

Editorial/Uvodnik

When artificial intelligence enters the editorial office: Responsibility, integrity, and trust in scientific publishing

Ko umetna inteligenca vstopi v uredništvo: odgovornost, integriteta in zaupanje v znanstvenem publiciranju

Mirko Prosen^{1,*}

Artificial intelligence (AI), particularly generative AI and large language models (LLMs), has rapidly become part of scientific writing and publishing. For journal editors, however, this development is no longer an abstract debate about what AI might do in the future; it is already evident in daily editorial practice. Submitted manuscripts contain polished language but surprisingly shallow argumentation, references that do not support the claims for which they are cited, citations of publications that cannot be found, methodological statements that sound authoritative but make little sense in the context of the study, and passages that seem unrelated to the reported research.

It is difficult to reliably determine the extent of AI involvement in scientific publishing. Current AI-detection approaches carry significant limitations, including false identification of human-written text and difficulty distinguishing between text that has been linguistically edited using AI and text that has been substantially generated by AI (Naddaf, 2026). Nevertheless, emerging evidence suggests that AI-generated text is increasingly common in both manuscripts and peer-review reports. A recent analysis discussed by Naddaf (2026), for example, examined approximately 7,000 manuscript abstracts and 8,000 peer-review reports and estimated that over 30% of peer-review reports contained some AI-generated text.

Editors do not necessarily need an AI detector to recognise potential flaws in a manuscript. At times, issues may be evident in the scientific content itself. The challenge is not simply that an LLM may "hallucinate". A more concerning ethical problem is that an author may submit hallucinated information without prior verification. References provide a particularly troubling example. LLMs can generate citations that look entirely plausible: the authors may be real, the journal appropriate, the title convincing,

and the publication year credible—and yes, the em dash is deliberate; editors may have noticed how remarkably fashionable it has become in the age of generative AI—yet the cited article does not exist.

This is no longer a hypothetical concern. Resnik & Hosseini (in press) cite striking examples from the published literature: one paper contained 19 hallucinated citations among 29 references, another 18 among 76, and a third 12 among only 14 references. They also report an experimental study in mental health, in which 56% of GPT-4o-generated citations contained errors and approximately one in five was hallucinated. Similarly, in an audit of approximately 2.5 million biomedical papers, Topaz et al. (2026) identified fabricated references already embedded in the peer-reviewed literature, and reported a marked increase over time. The proportion of papers containing at least one fabricated reference rose from approximately one in 2,828 papers in 2023 to one in 458 in 2025 and one in 277 during the first seven weeks of 2026.

These are not harmless technical errors. A reference is a scholarly claim that a source exists and provides evidence that can be independently examined. Resnik & Hosseini (in press) make an important distinction between bibliometrically incorrect citations, contextually inaccurate citations, and hallucinated citations – the latter referring to scholarly works that simply do not exist. More provocatively, they argue that hallucinated citations may, under particular circumstances, constitute academic misconduct. Where citations themselves function as research data, such as in systematic reviews or bibliometric research, unchecked hallucinated references may amount to data fabrication if the researcher has demonstrated indifference to the known risk of fabrication (Resnik & Hosseini, in press).

This reframes the ethical question considerably. The fundamental problem is not that the machine

¹ University of Primorska, Faculty of Health Sciences, Polje 42, 6310 Izola, Slovenia

* Corresponding author/Korespondenčni avtor: mirko.prosen@fvz.upr.si

Received/Prejeto: 23. 8. 2026

Accepted/Sprejeto: 29. 8. 2026



made an error; LLMs are known to hallucinate. The problem is that the human researcher failed to verify the machine-generated content. Responsible AI use requires, at a minimum, human verification of the veracity, accuracy, and validity of its output (Resnik & Hosseini, in press). This distinction between AI error and human responsibility should be central to our discussion of AI use in scientific publishing.

