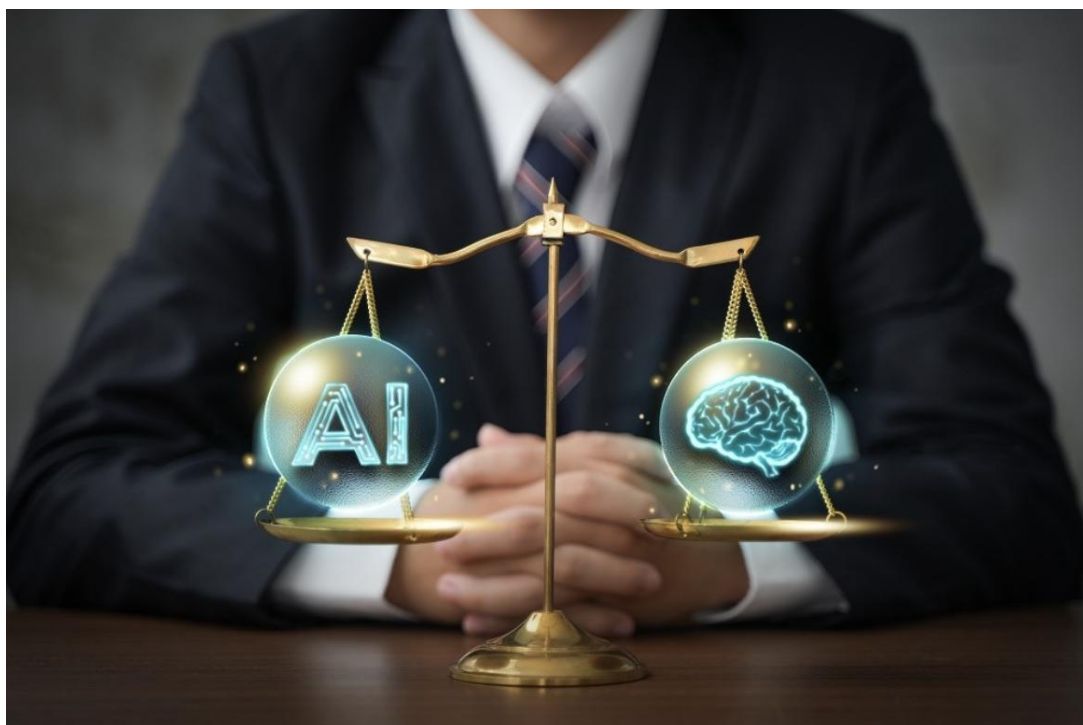




# SMERNICE ZA ODGOVORNO IN KRITIČNO UPORABO UMETNE INTELIGENCE PRI PRAVNEM RAZISKOVANJU

Smernice pripravljene v okviru projekta LexDigital – smernice za odgovorno rabo AI



*Operacija se izvaja v okviru javnega razpisa Problemsko učenje študentov v delovno okolje 2024–2027 (PUŠ v*

*Slika1: Uporaba umetne inteligence pri pravnem raziskovanju*

*delovno okolje 2024–2027). Operacijo sofinancirata Republika Slovenija, Ministrstvo za izobraževanje, znanost in mladino ter Evropska unija iz Evropskega socialnega sklada plus (ESS+).*

Ljubljana, junij 2026



NOVA  
UNIVERZA

## Kazalo vsebine

|                                                                        |           |
|------------------------------------------------------------------------|-----------|
| <b>1. UVOD .....</b>                                                   | <b>4</b>  |
| 1.1. Namen in cilji smernic.....                                       | 4         |
| 1.2. Definicije izrazov v smernicah .....                              | 4         |
| 1.3. Metodologija.....                                                 | 4         |
| <b>2. TEMELJNI POJMI.....</b>                                          | <b>5</b>  |
| 2.1. Kaj je umetna inteligenca (UI)? .....                             | 5         |
| 2.1.1 Kratek zgodovinski pregled.....                                  | 6         |
| 2.1.2 Strojno učenje .....                                             | 6         |
| 2.1.3 Globoko učenje.....                                              | 7         |
| 2.2 Generativni modeli umetne inteligence .....                        | 8         |
| 2.3 Pozitivne lastnosti umetne inteligence .....                       | 8         |
| 2.4 Negativne lastnosti umetne inteligence.....                        | 9         |
| <b>3. UPORABA UI V PRAVNEM RAZISKOVANJU.....</b>                       | <b>10</b> |
| 3.1. Pozitivna funkcija uporabe UI v pravu .....                       | 10        |
| 3.2. Meje dopustne in odgovorne uporabe UI v pravu .....               | 11        |
| 3.3. Tveganja uporabe UI v pravnem raziskovanju.....                   | 11        |
| 3.3.1. Zasebnost in varstvo osebnih podatkov .....                     | 12        |
| 3.3.2. Halucinacije .....                                              | 12        |
| 3.3.3. Pristranskost .....                                             | 12        |
| 3.3.4. Pomanjkanje preglednosti .....                                  | 13        |
| 3.3.5. Odgovornost za napake UI .....                                  | 13        |
| 3.3.6. Intelektualna lastnina .....                                    | 13        |
| 3.3.7. Varstvo demokracije.....                                        | 13        |
| 3.3.8. Kibernetika varnost .....                                       | 13        |
| <b>4. PRIMER SLABE UPORABE UMETNE INTELIGENCE PRAVNI PRAKSI.....</b>   | <b>13</b> |
| 4.1. Ozadje primera .....                                              | 14        |
| 4.2. Potek dogajanja in stopnjevanje napak .....                       | 14        |
| 4.3. Odziv sodnika in pravna presoja .....                             | 15        |
| 4.4. Širši kontekst .....                                              | 15        |
| 4.5. NAUK IZ FOURTEJEVEGA PRIMERA.....                                 | 15        |
| <b>5. PRAVNI IN ETIČNI VIDIKI ODGOVORNE UPORABE UMETNE INTELIGENCE</b> |           |
| 15                                                                     |           |
| 5.1. Odgovornost uporabe UI .....                                      | 15        |
| 5.2. Skrbnost uporabe UI .....                                         | 16        |
| 5.3. Varstvo podatkov GDPR .....                                       | 16        |

|             |                                                  |           |
|-------------|--------------------------------------------------|-----------|
| <b>5.4.</b> | <b>Poštenost, spoštovanje in pravičnost.....</b> | <b>17</b> |
| <b>5.5.</b> | <b>ODGOVORNA UPORABA UI V PRAVU .....</b>        | <b>17</b> |
| 5.5.1.      | Človeški nadzor .....                            | 18        |
| 5.5.2.      | Preverljivost rezultatov.....                    | 18        |
| 5.5.3.      | Načelo sorazmernosti .....                       | 18        |
| 5.5.4.      | Transparentnost uporabe UI .....                 | 19        |
| 5.5.5.      | Nenehno izobraževanje uporabnikov .....          | 19        |
| <b>6.</b>   | <b><i>VIRI IN LITERATURA.....</i></b>            | <b>20</b> |

## **Kazalo slik**

|                                                                    |   |
|--------------------------------------------------------------------|---|
| Figure 1: Uporaba umetne inteligence pri pravnem raziskovanju..... | 1 |
|--------------------------------------------------------------------|---|

## 1. UVOD

### 1.1. Namen in cilji smernic

Umetna inteligenca (v nadaljevanju: UI) postaja vse pomembnejše orodje tudi na področju prava, zlasti pri pravnem raziskovanju, iskanju sodne prakse, analizi pravnih virov ter pripravi osnutkov pravnih besedil. Njena uporaba lahko prispeva k večji učinkovitosti, hitrejšemu dostopu do informacij in lažji obdelavi obsežnih podatkovnih zbirk. Hkrati pa uporaba UI odpira številna pravna in etična vprašanja, povezana z zanesljivostjo rezultatov, varstvom podatkov, ter ohranjanjem kakovosti pravnega dela.

Namen teh smernic je opredeliti priporočila za odgovorno, etično in kritično uporabo UI pri pravnem raziskovanju. Smernice analizirajo tveganja uporabe UI, opozarjajo na omejitve tovrstnih sistemov ter poudarjajo pomen preverjanja informacij in uporabe verodostojnih pravnih virov. Posebna pozornost je namenjena tudi varstvu osebnih podatkov in zaupnosti informacij, ki jih uporabniki vnašajo v sisteme UI.

Cilj smernic je uporabnikom pomagati pri varni, učinkoviti in strokovno ustrezni uporabi UI pri pravnem raziskovanju. Temeljno izhodišče smernic je, da UI predstavlja podporno orodje pri delu pravnika, ne more pa nadomestiti njegove strokovne presoje, odgovornosti in njegove končne odločitve.

