



Marek Świerczyński^{a)}

 <https://orcid.org/0000-0002-4079-0487>

Duże modele językowe a prawo prywatne międzynarodowe

Abstract: This paper examines the functioning of large language models (LLMs) in the field of private international law. Its primary objective is to assess their effectiveness in determining the applicable law, particularly in the analysis of choice-of-law clauses, and to identify structural limitations arising from the architecture of these models. The study applies a dogmatic-legal and functional methodology, drawing on examples from legal practice, including due diligence processes and the use of AI systems in law firms. The central thesis is that the challenges associated with LLMs in private international law are structural in nature and stem from their probabilistic design, including tokenization, hallucinations, and semantic averaging. These features create a risk of misidentifying conflict-of-law rules, misinterpreting choice-of-law clauses, and introducing systemic bias toward dominant legal systems. The conclusions indicate that despite significant technological progress, particularly in reasoning models and retrieval-augmented generation (RAG), LLMs cannot replace legal professionals in applying conflict-of-law rules. Their use requires continuous human oversight, and the development of specialized, locally trained models may provide a more suitable path for private international law practice.

Keywords: large language models — private international law — applicable law — choice-of-law clauses — artificial intelligence — hallucinations — tokenization — retrieval-augmented generation — human oversight

^{a)} Prof. dr hab., Uniwersytet Kardynała Stefana Wyszyńskiego w Warszawie, m.swierczynski@uksw.edu.pl.

1. Wstęp

Istotną częścią moich zadań jako młodego prawnika pracującego w międzynarodowej kancelarii prawnej był udział w złożonych projektach *due diligence*. Ich celem było ustalenie ryzyka prawnego w oparciu o badanie dokumentacji prawnej działających w skali transgranicznej spółek, np. przed dokonaniem ich sprzedaży. W praktyce oznaczało to konieczność wydrukowania tysięcy stron umów, protokołów i innych dokumentów korporacyjnych, a następnie ich uważne przeczytanie, strona po stronie, z zakreślaczem w dłoni, starając się ustalić główne ryzyka prawne. Kluczowe było zidentyfikowanie klauzul wyboru prawa właściwego, odesłań do różnych systemów prawnych, klauzul jurysdykcyjnych i arbitrażowych. Przeoczenie takich postanowień, zwłaszcza w złożonych strukturach holdingowych, mogło mieć poważne konsekwencje dla powodzenia transakcji, począwszy od niespójności regulacyjnej po istotne ryzyka procesowe w kilku państwach jednocześnie.

Z czasem w projektach *due diligence* zaczęto wprowadzać tzw. wirtualne przestrzenie danych (ang. *virtual data rooms*). Nie trzeba już było drukować dokumentów, lecz nadal należało je czytać na ekranie. Jednak ponieważ dane te miały postać cyfrową, zaczęto stopniowo stosować narzędzia służące automatyzacji, w tym wyszukiwarki słów kluczowych czy klasyfikatory dokumentów¹. Każdy kolejny etap takiej cyfryzacji zmniejszał fizyczny trud pracy, ale istota pozostawała ta sama — to prawnik czytał dokumenty, analizował je i ocenił ryzyka prawne.

Obecnie coraz częściej powyższe zadania wykonują systemy sztucznej inteligencji oparte na dużych modelach językowych (ang. *Large Language Models* — LLM)². Prawnika, który nadal ponosi pełną

¹ Postawione przez P. Machnikowskiego 10 lat temu przewidywania okazały się trafne: P. Machnikowski, *Prawo zobowiązań w 2025 roku. Nowe technologie, nowe wyzwania*, w: *Współczesne problemy prawa zobowiązań*, red. A. Olejniczak et al., Wolters Kluwer, Warszawa 2015, s. 380 i n.

² Model językowy jest modelem zdolności ludzkiego mózgu do tworzenia języka naturalnego. Modele językowe są przydatne w przypadku wielu zadań, w tym do rozpoznawania mowy, tłumaczenia maszynowego, generowania języka naturalnego (generowania tekstu przypominającego tekst pisany przez człowieka), optycznego rozpoznawania znaków (OCR) czy wyszukiwania informacji. Duże modele językowe (LLM) to obecnie najbardziej zaawansowana forma modeli językowych. Są to modele językowe wyszkolone na dużej ilości danych, umożliwiające generowanie tekstu oraz realizację zadań związanych z przetwarzaniem języka naturalnego. Przykładami dużych modeli językowych są modele z serii GPT zbudowane przez OpenAI, używane w chatbotach ChataGPT

odpowiedzialność za wynik przeprowadzonej analizy, sprawuje nad nimi jednak coraz mniejszą realną kontrolę³. Dochodzimy do momentu, gdy tzw. systemy agentowe sztucznej inteligencji (ang. *agentic AI*)⁴ są w stanie samodzielnie wyszukiwać dokumenty w bazach danych badanej spółki, klasyfikować je według rodzaju i prawa właściwego, identyfikować klauzule ryzyka oraz generować raporty. Rodzi to pytania nie tylko o jakość analizy prawnej, lecz także o cyberbezpieczeństwo i ochronę poufności⁵. Z perspektywy prawa prywatnego międzynarodowego pojawia się natomiast pytanie, czy systemy te docierają do właściwych źródeł, czy potrafią prawidłowo ustalić kryteria wpływające na właściwość prawa danego państwa, w szczególności ocenić skuteczność klauzul wyboru prawa, a także ich konsekwencje, w tym ryzyka prawne.

Celem niniejszego opracowania jest wyjaśnienie, w jaki sposób działają duże modele językowe w kontekście prawa prywatnego międzynarodowego, jakie mają ograniczenia, w szczególności przy przeprowadzaniu analizy klauzul wyboru prawa, oraz jakie wnioski z tego wynikają nie tylko dla praktyki, lecz także dla dalszego rozwoju prawa prywatnego międzynarodowego⁶. Jest to tym bardziej istotne, ponieważ z narzędzi tych zaczynają korzystać również sędziowie⁷, a nawet

i Microsoft Copilot, a także modele Llama zbudowane przez Meta Platforms. Istnieją również chińskie modele, jak DeepSeek, czy polskie, jak Bielik i PLLuM. Zob. szerzej: P. Homoki, Z. Zsolt, *Large Language Models and Their Possible Uses in Law*, „Hungarian Journal of Legal Studies” 2023, vol. 64, no. 3, s. 436–439; M. Dahl et al., *Large Legal Fictions: Profiling Legal Hallucinations in Large Language Models*, „The Journal of Legal Analysis” 2024, vol. 16, s. 66.

³ Por. M. Beane, *The Skill Code: How to Save Human Ability in an Age of Intelligent Machines*, Harper Collins, New York 2024, s. 16.

⁴ W uproszczeniu różnica między chatbotem a agentem AI jest taka, że chatbot daje użytkownikowi odpowiedź, a agent działa na jego rzecz. H. Naveed et al., *A Comprehensive Overview of Large Language Models*, „ACM Transactions on Intelligent Systems and Technology” 2025, vol. 16 (5), s. 32–33.

⁵ Badania na autonomicznych agentach wykazały szereg luk w zabezpieczeniach, w tym nieautoryzowaną zgodność z instrukcjami osób trzecich oraz przypadkowe ujawnianie danych wrażliwych osadzonych w korespondencji. Zob. głośny raport: N. Shapira et al., *Agents of Chaos. Report, Study Timeline Feb 2–22, 2026*, <https://agentsofchaos.baulab.info/report.html> [dostęp: 1.06.2026].

⁶ Do przygotowania tego tekstu zainspirowała mnie znana publikacja H. Surdena pt. *ChatGPT, Large Language Models, and Law*, „Fordham Law Review” 2024, vol. 92 (5), s. 1942–1972. Autor odwołuje się w niej m.in. do zagadnień kolizyjnych analizowanych przez systemy sztucznej inteligencji. Od czasu publikacji tego artykułu Surdena nastąpiły przełomowe zmiany technologiczne, bezpośrednio wpływające na sposób działania tych systemów, w szczególności rozwój funkcji rozumowania.