Authors are only one side of the story

Editors should avoid reducing AI use solely to a matter of authorship. AI has also entered peer review, which may raise even more difficult ethical questions. Authors submit their work with the reasonable expectation that it will be evaluated by qualified peers exercising sound professional and scientific judgement. Yet, both as an editor and as an author, I have encountered reviews containing generic, irrelevant, or methodologically inappropriate recommendations that appear disconnected from the manuscript under review. An author may be asked to justify a methodological choice that was never made, conduct an analysis inappropriate for the study design, discuss a concept unrelated to the research question, or address a limitation that does not, in fact, exist.

An LLM can produce a review that sounds impressively scholarly while being scientifically wrong. If a reviewer uncritically incorporates such output into a peer-review report, the author may be required to spend time responding to comments with little methodological or substantive relevance to the study. This is not merely inefficient peer review; it raises questions about fairness, professional responsibility, and the integrity of the editorial process.

Importantly, this concern is no longer hypothetical. AI-generated content has been detected not only in submitted and published manuscripts but also in peer-review reports (Naddaf, 2026). In response, the International Committee of Medical Journal Editors (ICMJE, 2026) has issued guidelines on AI use by authors, reviewers, and editors, identifying four central concerns: validity, confidentiality, transparency, and accountability. The issue of confidentiality deserves particular attention. Submitted manuscripts are characterised as privileged communications. The ICMJE (2026) therefore states that editors, reviewers, and publishers should not upload submitted manuscripts to AI systems in which confidentiality cannot be assured without the authors' explicit permission. Similarly, current JAMA Network guidelines prohibit the use of external AI tools in peer review, as entering an author's manuscript – or even a reviewer's summary – into such systems may breach confidentiality (Flanagin et al., 2026).

The ethical question, therefore, is not simply whether reviewers should be permitted to use AI. It is what happens when assistance becomes substitution: when

a reviewer delegates the intellectual responsibility of peer review to a system incapable of being accountable to either the editor or the author.

From detection to responsibility

As an editor, I closely follow email discussions among members of the International Academy of Nursing Editors (INANE). These exchanges offer a revealing glimpse into how rapidly the challenges surrounding AI in scholarly publishing are evolving. The experiences editors share from practice are diverse and, at times, genuinely astonishing. While there appears to be broad agreement on the importance of transparency, verification, integrity, and human accountability, there is far less consensus on where exactly the boundaries of acceptable AI use should be drawn. This is unsurprising, since 'AI use' is not a single, uniform activity.

Using an LLM to improve the readability of a sentence is fundamentally different from prompting it to generate a discussion section. Translation differs from interpreting results. Brainstorming potential headings differs from generating references. Similarly, using an AI tool to improve the language of a review differs from asking it to evaluate a manuscript and then submitting its output as one's own professional judgement.

Evidence shows that undeclared AI-generated material has already appeared in articles published in established scientific journals. Strzelecki (2025) identified characteristic ChatGPT-generated passages in papers published by recognised publishers, including those indexed in Web of Science and Scopus, and demonstrated that some of them were subsequently cited in other scientific publications. His findings raise an uncomfortable question not only about authorship practices but also about whether existing reviewing and editorial processes are adequately equipped to detect such content (Strzelecki, 2025).

This is why an editorial response based primarily on detecting AI-generated language is unlikely to be sufficient. Detection itself remains unreliable, and distinguishing legitimate AI-assisted editing from substantive AI-generated content is particularly difficult. The more durable question is therefore not "Was AI used?" but rather: "How was AI used, was its use appropriate and transparent, was its output verified, and who is accountable for the final scholarly product?"