### 1.2. Definicije izrazov v smernicah

Za potrebe teh smernic imajo uporabljeni izrazi naslednji pomen:

- **Umetna inteligenca (UI)** pomeni računalniške sisteme, ki so sposobni izvajati naloge, ki običajno zahtevajo človeško inteligenco, kot so analiza podatkov, obdelava naravnega jezika in generiranje vsebin.<sup>1</sup>
- **Pravno raziskovanje** pomeni postopek iskanja, analize in vrednotenja pravnih virov, kot so zakonodaja, sodna praksa, pravna literatura in drugi relevantni pravni dokumenti.<sup>2</sup>
- **Halucinacija UI** pomeni primer, ko sistem umetne inteligence ustvari napačno, zavajajočo ali neobstoječo informacijo, ki je predstavljena kot dejstvo.<sup>3</sup>
- **Pravni viri** so zakonodaja, sodna praksa, pravna literatura in drugi uradni ali strokovno priznani viri pravnih informacij.<sup>4</sup>

### 1.3. Metodologija

Pri pripravi smernic je bila uporabljena deskriptivna in analitična metoda raziskovanja. Dokument se opira na samostojno empirično raziskavo UI orodja ModernLegal, predvsem pa temelji na sistematičnem pregledu in analizi zanesljivih domačih in tujih

---

<sup>1</sup> <https://www.europarl.europa.eu/topics/sl/article/20200827STO85804/kaj-je-umetna-inteligenca-in-kako-se-uporablja-v-praksi>

<sup>2</sup> <https://www.europarl.europa.eu/topics/sl/article/20200827STO85804/kaj-je-umetna-inteligenca-in-kako-se-uporablja-v-praksi>

<sup>3</sup> <https://www.europarl.europa.eu/topics/sl/article/20200827STO85804/kaj-je-umetna-inteligenca-in-kako-se-uporablja-v-praksi>

<sup>4</sup> <https://www.europarl.europa.eu/topics/sl/article/20200827STO85804/kaj-je-umetna-inteligenca-in-kako-se-uporablja-v-praksi>

virov. Uporabljeni so bili strokovna in znanstvena literatura, pravni predpisi, smernice mednarodnih organizacij, uradni dokumenti evropskih institucij ter relevantni prispevki s področja umetne inteligence, prava in varstva osebnih podatkov.

Na podlagi pregleda navedenih virov so bile ugotovljene ključne prednosti, omejitve in tveganja uporabe umetne inteligence pri pravnem raziskovanju. Ugotovitve so bile nato sintetizirane v obliki praktičnih priporočil in smernic, namenjenih uporabnikom umetne inteligence v pravnem okolju.

Ugotovitve raziskave predstavljajo podlago za oblikovanje priporočil za varno in kritično uporabo UI pri pravnem raziskovanju.

## **2. TEMELJNI POJMI**

### **2.1. Kaj je umetna inteligenca (UI)?**

Umetna inteligenca (v nadaljevanju UI), je definirana kot inteligenca strojev oz. Inteligenca predstavljena s stroji v kontrastu z naravno inteligenco ponazorjeno z ljudmi. Termin UI je pogosto opisuje stroje, naprave, ki posnemajo človeške kognitivne funkcije, kot so učenje, razumevanje, reševanje problemov, kompleksno razmišljanje.<sup>5</sup> Je hitro razvijajoča se družina tehnologij, ki prispeva k širokemu naboru gospodarskih, okoljskih in družbenih koristi v celotnem spektru panog in družbenih dejavnosti. Z izboljšanjem napovedovanja, optimizacijo poslovanja in dodeljevanja virov, ki so na voljo posameznikom in organizacijam lahko uporaba umetne inteligence zagotovi ključne konkurenčne prednosti podjetjem in podpre družbeno in okoljsko koristne rezultate (npr. v zdravstvu, kmetijstvu, izobraževanju, varnosti hrane, usposabljanju medijev, športu, kulturi, energetiki, itd.).<sup>6</sup>

Hkrati lahko umetna inteligenca, odvisno od okoliščin v zvezi z njeno specifično uporabo in stopnjo tehnološkega razvoja, ustvari tveganja in povzroči škodo javnim interesom in temeljnim pravicam, ki jih varuje pravo Unije. Takšna škoda je lahko materialna ali nematerialna, vključno s fizično, psihološko, družbeno ali gospodarsko škodo.<sup>7</sup>

Pojem 'sistem UI' po Uredbi (EU) 2024/1689 mora biti jasno opredeljen in tesno usklajen z delom mednarodnih organizacij, ki delajo na področju UI, da zagotovijo pravno gotovost, olajšajo mednarodno konvergenco in široko sprejetost ter hkrati zagotovijo prilagodljivost za hitro tehnološko rast na tem področju. Poleg tega bi morala definicija temeljiti na ključnih značilnostih sistemov UI, ki jih ločijo od enostavnejših tradicionalnih programskih sistemov.<sup>8</sup> Ključna značilnost sistemov UI je njihova sposobnost sklepanja. Ta sposobnost sklepanja se nanaša na proces pridobivanja izhodov, kot so napovedi, vsebina, priporočila ali odločitve, ki lahko

---

<sup>5</sup>EU komisija (AI Watch): Historical Evolution of Artificial Intelligence, [https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence\\_en](https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence_en), str. 5

<sup>6</sup> Uredba (EU) 2024/1689 evropskega parlamenta in sveta, <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng> 5. člen

<sup>7</sup> Uredba (EU) 2024/1689 evropskega parlamenta in sveta, <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng> 12. člen

<sup>8</sup> Uredba (EU) 2024/1689 evropskega parlamenta in sveta, <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng> 12. člen (prvi odst.)

vplivajo na fizična in virtualna okolja, ter na sposobnost sistemov umetne inteligence, da iz vhodov ali podatkov izpeljejo modele ali algoritme ali oboje.<sup>9</sup>

Umetna inteligenca (v nadaljevanju UI), je krovni pojem, ki zajema širok nabor tehnologij, vključno s strojnim učenjem, globokim učenjem in obdelavo naravnega jezika (NLP).<sup>10</sup> Čeprav se izraz pogosto uporablja za opis različnih tehnologij, se mnogi ne strinjajo ali te dejansko predstavljajo UI, namreč trdijo, da velik del tehnologije dane v uporabi predstavlja napredno strojno učenje, ki je zgolj prvi korak k pravi umetni inteligenci oz. splošni umetni inteligenci (GAI). Splošna umetna inteligenca se nanaša na teoretično stanje, v katerem računalniški sistemi lahko dosežejo ali presežejo človeško inteligenco. Drugače povedano GAI je »prava« umetna inteligenca, kot je prikazano v spletnih medijih in filmskih platnih.<sup>11</sup>

### 2.1.1 Kratek zgodovinski pregled

Po zgodovini razvoja delimo evolucijo UI na tri obdobja. Prva »UI doba« se je začela s konferenco Dartmouth leta 1956, kjer je UI (umetna inteligenca) dobila svoje ime in pomen. McCarthy je upodobil izraz umetna inteligenca, ki je postal znanstveni termin.<sup>6</sup> To obdobje (1950–1970) je prineslo temelje, kot so Turingov test, perceptron, prve igre in robote — a se je razbilo ob kombinatorični eksploziji in nerealiziranih napovedih, kar je sprožilo prvo umetno zimo.<sup>12</sup> Drugo obdobje (1970–1990) je stavilo na ekspertne sisteme in simbolno kodiranje človeškega znanja; kljub začetnim uspehom (XCON, milijardne naložbe Japonske in EU) so se sistemi zlomili ob težavnosti pridobivanja znanja, počasnosti simbolnih jezikov in preseženi zmogljivosti poceni PC-jev — sledila je druga zima.<sup>13</sup> Tretje obdobje oz. sodobno obdobje (1990–danes) je temeljito spremenilo paradigmo: namesto ročno kodiranih pravil so se stroji začeli učiti iz podatkov, razvilo se je strojno učenje, posledica nadaljnje nadgradnje tega pa se je pojavilo globoko učenje z eksplozijo podatkov (ImageNet) in računske moči, ki je premagalo človeka v šahu, Go-ju, pokru in kompleksnih videoigrah, čeprav ostajata razumevanje vzročnosti in posploševanje med domenami še vedno nedosegljivi.<sup>14</sup>

### 2.1.2 Strojno učenje

Neposredno pod umetno inteligenco imamo strojno učenje, ki vključuje ustvarjanje modelov z učenjem algoritma za napovedovanje ali odločitve na podlagi podatkov. Zajema širok nabor tehnik, ki računalnikom omogočajo, da se učijo iz podatkov in na podlagi njih sklepajo, ne da bi bili izrecno programirani za specifične naloge.<sup>15</sup>

Tehnike, ki omogočajo sklepanje pri gradnji sistema umetne inteligence, vključujejo pristope strojnega učenja, ki se iz podatkov učijo, kako doseči določene cilje, ter logično in znanje-temelječe pristope, ki sklepajo iz kodiranega znanja ali simbolne

---

<sup>9</sup> Uredba (EU) 2024/1689 evropskega parlamenta in sveta, <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng> 12. člen

<sup>10</sup> <https://www.coursera.org/articles/what-is-artificial-intelligence?msocid=26c0a2d04616618637fcad2c478a6017>