⁷ M. Świerczyński, *Generatywna sztuczna inteligencja w służbie wymiaru sprawiedliwości w świetle wytycznych Rady Europy (AIAB CEPEJ)*, w: *Ius Scientiam*

stają się one częścią rządowych planów reformy wymiaru sprawiedliwości, również w Polsce⁸.

Teza opracowania jest następująca. Problemy związane z wykorzystaniem sztucznej inteligencji w dziedzinie prawa prywatnego międzynarodowego nie wynikają jedynie z niewłaściwego doboru danych treningowych⁹, ale z samej natury działania modeli językowych, w tym tzw. tokenizacji, czyli dzielenia wyrazów na mniejsze jednostki¹⁰, oraz tzw. halucynacji (czyli błędów). Ponadto obecne modele językowe opierają się na danych pochodzących przede wszystkim z dominujących w rozwoju technologicznym państw (USA oraz Chiny), co może prowadzić do ukrytej dyskryminacji mniej reprezentowanych w danych treningowych systemów prawnych¹¹, w tym europejskich porządków prawnych. Rosnący stopień wykorzystania sztucznej inteligencji w praktyce prawniczej, także w sądach, może mieć wpływ na zmianę kierunku rozwoju prawa prywatnego międzynarodowego¹², powodując, bardziej lub mniej ukryte, wypieranie europejskich koncepcji kolizyjnoprawnych przez rozwiązania amerykańskie bądź chińskie¹³. Jest to zagadnienie

Persequitur. Księga jubileuszowa z okazji XXV-lecia Wydziału Prawa i Administracji Uniwersytetu Kardynała Stefana Wyszyńskiego w Warszawie, red. M. Badowski, J. Wesserling, L. Karski, Wolters Kluwer, Warszawa 2025, s. 287—302. O wykorzystywaniu dużych modeli językowych przez chińskich sędziów: J. Liu, L. Xueyao, *How Do Judges Use Large Language Models? Evidence from Shenzhen*, „Journal of Legal Analysis” 2024, vol. 16, s. 235—262. Zob. także UNESCO, *AI Essentials for Judges*, 2026, CI/DIT/AI/Judges2026, <https://unesdoc.unesco.org/ark:/48223/pf0000396991> [dostęp: 25.07.2026].

⁸ Polityka rozwoju sztucznej inteligencji w Polsce do 2030 roku, <https://www.gov.pl/web/cyfrizacja/polityka-rozwoju-sztucznej-inteligencji-w-polsce-do-2030-roku--bezpieczne-i-odpowiedzialne-wykorzystanie-sztucznej-inteligencji> [dostęp: 1.06.2026].

⁹ E. Mik, *Caveat Lector: Large Language Models in Legal Practice*, „Rutgers Business Law Review” 2024, vol. 19, no. 2, s. 123—125.

¹⁰ W uproszczeniu jest to podstawowa jednostka obliczeniowo-rozliczeniowa dużych modeli językowych. AI tworzy je, wykonując obliczenia na akceleratorach, gdy zadajemy jej pytanie lub wydajemy polecenie. Gdy mówimy o generowaniu tekstu, jeden token to średnio około czterech znaków, czyli odpowiada niewielkim ilościom danych.

¹¹ Bias (skrzywienie algorytmiczne): R. Ghoshal, *Cross-Border AI Governance for Legal Tech: Standardizing Ethical and Legal Norms in Access to Justice*, „International Journal of Law, Justice and Jurisprudence” 2025, vol. 5 (1), s. 187. Zob. także M. Coeckelbergh, *AI Ethics*, MIT Press, Cambridge 2020. Jak zauważa autor, mechanizm tworzenia systemów AI wynika nie tylko z zastosowanej technologii, lecz także z istniejących relacji władzy oraz asymetrii w dostępie do danych.

¹² I szerzej, prawa międzynarodowego publicznego: A. Deeks, D. Hollis, *Large Language Models and International Law*, „Chicago Journal of International Law” 2025, vol. 26, no. 1, s. 117—144.

¹³ Zob. szerzej M. Świerczyński, Z. Więckowski, *Sztuczna inteligencja w prawie międzynarodowym. Rekomendacje wybranych rozwiązań*, Difin, Warszawa 2021.

istotne także w kontekście wymaganego od prawników (kolizjonistów, sędziów, pełnomocników procesowych) stopnia ludzkiego nadzoru przy korzystaniu z systemów sztucznej inteligencji¹⁴. Obowiązek sprawowania takiego nadzoru (ang. *human oversight*) wynika z przepisów unijnego AI Akt¹⁵ oraz konwencji ramowej Rady Europy o sztucznej inteligencji¹⁶.

2. Jak działają duże modele językowe?

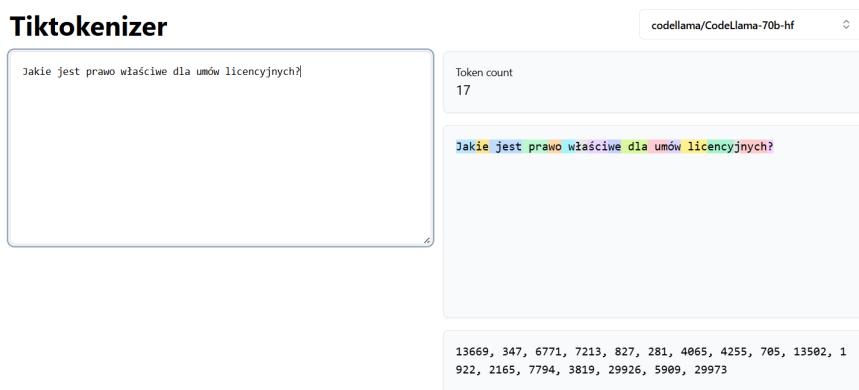
Duże modele językowe, takie jak powszechnie wykorzystywane GPT, Claude czy Gemini, działają w istocie jako bardzo złożone systemy przewidywania następnego słowa. Użytkownik wprowadza zapytanie (tzw. prompt) w rodzaju: „Jakie prawo jest właściwe dla umów licencyjnych?” (wgrywając dokument umowy do systemu AI), a model generuje odpowiedź, przewidując jej najbardziej prawdopodobne brzmienie na podstawie wzorców wyuczonych uprzednio na ogromnych zbiorach tekstów. Ponadto systemy te operują tokenami, czyli fragmentami słów. Może to prowadzić do sytuacji, w której model „rozpoznaje” poszczególne człony nazwy, ale traci ich unikalne znaczenie prawne, co jest szczególnie problematyczne przy analizie kolizyjnoprawnej.

¹⁴ A. Wiśniewski, *Sztuczna inteligencja i prawa człowieka w kontekście prawa międzynarodowego*, „Prawo i Więź” 2024, t. 47, nr 4, s. 35; B. Wang, *Prompts and Large Language Models: A New Tool for Drafting, Reviewing and Interpreting Contracts?*, „Law, Technology and Humans” 2024, vol. 6, no. 2, s. 102. Na temat zjawiska „shadow learning” oraz utraty więzi uczeń—mistrz w dobie automatyzacji: M. Beane, *The Skill Code...*, s. 116—120.

¹⁵ Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonized rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act), PE/24/2024/REV/1.

¹⁶ Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, otwarta do podpisu 5 września 2024 r., Council of Europe Treaty Series — No. 225. Konwencja jest pierwszym wiążącym traktatem międzynarodowym w dziedzinie sztucznej inteligencji.

Tiktokenizer



Rys.1. Zapytanie użytkownika jest rozbijane na fragmenty (tokeny) w formacie zrozumiałym dla systemu AI. Na ich podstawie system przewiduje statystycznie najbardziej prawdopodobną odpowiedź

Źródło: opracowanie własne przy wykorzystaniu narzędzia tiktokenizer.

Model nie „rozumie”, czym jest prawo właściwe albo czym jest norma kolizyjna. Na podstawie przeprowadzonej wcześniej w ramach treningu analizy miliardów dokumentów uczy się, że po pewnych sekwencjach słów („prawo właściwe”, „rozporządzenie Rzym I”, „umowa konsumentska”) zwykle pojawiają się inne sekwencje („art. 6”, „miejsce zwykłego pobytu konsumenta”, „ograniczenia autonomii woli stron”). Jego „inteligencja” polega na statystycznym dopasowywaniu, nie na rozumieniu w sensie kognitywnym¹⁷. Fundamentem jest semantyka dystrybucyjna, czyli założenie, że słowo definiuje jego sąsiedztwo¹⁸. Modele te nie odróżniają prawdy od prawdopodobieństwa statystycznego. Mogły nie zapoznać się w czasie treningu z klasycznym opracowaniem Henryka Trammera o strukturze normy kolizyjnej prawa prywatnego międzynarodowego¹⁹ (ponieważ nie jest on łatwo dostępny w Internecie), a nawet jeśli artykuł ten stanowił część korpusu treningowego, to i tak nic z niego nie zrozumiały.

