne sviđa. Tako je britanska zvijezda **Brian Cox** javno progovorio na **Sky Newsu** o upotrebi umjetne inteligencije na televiziji i u filmskoj industriji. Iako nije precizirao što ga točno zabrinjava, istaknuo je kako bi u budućnosti novi glumci mogli „prodavati“ prava na svoj izgled, a upravo je to ono što **SAG-AFTRA**, udruženje glumaca, želi spriječiti. Mlađi glumci dovedeni su u situaciju da im se kaže da to moraju učiniti, rekao je Cox. „Mislim da je umjetna inteligencija pitanje ljudskih prava. To nije samo problem sindikata. To je zapravo krađa identiteta. I to je vrlo, vrlo rašireno u ovom trenutku.“ Cox je već ranije izrazio svoje nezadovoljstvo korištenjem umjetne inteligencije u filmskoj industriji. Čak je pozvao svoje kolege glumce da zaustave tehnologiju. „Ne postoji crta povučena niz sredinu Atlantskog oceana koja kaže da problem tu prestaje“, rekao je Cox na skupu solidarnosti SAG-a. „Ovo je zajednička borba svih sindikata i radničke klase protiv zlouporabe tehnologije i automatizacije.“

Jedna od onih koja je poduzela prve korake jest poznata glumica **Scarlett Johansson**, koja je podigla tužbu protiv tvrtke koja razvija aplikaciju pogonjenu umjetnom inteligencijom, a u reklami koristi Scarlettin glas i lik. Riječ je o oglasu na platformi X pod nazivom „**Lisa AI: 90s Yearbook & Avatar**“, koji je objavila tvrtka iza aplikacije koja korisnicima omogućuje stvaranje AI-avatara ili primjenu umjetničkih filtara na slike. Oglas prikazuje lice Scarlett Johansson na nekoliko slika koje je generirala umjetna inteligencija, a oponaša i njezin glas. Iako je oglas ubrzo uklonjen, očito je da tvrtka koja je objavila oglas nije dobila odobrenje od Johansson.

TUŽBU PODNIJELI I KNJIŽEVNICI

Već prije protiv AI-ja ustao je ceh književnika koji predstavlja svjetski poznate autore poput **Jodi Picoult** i **Georgea R. R. Martina**, te autori publicistike, koji u grupnoj tužbi navode kako je OpenAI kopirao desetke tisuća publicističkih knjiga bez pristanka autora kako bi obučio svoj chat-bot za odgovaranje na upite. Hoće li Microsoftu i drugim tuženim kompanijama na sudu proći tvrdnja da su kopirali u cilju istraživanja i razvoja tehnologije generativne umjetne inteligencije koja donosi koristi društvu i podržava koncept informirane javnosti – tek će se vidjeti. Ukoliko bi im sud priznao to načelo, to bi bio presedan koji bi u potpunosti srušio zaštitu autorskih prava u cijelom svijetu.

Možda je odnos AI-ja i razmišljanja najbolje sažeo jedan od najvećih umova 21. stoljeća, **Noam Chomsky**, za **NYT** u ožujku 2023. godine: „Ljudski um nije statistički stroj poput ChatGPT-a i njemu sličnih, koji pohlepno guta stotine terabajta podataka kako bi došao do najvjerojatnijeg odgovora na razgovor ili najvjerojatnijeg odgovora na znanstveno pitanje.“ Naprotiv, kaže Chomsky, „ljudski um je iznenađujuće učinkovit i elegantan sustav koji radi s ograničenom količinom informacija. Ne pokušava pokrpati korelacije između podataka, već pokušava stvoriti objašnjenja.“ „Prestanimo to nazivati ‘umjetnom inteligencijom’ i nazovimo je pravim imenom – ‘softver za plagijat’ – jer on ne stvara ništa, već kopira postojeća djela postojećih umjetnika i modificira ih na takav način da se mogu izbjeći autorska prava.“ „Ovo je najveća krađa intelektualnog vlasništva ikad zabilježena otkako su europski kolonizatori stigli u indijanske zajednice“, zaključuje Chomsky.



Članak je objavljen uz financijsku potporu Agencije za elektroničke medije iz Programa poticanja novinarske izvrsnosti

[Hrvatski stručnjak koji se bavi razvojem AI tehnologija: "Ljudi danas ne čitaju vijesti, oni ih samo skimaju"](#)

[Direktor i suosnivač dvije startup tvrtke: "62% mladih u Hrvatskoj nikad nije učilo o medijskoj pismenosti – to je točno ono što ne želimo"](#)

[Od beta kazeta do umjetne inteligencije: Kako se novinarstvo promijenilo u posljednjih 20-ak godina](#)

[Mnogi smatraju ovog čovjeka najboljim podcasterom u regiji. Otkrio nam je zanimljive detalje iz svog života: "Upadao sam u probleme i pokušavao ih rješavati riječima, češće bezuspješno nego uspješno"](#)

["Dezinformacije su stare koliko i čovječanstvo, ali danas cvjetaju na društvenim mrežama"](#)

[Božo Skoko: "Lažne vijesti, senzacionalizam i pad profesionalnih standarda srozali su povjerenje u medije"](#)

Moja reakcija na članak je...

Ljubav

1

Haha

0

Nice

0

What?

0

Laž

0

Sad

0

Mad

1

NAPIŠI, DOJAVI, SLIKAJ,
UKAŽI... BUDI DIO NAS.