<sup>11</sup> <https://www.coursera.org/articles/what-is-artificial-intelligence?msocid=26c0a2d04616618637fcad2c478a6017>

<sup>12</sup> EU komisija (AI Watch): Historical Evolution of Artificial Intelligence, [https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence\\_en](https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence_en), str. 7-8

<sup>13</sup> EU komisija (AI Watch): Historical Evolution of Artificial Intelligence, [https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence\\_en](https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence_en), str. 9-10

<sup>14</sup> EU komisija (AI Watch): Historical Evolution of Artificial Intelligence, [https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence\\_en](https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence_en), str. 11-17

<sup>15</sup> <https://www.ibm.com/think/topics/artificial-intelligence>

predstavitve naloge, ki jo želimo rešiti. Sposobnost sistema UI za sklepanje presega osnovno obdelavo podatkov z omogočanjem učenja, sklepanja ali modeliranja. Izraz 'strojno osnovano' se nanaša na dejstvo, da sistemi umetne inteligence delujejo na strojih. Sklicevanje na eksplicitne ali implicitne cilje poudarja, da lahko sistemi UI delujejo v skladu z izrecno opredeljenimi cilji ali implicitnimi cilji.<sup>16</sup> 12 odstavek uredbe

Ena najbolj priljubljenih vrst algoritmov strojnega učenja pa se imenuje nevronska mreža (ali umetna nevronska mreža). Nevronska mreža je sestavljena iz medsebojno povezanih plasti vozlišč (analognih nevronom), ki sodelujejo pri obdelavi in analizi kompleksnih podatkov.<sup>17</sup>

### 2.1.3 Globoko učenje

Globoko učenje je podskupina strojnega učenja in umetne inteligence, ki uporablja večplastne nevronske mreže, ki se predstave podatkov uči z različnimi stopnjami abstrakcije. Te večplastni omogočajo nenadzorovano učenje: lahko avtomatizirajo pridobivanje značilnosti iz velikih, neoznačenih in nestrukturiranih podatkovnih nizov ter naredijo lastne napovedi o tem, kaj podatki predstavljajo, v primerjavi s klasičnimi algoritmi, ki imajo proces počasnejši.<sup>18</sup>

Arhitektura nevronske mreže globokega učenja vključuje:

- Globoke nevronske mreže, ki se učijo podatkovnih reprezentacij na več ravneh abstrakcije in so osnova večine sodobnih sistemov UI.
- Globoke verske mreže, so generativni modeli, ki se učijo verjetnostnih porazdelitev nad podatki in so sposobne odkrivati skrite vzorce v neoznačenih zbirkah podatkov.
- Ponavljajoče nevronske mreže so zasnovane za obdelavo zaporednih podatkov, kot so besedilo, govor ali časovne vrste, saj si med posameznimi koraki obdelave »zapomnijo« prejšnje stanje — njihova nadgrajena različica, dolgi kratkoročni spomin, pa rešuje problem izgube informacij pri dolgih zaporedjih.
- Konvolucijske nevronske mreže (CNN) so specializirane za obdelavo prostorskih podatkov, zlasti slik, kjer z drsečimi filtri samodejno zaznavajo lokalne vzorce, kot so robovi, oblike in texture, ter jih hierarhično združujejo v vse bolj abstraktne reprezentacije.<sup>19</sup>

Ponazorimo na primeru, kako bi globoko učenje in drugi klasični algoritmi rešili problem prepoznavanja mačk na sliki. Klasični algoritmi bi sliko obdelali z različnimi metodami, da bi identificirali obe očesi, nos, tačke in številne druge dele mačjega telesa. Te metode bi zahtevale zapletene in obsežne programe, medtem ko je pri globokem učenju postopek prepoznavanja slik po izbiri arhitekturnega modele, ki bi v tem primeru najverjetneje bile konvolucijske nevronske mreže, skoraj popolnoma

---

<sup>16</sup> Uredba (EU) 2024/1689 evropskega parlamenta in sveta, <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng> 12. člen

<sup>17</sup> <https://www.coursera.org/articles/what-is-artificial-intelligence?msocid=26c0a2d04616618637fcad2c478a6017>

<sup>18</sup> <https://www.ibm.com/think/topics/artificial-intelligence>

<sup>19</sup> EU komisija (AI Watch): Historical Evolution of Artificial Intelligence, [https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence\\_en](https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence_en), str. 11-12

samodejen. Globoko učenje se prepoznavanja mačk nauči samodejno, potem ko opazuje veliko primerov slik z mačkami in brez njih.<sup>20</sup>

## 2.2 Generativni modeli umetne inteligence

Generativna umetna inteligenca, včasih imenovana tudi " GenAI", se nanaša na modele globokega učenja, ki lahko ustvarijo kompleksno izvirno vsebino. To so lahko oblike besedil, audio, slike in video, ki prej niso obstajali v taki obliki. Sposobnost generacije novega materiala razlikuje GenAI od ostalih modelov UI, ki lahko le sortirajo, iščejo ali filtrirajo informacije brez ustvarjanja nove vsebine.<sup>21</sup>

Najpogosteje uporabljeni generativni modeli UI na področju prava so besedilno naravnani modeli, poznani kot široki jezikovni modeli (angl. Large language models - LLM). Primeri javno dostopnih LLM-ov so Copilot od Microsofta, OpenAI-jev ChatGPT, CLAUDE (Anthropic) in Googlov pomočnik Gemini. Te modeli so denificirani kot UI sistemi za splošni pomen in zaradi njihove narave so lahko uporabljeno preko različnih poslovnih sektorjev za splošno namembnost.<sup>22</sup>

Na splošni ravni generativni modeli kodirajo poenostavljeno predstavitev svojih učnih podatkov in nato iz te predstavitev črpajo ustvarjanje novega dela, ki je podobno a ne identično izvirnim podatkom. Modeli generativne UI se že leta uporabljajo v statistiki za analizo in generiranje bolj kompleksnih podatkovnih tipov. Razvili so se sočasno s pojavom treh prefinjenih vrst modelov globokega učenja:<sup>23</sup>

- Variacijski avto kodirniki ali VAE, ki so bili predstavljeni leta 2013, so omogočili modele, ki so lahko generirali različice vsebine kot odziv na poziv ali ukaz.

Difuzijski modeli, prvič predstavljeni leta 2014, dodajajo "šum" slikam, dokler niso več prepoznavne, nato pa šum odstranijo, da ustvarijo izvirne slike.

- Transformatorji (ti. transformatorski modeli), ki so usposobljeni na sekvenciranih podatkih za generiranje razširjenih zaporedij vsebin (npr. besede v stavkih, okvirji videa, oblika slike ali ukazi v programski kodi). Te so v središču večine sedanjih generativnih orodij za generativno UI, vključno s ChatGPT, GPT-4, Copilot, Gemini, Claude, Bard, BERT in Midjourney.<sup>24</sup>

## 2.3 Pozitivne lastnosti umetne inteligence

UI prodira na vsa področja življenja in prinaša izjemne možnosti — od napredka v medicini, kmetijstvu in izobraževanju do varnejšega prometa ter boljšega upravljanja javnih storitev. (uvodna izjava 1) Namen harmoniziranega pravnega okvira ni zaviranje UI, temveč spodbujanje uvajanja človeško usmerjene in zaupanja vredne umetne inteligence ob hkratni visoki ravni varstva temeljnih pravic.<sup>25</sup> Sisteme umetne inteligence je mogoče enostavno uporabiti v številnih sektorjih gospodarstva in v

---

<sup>20</sup> EU komisija (AI Watch): Historical Evolution of Artificial Intelligence, [https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence\\_en](https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence_en), str. 12 (prvi odst.)

<sup>21</sup> Law Society of Ireland, Guidelines for the Use of Generative Artificial Intelligence by the Legal Profession in Ireland, 2025, str. 2 (3. odst.)

<sup>22</sup> Law Society of Ireland, Guidelines for the Use of Generative Artificial Intelligence by the Legal Profession in Ireland, 2025, str. 2 (4. odst.)