¹⁷ Modele mogą przeprowadzać symulację rozumowania, ponieważ język naturalny sam w sobie koduje struktury logiczne i algorytmiczne (np. relacje „starszy niż”, „jeśli..., to...”). Modele uczą się imitować te wzorce, co sprawia, że ich odpowiedzi wyglądają na efekt procesów poznawczych, podczas gdy w rzeczywistości są jedynie statystycznym odzwierciedleniem struktur zawartych w ogromnych korpusach tekstowych wykorzystywanych do ich treningu.

¹⁸ Semantyka dystrybucyjna to dziedzina zajmująca się określaniem znaczenia słów na podstawie ich użycia w kontekście leksykalnym i gramatycznym. Opiera się na hipotezie, że słowa występujące w podobnym otoczeniu w dużych zbiorach danych tekstowych mają podobne znaczenie, reprezentowane najczęściej przez wielowymiarowe wektory.

¹⁹ H. Trammer, *Z rozważań nad strukturą normy kolizyjnej prawa prywatnego międzynarodowego*, „Studia Cywilistyczne” 1969, t. XII—XIV.

Dlaczego zatem w wielu przypadkach odpowiedzi udzielane prawnikom przez systemy sztucznej inteligencji wydają się tak trafne?²⁰ Przełomem w rozwoju modeli językowych były wektory słów (ang. *word embeddings*)²¹, czyli numeryczne reprezentacje znaczenia wyrazów w wielowymiarowej przestrzeni. Słowa o podobnym czy wręcz takim samym znaczeniu, jak choćby „prawo właściwe” i „applicable law”, są w tej przestrzeni blisko siebie, co pozwala modelowi skutecznie skojarzyć różne wyrażenia odnoszące się do podobnych pojęć. W prawie prywatnym międzynarodowym rodzi to jednak istotny problem. Wiele z nich ma bowiem różne znaczenie w poszczególnych systemach prawnych. Domicyl w prawie angielskim, francuskim i amerykańskim obejmuje odmienne konstrukcje. Siedziba spółki może być inaczej rozumiana w prawie unijnym (np. w rozporządzeniu Rzym I) niż w prawie krajowym. Model uznaje „miejsce zamieszkania” oraz „zwykły pobyt” za bliskie wektory, ponieważ często występują w podobnych kontekstach, mimo że ich skutki kolizyjnoprawne są różne. Model, który „uśrednia” znaczenia na podstawie rozproszonych tekstów, nie rozpoznaje tych subtelnych różnic²², a to właśnie one stają się kluczowe dla prawidłowego zastosowania norm kolizyjnych. Zatem odpowiedzi udzielane przez modele językowe jedynie początkowo wydają się prawidłowe, a dopiero uważna analiza umożliwia dostrzeżenie nieprawidłowości.

Wraz z wprowadzaniem nowych, coraz bardziej złożonych modeli językowych odpowiedzi udzielane przez takie systemy stają się bardziej zniuansowane, z odniesieniami do przepisów, orzecznictwa, a nawet doktryny (przynajmniej tej znajdującej się w otwartym dostępie, czego przykładem mogą być publikacje zawarte w „Problemach Prawa Prywatnego Międzynarodowego”). Takie zmiany umożliwiło wprowadzenie architektury Transformer, która jest oparta na mechanizmie uwagi (ang. *self-attention*)²³. Pozwala ona modelowi identyfikować, które słowa i fragmenty tekstu są kluczowe dla rozwiązania danego problemu²⁴. Przykładowo, przy zapytaniu: „Umowa została zawarta w Warszawie przez polską spółkę z niemieckim kontrahentem. Umowa zawiera klauzulę wyboru prawa niemieckiego. Jakie jest prawo właściwe dla zobowiązań pozaumownych związanych z tą transakcją?”, model powinien „zwrócić uwagę” m.in. na wyrażenia: „zobowiązania pozaumowne”, „polską spółkę” czy „niemieckim kontrahentem”. Od tego, jak powiąże je

²⁰ Autorka podkreśla: „LLMs may therefore appear more capable than they actually are”. E. Mik, *Caveat...*, s. 70.

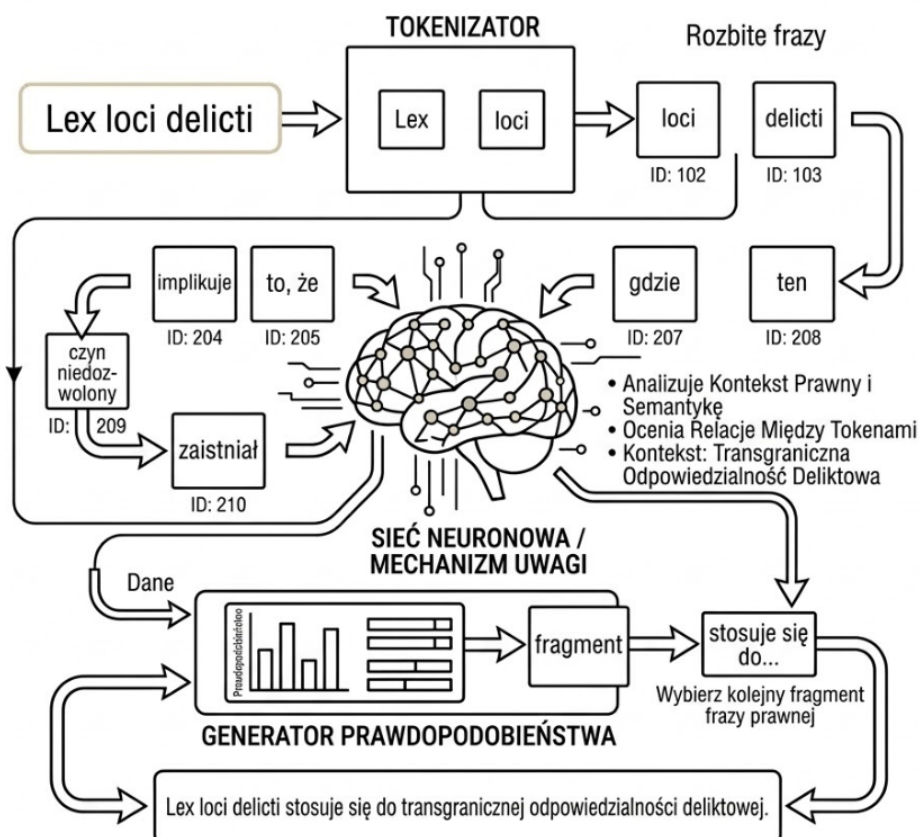
²¹ H. Surden, *ChatGPT...*, s. 1965—1966.

²² E. Mik, *Caveat...*, s. 70 i n.

²³ H. Surden, *ChatGPT...*, s. 1962—1965.

²⁴ B. Wang, *Prompts...*, s. 90.

z rozporządzeniem Rzym II i orzecznictwem TSUE, zależy kompletność odpowiedzi.



Rys. 2. Ilustracja działania sieci neuronowej z mechanizmem uwagi na przykładzie zapytania dotyczącego odpowiedzialności pozaumownej

Źródło: opracowanie własne.

Z najnowszych analiz architektury Transformer wynika jednak, że przy bardzo długich i złożonych dokumentach dochodzi do zjawiska „rozmycia uwagi” (ang. *attention dilution*). Model, starając się powiązać różne okoliczności, może błędnie przypisać skutek prawny z jednej klauzuli do stanu faktycznego opisanego w innym miejscu dokumentu. W prawie prywatnym międzynarodowym, gdzie jedno zdanie odnoszące się do siedziby strony umowy może mieć kluczowe znaczenie, taki błąd statystyczny staje się krytycznym błędem prawnym. Modele te mogą podać błędny wynik nawet przy prostych zadaniach, jeśli wzorzec statystyczny w danych treningowych przeważa nad logiką problemu.