AI as assistance, not substitution

When used responsibly, AI can be genuinely valuable, including in supporting authors in scientific writing (Watson, 2024). It may improve language, readability, organisation and translation, while increasing overall efficiency. This can be particularly beneficial for researchers writing in a language other

Table 1: *Principles of responsible AI use in scientific publishing*

<i>Area</i>	<i>Responsible use of AI</i>	<i>Key ethical requirement</i>
Authorship and accountability	Authors may use AI-assisted technologies in accordance with journal policy but should disclose how AI was used when required. AI tools must not be listed as authors.	Human authors remain fully responsible for the accuracy, integrity, originality, and appropriate attribution of all submitted content.
Manuscript preparation	AI may assist with language, clarity, and manuscript preparation, where permitted by journal policy.	Relevant use should be disclosed when required; authors remain responsible for the entire manuscript.
Grammar and spelling	Generally acceptable as language assistance.	Human verification remains necessary.
Translation	AI may support the translation of scholarly text.	Accuracy, confidentiality, and appropriate disclosure must be ensured.
Literature searching	AI may assist in identifying potentially relevant literature.	Sources and claims must be independently verified.
References	Reliance on AI-generated references should be avoided.	Every reference must be verified against the original source; hallucinated references threaten research integrity.
Research and data	AI use may be appropriate for specified research tasks where methodologically justified.	Use must be transparently reported and outputs validated.
Peer review	Use of AI in peer review should comply with journal policy. Manuscript content should not be uploaded to external AI tools where confidentiality cannot be assured.	Reviewers must maintain confidentiality, disclose AI use where required, verify the appropriateness and validity of AI-assisted content, and remain fully responsible for their review.
Editorial and publishing processes	AI may support editorial, production, and ancillary activities where permitted by journal policy and where appropriate safeguards are in place.	Editors remain responsible for the accuracy and validity of AI-assisted content; confidentiality and transparency must be ensured.

Legend: AI – artificial intelligence

Note: Author's synthesis based on the principles and recommendations of the ICMJE (2026) and JAMA (Flanagin et al., 2026).

than their first, and may reduce linguistic barriers to international scientific communication.

Current JAMA Network guidelines similarly recognise the potential value of AI use while placing responsibility firmly on authors. AI may be used for several purposes, including manuscript preparation, literature search, translation, and certain aspects of research, provided its use is transparently reported and its outputs verified, with authors retaining full responsibility for the content. These guidelines do not require disclosure of basic grammar and spelling checks. In contrast, using AI to generate bibliographic references is discouraged due to the risks to accuracy and integrity (Flanagin et al., 2026). This position is particularly important in light of editorial experience with hallucinated references. JAMA Network editors report receiving manuscripts containing multiple inaccurate or non-existent references. They caution that even advanced AI models can generate seemingly realistic references, with real authors and plausible titles, that nevertheless do not exist (Flanagin et al., 2026).

The relevant distinction is therefore not between using and not using AI, but between assistance and substitution. While AI may assist in scholarly work, it must not substitute for scholarly judgement (Kendall & Teixeira da Silva, 2024).

We can delegate tasks, but not responsibility

This principle applies equally to authors, reviewers, and editors: responsibility for scholarly content, peer review, and editorial decisions rests with humans. It also lies at the centre of current international guidelines. The ICMJE (2026) emphasises that AI cannot be listed as an author; that humans remain responsible for appropriate attribution and for avoiding plagiarism; that the confidentiality of submitted manuscripts must be protected; and that authors, reviewers, editors, and publishers should be transparent about AI use throughout manuscript preparation and the editorial process. Journals, in turn, should establish explicit AI policies and communicate them to all participants in the publication process. While transparency is necessary, it is not sufficient. Declaring that an AI tool was used does not excuse an inaccurate analysis, a fabricated reference, an inappropriate review, or a breach of confidentiality.

These principles should inform the development of practical recommendations for responsible AI use in scientific publishing (Table 1). Their ultimate purpose is to safeguard the very foundation of scientific publishing: trust.

AI can help us write more clearly, work more efficiently, and perhaps even think differently. But it cannot accept responsibility for fabricated findings, flawed interpretations, hallucinated references, or

misleading conclusions. It cannot respond to an author who has received an inappropriate peer-review request. It cannot assume responsibility for the integrity of the scientific record, nor can it be held accountable for editorial decisions. These responsibilities remain ours. We may delegate tasks to AI, but we cannot delegate scholarly responsibility.