<sup>23</sup> <https://www.ibm.com/think/topics/artificial-intelligence>

<sup>24</sup> <https://www.ibm.com/think/topics/artificial-intelligence>

<sup>25</sup> Uredba (EU) 2024/1689 evropskega parlamenta in sveta, <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng> 1. člen

mnogih delih družbe, tudi čez meje, ter se lahko enostavno širijo po vsej Uniji. Nekatere države članice so že preučile možnost sprejetja nacionalnih pravil, da bi zagotovile, da je UI zaupanja vredna in varna ter da je razvita in uporabljena v skladu z obveznostmi temeljnih pravic.<sup>26</sup>

Umetna inteligenca prinaša številne prednosti, ki pozitivno vplivajo na posameznika, organizacije in celotno družbo:

- Visoka učinkovitost in hitrost: UI obdeluje ogromne količine podatkov v delčku sekunde, kar bistveno presega človekove zmožnosti, omogoča natančnejše napovedi ter zanesljive, na podatkih temelječe odločitve.
- Razpoložljivost 24/7: deluje neprekinjeno, brez potrebe po počitku ali odmoru.
- Natančnost in zmanjšanje napak: pri ponavljajočih se ali podatkovnih nalogah je delež napak bistveno nižji kot pri človeku.
- Avtomatizacija rutinskih nalog: sprošča človeški čas za ustvarjalne in strateške dejavnosti, tako da avtomatizira ponavljajoča in pogosto dolgačasna opravila vključno z digitalnimi nalogami, kot so zbiranje podatkov, vnos in preobdelava ter fizičnimi nalogami, kot so prevzem zalog v skladišču in proizvodni procesi.
- Personalizacija: sistemi UI se prilagajajo posamezniku (npr. priporočila vsebin, prilagojeno izobraževanje).
- Napredek v medicini in znanosti: UI pospešuje odkrivanje zdravil, diagnostiko in obdelavo znanstvenih podatkov.
- Dostopnost storitev: demokratizira dostop do informacij in storitev, ki so bile prej rezervirane za privilegirane.<sup>27</sup>

## 2.4 Negativne lastnosti umetne inteligence

Umetna inteligenca ima negativne plati in njena uporaba lahko ustvari nova tveganja. Številne pomanjkljivosti UI, vključno s kakovostjo podatkov in njihovo reprezentativnostjo, lahko vodijo do pristranske, nepregledne in slabo delujoče umetne inteligence. Zlonamerni akterji lahko UI izkoristijo za ustvarjanje lažnih novic, deepfake vsebin, izvajanje visoko avtomatiziranih kibernetских napadov in podobno. Tveganjem, ki jih prinaša UI, je treba resno pristopiti. Številni etični odbori za umetno inteligenco se ukvarjajo z vprašanji pravičnosti, odgovornosti, preglednosti in razložljivosti UI ter s tem tlakujejo pot za regulacijo na ravni držav in mednarodnih institucij.<sup>28</sup>

Izzivi in tveganja UI modelov:

- Izguba delovnih mest: avtomatizacija ogroža številna poklicna področja, zlasti rutinska in administrativna dela.
- Pristranskost in diskriminacija: modeli so pogosto naučeni na pristranskih podatkih, kar vodi v diskriminatorne odločitve (npr. pri zaposlovanju, kreditnih ocenah).
- Deepfakes: ustvarjanje umetnih vsebin, ki predpostavlja nujnost večje pozornosti in nezaupanja v spletne vsebine. Uporabniki UI modelov za

---

<sup>26</sup> Uredba (EU) 2024/1689 evropskega parlamenta in sveta, <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng> 3. člen

<sup>27</sup> <https://www.coursera.org/articles/what-is-artificial-intelligence?msocid=26c0a2d04616618637fcad2c478a6017>

<sup>28</sup> Uredba (EU) 2024/1689 evropskega parlamenta in sveta, <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>

generiranje ali manipulacijo slikovnih ali avdio ali video vsebin, morajo jasno in razločljivo razkriti, da je bila vsebina umetno ustvarjena ali manipulirana, tako da ustrezno označijo izhod umetne inteligence in razkrijejo njegov umetni izvor.<sup>29</sup>

- Zasebnost in nadzor: obdelava ogromnih količin osebnih podatkov odpira vprašanja varnosti in množičnega nadzora.
- Pomanjkanje transparentnosti: Trenutno pri skoraj vseh generativnih sistemih UI zasledimo t. i. pojav »črne skrinjice«, pri katerem so notranji procesi sklepanja nepregledni in težko razložljivi. Uporabnik na primer ne more vprašati velikega jezikovnega modela, zakaj je podal določen odgovor, saj model svojih izhodnih podatkov dejansko ne »sklepa« niti ne »razume«.<sup>30</sup>
- Odvisnost od tehnologije: prekomerno zanašanje na UI zmanjšuje kritično mišljenje in človeško presojo.
- Varnostna tveganja: UI se lahko zlorabi v kibernetških napadih, vojskovanju ali za manipulacijo javnega mnenja.
- Okoljski odtis: treniranje velikih modelov zahteva ogromne količine energije in vode.<sup>31</sup>

### 3. UPORABA UI V PRAVNEM RAZISKOVANJU

#### 3.1. Pozitivna funkcija uporabe UI v pravu

Sodobno družbo danes zaznamuje hiter tehnološki napredek, zaradi česar je ena izmed najbolj perečih razprav trenutno umetna inteligenca.<sup>32</sup> Sistemi UI so danes že osrednji del vsakega vidika življenja ljudi, razvijajo pa se s svetlobno hitrostjo.<sup>33</sup> Ker UI vpliva tako na življenja posameznih ljudi, kot tudi družbo kot celoto, je eno izmed področij, ki se sooča s tovrstnimi izzivi tudi pravo.<sup>34</sup>

Uporaba orodij UI je lahko v pravnem raziskovanju tudi pozitivna, če so njeni uporabniki dovolj ozaveščeni o uporabi in omejitvah.<sup>35</sup> To velja zlasti za področja povzemanja, prevajanja, pisanja, »brainstorming« idej in načrtovanja strategij.<sup>36</sup> Primerna pa je tudi pri ustvarjanju kontrolnih seznamov, povzemanju gradiva, izboljšanju prepričevanja in preoblikovanju tona.<sup>37</sup> Uporaba UI v pravnem raziskovanju ne prinaša le tveganj, ampak tudi dolg seznam potencialnih koristi pri varni in pravilni uporabi.<sup>38</sup> Avtomatizirano ustvarjanje dokumentov, hitra analiza velikih količin podatkov, enostavnejša komunikacija s strankami, izboljšane pravne raziskave, boljša

---

<sup>29</sup> Uredba (EU) 2024/1689 evropskega parlamenta in sveta, <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng> 134. člen

<sup>30</sup> Law Society of Ireland, Guidelines for the Use of Generative Artificial Intelligence by the Legal Profession in Ireland, 2025, str. 5 (7.odst.)

<sup>31</sup> <https://www.coursera.org/articles/what-is-artificial-intelligence?msoclid=26c0a2d04616618637fcad2c478a6017>

<sup>32</sup> Nemas, 2023, str. 6.

<sup>33</sup> Prav tam.

<sup>34</sup> Prav tam.

<sup>35</sup> Law Society of Ireland, Guidelines for the Use of Generative Artificial Intelligence by the Legal Profession in Ireland, 2025, str. 4.

<sup>36</sup> Prav tam, str. 7.

<sup>37</sup> Prav tam, str. 7.

<sup>38</sup> CCBE guide on the use of generative AI by lawyers, 2025, str. 7.

kakovost dela (zmanjšanje napak in preverjanje skladnosti) so le nekatere izmed koristi, ki jih ponuja uporaba UI v pravu in lahko rezultira v potencialnem prihranku stroškov, hitrejši obdelavi zadev, večjem poudarku na kvantitativnih, ne rutinskih nalogah in izboljšanju storitev za stranke (npr. s hitrejšim odzivnim časom).<sup>39</sup> Orodja UI številnim odvetniškim pisarnam in drugim pravnim službam ponujajo večjo produktivnost, nove zmogljivosti in izboljšanje rezultatov za stranke.<sup>40</sup> Poleg tega pa lahko njihova uporaba omogoča tudi hitrejša pravna raziskave (UI lahko namreč v nekaj sekundah pregleda na tisoče primerov in poudari ustrezne precedense za hitrejša odločanja), avtomatiziran pregled pogodb, zmanjšanje napak (izboljšanje natančnosti), zmanjšanje administrativnega dela in drugo.<sup>41</sup>

### **3.2. Meje dopustne in odgovorne uporabe UI v pravu**

Pravo je temelj vsake družbe, saj ureja odnose med ljudmi, varuje njihove pravice in določa pravila vedenja.<sup>42</sup> Njegov cilj je zagotoviti red, pravičnost in varnost ter zaščititi pravice posameznikov ter organizacij.<sup>43</sup> Dolžnost pravnikov pa je uporaba prava v dobrobit celotne družbe na način, da predstavljajo svoje pravne argumente in teze, skladno s spoštovanjem človekovih pravic in temeljnimi načeli pravne države.<sup>44</sup>

Čeprav pravna tehnologija predstavlja pomembno prelomnico, ni povsem brez pomanjkljivosti in napak, zaradi česar je ključno prilagajanje, učenje in razvijanje kritičnega pogleda.<sup>45</sup> Dopustna in odgovorna uporaba UI v pravnem raziskovanju je lahko zagotovljena ob zavedanju, da so uporabniki odgovorni za informacije, ki jih posredujejo orodju UI (kar lahko krši njihove poklicne obveznosti) in se vzdržijo vnosa osebnih, zaupnih in drugih pomembnih podatkov.<sup>46</sup> Orodja UI lahko pozitivno pripomorejo pravnemu raziskovanju, vendar mora njen uporabnik (npr. odvetnik) opraviti oceno tveganja za pravilno uporabo orodja in odgovorno uporabo v skladu z GDPR, spoštujoč etične in profesionalne standarde.<sup>47</sup>

### **3.3. Tveganja uporabe UI v pravnem raziskovanju**

Kljub temu orodja umetne inteligence na splošno izboljšujejo učinkovitost pravnega svetovanja in storitev, obstajajo tudi številna tveganja pri njihovi uporabi.<sup>48</sup> Takšna tveganja lahko povzročijo tudi negativne posledice pri opravljanju številnih poklicnih

---

<sup>39</sup> CCBE guide on the use of generative AI by lawyers, 2025, str. 3.