3. Dlaczego pierwsze modele językowe nie radziły sobie ze złożonością prawa prywatnego międzynarodowego?

Prawo prywatne międzynarodowe jest dziedziną wyspecjalizowaną, o relatywnie niewielkim wolumenie tekstów w porównaniu do prawa materialnego czy procesowego. W ogólnym dostępie dominują przy tym źródła anglojęzyczne. Dane dotyczące rozporządzeń Rzym I i Rzym II, konwencji haskich czy krajowych regulacji kolizyjnych (jak polska ustawa z 2011 r.) stanowią niewielki ułamek w globalnych korpusach treningowych wykorzystywanych na potrzeby rozwoju sztucznej inteligencji.

W konsekwencji wczesne modele językowe (np. przełomowy dla świata GPT 3.5)²⁵ miały bardzo ograniczoną ekspozycję na typowe zadania z zakresu prawa prywatnego międzynarodowego, jak choćby ocenę klauzul wyboru prawa, samo rozróżnienie jurysdykcji krajowej i wyboru prawa właściwego, stosowanie odesłania, identyfikację przepisów wymuszających swoją właściwość itp. Skutkowało to odpowiedziami udzielanymi użytkownikom, które były po prostu błędne lub oparte na uproszczonych skojarzeniach pochodzących z systemów *common law*²⁶. Jeszcze parę lat temu na zajęciach z prawa prywatnego międzynarodowego pokazywałem studentom te odpowiedzi, aby wykazać, jak dalece systemy sztucznej inteligencji nie radzą sobie ze złożonością kasusów kolizyjnoprawnych.

W literaturze z tego okresu wady tych modeli językowych podkreślano za pomocą barwnego sformułowania o „stochastycznych papugach”²⁷, które tworzą pozór rozumienia poprzez, często przypadkową, kompilację różnych fragmentów tekstu. Zapytanie w rodzaju: „Wyjaśnij skutki klauzuli »This agreement shall be governed by the laws of Switzerland« dla załączonej umowy zawartej pomiędzy polskim i niemieckim przedsiębiorcą” przekraczało możliwości ówczesnych systemów. Odpowiedź mogła zawierać ogólniki np. o stabilności prawa szwajcarskiego, całkowicie ignorując wymiar kolizyjnoprawny²⁸. Koncepcja „stochastycznych

²⁵ A. Bent, *Large Language Models: AI's Legal Revolution*, „Pace Law Review” 2023, vol. 44, no. 1, s. 120.

²⁶ O powodach szerzej: B. Shneiderman, *Human-Centered AI*, Oxford University Press, Oxford 2022, s. 53–68.

²⁷ H. Surden, *ChatGPT...*, s. 1948; B. Wang, *Prompts...*, s. 90; Niewątpliwie nawiązuje to do spostrzeżenia D. Diderota: „S'il se trouvait un perroquet qui répondit à tout, je prononcerais sans balancer que c'est un être pensant”. D. Diderot, *Œuvres complètes*, éd. J. Assézat, M. Tourneux, Garnier, Paris 1876, s. 204.

²⁸ Podobnie proste odpowiedzi możemy uzyskać też obecnie, gdy będziemy korzystać z lokalnych modeli językowych zainstalowanych na naszym urządzeniu.

papug” znajduje odzwierciedlenie w zjawisku tzw. długodystansowych zależności (ang. *long-range dependencies*)²⁹. Poprzednie modele językowe napotykały trudności w utrzymaniu spójności w tekście, gdzie klauzula wyboru prawa zamieszczona na końcu umowy wywierała wpływ na obowiązki stron opisane w treści całego dokumentu. Dopiero nowa architektura modeli językowych pozwoliła na „ważenie” relacji między odległymi fragmentami tekstu, co jest kluczowe dla wieloetapowej analizy kolizyjnej.

4. Postęp po 2025 r. — obecne modele językowe a prawo prywatne międzynarodowe

Od przełomu związanego z wprowadzeniem pierwszych modeli językowych nastąpił gwałtowny wzrost zdolności „rozumowania” (ang. *chain-of-thought reasoning*, czyli rozumowanie krok po kroku). Ewolucja ta stanowi przejście od intuicyjnego „systemu 1” (myślenia szybkiego) do „systemu 2” (myślenia wolnego/analizycznego), co pozwala już na imitację wieloetapowego wnioskowania prawniczego. Takie systemy zaczęły być powszechnie stosowane w sektorze prawnym, począwszy od 2025 r. Tworzone są wyspecjalizowane systemy prawniczej sztucznej inteligencji, czego przykładem w Polsce jest Libra bądź Lex Expert AI. Dotychczasowe systemy informacji prawnej, takie jak Lex czy Legalis, muszą być uzupełniane przez komponenty sztucznej inteligencji, aby nie utracić swojej pozycji rynkowej.

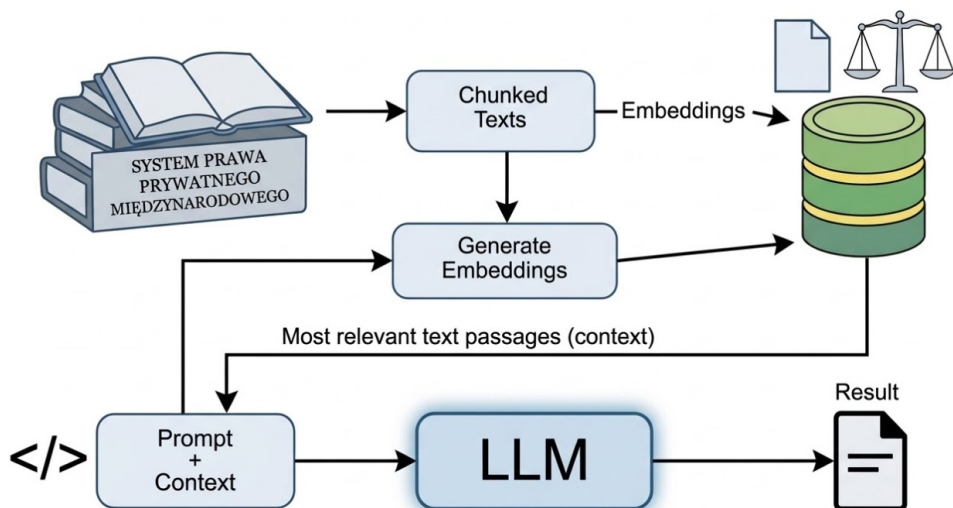
Obecnie wykorzystywane modele językowe potrafią, przynajmniej dla wielu typowych scenariuszy, prawidłowo odwołać się do przepisów kolizyjnoprawnych, a także orzecznictwa i doktryny. Kluczową innowacją w zastosowaniach prawnych jest technika *Retrieval-Augmented Generation* (RAG)³⁰. Polega ona na tym, że przed wygenerowaniem odpowiedzi model otrzymuje dodatkowo fragmenty tekstów wyszukanych w zewnętrznej prawniczej bazie danych bądź przedstawionych bezpośrednio przez użytkownika³¹. Model „widzi” więc w oknie kontekstowym nie tylko pytanie, lecz także właściwe przepisy mające zastosowanie

²⁹ M. Beane, *The Skill Code...*, s. 92.

³⁰ E. Mik, *Caveat...*, s. 113.

³¹ D. Schwarcz et al., *AI-Powered Lawyering: AI Reasoning Models, Retrieval Augmented Generation, and the Future of Legal Practice*, „Journal of Law and Empirical Analysis” 2026, vol. 3 (1), s. 220—250, <https://dx.doi.org/10.2139/ssrn.5162111>.

w danej sprawie, orzecznictwo oraz doktrynę. Budzi to skojarzenie z obowiązkiem sędziego, który stosując prawo obce, musi uwzględnić nie tylko brzmienie przepisu, ale też obce orzecznictwo oraz poglądy obcej doktryny. Rodzi to pokusę sięgnięcia do modelu językowego zamiast powołania kosztownego biegłego (z czym mamy już do czynienia również w Polsce)³².