Slovenian translation/Prevod v slovenščino

Umetna inteligenca (UI), zlasti v obliki generativne UI in velikih jezikovnih modelov (VJM), je hitro postala del znanstvenega pisanja in publiciranja. Za urednike revij ta sprememba ne predstavlja zgolj abstraktne razprave o morebitnih prihodnjih učinkih UI, temveč njeno aktualno in vidno prisotnost v vsakodnevni uredniški praksi. Avtorski prispevki se morda ponašajo z izpiljenim jezikom, vendar vsebujejo presenetljivo površno argumentacijo, navedbe nerelevantnih ali celo neobstoječih publikacij, metodološke izjave, ki sicer zvenijo avtoritativno, a v kontekstu raziskave nimajo veliko smisla, ter odlomke, ki se zdijo nepovezani z raziskavo.

O natančnem obsegu prodora UI v znanstveno publiciranje je težko zanesljivo soditi. Pomembne omejitve sedanjih pristopov za detekcijo besedil, ustvarjenih s pomočjo UI, vključujejo napačno prepoznavo avtorskega besedila ter težave pri razlikovanju med besedilom, ki je bilo jezikovno urejeno z UI, in besedilom, ki ga je v celoti ustvarila UI (Naddaf, 2026). Najnovejše ugotovitve kažejo, da so besedila, ki jih je ustvarila UI, vse pogostejša ne le med prispevki, temveč tudi med recenzijami. Kot poroča Naddaf (2026), so rezultati ene od nedavnih raziskav, ki se je posvetila približno 7.000 povzetkom člankov in 8.000 strokovnim recenzijam, pokazali, da je več kot 30 % recenzij vsebovalo določen delež besedila, ki ga je ustvarila UI.

Uredniki ne potrebujejo nujno detektorja UI, saj se potencialna neustreznost prispevka pogosto razkrije že skozi znanstveno vsebino. Težava ni le v tem, da VJM morda »halucinira«. Veliko bolj pereč etični problem je, da avtor sploh predloži »halucinirane« informacije, ne da bi jih predhodno preveril. Navedbe virov so še posebej zaskrbljujoč primer. VJM lahko ustvarijo navidez povsem legitimne navedbe virov z resničnimi avtorji, ustrezno revijo, prepričljivim naslovom in verodostojnim letom objave – (in da, dolg pomišljaj je na tem mestu nameren; uredniki so morda opazili, kako izjemno moden je postal v dobi generativne UI) –, vendar citirani članek ne obstaja.

Ne gre zgolj za hipotetični pomislek. Resnik & Hosseini (in press) navajata presenetljive primere iz znanstvenih objav: v enem članku je bilo izmišljenih 19 od 29 citatov, v drugem 18 od 76, v tretjem pa 12 od skupno 14. Avtorja poročata tudi o eksperimentalni raziskavi na področju duševnega zdravja, v kateri je

56 % navedb virov, ki jih je ustvaril GPT-4o, vsebovalo napake, približno ena od petih pa je bila izmišljena. Podobno so Topaz et al. (2026) pri pregledu približno 2,5 milijona člankov s področja biomedicine med strokovno recenzirano literaturo odkrili izrazito povečanje izmišljenih bibliografskih navedb skozi čas. Delež člankov z vsaj eno izmišljeno navedbo se je povečal s približno enega na 2.828 člankov v letu 2023 na enega na 458 v letu 2025 ter enega na 277 v prvih sedmih tednih leta 2026.

To niso zanemarljive tehnične napake. Navedba vira je znanstvena trditev, da vir obstaja in zagotavlja dokaze, ki jih je mogoče neodvisno preveriti. Resnik & Hosseini (in press) postavljata pomembno ločnico med bibliometrično napačnimi, kontekstualno netočnimi in izmišljenimi navedbami virov – pri čemer se slednje nanašajo na znanstvena dela, ki preprosto ne obstajajo. Še več, trdita, da lahko navajanje izmišljenih virov v določenih okoliščinah predstavlja kršenje raziskovalne etike. Kadar navedbe virov predstavljajo raziskovalne podatke, kot na primer v sistematičnih pregledih literature ali bibliometričnih raziskavah, lahko nekritično navajanje haluciniranih virov ob zanemarjanju znanega tveganja potvarjanja pomeni ponarejanje podatkov (Resnik & Hosseini, in press).

