

Stosowanie systemów sztucznej inteligencji w e-legislacji. Wprowadzenie do problematyki²

1. Zakres artykułu

Celem niniejszego artykułu jest zbudowanie u czytelnika świadomości szans i problemów związanych ze stosowaniem

-
- 1 Michał Araszkiewicz jest autorem części: Zakres artykułu, Pojęcie sztucznej inteligencji (*Artificial Intelligence*), Podejście oparte na rozwiązywaniu problemów, Wyszukiwanie, klasyfikacja i prezentacja informacji, Sprawdzanie jakości tworzonych tekstów, Generowanie dokumentów spełniających określone kryteria; Patryk Ciurak jest autorem części: Potrzeby osób zaangażowanych w proces legislacyjny, Dane, Wykorzystanie danych, Efektywność legislacji.
 - 2 Artykuł stanowi dostosowaną do wymogów publikacyjnych wersję ekspertyzy stworzonej na zlecenie Kancelarii Prezesa Rady Ministrów w ramach prac zespołu „e-Legislacja”.

systemów sztucznej inteligencji w celu wsparcia wykonywania różnorodnych zadań w procesie legislacyjnym. W szczególności diskutowane są istotne prerekwizyty efektywnego wdrażania takich systemów, w tym dotyczące dostępności odpowiednio przygotowanych baz danych.

W artykule przedstawiono metodę dokonywania oceny, czy włączenie w prace legislacyjne rozwiązań wykorzystujących sztuczną inteligencję jest uzasadnione z punktu widzenia problemów zidentyfikowanych w procesie legislacyjnym. Kluczowe znaczenie w tym zakresie mają dane opisujące poszczególne etapy procesu oraz wymogi, którym te dane powinny odpowiadać. Dysponując wiarygodnym zbiorem danych, można np. monitorować i dokonywać oceny wydajności procesu legislacyjnego przy użyciu podstawowych metod statystycznych. W dalszej części artykułu zaprezentowano spektrum szans wykorzystania narzędzi opartych o sztuczną inteligencję, poczynając od zaawansowanego wyszukiwania, klasyfikacji i grupowania informacji, poprzez eliminację błędów (o różnym stopniu skomplikowania), a kończąc na generowaniu tekstów na podstawie wytycznych zadanych przez legislatora. Rozważania prowadzone są przy założeniu racjonalności koncepcji tworzenia ekosystemu mikrousług informatycznych mających wspierać rozwiązywanie różnych problemów aktualizujących się w odniesieniu do zarządzania procesami tworzenia prawa oraz tworzenia treści w ramach tych procesów, w tym projektów aktów normatywnych. W konsekwencji celem niniejszego opracowania nie jest postulowanie stworzenia wielokomponentowego systemu, który miałby wspierać wszystkie zadania właściwe dla procesu tworzenia prawa. Artykuł nie porusza również zagadnień związanych z prawnymi, etycznymi oraz technicznymi (w tym dotyczącymi się cyberbezpieczeństwa) aspektami wdrażania rozwiązań opartych na AI.

2. Pojęcie sztucznej inteligencji (*Artificial Intelligence*)

Pojęcie sztucznej inteligencji (ang. *Artificial Intelligence*, dalej jako AI), pomimo jego rozpowszechnienia, jest notoryjnie dyskusyjne³. Normatywne pojęcie systemu AI zostało zdefiniowane w art. 3 pkt 1 rozporządzenia Parlamentu Europejskiego i Rady (UE) 2024/1689 z dnia 13 czerwca 2024 r. w sprawie ustanowienia zharmonizowanych przepisów dotyczących sztucznej inteligencji oraz zmiany rozporządzeń (WE) nr 300/2008, (UE) nr 167/2013, (UE) nr 168/2013, (UE) 2018/858, (UE) 2018/1139 i (UE) 2019/2144 oraz dyrektyw 2014/90/UE, (UE) 2016/797 i (UE) 2020/1828 (akt w sprawie sztucznej inteligencji)⁴ i obejmuje następujące elementy⁵:

1. system maszynowy,
2. który został zaprojektowany do działania z różnym poziomem autonomii po jego wdrożeniu oraz
3. który może wykazywać zdolność adaptacji po jego wdrożeniu,
4. a także który – na potrzeby wyraźnych lub dorozumianych celów – wnioskuje, jak generować na podstawie otrzymanych danych wejściowych wyniki, takie jak predykcje, treści, zalecenia lub decyzje, które mogą wpływać na środowisko fizyczne lub wirtualne.

Powyższe cechy są spełnione w wysokim stopniu przede wszystkim przez systemy określane jako oparte na tzw. inteligencji

3 S. Russell, P. Norvig, *Artificial Intelligence. A Modern Approach*, 3rd ed., Upper Saddle River 2016.

4 Dz. Urz. UE L z 2024 r. poz. 1689.

5 Elementy przejęte są z definicji systemu AI proponowanej przez OECD; por. OECD, *Explanatory memorandum on the updated OECD definition of an AI system*, „OECD Artificial Intelligence Papers” 2024, nr 8, <https://doi.org/10.1787/623da898-en> [dostęp: 25.05.2025 r.].

obliczeniowej (*Computational Intelligence*)⁶. Różnorodne modele zaliczane do tej kategorii (np. sztuczne sieci neuronowe, obliczenia ewolucyjne, maszyny wektorów nośnych) charakteryzują się m.in. następującymi właściwościami:

- analizują dane wyrażone numerycznie, a wykonywane przez nie operacje myślowe mają charakter matematyczny;
- ich model wiedzy jest implicytny, wyuczony na podstawie danych i reprezentowany w parametrach modelu;
- istotą ich działania jest stopniowa zmiana zachowania w ekspozycji na dane (uczenie maszynowe, w tym uczenie nadzorowane, nienadzorowane oraz uczenie przez wzmocnienie⁷);
- w konsekwencji wykazują znaczną adaptacyjność, dostosowując się do zmieniających się danych, kontekstów, środowisk; mogą wykonywać operacje i generować rezultaty zaskakujące nie tylko użytkowników, ale i twórców oraz ekspertów;
- ogólna zasada funkcjonowania systemu jest oczywiście znana, ale jego sposób działania w danej sytuacji może być niezrozumiały także dla specjalistów; systemy cechują się więc zróżnicowaną, ale ograniczoną lub wręcz znikomą wyjaśnialnością/objaśnialnością (ang. *explainability*)⁸ – bywają określane jako tzw. czarne skrzynki (ang. *black box AI*); w związku z tym zjawiskiem powstał dział badań i technologii określany jako XAI (ang. *eXplainable AI*)⁹ nakierowany na tworzenie rozwiązań wspierających objaśnianie działania złożonych systemów AI;

6 M. Flasiński, *Introduction to Artificial Intelligence*, Cham 2016.

7 M. Kubat, *An Introduction to Machine Learning*, Cham 2017.

8 A.B. Arrieta et al., *Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI*, „Information Fusion” 2020, nr 58, s. 82–115.