Rys. 3. Ilustracja obrazująca wykorzystanie danych kontekstowych przez system sztucznej inteligencji. Oznaczenie </> ilustruje prompt (instrukcję) od użytkownika

Źródło: opracowanie własne.

³² Sposób wykorzystania systemu sztucznej inteligencji przez sędziego został wyjaśniony w postanowieniu z 23 stycznia 2024 r. (sygn. akt XXII GW 219/20). W uzasadnieniu sąd wskazał: „Sąd z urzędu ustala i stosuje jedynie właściwe prawo obce (art. 51a § 1 zd. 1 u.s.p.), czyli zobowiązany jest pozyskać tekst obcych aktów normatywnych. Tekst prawa Chile został ustalony na podstawie innych środków dowodowych. Przewodniczący zapoznał się z tekstem ustawy, który zamieszczono pod linkiem przekazanym przez Ministerstwo Sprawiedliwości, a następnie zlecił przetłumaczenie go za pośrednictwem ChatGPT 4.0 (k. 2366). Tłumaczenie jest zgodne z tekstem ustawy, który powołał pełnomocnik powoda w piśmie z 20 marca 2023 r. Zarazem wiedzą powszechną jest, że ChatGPT jest urządzeniem, które bardzo dobrze radzi sobie z tłumaczeniem tekstu. Wykorzystanie tego instrumentu znacznie przyspieszyło rozpoznanie sprawy i ograniczyło koszty. Otwarty katalog możliwości dowodzenia obcego prawa wynika z art. 51a § 3 u.s.p., zgodnie z którym celem ustalenia treści prawa obcego lub obcej praktyki sądowej albo istnienia wzajemności sąd może zastosować także inne środki, w tym zasięgnąć opinii biegłych”.

W rezultacie odpowiedź nie opiera się wyłącznie na wytrenowanej uprzednio „pamięci” modelu, lecz na aktualnej wykładni przepisów kolizyjnych. Warunkiem skuteczności jest jednak odpowiednia jakość i kompletność repozytoriów danych, które model wykorzystuje³³. Jeżeli repozytorium jest zdominowane przez prawo amerykańskie czy chińskie, a nie zawiera np. polskiej ustawy o prawie prywatnym międzynarodowym, polskiego orzecznictwa oraz polskiej literatury, system może próbować „dopowiadać” treść norm kolizyjnych oraz ich interpretacji na podstawie ogólnych skojarzeń³⁴.

Dochodzą do tego inne problemy. Przykładem jest zjawisko tzw. szumu w wyszukiwaniu³⁵. Nawet jeśli system AI ma dostęp do bazy aktów prawnych (typu Lex bądź Legalis), może nadać większą wagę nieprecyzyjnemu komentarzowi z Internetu niż treści rozporządzenia Rzym I, o ile ten pierwszy jest sformułowany w sposób statystycznie bardziej „prawdopodobny” dla architektury wykorzystywanej przez ten system. Poza tym, ze względów czysto biznesowych, współczesne modele dostraja się do preferencji użytkowników. Promuje to odpowiedzi, które brzmią uprzejmie i przekonująco, ale nie zawsze są zgodne z prawdą. Dla prawnika jest to pułapka. System może wygenerować niezwykle przekonujący wywód o właściwości prawa szwajcarskiego (bo tak brzmi „typowy” wzorzec bezpiecznej porady), ignorując specyficzne właściwości danej umowy, których „nie lubi” statystyka modelu. W istocie modele nazywane obecnie (dosyć myląco) „rozumowymi”, choć pozwalają na „rozłożenie” problemu kolizyjnoprawnego na mniejsze etapy, imitując pracę prawnika, nadal wykazują tendencję do zgadywania końcowego wyniku (np. właściwości prawa szwajcarskiego) i dopasowywania do niego błędnej argumentacji. Przykładem może być wymyślenie (halucynowanie) rzekomego brzmienia art. 4 rozporządzenia Rzym I, gdy pytanie dotyczy prawa właściwego dla umów licencyjnych, przy wykorzystaniu treści szwajcarskiej ustawy o prawie prywatnym międzynarodowym, która taki przepis rzeczywiście zawiera.

³³ M. Dahl et al., *Large Legal Fictions...*, s. 88.

³⁴ E. Mik, *Caveat...*, s. 114.

³⁵ Jest to zjawisko polegające na obecności nieistotnych, mylących lub błędnych informacji w danych wejściowych (treningowych) lub wynikach generowanych przez sztuczną inteligencję, co utrudnia wyodrębnienie prawdziwych i istotnych treści.

5. Systemy agentowe oraz lokalne modele językowe

Obecnie, w 2026 r., widoczne są dwie silne tendencje rozwoju systemów AI. Są to tzw. agentowe systemy sztucznej inteligencji (ang. *agentic AI*) oraz coraz bardziej dostępna możliwość tworzenia własnej infrastruktury w oparciu o lokalne modele sztucznej inteligencji. Na przykład w kancelarii prawnej możemy samodzielnie dostroić system sztucznej inteligencji do potrzeb naszej specjalizacji, nawet jeśli dotyczy ona tak szczególnej dziedziny jak prawo prywatne międzynarodowe.

Te pierwsze to systemy, które nie tylko generują odpowiedzi, lecz także potrafią samodzielnie planować zadania, dzielić je na kroki, wywoływać inne narzędzia (np. wyszukiwarki, skrzynki pocztowe, kancelaryjne bazy danych), a następnie integrować wyniki. W praktyce prawniczej oznacza to system, który samodzielnie przeszukuje zasoby prawne, identyfikuje klauzule prawa właściwego, jurysdykcji i arbitrażu, klasyfikuje umowy według państwa i gałęzi prawa, generuje wstępny raport ryzyka prawnego itd.

Prognozy dotyczące rynku technologii prawnych wskazują, że w perspektywie kilku lat systemy agentowe będą wbudowaną, nieodłączną częścią aplikacji prawniczych. Ewolucja w stronę tych systemów (z dotychczasowych prostych chatbotów) opiera się na łączeniu modeli językowych z zewnętrznymi narzędziami i planowaniem (ang. *orchestration*). Pozwala to na interakcję z bazami danych w czasie rzeczywistym, co teoretycznie rozwiązuje problem nieaktualności prawa. Systemy te jednak nadal bazują na probabilistyce. Agent może „zaplanować” sprawdzenie danych zewnętrznych dotyczących konwencji haskich, ale błędnie zinterpretować wynik wyszukiwania ze względu na brak kognitywnego rozumienia hierarchii źródeł prawa. Metaforycznie system agentowy to następująca sytuacja. Wyjeżdżamy na dwutygodniowy urlop i przekazujemy młodemu prawnikowi podręcznik prof. Maksymiliana Pazdana, aby nauczył się prawa prywatnego międzynarodowego i odpowiadał klientom na wszelkie pytania kolizyjnoprawne, dając mu do dyspozycji dostęp do naszej skrzynki pocztowej, zasobów kancelaryjnych, zewnętrznych baz danych itp. Oczywiście w przypadku świata realnego działanie takie należałoby uznać za nieracjonalne, choćby z tego powodu, że niekontrolowany dostęp do zasobów może skutkować naruszeniem poufności i bezpieczeństwa. Natomiast w przypadku systemów agentowych próby takie rzeczywiście są podejmowane, tak jakby nawet doświadczeni eksperci nagle tracili rozsądek w zetknięciu z nowinkami technologicznymi.