To v precejšnji meri spreminja naravo tega etičnega vprašanja. Osnovni problem namreč ni v tem, da je »stroj« naredil napako, saj je pri rabi VJM tveganje haluciniranja znano. Težava je torej v tem, da raziskovalec ni preveril, kaj je »stroj« ustvaril. Odgovorna raba UI zahteva vsaj človeško presojo resničnosti, točnosti in veljavnosti izhodnih podatkov (Resnik & Hosseini, in press). Zato bi morala razlika med napako UI in človeško odgovornostjo zavzemati osrednje mesto razprave o UI v znanstvenem publiciranju.

Avtorji predstavljajo le eno plat zgodbe

Uredniki se morajo izogibati obravnavi rabe UI zgolj v kontekstu avtorstva. Dejstvo, da UI prodira tudi na področje strokovne recenzije, sproža morda še bolj pereča etična vprašanja. Avtorji svoja dela oddajo z razumnim pričakovanjem, da jih bodo ocenili usposobljeni strokovnjaki na podlagi lastne strokovne in znanstvene presoje. Vendar sem tako kot urednik kot tudi kot avtor že naletel na recenzije, ki vsebujejo splošna, nebitvena ali metodološko neustrezna priporočila, ki se zdijo nepovezana s samim prispevkom. Recenzenti morda pozivajo k utemeljitvi metodološke odločitve, ki je avtor sploh ni navedel, izvedbi analize, ki ni primerna glede na zasnovo raziskave, obravnavi koncepta, ki ni povezan z raziskovalnim vprašanjem, ali naslovitvi omejitve, ki v resnici ne obstaja.

VJM lahko ustvari recenzijo, ki zveni izjemno znanstveno, vendar je znanstveno neutemeljena. Če recenzent takšno izhodno besedilo orodja UI nekritično prenese v recenzijo, lahko od avtorja zahteva,

da porabi čas za odgovarjanje na komentarje, ki za raziskavo nimajo zadostne metodološke ali vsebinske relevantnosti. Takšna recenzija ni le neučinkovita, temveč postavlja pod vprašaj pravičnost, strokovno odgovornost in integriteto uredniškega procesa.

Pomembno se je zavedati, da ta pomislek ni več zgolj hipotetičen, saj analize že potrjujejo prisotnost vsebin, ustvarjenih z UI, tako v predloženih in objavljenih prispevkih kot tudi strokovnih recenzijah (Naddaf, 2026). Mednarodni odbor urednikov medicinskih revij (ICMJE, 2026) zato uporabo UI obravnava ne le z vidika avtorjev, ampak tudi recenzentov in urednikov ter pri tem zastavlja štiri osrednja vprašanja: o veljavnosti, zaupnosti, preglednosti in odgovornosti. V tem kontekstu velja posebno pozornost posvetiti vprašanju zaupnosti, saj so predloženi prispevki zaupne narave. ICMJE (2026) zato navaja, da jih uredniki, recenzenti in založniki brez izrecnega dovoljenja avtorjev ne smejo nalagati v sisteme UI, ki ne zagotavljajo zaupnosti. Podobno veljavna navodila mreže JAMA Network ne dovoljujejo uporabe zunanjih orodij UI za strokovno recenzijo, saj lahko vnos avtorjevega prispevka ali celo recenzentovega povzetka le-tega pomeni kršitev načela zaupnosti (Flanagin et al., 2026).

Etično dilemo zato ne predstavlja le vprašanje, ali recenzentom dovoliti rabo UI. Gre za to, kaj se zgodi, ko od pomoči pri pisanju preidemo k nadomeščanju pisanja – ko recenzent intelektualno odgovornost strokovnega pregleda prenese na sistem, ki ne more biti odgovoren niti uredniku niti avtorju.