<sup>40</sup> Couture, R.J., The Impact of Artificial Intelligence on Law Firms' Business Models, Center on the Legal Profession, 2025.

<sup>41</sup> How Is AI Impacting the Legal Profession? Vanderbilt University, Law School, e-vir.

<sup>42</sup> Osnove prava: Kaj morate vedeti?, 2025, e-vir.

<sup>43</sup> Prav tam.

<sup>44</sup> Letnar Černič, 2009, e-vir.

<sup>45</sup> Kodrin, Strajnar, 2024, e-vir.

<sup>46</sup> Law Society of Ireland, Guidelines for the Use of Generative Artificial Intelligence by the Legal Profession in Ireland, 2025, str. 10.

<sup>47</sup> Prav tam, str. 15.

<sup>48</sup> CCBE guide on the use of generative AI by lawyers, 2025, str. 13.

obveznosti (npr. pri izvrševanju odvetniškega poklica), tovrstna tveganja pa so odvisna od primera do primera in so lahko manjša ali večja.<sup>49</sup>

### **3.3.1. Zasebnost in varstvo osebnih podatkov**

Eno izmed pomembnih tveganj pri uporabi orodij UI predstavlja neustrezna uporaba teh orodij, ki lahko povzroči posege v zasebnost in varstvo osebnih podatkov. Ta področja so v Republiki Sloveniji zagotovljena že s 35. členom Ustave Republike Slovenije, ki zagotavlja nedotakljivost človekove telesne in duševne celovitosti, njegove zasebnosti ter osebnostih pravic.<sup>50</sup> Posameznikova zasebnost pa je varovana tudi s 8. členom EKČP, ki posamezniku omogoča sprožitev postopka proti državi, ki je kršila to pravico pred Evropskim sodiščem za človekove pravice.<sup>51</sup> Ker se orodja UI učijo na ogromnih naborih podatkov, v katere uporabniki vnašajo različne informacije (t.i. vhodni podatki), lahko neustrezna uporaba teh orodij rezultira v razkritju osebnih ali občutljivih podatkov.<sup>52</sup> Uporabniki orodja UI bi lahko, brez jasnih razkritij upravljalcev, pri vnosu informacij nezavedno razkrili občutljive informacije, ki jih sicer niso želeli.<sup>53</sup> Dodatno tveganje predstavlja tudi omogočen dostop upravljalcev orodja do vhodnih in izhodnih podatkov, osebne informacije pa so lahko tudi zavestno ali nevede vključene v nabor podatkov za učenje sistema.<sup>54</sup>

### **3.3.2. Halucinacije**

Halucinacije so rezultat netočnega oziroma nelogičnega odgovora sistema umetne inteligence, ki izvira iz več dejavnikov (npr. omejitve učnih podatkov, verjetnostna narava modelov umetne inteligence, nerazumevanje konteksta in drugo).<sup>55</sup> Orodja UI lahko ustvarijo tudi izmišljeno sodno prakso, sodna mnenja in neobstoječe sodne primere ter oblikujejo povsem izmišljene pravne argumente.<sup>56</sup> Halucinacije se pogosto pojavljajo tudi pri ustvarjanju zakonske razlage z izmišljenimi pravnimi načeli in napačnimi povezavami med pravnimi koncepti.<sup>57</sup> Orodja UI lahko ustvarjajo tudi sklice na znanstvene članke, ki ne obstajajo ali v povzetku navajajo primere, ki v resnici ne obstajajo (npr. leta 2023 je newyorški odvetnik uporabil pravni povzetek s primerom, ki si ga je ChatGBT izmisli).<sup>58</sup>

### **3.3.3. Pristranskost**

Obstaja tveganje, da umetna inteligenca daje prednost ustvarjanju sprejemljivih odzivov pred zagotavljanjem natančnih in kritičnih informacij, kar lahko rezultira v zavajajočih in neuravnoteženih rezultatih.<sup>59</sup>

---

<sup>49</sup> Prav tam.

<sup>50</sup> Pustovrh, 2013, str. 12.

<sup>51</sup> Prav tam.

<sup>52</sup> CCBE guide on the use of generative AI by lawyers, 2025, str. 13.

<sup>53</sup> Prav tam.

<sup>54</sup> Prav tam.

<sup>55</sup> Prav tam, str. 13-14.

<sup>56</sup> Prav tam.

<sup>57</sup> Prav tam.

<sup>58</sup> Choi, Xiaoying Mei, 2025, e-vir.

<sup>59</sup> CCBE guide on the use of generative AI by lawyers, 2025, str. 14.

### **3.3.4. Pomanjkanje preglednosti**

Vsa orodja UI trenutno delujejo na principu tako imenovane »črne skrinjice« kjer so notranji procesi sklepanja nepregledni in težko pojasnljivi.<sup>60</sup> Sistemi UI so mnogokrat nepregledni in je težko vedeti, kako sprejemajo odločitve.<sup>61</sup> Tveganje uporabe UI v pravnem raziskovanju predstavlja pomanjkanje preglednosti, saj uporabnik pogosto ne ve, kako je orodje UI prišlo do določenega pravnega zaključka.<sup>62</sup>

### **3.3.5. Odgovornost za napake UI**

Ena izmed ključnih skrbi pri uporabi UI je vprašanje odgovornosti za morebitne napake in posledice, ki izhajajo iz njenih odločitev.<sup>63</sup> Ko UI povzroči napake, se vprašanje odgovornosti razporedi med različne akterje, vključno z uporabniki, razvijalci in ponudniki storitev.<sup>64</sup> Odgovornost za napake UI na koncu še vedno nosijo človeški akterji (razvijalci in upravljalci sistema), ki prevzamejo pravno in moralno odgovornost, kadar sistem odpove.<sup>65</sup>

### **3.3.6. Intelktualna lastnina**

En izmed aktualnih problematik na področju umetne inteligence je tudi učenje orodij na podlagi podatkov, zaščiteneh z avtorskimi pravicami, ki so lahko tudi kršene, če vhodni podatki, ki vsebujejo avtorska dela, povzročijo prepoznavne izhode.<sup>66</sup>

### **3.3.7. Varstvo demokracije**

UI lahko z ustvarjanjem realističnih lažnih avdio in video vsebin (deepfake) širi dezinformacije, ki lahko ogrozijo volilni proces in demokratični red, zato je nujna uvedba odgovornosti, nadzora in sankcij za njihovo ustvarjanje.<sup>67</sup>

### **3.3.8. Kibernetška varnost**

Danes se gospodarstva preusmerjajo na digitalne in spletne modele, kar pozroča grožnje tradicionalnim pristopom k varnosti podatkov.<sup>68</sup> Orodja UI lahko uvedejo številne nove metode napadov na kibernetško varnost, kot so poškodba podatkov, virov in tako imenovana zastrupitev modelov.<sup>69</sup>

## **4. PRIMER SLABE UPORABE UMETNE INTELIGENCE PRAVNI PRAKSI**

Na pravnem področju se je v zadnjih letih uporaba umetne inteligence kot pomočnika ali pametnega iskalnika razširila. Z vsemi razpoložljivimi podatki so generativni jezikovni modeli postali do te mere sofisticirani, da najdejo vsaj aproksimacijo

---

<sup>60</sup> Prav tam.

<sup>61</sup> Koristi in tveganja umetne inteligence, e-vir.

<sup>62</sup> Izzivi in tveganja pri uporabi umetne inteligence, 2025, e-vir.

<sup>63</sup> Umetna inteligenca v praksi: Kdo je odgovoren za napake, ki jih povzroči umetna inteligenca?, 2025, e-vir.

<sup>64</sup> Prav tam.

<sup>65</sup> Tiwari, 2025, e-vir.

<sup>66</sup> Prav tam, str. 14-15.

<sup>67</sup> Gates, 2023, e-vir.

<sup>68</sup> Keck, Gillani, Dermish, Grossman, Rühmann, 2022, e-vir.

<sup>69</sup> CCBE guide on the use of generative AI by lawyers, 2025

pravega odgovora na pravno vprašanje ali na iskano pravno prakso poiščejo primere sodb in zadev.

Vendar je dolžnost vsakega pravnika in odvetnika, ki to orodje uporablja, da temeljito preveri vir in navedbe za resnične. UI modeli namreč niso še na taki stopnji razvoja, da bi vedno podali pravilen odgovor, ponekod jim podatkov zmanjka in začnejo halucinirati, kar pomeni, da podajajo neresnične, izmišljene rešitve, ki na zunaj zvenijo prepričljivo in strokovno.