Skoro, jak się wydaje, nie jesteśmy w stanie uciec przed wprowadzaniem sztucznej inteligencji do praktyki prawniczej, bezpieczniejszą alternatywą staje się rozwój lokalnych systemów (opartych na własnej infrastrukturze sprzętowej)³⁶. Mogą one bazować na polskich modelach językowych Bielik oraz PLLuM, o ile te jeszcze bardziej się rozwijają, ewentualnie na wyspecjalizowanym (wytrenowanym na odpowiednich zasobach) prawniczym modelu³⁷ bądź na zastosowaniu podejścia mieszanego (hybrydowego)³⁸. Kancelarie prawne mogą już teraz rozpocząć wdrażanie sztucznej inteligencji w sposób kontrolowany i w ramach własnej infrastruktury sprzętowej, bez przekazywania dokumentów klientów do zewnętrznego dostawcy. Mniejsze, wyspecjalizowane dla potrzeb prawnych modele (tzw. SLM) okazują się bardziej użyteczne niż systemy komercyjne pochodzące od zagranicznych globalnych korporacji technologicznych³⁹. W kontekście prawa prywatnego międzynarodowego oznacza to możliwość opracowania dedykowanego modelu typu „LLM PPM”, wytrenowanego na zestawach pytań i odpowiedzi dotyczących prawa właściwego i jurysdykcji, przykładach przeprowadzonej analizy klauzul wyboru prawa, przepisach kolizyjnych, orzecznictwie i doktrynie. W literaturze opisano już wyspecjalizowane dla potrzeb prawa międzynarodowego publicznego modele, które dzięki takim działaniom osiągają lepsze wyniki⁴⁰. Podobne rozwiązania mogą zostać zastosowane w dziedzinie prawa prywatnego międzynarodowego.

³⁶ R. Bhambhoria et al., *Evaluating AI for Law: Bridging the Gap with Open-Source Solutions*, <https://doi.org/10.48550/arXiv.2404.12349>.

³⁷ P. Colombo et al., *SaulLM-7B: A Pioneering Large Language Model for Law*, <https://doi.org/10.48550/arXiv.2403.03883>; P. Colombo et al., *SaulLM-54B & SaulLM-141B: Scaling Up Domain Adaptation for the Legal Domain*, <https://doi.org/10.48550/arXiv.2407.19584>.

³⁸ Wybór modelu powinien uwzględniać gradację między modelami typu Frontier, LLM i SLM. Podczas gdy modele Frontier (np. najnowsze modele GPT, Claude czy Gemini) oferują najwyższy poziom rozumowania w złożonych, wieloetapowych zadaniach, małe modele językowe (SLM) mogą przewyższać je w wąsko zdefiniowanych zadaniach, takich jak klasyfikacja dokumentów pod kątem prawa właściwego, oferując przy tym wyższą szybkość i niższe koszty operacyjne.

³⁹ A. Bent, *Large Language Models...*, s. 129–131.

⁴⁰ J. Rochel, *Encoded Normativity: Designing Large Language Models for International Law*, „Nordic Journal of International Law” 2025, vol. 95, s. 1–32.

6. Halucynacje, stroniczość i inne zagrożenia dla prawa prywatnego międzynarodowego

Najpoważniejszym obecnie problemem w wykorzystaniu systemów sztucznej inteligencji w obszarze prawnym jest generowanie nieprawdziwych, ale wiarygodnie brzmiących informacji, tzw. halucynacji (czy też raczej konfabulacji)⁴¹. W przypadku prawa prywatnego międzynarodowego przykładem będzie powoływanie się przez model językowy na nieistniejące artykuły rozporządzeń Rzym I czy Rzym II, fikcyjne orzeczenia bądź spreparowane w tym celu tytuły opracowań naukowych. W praktyce znane są już sprawy, w których prawnicy powoływali się w pismach procesowych na nieistniejące orzeczenia wygenerowane przez ChatGPT. W sporach międzynarodowych skala ryzyka jest większa, ponieważ sąd lub druga strona mogą nie mieć łatwego dostępu do źródeł obcego prawa.

Halucynacje są błędem, którego nie można łatwo wyeliminować w kolejnych wersjach modeli językowych, ponieważ wynikają one z fundamentalnej niemożności odróżnienia przez model faktu od przekonania⁴². Dla modelu prawda nie istnieje jako kategoria obiektywna. Liczy się jedynie prawdopodobieństwo wystąpienia danego słowa po innym, co prowadzi do generowania brzmiących przekonująco nonsensów, czyli odpowiedzi językowo poprawnych, ale całkowicie błędnych merytorycznie⁴³.

Zjawisko to wiąże się z tzw. problemem zagubienia w środku, opisywanym w najnowszych badaniach. Modele, mimo dużych okien kontekstowych (tj. podania im bardzo precyzyjnych i szczegółowych danych wejściowych, jak np. kompletu dokumentów dotyczących danego sporu transgranicznego), wykazują tendencję do lepszego zapamiętywania informacji z początku i końca wprowadzonego tekstu. W przypadku prawa prywatnego międzynarodowego może to prowadzić do sytuacji, w której system prawidłowo zidentyfikuje strony umowy (początek) i klauzulę prawa właściwego (koniec), ale całkowicie pominie ukryte w środku dokumentu zastrzeżenia.

⁴¹ M. Dahl et al., *Large Legal Fictions...*, s. 64–93.

⁴² Zob. szerzej: B. Shneiderman, *Human-Centered AI...*, s. 151–163; Z. Xu, S. Jain, M. Kankanhalli, *Hallucination is Inevitable: An Innate Limitation of Large Language Models*, <https://doi.org/10.48550/arXiv.2401.11817>.

⁴³ E. Mik, *Caveat...*, s. 108–109.

Modele są dostosowywane do tego, aby ich odpowiedzi były przekonujące i pomocne dla użytkownika, co jednak sprzyja zjawisku tzw. pochlebstwa algorytmicznego (sykofancji; gr. *συκοφαντία/sykophantía*)⁴⁴. System jest więc skłonny generować argumentację wspierającą błędną tezę postawioną przez prawnika w zapytaniu, np. o właściwości prawa danego państwa, zamiast zwrócić uwagę na obiektywne ograniczenia autonomii woli.

Najtrudniejsze dla modeli językowych pozostają klauzule generalne, takie jak „porządek publiczny”. Ich zastosowanie wymaga wartościowania, czego system nie może samodzielnie dokonać. System sztucznej inteligencji może jedynie przytaczać poglądy pochodzące z doktryny i orzecznictwo, lecz nie jest w stanie zastąpić prawnika przy ocenie ich znaczenia w danej sprawie.

Globalne modele językowe trenowane są głównie na materiałach anglojęzycznych, z przewagą prawa amerykańskiego i brytyjskiego. Może to powodować domyślne stosowanie konstrukcji *common law* (np. *proper law of the contract*), marginalizowanie krajowych kodyfikacji kolizyjnych, mniej obecnych w otwartych źródłach, oraz „przesunięcie” metodologii w kierunku kazuistycznego, precedensowego myślenia. Systemy, operując na wektorach słów, dążą do uśredniania znaczeń na podstawie najliczniejszych wzorców w danych treningowych. W praktyce oznacza to ryzyko statystycznego „wchłonięcia” rzadkich norm kolizyjnych przez dominujące konstrukcje *common law*⁴⁵, co nakłada na prawnika obowiązek szczególnej weryfikacji wyników pod kątem autonomii pojęciowej prawa unijnego i krajowego.

Prawo prywatne międzynarodowe funkcjonuje w środowisku wielojęzycznym. Nawet w obrębie UE to samo rozporządzenie rzymskie czy brukselskie istnieje w wielu wersjach językowych, które nie zawsze są w pełni ekwiwalentne. Modele językowe mają trudność z zachowaniem pełnej spójności między tymi wersjami, a w praktyce mogą mieszać terminologię pochodzącą z różnych języków krajowych. Problem ten

⁴⁴ Pojęcie określające zachowanie polegające na nieszczerym, przesadnym schlebianiu komuś lub potwierdzaniu cudzych poglądów. W języku potocznym bywa łączone z lizusostwem i serwilizmem. Stosuje się je obecnie w odniesieniu do skłonności modeli językowych do bezkrytycznego potwierdzania tez użytkownika, wzmacniania jego przekonań i schlebiania mu niezależnie od poprawności faktów, co wiąże się ze sposobem ich trenowania, zwłaszcza z etapami wykorzystującymi preferencje ludzkie. W trakcie tego procesu ludzie oceniają odpowiedzi modeli, częściej nagradzając wypowiedzi uprzejme, zgadzające się z rozmówcą i utrzymujące przyjazny ton niż otwarcie korygujące lub polemiczne. Zjawisko to sprzyja wzmacnianiu takich wzorców zachowania modelu, które dają wysokie oceny satysfakcji użytkownika, ale niekoniecznie zapewniają prawdziwość i pełność udzielanych informacji.