Od detekcije do odgovornosti

Kot urednik pozorno spremljam aktualne dopisovalne razprave med člani Mednarodne akademije urednikov na področju zdravstvene nege (angl. *International Academy of Nursing Editors - INANE*). V izmenjavah mnenj, ki ponujajo zanimiv vpogled v sodobne izzive, povezane z rabo UI v znanstvenem publiciranju, uredniki delijo raznolike in mestoma nepričakovane izkušnje iz uredniške prakse. Čeprav se zdi, da obstaja splošen konsenz o pomenu preglednosti, preverljivosti, integritete in človeške odgovornosti, pa stališča glede natančnih meja sprejemljive rabe UI niso enotna. To je razumljivo, saj tudi izraz »raba UI« ne označuje enotne dejavnosti.

Raba VJM za izboljšanje berljivosti stavka se bistveno razlikuje od generiranja celotnega poglavja diskusije. Prevajanje se razlikuje od interpretacije rezultatov. Zbiranje predlogov potencialnih naslovov se razlikuje od ustvarjanja seznama bibliografskih virov. Podobno se uporaba orodij UI za izboljšanje jezikovnega sloga recenzije razlikuje od poziva sistemu k strokovni oceni prispevka in nato predložitve generiranega besedila kot lastne recenzije.

Raziskave kažejo, da je prikrita raba gradiva, generiranega z UI, že prisotna v znanstvenih objavah uveljavljenih znanstvenih revij. Strzelecki (2025)

poroča, da znanstvene objave v revijah priznanih založnikov (tudi v bazah Web of Science in Scopus) vsebujejo značilne odlomke, ustvarjene z orodjem ChatGPT, nekatere med njimi pa se celo navajajo v drugih znanstvenih publikacijah. Ugotovitve te raziskave odpirajo neprijetno vprašanje ne le o avtorskih praksah, temveč tudi o ustrezni opremljenosti obstoječih recenzijskih in uredniških postopkov za prepoznavanje tovrstnih vsebin (Strzelecki, 2025).

Zato je malo verjetno, da bo uredniški odziv, ki temelji predvsem na detekciji besedil, ustvarjenih z UI, zadosten. Sama detekcija tovrstnih besedil ostaja nezanesljiva, razlikovanje med legitimnim jezikovnim urejanjem besedila s pomočjo UI in generiranjem vsebinskih delov besedila z UI pa je še posebej težavno. Temeljno vprašanje se zato ne glasi: »Ali je bila uporabljena UI?«, temveč: »Na kakšen način je bila uporabljena UI, ali je bila njena uporaba ustrezna in pregledna, ali je bilo izhodno besedilo preverjeno in kdo je odgovoren za končni znanstveni izdelek?«

UI kot pomoč, ne kot nadomestilo

Ob odgovorni rabi lahko UI avtorjem nudi neprecenljivo pomoč pri pisanju znanstvenih besedil (Watson, 2024). Izboljša lahko njihovo jezikovno ustreznost, berljivost, strukturo, prevod in učinkovitost besedila. V takšni funkciji je lahko še posebej dragocena za raziskovalce, ki pišejo v jeziku, ki ni njihov materni, saj lahko zmanjša nekatere jezikovne ovire v mednarodni znanstveni komunikaciji.

Aktualna navodila mreže JAMA Network priznavajo potencialno vrednost UI, vendar odgovornost odločno pripisujejo avtorjem. UI se lahko uporablja za različne namene, vključno s pripravo prispevka, iskanjem literature, prevajanjem in nekaterimi vidiki raziskovanja, pod pogojem, da se načini njene uporabe pregledno navedejo ter da avtorji preverijo rezultate in ohranijo polno odgovornost za vsebino. Omenjena navodila ne zahtevajo razkritja rabe UI za osnovni pregled slovnice in pravopisa, odsvetujejo pa zanašanje nanjo pri sestavljanju seznamov uporabljenih literature zaradi tveganj, povezanih s točnostjo in celovitostjo (Flanagin et al., 2026). To stališče je še posebej pomembno glede na uredniške izkušnje s pojavljanjem izmišljenih virov. Uredniki mreže JAMA Network poročajo o prispevkih s številnimi netočnimi citati ali navedbami neobstoječih virov ter opozarjajo, da lahko tudi napredni modeli UI ustvarijo navidez legitimne bibliografske vire, ki vključujejo resnične avtorje in verjetne naslove, a kljub temu ne obstajajo (Flanagin et al., 2026).