Primer slabe uporabe in neopravljene dolžnosti odvetnika je razvpiti ameriški sodni primer, v katerem je bil odvetnik Michael Fourte obtožen uporabe umetne inteligence, ki mu je posredovala napačne in izmišljene informacije. (FUTURISM)

#### 4.1. Ozadje primera

Primer se je odvijal pred Vrhovnim sodiščem države New York. Obrambni odvetnik Michael Fourte je v sodnih dokumentih, ki so se nanašali na sporno posojilo, pustil halucinacijske citate, kar pomeni napačne ali popolnoma izmišljene pravne reference, ki jih je ustvarila umetna inteligenca. Napake je prva odkrila pravna ekipa tožnika, ki je zahtevala, da sodnik Fourteju izreče sankcijo.<sup>70</sup>

Zadeva je pritegnila pozornost sodnika vrhovnega sodišča New Yorka Joela Cohena, ki je primer označil kot "še eno nesrečno poglavje zgodbe o zlorabi umetne inteligence v pravni stroki."<sup>71</sup>

#### 4.2. Potek dogajanja in stopnjevanje napak

Dogajanje v primeru je potekalo v več zaporednih fazah:

- Prva vložitev: Fourtejev obrambni dokument je vseboval netočne ali izmišljene citate in pravne reference, ki jih je generirala umetna inteligenca.<sup>72</sup>
- Povzetek z umetno inteligenco: Ko so ga ujeli pri napaki, je Fourte predložil pojasnilo o svoji uporabi umetne inteligence, ki je bil prav tako napisan z jezikovnim modelom.
- Stopnjevanje: V nasprotovanju predlogu za sankcije je Fourtejev dokument vseboval "več kot dvakrat izmišljenih ali napačnih navedb" kot prvič, je poudaril sodnik Cohen.<sup>73</sup>
- Zanihanje: Fourte sprva ni priznal niti zanihal uporabe umetne inteligence, temveč je napačne navedbe poskušal predstaviti kot "nedolžne parafraze natančnih pravnih načel."
- Delno priznanje: Po nadaljnjem zaslišanju je Fourte priznal uporabo umetne inteligence, a je hkrati skušal del odgovornosti prevaliti na druge sodelavce.
- Zameglitev: Nazadnje je izjavil: "Nikoli nisem rekel, da ne uporabljam umetne inteligence. Rekel sem, da nisem uporabljal nepreverjene umetne inteligence."<sup>74</sup>

---

<sup>70</sup> <https://www.404media.co/lawyer-using-ai-fake-citations/>

<sup>71</sup> <https://www.404media.co/lawyer-using-ai-fake-citations/>

<sup>72</sup> <https://futurism.com/artificial-intelligence/lawyer-caught-using-ai-responds>

<sup>73</sup> <https://futurism.com/artificial-intelligence/lawyer-caught-using-ai-responds>

<sup>74</sup> <https://futurism.com/artificial-intelligence/lawyer-caught-using-ai-responds>

### 4.3. Odziv sodnika in pravna presoja

Sodnik Joel Cohen je Fourtejev argument zavrnil z jasno logično presojo: „Če vključujete citate, ki ne obstajajo, obstaja samo ena razlaga za to. Gre za to, da vam je umetna inteligenca dala kazni, vi pa jih niste preverili. To je definicija nepreverjene umetne inteligence.“<sup>75</sup>

Na koncu je sodnik ugodil tožnikovemu predlogu za sankcije in Fourteju izrekel disciplinske ukrepe.

### 4.4. Širši kontekst

Fourtejev primer ni edinstven. Več deset odvetnikov je bilo ujetih pri podobnih napakah, pri predložitvi napačne ali izmišljene sodne prakse. Med njimi so bili posamezniki, ki so uporabljali splošna orodja, kot je ChatGPT, pa tudi specializirana pravna orodja umetne inteligence, kar kaže, da je problem tehnološko širši in ne vezan le na eno platformo.<sup>76</sup>

Znano je tudi, da je odvetniška pisarna Morgan & Morgan poslala interno opozorilno sporočilo vsem zaposlenim, ko sta bila dva njena odvetnika sankcioniran a zaradi navajanja halucinacij umetne inteligence. Sodniki so medtem opozorili na resnost problema in naredili vzorčne primere iz odvetnikov, ki so bili premalo pazljivi pri zanašanju na jezikovne modele.<sup>77</sup>

### 4.5. NAUK IZ FOURTEJEVEGA PRIMERA

- Odgovornost ostaja pri odvetniku: Umetna inteligenca je orodje, ne porok točnosti. Vsak pravni dokument mora biti neodvisno preverjen iz primarnih virov.
- Halucinacije so resnično tveganje: Jezikovni modeli lahko s prepričljivim slogom podajajo izmišljene sodne zadeve, datume in citate.
- Prikrivanje napak poslabša položaj: Fourtejev primer kaže, da poskusi zanikanja in zamegljevanja vodijo le v večje sankcije.
- Protokoli preverjanja so nujni: Vsaka pravna pisarna mora vzpostaviti jasne postopke nadzora pri uporabi umetne inteligence.

## 5. PRAVNI IN ETIČNI VIDIKI ODGOVORNE UPORABE UMETNE INTELIGENCE

### 5.1. Odgovornost uporabe UI

Odgovorna uporaba pomeni obvladovanje tveganj UI, ki je lahko doseženo s kontrolnim seznamom, ki bi uporabnikom orodja UI v pravnem raziskovanju nudil strukturiran postopek pregleda za zagotovitev natančnega postopka, podprtega z UI.<sup>78</sup> Tako kot se piloti zanašajo na varnostne kontrolne sezname, je tudi v pravnem raziskovanju priporočena uporaba kontrolnega seznama za pregled vsebine, ustvarjene z umetno inteligenco, preden se ta deli bodisi s strankami bodisi z ostalimi uporabniki orodja.<sup>79</sup> Nekatere pomembne točke seznama za odgovorno rabo so:

---

<sup>75</sup> <https://futurism.com/artificial-intelligence/lawyer-caught-using-ai-responds>

<sup>76</sup> <https://www.404media.co/lawyer-using-ai-fake-citations/>

<sup>77</sup> <https://www.404media.co/lawyer-using-ai-fake-citations/>

<sup>78</sup> Keck, Gillani, Dermish, Grossman, Rühmann, 2022, e-vir.

<sup>79</sup> A Practical Checklist for Using AI Responsibly in Your Law Firm, 2026, e-vir.

- uporaba preverjenih orodij UI (npr. odobrenih s strani odvetniške pisarne),
- preverba varnosti in zaupnosti orodja (zagotavljanje ničelne hrambe podatkov),
- preverba dejanske točnosti (ravnanje z rezultati orodja UI kot z osnutkom sodelavca)
- preverba virov (potrditev citiranih primerov, zakonov in predpisov v preverjenih bazah podatkov),
- analiza kakovosti sklepanja (potrebna je tudi analiza logike in ne samo zaključka),
- analiza pristranskosti,
- analiza oblikovanja (pravilno oblikovanje naslovov, napisov, podpisov in drugih pravil).<sup>80</sup>

Upoštevati je treba, da niso vsa orodja UI enaka, zato se tudi tveganja razlikujejo in terjajo različen pristop.<sup>81</sup> Kljub temu da nekatera orodja UI zagotavljajo zaupanja vredne podatkovne baze pravnih raziskav in izpolnjujejo zahteve glede varnosti, preglednosti in natančnosti, je človeški pregled še vedno pomemben, saj so pravniki (in ostali uporabniki UI) tisti, ki ostajajo odgovorni za natančnost in etično skladnost svojega dela.<sup>82</sup>

## **5.2. Skrbnost uporabe UI**

Udeleženci v obligacijskem razmerju morajo pri izpolnjevanju obveznosti iz svoje poklicne dejavnosti ravnati z večjo skrbnostjo, po pravilih stroke in običajih (skrbnost dobrega strokovnjaka).<sup>83</sup> Pravniki, odvetniki, sodniki, notarji in drugi zaposleni v pravni stroki morajo pri izvajanju svojega poklica ravnati z večjo skrbnostjo, zaradi česar se pričakuje tudi skrbna in odgovorna uporaba UI pri njihovem poklicu. Skrbna uporaba pomeni tudi zagotovitev pismenosti na področju UI, ki se ne nanaša le na ponudnike in uvajalce visokotveganih sistemov, temveč tudi na ponudnike in uvajalce vseh sistemov UI.<sup>84</sup> Da lahko uporabniki orodja UI in zaposleni, ki pri pravnem raziskovanju to orodje uporabljajo, sploh lahko ravnajo skrbno, je najprej potrebno informiranje posameznikov o sistemih UI, njihovih priložnostih in tveganjih ter o škodi, ki jo lahko povzročijo.<sup>85</sup>

## **5.3. Varstvo podatkov GDPR**

GDPR, ki opredeljuje pravila o varstvu posameznikov pri obdelavi osebnih podatkov in pravila o prostem pretoku osebnih podatkov, se uporablja ne glede na uporabljeno tehnologijo, če gre za obdelavo osebnih podatkov.<sup>86</sup> Pri uporabi storitev UI sta za

---

<sup>80</sup> Prav tam.