9 *Ibidem*.

- systemy mogą z istoty swojej generować błędne rezultaty, a prawidłowo zaprojektowany proces maszynowego uczenia ma zapewniać zmniejszanie liczby błędów; wymaga podkreślenia, że zjawisko to ma miejsce również wtedy, gdy zbiór danych treningowych jest dobrze przygotowany i nie zawiera błędnych danych; trenowanie systemu na błędnych czy źle zbalansowanych danych oczywiście tworzy znaczące ryzyko generowania niewłaściwych rezultatów; odpowiednie przygotowanie danych treningowych i testowych może być zadaniem kosztownym i prowadzącym do sporów wśród ekspertów odpowiedzialnych za te zadania, dlatego istotne jest dokładne dokumentowanie tych procesów;
- systemy są względnie dobrze skalowalne; z powodzeniem stosowane są w odniesieniu do działań na dużych zbiorach danych i w związku z rozwiązywaniem wielu złożonych problemów;
- typowe obszary zastosowań systemów opartych na inteligencji obliczeniowej to np. rozpoznawanie wzorców, przetwarzanie języka naturalnego, analityka danych czy predykcja; bardzo popularne w ostatnich latach systemy generatywnej AI oparte na tzw. dużych modelach językowych również zaliczane są do systemów inteligencji obliczeniowej w przyjętym tu rozumieniu.

Normatywna definicja systemów AI w mniejszym stopniu dotyczy metod i technik określanych tradycyjnie jako symboliczna AI¹⁰. Systemy zaliczane do tej kategorii (w szczególności oparte na logice, regułach oraz tzw. strukturalnych modelach reprezentacji wiedzy) można scharakteryzować w sposób następujący:

10 M. Flasiński, *Introduction to..., op. cit.*

- dla swojego działania wymagają bazy wiedzy wyrażonej w języku formalnym; wykonywane przez nie operacje polegają na stosowaniu schematów rozumowania, w szczególności reguł logiki;
- ich model wiedzy ma charakter eksplicytny; poszczególnym elementom bazy wiedzy może zostać przypisane znaczenie;
- istotą ich działania jest wyprowadzanie wniosków na podstawie danych wyrażonych w sposób sformalizowany; wykonują jedynie te operacje, które zostały uprzednio przewidziane i określone w postaci wyraźnych instrukcji;
- w konsekwencji ich adaptacyjność jest ograniczona, chociaż istnieją modele maszynowego uczenia właściwe dla systemów symbolicznej AI;
- zarówno ogólna zasada działania systemu, jak i poszczególne wykonywane przezeń kroki są zrozumiałe dla osoby posiadającej odpowiednie przygotowanie; działanie systemu jest transparentne; wiele systemów opartych na założeniach symbolicznej AI jest wyposażanych w tzw. moduły wyjaśniające, które prezentują działanie systemu krok po kroku;
- ich działanie jest w zasadzie niezawodne, z zastrzeżeniem ewentualnych błędów w bazie wiedzy oraz błędnego sformułowania instrukcji działania programu; co jednak istotne, wyeliminowanie błędu wymaga co do zasady interwencji, w tym manualnego wprowadzania korekt;
- systemy są zazwyczaj względnie słabo skalowalne, ponieważ rozwiązywanie nowych problemów wymaga rozwijania bazy wiedzy oraz w razie potrzeby tworzenia nowych instrukcji; przygotowanie takiego systemu jest kosztowne i czasochłonne z uwagi na konieczność manualnego wprowadzania poszczególnych elementów bazy wiedzy, co zazwyczaj wymaga współpracy pomiędzy ekspertami w zakresie symbolicznej AI

oraz ekspertami dziedzinowymi; sprawdzanie poprawności bazy wiedzy jest również zadaniem żmudnym i cechującym się podatnością na błędy;

- typowe obszary zastosowań systemów tego rodzaju to formułowanie propozycji rozwiązań problemów w obszarach, gdzie eksplicytna wiedza jest stosunkowo łatwa do ustrukturyzowania, a od efektywności i skali działania ważniejsze są dokładność oraz możliwość prześledzenia rozumowania (wiarogodne uzasadnienie odpowiedzi, a nie tylko sam rezultat), i gdzie np. ilość potrzebnych informacji lub stopień skomplikowania rozumowania są na tyle duże, że poprawne wykonanie zadania jest trudne dla człowieka; taką funkcję spełniały i nadal spełniają tzw. systemy ekspertowe, w szczególności tworzone w obszarze zastosowań prawniczych; prawnicze zastosowania zarówno systemów inteligencji obliczeniowej, jak i symbolicznej AI są od ponad czterech dekad intensywnie badane w ramach nurtu określanego jako „sztuczna inteligencja i prawo”¹¹.

Należy dodać, że od ponad ćwierć wieku rozwijany jest system standardów Sieci Semantycznej, który jest nakierowany na uzyskanie efektu przetwarzania znaczenia danych dostępnych w sieci przez maszyny¹². Zbiór standardów obejmuje m.in.:

- Unicode – standard pozwalający na wyrażenie w sposób odczytywalny dla maszyn dowolnych symboli języka pisanego;
- URI – standard identyfikatorów zasobów w sieci;

11 M. Araszkiewicz, *Algorytmizacja myślenia prawniczego. Modele, możliwości, ograniczenia* [w:] *Legal tech. Czyli jak bezpiecznie korzystać z narzędzi informatycznych w organizacji, w tym w kancelarii oraz dziale prawnym*, red. D. Szostek, Warszawa 2021, s. 57–87.