⁴⁵ Por. R. Ghoshal, *Cross-Border...*, s. 191.

znajduje techniczne wyjaśnienie w badaniach nad *Cross-Lingual Word Embeddings*. Mapowanie pojęć między językami (np. polskim miejscem zamieszkania a angielskim domicylem) nie jest idealne. Model trenowany głównie na prawie amerykańskim może „przemycać” cechy danej instytucji, ponieważ w przestrzeni wektorowej te pojęcia znajdują się zbyt blisko siebie.

Co więcej, empiryczne testy potwierdzają, że uwarunkowania polityczne oraz ukryte założenia wprowadzone do modeli językowych mogą blokować dostęp do określonych informacji lub zniekształcać odpowiedzi⁴⁶, co stanowi realne zagrożenie dla obiektywności analizy w sprawach transgranicznych. Przykładowo, wystarczy zapytać jakiegokolwiek chiński model językowy, czy Tajwan jest odrębnym państwem.

7. Obowiązki wynikające z nowych przepisów

Od 2026 r. działamy jako prawnicy (kolizjoniści, sędziowie, profesjonaliści pełnomocnicy) w warunkach dwóch kluczowych regulacji unijnego rozporządzenia o sztucznej inteligencji (AI Akt) oraz konwencji ramowej Rady Europy o sztucznej inteligencji, prawach człowieka, demokracji i praworządności. Oba te akty prawne nakładają na użytkowników systemów AI, w tym prawników, istotne obowiązki w zakresie bezpieczeństwa, przejrzystości, nadzoru ludzkiego oraz odpowiedzialnego stosowania technologii⁴⁷. Na uwagę zasługuje też art. 4 AI Akt, wprowadzający obowiązek rozwoju kompetencji osób korzystających z systemów sztucznej inteligencji (ang. *AI literacy*)⁴⁸.

⁴⁶ J. Pan, X. Xu, *Political Censorship in Large Language Models Originating from China*, „PNAS Nexus” 2026, vol. 5 (2), s. 1–8, <https://doi.org/10.1093/pnasnexus/pgag013>.

⁴⁷ M. Linera, A. Meuwese, *Regulating AI from Europe: A Joint Analysis of the AI Act and the Framework Convention on AI*, „The Theory and Practice of Legislation” 2025, vol. 13 (3), s. 292–311; R. Jokubauskas, M. Świerczyński, *Relationship between the Council of Europe Framework Convention on Artificial Intelligence, Human Rights, Democracy and the Rule of Law and the European Union’s Artificial Intelligence Act (AI Act)*, w: *Wpływ Aktu w sprawie sztucznej inteligencji na funkcjonowanie administracji publicznej*, red. G. Sibiga, C.H. Beck, Warszawa 2025, s. 22–36; R. Ghoshal, *Cross-Border...*, s. 187–188.

⁴⁸ K. Strzępek, *Convergence or Divergence? Balancing Regulatory Approaches in the Council of Europe AI Convention and the EU AI Act*, „Rocznik Administracji Publicznej” 2025, nr 11 (2), s. 359.

Dla prawników oznacza to konieczność zrozumienia, jak działają modele językowe używane w narzędziach prawniczych, uwzględniania ryzyka związanego z błędami i halucynacjami w sprawach transgranicznych, dokumentowania sposobu wykorzystania sztucznej inteligencji w obsłudze klientów i przygotowania się na sytuacje, w których użycie danego narzędzia przez stronę lub sąd stanie się przedmiotem sporu (np. zarzut zawarcia halucynacji systemu w piśmie procesowym)⁴⁹.

Ze względu na zjawisko antropomorfizacji systemów sztucznej inteligencji, czyli tendencję do traktowania ich jako racjonalnych umysłów (co jest ewolucyjnie uwarunkowane u ludzi), prawnicy mogą nadmiernie wierzyć generowanym przez nie odpowiedziom⁵⁰. Tymczasem systemy te należy traktować wyłącznie jako narzędzie wspomagające, a nie zastępujące ludzką inteligencję, oraz wymagające stałego nadzoru, skoro nie znają one koncepcji czasu, odpowiedzialności ani nie posiadają świadomości świata zewnętrznego⁵¹.

W świetle powyższego prawnicy specjalizujący się w dziedzinie prawa prywatnego międzynarodowego, zamierzający korzystać z narzędzi sztucznej inteligencji, powinni zadać co najmniej następujące pytania:

- 1) Na jakich danych trenowano wykorzystywany model? Czy obejmuje one prawo unijne (Rzym I, Rzym II, Bruksela I bis), konwencje haskie i krajowe ustawy kolizyjne?
- 2) Czy system korzysta z zewnętrznych repozytoriów (oficjalnych baz aktów prawnych, orzeczeń, komentarzy)?
- 3) Czy system rozróżnia jurysdykcję i prawo właściwe, a także zakres zastosowania poszczególnych rozporządzeń i konwencji?
- 4) Czy dostawca systemu ujawnia wskaźniki dotyczące halucynacji i błędów w obszarze prawnym?
- 5) Czy istnieje możliwość lokalnego wdrożenia modelu (w ramach infrastruktury kancelarii) dla dokumentów poufnych?
- 6) Czy wyniki generowane przez system są zawsze weryfikowane przez prawnika, zanim zostaną przedstawione klientowi lub sądowi?
- 7) Jakie procedury wewnętrzne kancelarii regulują korzystanie z AI w kontekście obowiązków wynikających z obowiązujących przepisów?

⁴⁹ H. Surden, *ChatGPT...*, s. 1972; P. Homoki, Z. Zsolt, *Large...*, s. 452.

⁵⁰ Więcej o tym J. Ercilla García, *Justicia automatizada: entre las inteligencias artificiales que fingen y las que persuaden*, „Lex Social. Revista De Derechos Sociales” 2025, vol. 15 (1), s. 1—39.

⁵¹ Z. Więckowski, M. Świerczyński, *Analiza ryzyka dokonywana na podstawie konwencji ramowej Rady Europy o sztucznej inteligencji na przykładzie zastosowań w sektorze prawnym*, „Prawo i Więź” 2025, t. 54, nr 1, s. 409—428.

8. Zakończenie

Duże modele językowe oraz systemy agentowe wkraczają nieodwołalnie do praktyki prawniczej, w tym w obszar stosowania prawa prywatnego międzynarodowego⁵². W porównaniu z okresem sprzed 2025 r. osiągnęły one istotny postęp. Potrafią sprawniej analizować dokumenty, identyfikować klauzule wyboru prawa oraz generować coraz bardziej przekonujące odpowiedzi, także w złożonych stanach faktycznych o charakterze transgranicznym. Rozwój modeli typu *reasoning* oraz technik RAG znacząco zwiększył ich użyteczność operacyjną.

Postęp ten nie zmienia jednak ich fundamentalnej natury. Modele językowe pozostają systemami probabilistycznymi, operującymi na statystycznych zależnościach między tekstami, a nie na normach prawnych jako takich. W konsekwencji nie są zdolne do rzeczywistego stosowania prawa prywatnego międzynarodowego, które wymaga precyzyjnego różnicowania pojęć, rozumienia hierarchii źródeł prawa oraz dokonywania wartościowania kolizyjnego. Zjawiska takie jak halucynacje, semantyczne „uśrednianie” pojęć, błędy w mapowaniu wielojęzycznym czy stronniczość danych treningowych mają charakter strukturalny, a nie incydentalny.

Z tego względu sztuczna inteligencja może pełnić wyłącznie funkcję narzędzia wspierającego, a nie zastępującego prawnika. Kluczowe znaczenie ma świadomy i kompetentny nadzór człowieka (ang. *human oversight*), który w świetle regulacji europejskich nie powinien być traktowany jako formalny obowiązek, lecz jako podstawowa zasada bezpiecznego korzystania z tych technologii. Wymaga to od prawnika nie tylko wiedzy dogmatycznej, lecz także rozumienia technicznych ograniczeń modeli, w tym ich podatności na potwierdzanie błędnych założeń użytkownika (sykofania) czy na zniekształcenia wynikające z dominacji określonych systemów prawnych.