<sup>81</sup> Prav tam.

<sup>82</sup> Prav tam.

<sup>83</sup> OZ, 10. člen.

<sup>84</sup> Splošno o Aktu o umetni inteligenci – kaj je in kdaj se uporablja?, e-vir.

<sup>85</sup> Prav tam.

<sup>86</sup> Višič, 2023, str. 39-40.

skladnost z GDPR odgovorna upravljalec in obdelovalec osebnih podatkov.<sup>87</sup> GDPR in Akt o umetni inteligenci naj bi se med seboj dopolnjevala in zagotavljala regulativni okvir za storitve UI.<sup>88</sup> GDPR tako ureja obdelavo osebnih podatkov s strani obdelovalcev in upravljalcev osebnih podatkov ter si prizadeva za nadzor posameznikov nad njihovimi osebnimi podatki, namen Akta o UI pa je zagotavljanje spoštovanja temeljnih pravic, varnosti in etičnih načel.<sup>89</sup>

#### **5.4. Poštenost, spoštovanje in pravičnost**

Digitalna doba je z novimi algoritmi in UI povzročila preoblikovanje družbenih interakcij, kar je rezultiralo tudi v izzivih pri razumevanju in obravnavanju diskriminacije.<sup>90</sup> Slednja sedaj izhaja tako iz človeških kot tudi iz algoritmičnih pristranskosti, kar lahko rezultira v diskriminatornih odločitvah.<sup>91</sup> Pri stičišču varstva človekovih pravic in UI pa sta pomembna tudi inkriminacija kaznivega dejanja kršitve enakopravnosti po 131. členu KZ-1 in točka b) 5. člena Akta o UI, ki med prepovedane prakse UI uvršča tudi dajanje na trg, v uporabo ali v obratovanje take sisteme UI, ki izkoriščajo šibke točke neprivilegiranih oseb.<sup>92</sup> Te osebe Akt razlikuje po starosti, telesni in duševni invalidnosti in prepoveduje, da bi UI bistveno izkrivila vedenje takšnih oseb in s tem povzročila fizično in/ali psihično škodo njim ali drugim.<sup>93</sup> Akt o UI apelira tudi prepoved razvrščanja oseb s pomočjo UI po kriteriju, koliko so vredne zaupanje v določenem časovnem odboju na podlagi osebnostnih značilnosti ali družbenega vedenja.<sup>94</sup>

#### **5.5. ODGOVORNA UPORABA UI V PRAVU**

Na podlagi ugotovljenih prednosti, tveganj ter pravnih in etičnih vprašanj, predstavljenih v prejšnjih poglavjih, je mogoče oblikovati temeljna načela odgovorne uporabe umetne inteligence pri pravnem raziskovanju. Namen teh priporočil je uporabnikom ponuditi praktične smernice za varno in strokovno uporabo UI v pravni praksi.

Temeljno izhodišče teh smernic je, da umetna inteligenca predstavlja podporno orodje pri pravnem raziskovanju, ne pa nadomestila za pravnikovo strokovno znanje, presojo in odgovornost<sup>95</sup>. Končno odgovornost za pravilnost pravne analize, izbiro pravnih virov in sprejete odločitve vedno nosi uporabnik.

---

<sup>87</sup> GDPR in generativna umetna inteligenca, vodnik za javni sektor, 2024, e-vir.

<sup>88</sup> Prav tam.

<sup>89</sup> Prav tam.

<sup>90</sup> Paternolli, Dalla Preda, Giacobazzi, str.1.

<sup>91</sup> Prav tam.

<sup>92</sup> Ferlinc, 2024, str. 23.

<sup>93</sup> Prav tam.

<sup>94</sup> Prav tam.

<sup>95</sup> Sidra Nasir, Qasim Abbas, Shuo Bai in Riaz Ahmed Khan, *A Comprehensive Framework for Reliable Legal AI: Combining Specialized Expert Systems and Adaptive Refinement*, arXiv, 2024.

### 5.5.1. Človeški nadzor <sup>96</sup>

Pri uporabi UI mora biti vedno zagotovljen ustrezen človeški nadzor. Pravniki ne sme nekritično prevzemati rezultatov, ki jih generira sistem umetne inteligence, temveč jih mora presoditi z vidika njihove pravne pravilnosti, aktualnosti in relevantnosti za konkretni primer.

Rezultati našega testiranja so pokazali, da lahko UI sicer učinkovito identificira relevantna pravna področja in predlaga uporabne vire, vendar se pri zahtevnejših poizvedbah pojavljajo tudi napake, kot so nepopolni odgovori, napačna razlaga pravnih institutov ali pomanjkljiva izbira sodne prakse. Zaradi tega mora uporabnik vse bistvene ugotovitve samostojno preveriti.

Pravnik mora zato UI uporabljati kot pomoč pri raziskovanju in organizaciji informacij, ne pa kot samostojnega odločevalca. Človeški nadzor mora biti zagotovljen v vseh fazah pravnega raziskovanja – od oblikovanja poizvedbe do uporabe rezultatov pri pripravi pravnih dokumentov ali pravnih mnenj.

UI naj služi kot pomoč pri delu, končna pravna odločitev pa mora vedno ostati v pristojnosti pravnega strokovnjaka.

### 5.5.2. Preverljivost rezultatov<sup>97</sup>

Vsak rezultat, pridobljen s pomočjo UI, mora biti preverljiv. Uporabnik mora imeti možnost ugotoviti, na katerih virih temelji odgovor, in samostojno oceniti njihovo zanesljivost.

Pri pravnem raziskovanju je treba prednostno uporabljati uradne in verodostojne pravne vire, kot so zakonodaja, sodna praksa, uradne baze podatkov ter strokovna pravna literatura. Informacij, ki jih UI ne podpre z ustreznimi viri ali jih ni mogoče preveriti, ni dopustno uporabljati kot podlage za pravno argumentacijo.

Posebno previdnost zahtevajo primeri halucinacij, pri katerih sistem ustvari navidežno prepričljive, vendar dejansko napačne informacije. Takšne napake lahko vključujejo neobstoječe sodbe, napačne opravilne številke, nepravilne navedbe zakonov ali zavajajoče povzetke sodne prakse.

### 5.5.3. Načelo sorazmernosti <sup>98</sup>

Obseg uporabe UI mora biti sorazmeren z zahtevnostjo in pomembnostjo konkretne pravne naloge. Čim večji je vpliv posamezne odločitve na pravice in obveznosti posameznikov, tem višja mora biti stopnja previdnosti pri uporabi UI.

UI je lahko zelo koristna pri začetnem raziskovanju pravnega problema, prepoznavanju relevantnih pravnih vprašanj, iskanju zakonodaje ali sodne prakse ter

---

<sup>96</sup> [https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=OJ:L\\_202401689](https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=OJ:L_202401689)

<sup>97</sup> Sidra Nasir, Qasim Abbas, Shuo Bai in Riaz Ahmed Khan, *A Comprehensive Framework for Reliable Legal AI: Combining Specialized Expert Systems and Adaptive Refinement*, arXiv, 2024.

<sup>98</sup> Sidra Nasir, Qasim Abbas, Shuo Bai in Riaz Ahmed Khan, *A Comprehensive Framework for Reliable Legal AI: Combining Specialized Expert Systems and Adaptive Refinement*, arXiv, 2024.

pripravi osnutkov besedil. Pri pripravi končnih pravnih mnenj, sodnih odločb, pogodb ali drugih dokumentov, ki imajo neposredne pravne posledice, pa je nujna bistveno višja stopnja preverjanja in strokovne presoje.

Načelo sorazmernosti pomeni tudi, da mora uporabnik izbrati takšno uporabo UI, ki je primerna glede na namen obdelave in stopnjo tveganja<sup>99</sup>. Avtomatizirano generirane vsebine ne smejo postati edina podlaga za sprejemanje pomembnih pravnih odločitev.

Bolj kot je pravno vprašanje kompleksno ali občutljivo, manjša naj bo odvisnost od avtomatiziranih rezultatov in večji obseg samostojnega preverjanja.

#### **5.5.4. Transparentnost uporabe UI <sup>100</sup>**

Kadar je UI uporabljena pri pripravi pravnih analiz, osnutkov dokumentov ali raziskovalnega dela, mora biti njena uporaba transparentna. Uporabnik, in še posebej naslovnik, mora vedeti, katere naloge je opravil pripravljavec sam in pri katerih si je pomagal z UI.

Transparentnost prispeva k večji odgovornosti uporabnika, omogoča lažje preverjanje rezultatov in zmanjšuje tveganje nekritičnega prevzemanja avtomatiziranih vsebin. Posebej pomembna je v akademskem in raziskovalnem okolju, kjer je treba zagotoviti sledljivost uporabljenih metod dela.