12 L. Yu, *Introduction to the Semantic Web and Semantic Web Services*, Boca Raton 2019.

- uniwersalny język XML przeznaczony do reprezentowania różnych danych w ustrukturyzowany sposób;
- XML Schema – schematy wprowadzające ograniczenia typów i struktury danych;
- RDF – metoda opisu danych jako tzw. grafu skierowanego, obejmującego dwa wierzchołki i łączącą je krawędź;
- OWL – język służący do opisu tzw. ontologii, czyli formalnych reprezentacji pewnej dziedziny wiedzy, wyrażającej pojęcia oraz relacje pomiędzy nimi;
- systemy wnioskowania (ang. *reasoners*) wyprowadzające konkluzje ze zbioru podanych faktów w odniesieniu do bazy wiedzy systemu.

Warto podkreślić, że wiele systemów symbolicznej AI przygotowywanych dla celów wsparcia wykonywania zadań prawniczych opracowanych w ciągu ostatnich dwóch dekad korzysta właśnie ze standardów Sieci Semantycznej, co sprzyja weryfikacji stosowanych rozwiązań oraz ich uniwersalności¹³. W szczególności na tych standardach oparty jest system identyfikatorów ELI/ELI-DL, których pełne wdrożenie jest jednym z rekomendowanych rozwiązań w obszarze e-legislacji.

Pomimo faktu, że podejście określane jako symboliczna AI pozostaje zmarginalizowane z uwagi na bezprecedensowe sukcesy systemów opartych na inteligencji obliczeniowej (w tym w ostatnich latach na dużych modelach językowych), to w ostatnich latach na nowo dostrzega się jego zalety (zrozumiałość, przejrzystość) i eksploruje możliwości tworzenia tzw. modeli hybrydowych, łączących zalety podejścia symbolicznego i opartego na inteligencji

13 P. Casanovas et al., *Special Issue on the Semantic Web for the Legal Domain Guest Editors' Editorial: The Next Step*, „Semantic Web” 2016, nr 3, s. 213–227.

obliczeniowej¹⁴. Szczególnie należy zwrócić uwagę na trzy następujące, komplementarne podejścia:

- po pierwsze, stosowanie metod inteligencji obliczeniowej w celu wyszukania w zbiorach danych odpowiednich informacji (np. fragmentów dokumentów tekstowych – takich jak teksty aktów normatywnych), które następnie podlegają automatycznej strukturyzacji i stanowią dane wejściowe dla uprzednio przygotowanych systemów symbolicznej AI, wykonującej odpowiednie rozumowania¹⁵;
- po drugie, automatyczne ekstrahowanie ze zbiorów danych – np. z dokumentów tekstowych – struktur wiedzy, np. grafów wiedzy (ang. *knowledge graphs*) umożliwiających przejrzystą i precyzyjną reprezentację informacji i identyfikację ewentualnych problemów¹⁶;
- po trzecie, trenowanie systemów inteligencji obliczeniowej na bazie danych ustrukturyzowanych z wykorzystaniem elementów reprezentacji wiedzy właściwej dla wybranego modelu symbolicznej AI¹⁷.

3. Podejście oparte na rozwiązywaniu problemów

W zakresie ewentualnych wdrożeń systemów opartych na AI w e-legislacji należy przyjąć podejście oparte na rozwiązywaniu

14 K.D. Ashley, *Artificial Intelligence and Legal Analytics. New Tools for Legal Practice in the Digital Age*, Cambridge 2017.

15 *Ibidem*.

16 S. Servantez et al., *Computable Contracts by Extracting Obligation Logic Graphs* [w:] *ICAIL '23: Proceedings of the Nineteenth International Conference on Artificial Intelligence and Law*, red. M. Grabmair, New York 2023, s. 267–276.

17 L.K. Branting et al., *Scalable and explainable legal prediction*, „Artificial Intelligence and Law” 2021, nr 2, s. 213–238.

problemów (ang. *problem-solving*)¹⁸. Podejście zakłada, że podjęcie jakichkolwiek działań (tu: dobranie i wdrożenie odpowiednich rozwiązań technologicznych) powinno być poprzedzone próbą ścisłego zdefiniowania problemu, który te działania mają rozwiązać. Dotyczy to zwłaszcza podjęcia decyzji o automatyzacji danego działania, ponieważ wówczas konieczne jest ujęcie problemu rozwiązywanego przez owo działanie jako tzw. problemu dobrze zdefiniowanego (ang. *well-defined problem*). Problem jest dobrze zdefiniowany, jeżeli:

- stan wejściowy – nieformalnie zbiór okoliczności zastanych – jest precyzyjnie opisany;
- znany jest zbiór rozwiązań problemu (w szczególności cele, które mają być osiągnięte, są mierzalne albo co najmniej znana jest metodyka ewaluacji osiągniętego rozwiązania);
- znany jest zbiór operacji przeprowadzających stan wejściowy do poszczególnych rozwiązań.

Jeśli dysponujemy takim opisem problemu oraz odpowiednim zbiorem kryteriów, możliwe staje się porównywanie poszczególnych rozwiązań, np. pod względem kosztów ich realizacji, zaangażowania zasobów ludzkich, ryzyk związanych z każdym z możliwych rozwiązań itd. Dysponując precyzyjnym określeniem danego problemu, następnie można poddać analizie zasadność wdrożenia narzędzi technologicznych (w tym AI), kierując się następującym zestawem kroków, inspirowanym znanym testem regulacyjnym określonym w § 1 załącznika do rozporządzenie Prezesa Rady Ministrów z dnia 20 czerwca 2002 r. w sprawie „Zasad techniki prawodawczej”¹⁹, a w konsekwencji opartym na modelach racjonalnego podejmowania decyzji:

18 K. Dunbar, *Problem Solving* [w:] *A Companion to Cognitive Science*, red. W. Bechtel, G. Graham, Oxford 1998, s. 289–298.