Z perspektywy prawa prywatnego międzynarodowego szczególnie istotne jest ryzyko systemowego uprzywilejowania dominujących porządków prawnych, zwłaszcza *common law*, co może prowadzić do stopniowego wypierania europejskich konstrukcji kolizyjnoprawnych przez modele oparte na analizie kazuistycznej i precedensowej. Tendencja ta, wzmacniana przez strukturę danych treningowych, wymaga świadomego przeciwdziałania.

⁵² A. Deeks, D. Hollis, *Large...*, s. 117–144; M. Ramos-Maqueda, D.L. Chen, *The Data Revolution in Justice*, „World Development” 2025, vol. 186, s. 1–10, <https://doi.org/10.1016/j.worlddev.2024.106834>.

W konsekwencji dalszy rozwój sztucznej inteligencji w praktyce prawniczej powinien zmierzać w kierunku tworzenia wyspecjalizowanych, lokalnych modeli językowych, dostosowanych do konkretnych systemów prawnych i opartych na wysokiej jakości reprezentatywnych danych. Tylko takie podejście może ograniczyć ryzyka systemowe i zapewnić, że sztuczna inteligencja będzie narzędziem wspierającym, a nie deformującym rozwój prawa prywatnego międzynarodowego.

Ostatecznie bowiem, niezależnie od stopnia zaawansowania technologicznego, to nie algorytm, lecz prawnik pozostaje gwarantem prawidłowego stosowania norm kolizyjnych oraz sprawiedliwości proceduralnej w coraz bardziej cyfrowym środowisku prawnym.

Bibliografia

- Beane M., *The Skill Code: How to Save Human Ability in an Age of Intelligent Machines*, Harper Collins, New York 2024.
- Bent A., *Large Language Models: AI's Legal Revolution*, „Pace Law Review” 2023, vol. 44, no. 1.
- Bhambhoria R. et al., *Evaluating AI for Law: Bridging the Gap with Open-Source Solutions*, <https://doi.org/10.48550/arXiv.2404.12349>.
- Coeckelbergh M., *AI Ethics*, MIT Press, Cambridge 2020.
- Colombo P. et al., *SaulLM-7B: A Pioneering Large Language Model for Law*, <https://doi.org/10.48550/arXiv.2403.03883>.
- Colombo P. et al., *SaulLM-54B & SaulLM-141B: Scaling Up Domain Adaptation for the Legal Domain*, <https://doi.org/10.48550/arXiv.2407.19584>.
- Dahl M. et al., *Large Legal Fictions: Profiling Legal Hallucinations in Large Language Models*, „The Journal of Legal Analysis” 2024, vol. 16.
- Deeks A., Hollis D., *Large Language Models and International Law*, „Chicago Journal of International Law” 2025, vol. 26, no. 1.
- Diderot D., *Œuvres complètes*, éd. J. Assézat, M. Tourneux, Garnier, Paris 1876.
- Ercilla García J., *Justicia automatizada: entre las inteligencias artificiales que fingen y las que persuaden*, „Lex Social. Revista De Derechos Sociales” 2025, vol. 15 (1), s. 1—39.
- Ghoshal R., *Cross-Border AI Governance for Legal Tech: Standardizing Ethical and Legal Norms in Access to Justice*, „International Journal of Law, Justice and Jurisprudence” 2025, vol. 5 (1).
- Homoki P., Zsolt Z., *Large Language Models and Their Possible Uses in Law*, „Hungarian Journal of Legal Studies” 2023, vol. 64, no. 3.

- Jokubauskas R., Świerczyński M., *Relationship between the Council of Europe Framework Convention on Artificial Intelligence, Human Rights, Democracy and the Rule of Law and the European Union's Artificial Intelligence Act (AI Act)*, w: *Wpływ Aktu w sprawie sztucznej inteligencji na funkcjonowanie administracji publicznej*, red. G. Sibiga, C.H. Beck, Warszawa 2025, s. 22—36.
- Linera M., Meuwese A., *Regulating AI from Europe: A Joint Analysis of the AI Act and the Framework Convention on AI*, „The Theory and Practice of Legislation” 2025, vol. 13 (3), s. 292—311.
- Liu J., Xueyao L., *How Do Judges Use Large Language Models? Evidence from Shenzhen*, „Journal of Legal Analysis” 2024, vol. 16, s. 235—262.
- Machnikowski P., *Prawo zobowiązań w 2025 roku. Nowe technologie, nowe wyzwania*, w: *Współczesne problemy prawa zobowiązań*, red. A. Olejniczak et al., Wolters Kluwer, Warszawa 2015.
- Mik E., *Caveat Lector: Large Language Models in Legal Practice*, „Rutgers Business Law Review” 2024, vol. 19, no. 2.
- Naveed H. et al., *A Comprehensive Overview of Large Language Models*, „ACM Transactions on Intelligent Systems and Technology” 2025, vol. 16 (5), s. 1—72.
- Pan J., Xu X., *Political Censorship in Large Language Models Originating from China*, „PNAS Nexus” 2026, vol. 5 (2), s. 1—8, <https://doi.org/10.1093/pnas/nexus/pgag013>.
- Ramos-Maqueda M., Chen D.L., *The Data Revolution in Justice*, „World Development” 2025, vol. 186, s. 1—10, <https://doi.org/10.1016/j.worlddev.2024.106834>.
- Rochel J., *Encoded Normativity: Designing Large Language Models for International Law*, „Nordic Journal of International Law” 2025, vol. 95, s. 1—32.
- Schwarcz D. et al., *AI Reasoning Models, Retrieval Augmented Generation, and the Future of Legal Practice*, „Journal of Law and Empirical Analysis” 2026, vol. 3 (1), s. 220—250, <https://dx.doi.org/10.2139/ssrn.5162111>.
- Shapira N. et al., *Agents of Chaos. Report, Study Timeline Feb 2—22, 2026*, <https://agentsofchaos.baulab.info/report.html> [dostęp: 1.06.2026].
- Shneiderman B., *Human-Centered AI*, Oxford University Press, Oxford 2022.
- Strzępek K., *Convergence or Divergence? Balancing Regulatory Approaches in the Council of Europe AI Convention and the EU AI Act*, „Rocznik Administracji Publicznej” 2025, nr 11 (2), s. 351—368.
- Surden H., *ChatGPT, Large Language Models, and Law*, „Fordham Law Review” 2024, vol. 92 (5), s. 1942—1972.
- Świerczyński M., *Generatywna sztuczna inteligencja w służbie wymiaru sprawiedliwości w świetle wytycznych Rady Europy (AIAB CEPEJ)*, w: *Ius Scientiam Persequitur. Księga jubileuszowa z okazji XXV-lecia Wydziału Prawa i Administracji Uniwersytetu Kardynała Stefana Wyszyńskiego*

- w Warszawie, red. M. Badowski, J. Wesserling, L. Karski, Wolters Kluwer, Warszawa 2025, s. 287—302.
- Świerczyński M., Więckowski Z., *Sztuczna inteligencja w prawie międzynarodowym. Rekomendacje wybranych rozwiązań*, Difin, Warszawa 2021.
- Trammer H., *Z rozważań nad strukturą normy kolizyjnej prawa prywatnego międzynarodowego*, „Studia Cywilistyczne” 1969, t. XII—XIV.
- UNESCO, *AI Essentials for Judges*, 2026, CI/DIT/AI/Judges2026, <https://unesdoc.unesco.org/ark:/48223/pf0000396991> [dostęp: 25.07.2026].
- Wang B., *Prompts and Large Language Models: A New Tool for Drafting, Reviewing and Interpreting Contracts?*, „Law, Technology and Humans” 2024, vol. 6, no. 2, s. 88—106.
- Więckowski Z., Świerczyński M., *Analiza ryzyka dokonywana na podstawie konwencji ramowej Rady Europy o sztucznej inteligencji na przykładzie zastosowań w sektorze prawnym*, „Prawo i Więź” 2025, t. 54, nr 1, s. 409—428.
- Wiśniewski A., *Sztuczna inteligencja i prawa człowieka w kontekście prawa międzynarodowego*, „Prawo i Więź” 2024, t. 47, nr 4, s. 29—52.
- Xu Z., Jain S., Kankanhalli M., *Hallucination is Inevitable: An Innate Limitation of Large Language Models*, <https://doi.org/10.48550/arXiv.2401.11817>.