Zato ni toliko pomembna ločnica med uporabo in neuporabo UI, temveč med uporabo UI kot pomoči pri pisanju ter njeno uporabo namesto pisanja. UI lahko nudi pomoč pri oblikovanju znanstvenega dela, ne sme pa nadomestiti znanstvene presoje (Kendall & Teixeira da Silva, 2024).

Tabela 1: Načela odgovorne rabe UI v znanstvenem publiciranju

Področje	Odgovorna uporaba UI	Ključna etična zahteva
Avtorstvo in odgovornost	Avtorji lahko uporabljajo tehnologije, podprte z UI, v skladu z uredniško politiko revije, vendar morajo po potrebi razkriti, kako je bila UI uporabljena. Orodij UI se ne sme navajati kot avtorjev.	Človeški avtorji ostajajo v celoti odgovorni za točnost, celovitost, izvirnost in ustrezno navajanje virov vseh predloženih vsebin.
Priprava prispevka	UI lahko pomaga pri jezikovni obdelavi, jasnosti in pripravi rokopisa, če to dovoljuje uredniška politika revije.	Uporabo UI je treba po potrebi navesti; avtorji ostajajo odgovorni za celoten prispevek.
Slovnica in pravopis	Na splošno sprejemljiva kot jezikovna pomoč.	Človeški pregled ostaja nepogrešljiv.
Prevajanje	UI lahko podpira prevajanje znanstvenih besedil.	Zagotoviti je potrebno natančnost, zaupnost in ustrezno razkritje.
Iskanje literature	UI se lahko uporablja kot podpora pri iskanju potencialno relevantne literature.	Vire in trditve je potrebno neodvisno preveriti.
Navajanje virov	Izogibati se je potrebno zanašanju na navedbe virov, ki jih ustvari UI.	Vsako bibliografsko navedbo ali sklic je potrebno preveriti glede na izvirni vir; izmišljene bibliografske navedbe ogrožajo integriteto raziskave.
Raziskava in podatki	Uporaba UI je lahko primerna za določene raziskovalne naloge, kadar je to metodološko utemeljeno.	Uporabo UI je potrebno pregledno navesti, rezultate pa preveriti.
Strokovna recenzija	Uporaba UI pri recenziranju mora biti v skladu s politiko revije. Vsebin prispevka se ne sme nalogati v zunanja orodja UI, kadar zaupnost ni zajamčena.	Recenzenti morajo ohraniti zaupnost, po potrebi razkriti uporabo UI, preveriti ustreznost in veljavnost vsebine, pri kateri je bila uporabljena UI, ter ostati v celoti odgovorni za svojo recenzijo.
Uredniški in založniški postopki	UI lahko podpira uredniške, produkcijske ali pomožne dejavnosti, v kolikor to dovoljuje politika revije in so sprejeti ustrezni zaščitni ukrepi.	Uredniki ostajajo odgovorni za točnost in veljavnost vsebin, pri katerih je bila uporabljena UI; zagotovljeni morata biti zaupnost in transparentnost.

Legenda: UI – umetna inteligenca

Opomba: Avtorjev povzetek na podlagi pregleda načel in priporočil ICMJE (2026) ter JAMA (Flanagin et al., 2026).