19 T.j. Dz. U. z 2016 r. poz. 283; dalej: z.t.p.

1. przedstawienie precyzyjnej definicji problemu (może być zadaniem skomplikowanym i wymagać podziału problemu na podproblemy);
2. analiza zbioru środków potencjalnie służących rozwiązaniu problemu; w szczególności mogą być to środki o charakterze:
 - a) organizacyjnym (np. zmiana organizacji pracy zespołu) lub
 - b) prawnym (np. zmiana przepisów mająca umożliwić działanie do tej pory pozbawione podstawy prawnej) lub
 - c) edukacyjnym (np. szkolenie personelu w zakresie posługiwania się dostępną infrastrukturą techniczną) lub
 - d) technologicznym (np. wdrożenie nowego rozwiązania technicznego, w tym opartego na wybranym systemie AI);
3. porównanie diskutowanych rozwiązań pod względem przyjętych kryteriów;
4. wybór optymalnego rozwiązania;
5. ewaluacja rozwiązania po wdrożeniu.

W świetle powyższego szkicu metodyki postępowania każda decyzja dotycząca ewentualnego wdrożenia narzędzi opartych o AI w e-legislacji powinna być poprzedzona wnikliwą analizą dotyczącą tego, czy istotnie narzędzie tego rodzaju jest optymalnym środkiem do rozwiązania problemu. Przeprowadzenie takiej analizy wymaga dysponowania ekspercką wiedzą dotyczącą silnych cech oraz typowych ograniczeń systemów AI należących do poszczególnych kategorii (opisanych powyżej w sekcji 2), oraz w ramach owych szerokich kategorii, cech charakterystycznych poszczególnych modeli i wreszcie produktów informatycznych. Jest to niezbędne, by uniknąć błędów wynikających z prób wdrożenia systemów nienadających się do rozwiązywania problemów określonego rodzaju. Wiele problemów identyfikowanych przez aktorów procesu legislacyjnego może być adresowanych z wykorzystaniem rozwiązań technicznych innych niż systemy AI (tworzenie

API, ulepszanie procesów zarządzania danymi, strukturyzacja dokumentów, tworzenie modeli procesów) lub przez skupienie się na zagadnieniach innych niż techniczne, na przykład odpowiednich zmianach w regulacjach prawnych. Z kolei podkreślenia wymaga fakt, że podejścia te nie wykluczają się, lecz są komplementarne; jednak w każdym razie należy dysponować odpowiednio precyzyjnym opisem problemu, który ma podlegać rozwiązaniu. W obrocie prawnym (w tym związanym z tworzeniem prawa) często w punkcie wyjścia mamy do czynienia z tzw. problemami źle zdefiniowanymi, które charakteryzują się występowaniem co najmniej jednej z poniżej wymienionych cech²⁰:

1. niejasno sformułowane cele;
2. niedookreślone granice problemu;
3. rozbieżności pomiędzy osobami rozwiązującymi problem dotyczące podziału na fazę reprezentacji problemu i fazę jego rozwiązania;
4. brak powszechnie akceptowanych kryteriów oceny rozwiązań;
5. konieczność uzasadniania rozwiązań za pomocą środków retorycznych;
6. wrodzona sporność każdego rozwiązania;
7. brak wyraźnych kryteriów wskazujących na zakończenie procesu rozwiązywania problemu;
8. konieczność odwoływania się do stosunkowo obszernej bazy przesłanek rozumowania.

W rzeczywistości nie wszystkie problemy, z którymi mają do czynienia osoby tworzące prawo, dają się w sposób niesporny

20 J.F. Voss, *Toulmin's Model and the Solving of Ill-Structured Problems* [w:] *Arguing on the Toulmin Model. New Essays in Argument Analysis and Evaluation*, red. D. Hitchcock, B. Verheij, Dordrecht 2006, s. 303–311; C. Lynch *et al.*, *Concepts, structures, and goals: Redefining ill-definedness*, „International Journal of Artificial Intelligence in Education” 2009, nr 3, s. 253–266.

przedstawić jako dobrze zdefiniowane. Osoby odpowiedzialne za planowanie takich wdrożeń powinny więc mieć świadomość granic precyzyjnego definiowania problemów określonych rodzajów, np. w razie podjęcia decyzji o przygotowaniu nakładki na edytor tekstu mającej poprawiać komunikatywność tworzonych treści należy mieć na uwadze, że spornym może być, kiedy uzyskany tekst jest już dostatecznie komunikatywny. Przyjęcie określonych kryteriów pomiaru tej cechy może być kontestowane i osoba odpowiedzialna za wdrożenie takiego narzędzia musi zdawać sobie z tego sprawę.

Ponadto należy podkreślić, że wiele dostępnych systemów AI opartych o założenia inteligencji obliczeniowej (por. wyżej) stworzonych jest właśnie po to, by wspierać rozwiązywanie problemów, których precyzyjne zdefiniowanie okazuje się niemożliwe. Możemy np. dysponować pewnym wyobrażeniem na temat potencjalnych rozwiązań (lub przyczyn) problemu, ale sprecyzowanie przejść od stanu wejściowego do tych rozwiązań okazuje się skrajnie trudne. W takich wypadkach można np. posłużyć się modelami maszynowego uczenia nienadzorowanego w celu eksploracji danych, mając na uwadze wyszukanie informacji potencjalnie relewantnych dla rozwiązania problemu albo do wygenerowania propozycji rozwiązań, które być może będą miały co najmniej walory inspiracyjne. To, jaki stopień zdefiniowania rozwiązywanego problemu będzie potrzebny aktorowi odpowiedzialnemu za jego rozwiązanie, będzie zależało od tego, czy zamierzonym działaniem będzie dostarczenie jedynie wsparcia dla człowieka rozwiązującego problem manualnie, czy też częściowa lub całkowita automatyzacja procesu rozwiązania problemu. Oczywiście powyższe rozważania nie wyczerpują całokształtu zagadnień dotyczących wdrożeń AI w e-legislacji. Aktualnie tworzy się już bardziej szczegółowe metodyki dotyczące przygotowania wdrożeń AI w obszarze

prawa²¹, które mogą być dalej konkretyzowane w odniesieniu do szczególnego obszaru, jakim jest tworzenie prawa. Wskazane poniżej obszary problemowe dają asumpt do dokonywania prób takiej konkretyzacji.