#### **5.5.5. Nenehno izobraževanje uporabnikov <sup>101</sup>**

Odgovorna uporaba UI zahteva tudi ustrezno usposobljenost uporabnikov. Pravniki morajo poznati delovanje uporabljenih sistemov, njihove omejitve, tveganja in možnosti napak.

Ker se tehnologija umetne inteligence hitro razvija, je pomembno stalno spremljanje novih tehničnih, pravnih in etičnih vprašanj. Le ustrezno usposobljen uporabnik lahko kritično presoja kakovost rezultatov in prepreči njihovo nepravilno uporabo.

---

<sup>99</sup> [https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=OJ:L\\_202401689](https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=OJ:L_202401689)

<sup>100</sup> <https://rm.coe.int/1680afae3c>

<sup>101</sup> Sidra Nasir, Qasim Abbas, Shuo Bai in Riaz Ahmed Khan, *A Comprehensive Framework for Reliable Legal AI: Combining Specialized Expert Systems and Adaptive Refinement*, arXiv, 2024.

## **6. VIRI IN LITERATURA**

### **SAMOSTOJNE PUBLIKACIJE**

Council of Bars and Law Societies of Europe. (2025). CCBE guide on the use of generative AI by lawyers.

Dokument Sveta Evrope o UI in ČP. URL: <https://rm.coe.int/1680afae3c> 9.6.2026

Law Society of Ireland. (2025). Guidelines for the Use of Generative Artificial Intelligence by the Legal Profession in Ireland, Law Society of Ireland.

Nemac, L. (2023). Umetna inteligenca in delovno pravo (magistrsko delo). Kraj: Univerza v Ljubljani, Pravna fakulteta.

Pustovrh, T. (2013). Zasebnost in varstvo osebnih podatkov. Kraj: Evropska pravna fakulteta.

Sidra Nasir, Qasim Abbas, Shuo Bai in Riaz Ahmed Khan. (2024). A Comprehensive Framework for Reliable Legal AI: Combining Specialized Expert Systems and Adaptive Refinement

Evropska komisija. (1. 9. 2021). AI WATCH. Pridobljeno 12. 6. 2026 iz Historical Evolution of Artificial Intelligence: [https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence\\_en](https://ai-watch.ec.europa.eu/publications/historical-evolution-artificial-intelligence_en)

Evropski parlament. (13.6.2024). UREDBA (EU) 2024/1689 EVROPSKEGA PARLAMENTA IN SVETA. Pridobljeno 9. 6. 2026 iz Uredba (EU) 2024/1689 Evropskega parlamenta : <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>

### **ČLANKI V REVIJAH**

Couture, R.J. (2025). The Impact of Artificial Intelligence on Law Firms' Business Models. Center on the Legal Profession, 24.2.2025. URL: <https://clp.law.harvard.edu/knowledge-hub/insights/the-impact-of-artificial-intelligence-on-law-law-firms-business-models/>, 9.6.2026.

Ferlinc, A. (2024). Umetna inteligenca z vidika uporabe kazenskega prava. Pravna praksa, št. 13, str. 23.

Keck, M., Gillani, S., Dermish, A., Grossman, J., Rühmann, F. (2022). The role of cybersecurity and data security in the digital economy. Policy Accelerator, junij 2022. URL: <https://policyaccelerator.uncdf.org/all/brief-cybersecurity-digital-economy>, 9.6.2026.

Kudeikina, I., Kaija S. (2024). LIMITS OF THE USE OF ARTIFICIAL INTELLIGENCE IN LAW – ETHICAL AND LEGAL ASPECTS. Environment Technology Resources Proceedings of the International Scientific and Practical Conference 2, str. 188-191.

Višič, V. (2023). 1. konferenca Umetna inteligenca in pravo. Pravna praksa, št. 46-47, str. 39-40.

## **PRAVNI VIRI**

Obligacijski zakonik (OZ). Uradni list RS, št. 97/07 – uradno prečiščeno besedilo, 64/16 – odl. US in 20/18 – OROZ631.

## **SPLETNI VIRI**

A Practical Checklist for Using AI Responsibly in Your Law Firm. ABA Groups, Law Technology Today. URL:

[https://www.americanbar.org/groups/law\\_practice/resources/law-technology-today/2026/checklist-for-using-ai-responsibly-in-your-law-firm/](https://www.americanbar.org/groups/law_practice/resources/law-technology-today/2026/checklist-for-using-ai-responsibly-in-your-law-firm/), 9.6.2026.

Choi, A., Xiaoying Mei, K. What are AI hallucinations? Why AIs sometimes make things up. The conversation, Newsletters, 21.3.2025. URL:

<https://theconversation.com/what-are-ai-hallucinations-why-ais-sometimes-make-things-up-242896>, 9.6.2026.

Cole Stryker, E. K. (NZ). IBM-Kaj je umetna inteligenca'. Pridobljeno 6. 6 2026 iz <https://www.ibm.com/think/topics/artificial-intelligence>

Cole, S. (14. 10 2025). 404Media. Pridobljeno 6. 6 2026 iz <https://www.404media.co/lawyer-using-ai-fake-citations/>

Coursera. (14. 3 2026). Coursera. Pridobljeno 6. 6 2026 iz Kaj je umetna inteligenca? Definicija, uporaba in vrste: <https://www.coursera.org/articles/what-is-artificial-intelligence?msocid=26c0a2d04616618637fcad2c478a6017>

Gates, B. The risks of AI are real but manageable. GatesNotes, 11.7.2023. URL: <https://www.gatesnotes.com/meet-bill/tech-thinking/reader/the-risks-of-ai-are-real-but-manageable>, 11.6.2026.

GDPR in generativna umetna inteligenca, vodnik za javni sektor. Microsoft, april 2024. URL: [https://cdn-dynmedia-1.microsoft.com/is/content/microsoftcorp/microsoft/final/en-us/microsoft-brand/documents/GDPR%20and%20generative%20AI%20-%20A%20Guide%20for%20the%20Public%20Sector\\_SLV..pdf](https://cdn-dynmedia-1.microsoft.com/is/content/microsoftcorp/microsoft/final/en-us/microsoft-brand/documents/GDPR%20and%20generative%20AI%20-%20A%20Guide%20for%20the%20Public%20Sector_SLV..pdf), 9.6.2026.

How Is AI Impacting the Legal Profession? Vanderbilt University, Law School. URL: <https://law.vanderbilt.edu/master-legal-studies/articles/how-is-ai-impacting-the-legal-profession/>, 9.6.2026.

Izzivi in tveganja pri uporabi umetne inteligence. Kabi, 19.1.2025. URL: <https://www.kabi.info/Novice/izzivi-in-tveganja-pri-uporabi-umetne-inteligence/>, 11.6.2026.

Kodrin, L., Strajnar, B. Pravna tehnologija: trend ali nuja za študente prava? Revija Pamfil, 21.2.2024. URL: <https://revijapamfil.si/clanki/2024/2/21/pravna-tehnologija-trend-ali-nuja-za-tudente-prava>, 11.6.2026.

Koristi in tveganja umetne inteligence. Evropski svet, Svet Evropske unije. URL: <https://www.consilium.europa.eu/si/policies/benefits-and-risks-of-ai/#risks>, 11.6.2026.

Landymore, F. (15. 10 2025). Futurism. Pridobljeno 6. 6 2026 iz Odvetnik je ujet pri uporabi UI: <https://futurism.com/artificial-intelligence/lawyer-caught-using-ai-responds>

Letnar Černič, J. Pravo, družba in javna dobrobit. IUS-INFO, 30.10.2009. URL: <https://www.iusinfo.si/medijsko-sredisce/kolumne/48987>, 11.6.2026.

Osnove prava: Kaj morate vedeti? ODVETNIŠKA PISARNA POLIČ. URL: <https://op-polic.si/2025/04/16/in-hac-habitasse-platea-dictumst/>, 11.6.2026.

Paternolli, V., Dalla Preda, M., Giacobazzi, R. Fairness of AI Systems in the Legal Context. URL: <https://buildtrust.it/wp-content/uploads/2025/04/Fairness-of-AI-Systems-in-the-Legal-Context.pdf>, 9.6.2026.

Splošno o Aktu o umetni inteligenci – kaj je in kdaj se uporablja? Informacijski pooblaščenec. URL: <https://www.ip-rs.si/umetna-inteligenca/pogosta-vprašanja/>, 9.6.2026.

Tiwari, A.. Who Is Responsible When AI Makes a Mistake? RISKINFO.AI, 11.7.2025. URL: <https://www.riskinfo.ai/post/who-is-responsible-when-ai-makes-a-mistake>, 11.6.2026.

Umetna inteligenca v praksi: Kdo je odgovoren za napake, ki jih povzroči umetna inteligenca? Law & More attorneys, 2.10.2025. URL: <https://lawandmore.si/umetna-inteligenca-v-praksi--kdo-je-odgovoren-za-napake--ki-jih-povzroci-umetna-inteligenca/>