Prenesemo lahko naloge, ne pa odgovornosti

To v enaki meri velja za avtorje, recenzente in urednike: odgovornost za znanstveno vsebino, strokovno recenzijo in uredniške odločitve ostaja v rokah ljudi. Navedeno načelo predstavlja bistveni element veljavnih mednarodnih smernic. ICMJE (2026) poudarja, da se UI ne sme navajati kot avtorja, da odgovornost za ustrezno navajanje virov in preprečevanje plagiatstva nosijo ljudje, da je ključnega pomena zaščititi zaupnost predloženih prispevkov ter da morajo avtorji, recenzenti, uredniki in založniki zagotoviti transparentnost glede načina uporabe UI skozi celoten proces priprave prispevka in njegove uredniške obravnave. Revije morajo vzpostaviti jasna pravila glede rabe UI in z njimi seznaniti vse udeležence v procesu objavljanja. Pri tem je transparentnost sicer nujna, vendar ni zadostna. Sama izjava, da je bilo uporabljeno orodje UI, ne upravičuje netočne analize, izmišljenih virov, neprimerne recenzije ali kršitve zaupnosti.

Omenjena načela morajo služiti kot podlaga za oblikovanje praktičnih priporočil za odgovorno rabo UI v znanstvenem publiciranju (Tabela 1). Končni namen teh priporočil pa je varovanje zaupanja kot temeljnega principa znanstvenega publiciranja.

UI nam lahko pomaga pisati jasneje, delati učinkoviteje in morda celo razmišljati drugače. Vendar ne more prevzeti odgovornosti za izmišljene ugotovitve, napačne interpretacije, halucinirane navedbe virov ali zavajajoče zaključke. Ne more odgovoriti avtorju, ki je v strokovni recenziji prejel neprimerno zahtevo. Ne more prevzeti odgovornosti ne za integriteto znanstvenih zapisov ne za uredniške odločitve. Te odgovornosti ostajajo naše. Čeprav lahko na UI prenesemo določene naloge, nanjo ne moremo prenesti tudi znanstvene odgovornosti.

Conflict of interest

The author declares no conflicts of interest.

Literature

Flanagin, A., Perlis, R. H., & Bibbins-Domingo, K. (2026). Updated guidance for author use of AI in medical publication. *JAMA*. <https://doi.org/10.1001/jama.2026.16613>
PMCID:PMC12853283

International Committee of Medical Journal Editors (ICMJE). (2026). *Use of artificial intelligence in publishing*. <https://www.icmje.org/recommendations/browse/artificial-intelligence/>

Kendall, G., & Teixeira da Silva, J. A. (2024). Risks of abuse of large language models, like ChatGPT, in scientific publishing: Authorship, predatory publishing, and paper mills. *Learned Publishing*, 37(1), 55–62.

<https://doi.org/10.1002/leap.1578>

PMid:41833014; PMCID:PMC13051339

Naddaf, M. (2026, May 5). How much of the scientific literature is generated by AI? *Nature*. <https://doi.org/10.1038/d41586-025-03504-8>

Resnik, D. B., & Hosseini, M. (in press). Hallucinated citations produced by generative artificial intelligence may constitute research misconduct when citations function as data in scholarly papers. *Accountability in Research*.

<https://doi.org/10.1080/08989621.2026.2645390>

PMid:41833014; PMCID:PMC13051339

Strzelecki, A. (2025). 'As of my last knowledge update': How is content generated by ChatGPT infiltrating scientific papers published in premier journals? *Learned Publishing*, 38, Article e1650.

<https://doi.org/10.1002/leap.1650>

Topaz, M., Roguin, N., Gupta, P., Zhang, Z., & Peltonen, L.-M. (2026). Fabricated citations: An audit across 2.5 million biomedical papers. *The Lancet*, 407, 1779–1780.

[https://doi.org/10.1016/S0140-6736\(26\)00603-3](https://doi.org/10.1016/S0140-6736(26)00603-3)

PMid:42107362

Watson, R. (2024). The potential and pitfalls of artificial intelligence in nursing. *Obzornik zdravstvene nege*, 58(1), 4–6.

<https://doi.org/10.14528/snr.2024.58.1.3279>

Cite as/Citirajte kot:

Prosen, M. 2026. When artificial intelligence enters the editorial office: Responsibility, integrity, and trust in scientific publishing. *Obzornik zdravstvene nege*, 60(3), 268–274. <https://doi.org/10.14528/snr.2026.60.3.3400>