4. Potrzeby osób zaangażowanych w proces legislacyjny

Powstanie niniejszej publikacji poprzedziła seria rozmów z osobami bezpośrednio zaangażowanymi w przebieg procesu legislacyjnego po stronie Sejmu, Rady Ministrów oraz poszczególnych ministerstw. W ich trakcie rozmówcy sygnalizowali szereg potrzeb związanych ze swoją pracą, spośród których trzy zaznaczano niemalże za każdym razem:

1. agregacja informacji – konieczność zgromadzenia w jednym miejscu danych o przebiegu prac nad projektem aktu;
2. koordynacja prac nad projektem w ramach danej komórki organizacyjnej (np. departamentu), w ramach całego podmiotu (np. ministerstwa) oraz pomiędzy różnymi podmiotami;
3. rozliczalność działań – wymóg posiadania informacji na temat tego, kto, co i kiedy wykonał w trakcie procesu legislacyjnego.

Opisane wyżej potrzeby sygnalizowały przede wszystkim osoby na stanowiskach kierowniczych, zarządzające procesami tworzenia prawa. Osoby te wyrażały konieczność dostępu do aktualnej i wiarygodnej wiedzy o procesie. Jednakże, aby mogło to nastąpić, należy przetłumaczyć owe żądania na bardziej syste-

21 M. Araszkiewicz, *HOLAI: A Methodology for Holistic Legal AI, integrating interdisciplinary, multilayered, multifunctional, cross-domain, and adaptive legal intelligence* [w:] *Human-Centred AI for Good – AI, Ethics and Law, Proceedings of the 28th International Legal Informatics Symposium IRIS 2025*, red. E. Schweighofer et al., Berno 2025, s. 159–166.

może ujęcie. A zatem, żeby agregować informacje, koordynować prace oraz zapewnić rozliczalność działań, system informatyczny musi gwarantować automatyczne gromadzenie wiarygodnych danych, zapewniać ich wersjonowanie (tj. prezentować zmienność konkretnej danej w czasie) oraz archiwizację. W konsekwencji, by móc dostarczyć wymaganej wiedzy, konieczne jest opisanie procesu legislacyjnego w wymierny sposób, za pomocą danych. Zbiór surowych danych, bez umieszczenia ich w odpowiednim kontekście, oczywiście nie rozwiąże problemów sygnalizowanych przez rozmówców, bądź uczyni to w bardzo niewielkim stopniu. Natomiast będzie stanowić bazę dla obliczenia podstawowych statystyk opisowych (jak średnia, mediana, minimum i maksimum oraz kwartyle), określenia rozkładu danych czy przeprowadzenia bardziej zaawansowanych analiz statystycznych. W dalszej perspektywie może również posłużyć do przygotowania modeli będących podstawą rozwiązań wykorzystujących szeroko pojętą sztuczną inteligencję.

5. Dane

Dane opisujące proces legislacyjny muszą być gromadzone w pewnym konkretnym celu, zawierać w sobie odpowiedź na pytania, które można wyinterpretować z wymienionych wcześniej potrzeb. To, co można rejestrować, a więc z czego można wybierać potrzebne dane, to wszystkie kategorie zdarzeń, które mają miejsce w systemie: stworzenie, odczyt, modyfikacja lub usunięcie dokumentu (projektu), daty tych wydarzeń, unikalny identyfikator dokumentu, etapy procesu, przez które dokument przechodzi, czy wreszcie efekt końcowy procesu wraz z jego ewentualnym uzasadnieniem. Jeżeli zdarzenie listy nie zachodzi automatycznie, należy odnotować osobę, która owe zdarzenie zainicjowała. Pozwoli to ustalić

stopień zaangażowania zasobów ludzkich, a także zapewnić rozliczalność działań. Szczególną wagę należy przyłożyć do rejestrowania czasu, który upłynął pomiędzy kolejnymi zdarzeniami, jak też tego, ile czasu pochłonęło konkretne zdarzenie, np. jak długo projekt aktu znajdował się na etapie opiniowania. Zestawiając ze sobą zużyty czas oraz liczbę i kategorie osób zaangażowanych w pracę, można przygotować podsumowanie osobogodzin, które zostały zużyte w trakcie procesu legislacyjnego. Po uzupełnieniu o stawkę godzinową wynagrodzenia za pracę poszczególnych zaangażowanych osób można oszacować częściowy koszt prac nad projektem.

Gromadzenie danych powinno odbywać się także na poziomie samego dokumentu. W publicznej przestrzeni informacyjnej co pewien czas prezentowane są raporty wskazujące na rzekomy wzrost skomplikowania systemu prawa. W tym celu porównuje się np. liczbę stron aktów opublikowanych w Dzienniku Ustaw w poszczególnych latach. Takie działania należy postrzegać jako zabieg marketingowy, a nie badanie naukowe, które ma określić złożoność systemu prawa; trudno uznać liczbę stron za wiarygodny wyznacznik stopnia skomplikowania regulacji, podobnie jak objętość powieści nie świadczy o jej walorach artystycznych. Chcąc zaspokoić wskazane wyżej potrzeby, należy skupić się raczej na innych cechach tekstu dokumentu, jak i opisujących go metadanych. Podstawowe metadane takie jak tytuł, numer, rodzaj czy autor aktu będą służyć do ustalenia badanej próby na potrzeby analiz statystycznych. Istotne znaczenie należy przypisać datom opisującym dokument: wydania/uchwalenia, ogłoszenia oraz wejścia w życie bądź utraty mocy obowiązującej. Dwie ostatnie kategorie mogą ponadto obejmować tzw. datę główną (dotyczącą całego aktu) oraz daty częściowe stanowiące wyjątek od daty głównej i odnoszące się do wybranych jednostek redakcyjnych. Tym sa-

mym otrzymuje się zbiór cennych danych o przebiegu końcowych etapów procesu legislacyjnego, będących uzupełnieniem danych zbieranych na poziomie systemu informatycznego.

Istotne znaczenie będą miały dane na temat rodzaju i liczby przypadków występowania przepisów poszczególnych kategorii w tekście aktu. Na podstawie jednostkowych obserwacji (wymagających potwierdzenia naukowego) można przypuszczać, że im bardziej akt jest zróżnicowany pod tym względem, tym więcej zasobów (osobowych, materialnych, czasu) może pochłoniąć jego projektowanie i procedowanie. W szczególności tyczy się to metaprzepisów, tj. przepisów zmieniających, uchylających, wprowadzających oraz upoważnień ustawowych i odesłań. Stanowią one podstawę do wyodrębnienia specyficznej kategorii metadanych, czyli relacji danego aktu z innymi aktami²². Wśród relacji za podstawowe i najczęściej występujące można uznać relacje: „zmienia” (z różnymi jej wariantami), „uchyla” (z różnymi jej wariantami) i „wykonuje”. Inne relacje mogą występować rzadziej, lecz nie można ignorować ich wagi; mowa tu m.in. o relacjach „wprowadza”, „implementuje” a także „interpretuje”. Specyficznym przypadkiem jest odesłanie, które stanowi podstawę dla wyodrębnienia osobnej relacji tylko w niektórych systemach informacji prawnej²³. Choć temat wymaga pogłębionej analizy, należy spodziewać się, że relacje będą istotnym wskaźnikiem złożoności danej regulacji. W tym miejscu należy uczynić dwa istotne zastrzeżenia. Po pierwsze: relacje w systemach informacji prawnej są w istocie rekordami w bazie danych, a więc tworzone są po fakcie, tj. po ogłoszeniu aktu we właściwym dzienniku

22 G. Wierczyński, *Profesjonalne źródła informacji o prawie*, Warszawa 2024, s. 225.

23 Przykładowo, odesłanie jako osobna relacja łącząca akty prawne funkcjonuje w Internetowym Systemie Aktów Prawnych, podczas gdy systemy LEX czy Legalis ograniczają się jedynie do zamieszczenia hiperłącza w tekście.

urzędowym. W związku z tym w procesie legislacyjnym dla oceny złożoności projektowanej regulacji większe znaczenie będą miały liczba i rodzaj użytych w projekcie kategorii przepisów. Tworzone na ich podstawie relacje powinny mieć charakter roboczy i być wykorzystywane do modelowania w czasie rzeczywistym i *ex ante* skutków nowej regulacji. Identyfikacja i określenie „głębokości” zmiany w systemie prawa na skutek powstania danej relacji mogą być zadaniami powierzonymi systemom wykorzystującym sztuczną inteligencję. Z powyższym związane jest drugie zastrzeżenie: relacje pomiędzy aktami są efektem interpretacji przepisów dokonywanej przez redakcje poszczególnych systemów informacji prawnej²⁴. Dlatego też nie powinny być traktowane jako część aktu ogłoszonego we właściwym dzienniku urzędowym. Nie zmienia to faktu, że w ramach jednego systemu informatycznego siatka relacji stanowi bezcenny zbiór metadanych, który może służyć do przeprowadzania zaawansowanych analiz lub tworzenia nowych rozwiązań informatycznych.

Na szczególną uwagę zasługują wspomniane wcześniej odesłania. W chwili obecnej jedynie Internetowy System Aktów Prawnych odnotowuje odesłania występujące w tekście aktu i to wyłącznie odesłania wyraźne; odesłania generalne są pomijane. Obsługa odesłań przez systemy informacji prawnej LEX i Legalis jest nieregularna: sporadycznie nakładany jest na nie hiperlink

24 Tekst aktu ogłoszonego we właściwym dzienniku urzędowym zawiera konkretne przepisy (np. zmieniające, uchylające), które stanowią podstawę do utworzenia relacji, a sama relacja stanowi tylko wpis w bazie danych. Aby ów wpis utworzyć, potrzebna jest rekonstrukcja normy prawnej (np. zmieniającej, uchylającej) w drodze interpretacji wspomnianych przepisów, nawet jeżeli ich brzmienie zdaje się być jednoznaczne (*omnia sunt interpretanda*). Rekonstrukcja normy nie byłaby konieczna wyłącznie w wypadku, gdyby tekst opublikowanego aktu zawierał kod źródłowy, którego wykonanie spowoduje dodanie konkretnego wpisu do bazy danych. Omówienie prawnych i technicznych możliwości przyjęcia podobnego rozwiązania jednak znacznie przekracza ramy tego opracowania.

do wskazanego aktu (w wypadku odesłań wyraźnych) bądź grupy aktów (w wypadku odesłań generalnych w LEX). Tymczasem odesłania w tekście aktu należy traktować jako ważny wskaźnik stopnia skomplikowania systemu prawa oraz pracochłonności procesu legislacyjnego. Zgodnie z § 156 ust. 1 z.t.p. odesłaniami w akcie normatywnym można posłużyć się, jeżeli zachodzi potrzeba osiągnięcia skrótowości tekstu lub zapewnienia spójności regulowanych instytucji prawnych. W praktyce oznacza to konieczność wykonania dodatkowej pracy związanej z oceną, czy spójność z regulacją, do której prowadzi odesłanie, rzeczywiście jest zachowana. Co więcej, praca ta musi być ponawiana przy każdej nowelizacji aktu odsyłającego lub do którego się odsyła. Ponadto konieczne jest monitorowanie, czy odesłanie w dalszym ciągu jest prawidłowe. Dotyczy to sytuacji, gdy odesłanie prowadzi do aktu, który utracił moc obowiązującą.

Z punktu widzenia zrozumiałości projektowanych regulacji przydatne jest ustalenie, ile z głównych jednostek redakcyjnych projektu (artykułów, paragrafów) ma swoje wyjaśnienie w uzasadnieniu projektu (konkretnych jego akapitach) lub częściach oceny skutków regulacji. Aktualnie istnieje realne ryzyko, że pewna liczba projektów w swym oryginalnym brzmieniu może nie być pokryta w całości uzasadnieniem, tj. nie można znaleźć wyraźnego wytłumaczenia dla celu niektórych z planowanych regulacji. Dążąc do zapewnienia efektywności legislacji, należy zadbać, by liczba tego typu przypadków była jak najmniejsza.

Powyższy katalog danych jest jedynie przykładowym i bardzo ograniczonym wyliczeniem tego, co można pozyskać w trakcie prac legislacyjnych. Kluczowym jest uświadomienie sobie, że dysponując odpowiednim systemem informatycznym i wykonując swoje dotychczasowe zadania, równolegle można tworzyć niezwykle cenny zbiór danych.

6. Wykorzystanie danych

Tworząc zbiór danych, należy nieustannie pamiętać, do czego ów zbiór ma posłużyć. Wyżej przedstawiono potrzeby, które zidentyfikowano w trakcie rozmów z osobami zaangażowanymi w proces legislacyjny. Wyznaczają one jedynie ogólny kierunek działania, a tym samym są dość nieprecyzyjne. Aby uzyskać bardziej szczegółową wiedzę o przebiegu procesu legislacyjnego, wyartykułowane potrzeby należy przekształcić w bardziej precyzyjne pytania badawcze. Następnie należy poszukać na nie odpowiedzi przy wykorzystaniu zgromadzonych danych i metod statystycznych, co ostatecznie będzie podstawą do stworzenia rozwiązań opartych o sztuczną inteligencję.

W większości badań naukowych, w których zastosowano metody statystyczne, analizy rozpoczyna się od obliczenia tzw. podstawowych statystyk opisowych. Można do nich zaliczyć m.in. średnią (najczęściej średnią arytmetyczną)²⁵ wraz z odchyleniem standardowym, które wskazuje, jak bardzo poszczególne pomiary skupione są wokół średniej, a więc czy uzyskane wyniki są dość powtarzalne, czy też bardziej cechuje je większe zróżnicowanie²⁶. Średnia arytmetyczna i odchylenie standardowe będą użytecznymi miarami, gdy rozkład danych w próbie będzie mniej więcej symetryczny, o charakterystycznym kształcie przypominającym dzwon. O takim rozkładzie mówi się, że jest zbliżony do rozkładu normalnego²⁷. Wartość średniej jest wtedy zbliżona do wartości mediany, czyli wartości, która dzieli pomiary w próbie dokładnie na pół: 50% pomiarów leży powyżej tej wartości, a 50% poniżej²⁸.

25 M.O. Finkelstein, B. Levin, *Statistics for Lawyers*, New York 2015, s. 2.

26 *Ibidem*, s. 21.

27 *Ibidem*, s. 116–120.

28 *Ibidem*, s. 4.

Często rozkład danych będzie asymetryczny, gdzie większość pomiarów będzie znajdowała się bliżej minimum (najmniejszej wartości) bądź maksimum (wartości największej) w próbie. Wówczas na znaczeniu zyskuje właśnie mediana, która nie jest wrażliwa na wartości skrajne (w przeciwieństwie do średniej arytmetycznej) i zawsze będzie prezentowała dokładnie środek rozkładu danych. Mówiąc o medianie, należy również wspomnieć o kwartylach, a więc przedziałach, w których mieści się 25% pomiarów. Przedstawiając dane z zakresu drugiego i trzeciego kwartyla, prezentuje się 50% wszystkich pomiarów: 25% leżących powyżej i 25% leżących poniżej mediany²⁹. Na końcu należy wskazać na istnienie dominantry (zwanej też modą), czyli wartości najczęściej występującej w próbie³⁰.

Opierając się na samych podstawowych statystykach opisowych, można starać się sformułować pytania odpowiadające opisanym wcześniej potrzebom, dla przykładu:

- Ile wynoszą: średnia, mediana i moda dla liczby dni trwania procesu legislacyjnego, w podziale na ustawy i rozporządzenia?
- Ile wynoszą: średnia, mediana i moda dla liczby dni trwania poszczególnych etapów procesu legislacyjnego?
- Jaki kształt ma rozkład danych w powyższych przypadkach?
- Ile wynoszą: średnia, mediana i moda dla liczby osobogodzin zużytych w trakcie trwania całego procesu legislacyjnego oraz jego poszczególnych etapów?
- Ile wynoszą: średnia, mediana i moda dla kosztu wytworzenia nowego aktu prawnego, w podziale na ustawy i rozporządzenia?

29 *Ibidem*, s. 4, 25.

30 *Ibidem*, s. 4.

Podobne pytania badawcze mogą formułować zarówno osoby na stanowiskach kierowniczych, jak i obsługujące poszczególne etapy procesu. Należy spodziewać się, że zestawy pytań będą się różnić, podobnie jak inne są zadania i potrzeby tych dwóch grup. Jeżeli jednak pytania będą postawione w uczciwy sposób, oraz będą odnosiły się do działań, które rzeczywiście są kluczowe w procesie legislacyjnym, ułatwią identyfikację wąskich gardeł procesu i późniejsze monitorowanie wprowadzanych zmian. Co więcej, pozwolą na sformułowanie kryteriów sukcesu dla wdrożenia systemu (bądź systemów) obsługujących e-legislację.

Bazując na doświadczeniu osób zaangażowanych w powstawanie aktów prawnych, można skorzystać z bardziej zaawansowanych metod statystycznych i np. wskazać wartości, które mogą być ze sobą skorelowane. Kiedy istnienie korelacji zostanie potwierdzone i będzie ona dostatecznie duża, można posłużyć się analizą regresji (np. analizą regresji ujemnej dwumianowej) i stworzyć model prognozujący, jak zmieni się wartość jednej zmiennej (nazywanej zmienną zależną) w momencie zmiany wartości drugiej zmiennej (określanej jako zmienna niezależna)³¹. Tym samym można pozyskać narzędzie przydatne do planowania i koordynowania prac poszczególnych komórek organizacyjnych. Analiza regresji jest tematem bardzo obszernym, którego dokładne omówienie wykracza poza granice tego opracowania. Warto jednak zaznaczyć, że w polskiej nauce prawa istnieją badania wykorzystujące analizę regresji do prognozowania dynamiki pewnych zmian w systemie prawa³².

31 *Ibidem*, s. 30–34, 369.

32 P. Ciurak, *The relationship between the number of statutes coming into force and the number of regulations coming into force: Findings from empirical studies*, „Ruch Prawniczy, Ekonomiczny i Socjologiczny,” 2024, nr 4, s. 171–184.

Dysponując podstawowymi statystykami opisowymi, wiedząc o istnieniu korelacji między wybranymi zmiennymi i bazując na wiedzy osób uczestniczących, można próbować usprawnić proces legislacyjny: zmniejszyć koszty, skrócić czas, zwiększyć efektywność pracy (zagadnienie to poruszono niżej). Ocena, czy wprowadzone usprawnienia przyniosły zamierzony efekt, może być przeprowadzona z zastosowaniem metod statystycznych, które od wielu lat z powodzeniem stosowane są w naukach społecznych, medycznych czy przyrodniczych. Przykładowo, do porównania dwóch niezależnych prób sprzed i po wprowadzeniu zmiany w procesie mogą służyć testy t-Studenta bądź U Manna-Whitney'a³³. Dla porównania ze sobą średnich z wielu prób można posłużyć się analizą wariancji bądź testem Kruskala-Wallisa³⁴. Ich uzupełnieniem będą testy post-hoc, których zadaniem jest wskazać, które konkretnie próby różnią się między sobą³⁵.

Na sam koniec należy odnotować, że wymienione analizy mogą posłużyć do oceny stosowania technik legislacyjnych. Na przykład w badaniu opublikowanym w 2018 roku wykazano, że użycie klauzuli o terminowym utrzymaniu w mocy dotychczasowych aktów prawa miejscowego sprawia, że akty je zastępujące są wydawane szybciej niż w przypadku, gdy zastosowano klauzulę o bezterminowym utrzymaniu w mocy³⁶. W innym badaniu udowodniono, że istnieją istotne statystycznie różnice pomiędzy stosowaną długo-

33 M.O. Finkelstein, B. Levin, *Statistics for...*, op. cit., s. 231, 362.

34 M.O. Finkelstein, B. Levin, *Statistics for...*, op. cit., s. 241; A. Łomnicki, *Wprowadzenie do statystyki dla przyrodników*, Warszawa 2007, s. 163–165.

35 O.J. Dunn, *Multiple Comparisons Among Means*, „Journal of the American Statistical Association” 1961, nr 293, s. 52–64; O.J. Dunn, *Multiple Comparisons Using Rank Sums*, „Technometrics” 1964, s. 241–252.

36 P. Ciurak, *Wpływ czasowego utrzymania w mocy aktów wykonawczych na wypełnianie dyspozycji ustawowych przez gminy – wnioski z badań*, „Samorząd Terytorialny” 2018, nr 11, s. 41–51.

ścią *vacatio legis* w poszczególnych latach, tak dla grupy ustaw jak i rozporządzeń. Co więcej, dla aktów niesamodzielnych (głównie nowelizacji) prawodawca przyjmował istotnie krótszą *vacatio legis* niż w wypadku aktów samodzielnych³⁷. Przedstawione wnioski nie będą zaskakujące dla osób na co dzień zajmujących się legislacją. Jednakże pokazują, w jaki sposób można wykorzystać metody statystyczne, aby oceniać skuteczność i celowość stosowania określonych technik legislacyjnych, a także wnioskować o ich wpływie na funkcjonowanie systemu prawa.

7. Efektywność legislacji

Znamienne jest, że wśród potrzeb zgłaszanych przez osoby, z którymi prowadzono rozmowy, nie wskazano wprost konieczności zapewnienia efektywności (a zatem i jakości) tworzonego prawa. O przyczynach takiego stanu rzeczy można jedynie spekulować; efektywność legislacji może być utożsamiana przez rozmówców ze zbiorczym określeniem wyartykułowanych potrzeb. Innym wytłumaczeniem może być brak świadomości, że usprawnienie obsługi procesu legislacyjnego może przekładać się na zwiększenie efektywności. W tym miejscu należy dokonać ważnego rozróżnienia. Efektywność legislacji w ujęciu węższym rozumiana jest jako efektywność samego procesu legislacyjnego. W szerszym ujęciu, które obecnie jest szeroko reprezentowane w nauce prawa, o efektywności legislacji decydują również m.in. zrozumiałość użytego języka, identyfikacja problemów wymagających rozwiązania, zastosowane środki legislacyjne czy wreszcie rzeczywiste zaspokojenie potrzeb i rozwiązanie problemów sygnalizowanych

37 P. Ciurak, A. Głowacka, *Analiza empiryczna długości vacatio legis ustaw i rozporządzeń w praktyce legislacyjnej w latach 1992–2018*, „Państwo i Prawo” 2021, nr 10, s. 106–130.

przez obywateli³⁸. Szersze ujęcie do pewnego stopnia utożsamia efektywność prawa z jego jakością.

W 1981 roku Ronald Hedlund i Patricia Freeman opublikowali pracę zatytułowaną *A Strategy for Measuring the Performance of Legislatures in Processing Decisions*³⁹. Celem publikacji było zbadanie działania organów ustawodawczych w stanach Iowa i Wisconsin. Autorzy skoncentrowali się na ocenie, jak atrybuty organizacji – personel, technologia, struktura, podział prac oraz środowisko działania (obejmujące m.in. budżet, wpływy partii czy obłożenie pracą) wpływają na wymiary efektywności działania: szybkość, sprawność, wydajność i subiektywną ocenę pracy przez osoby, które wzięły udział w badaniu. Dane do analiz pochodziły zarówno ze źródeł oficjalnych, jak i wywiadów czy kwestionariuszy wypełnianych przez osoby zaangażowane bądź obserwujące proces legislacyjny. Rezultaty badań nie były zaskakujące – generalnie wykazano pozytywny związek między atrybutami organizacji a efektywnością jej pracy. Potwierdzono też, że założenia, miary i związki między nimi, które są właściwe dla teorii organizacji i dotyczą komercyjnych podmiotów prywatnych, można z powodzeniem zastosować do podmiotów publicznych. Wyjątkowość pracy polegała jednak na zastosowaniu regresji wielorakiej do określania rodzaju i siły związku między atrybutami organizacji a wymiarami efektywności działania. Dzięki temu możliwe było stworzenie ram do monitorowania, jak zmiany w atrybutach mogą wpływać na różne wymiary efektywności pracy: które z nich wzrosną, a które zostaną zmniejszone. Wskazano również, które ze związków między atrybutami a efektywnością mają największe znaczenie,

38 M. Mousmouti, *Operationalising Quality of Legislation through the Effectiveness Test*, „Legisprudence” 2012, nr 2, s. 191–205.

39 R. Hedlund, P. Freeman, *A Strategy for Measuring the Performance of Legislatures in Processing Decisions*, „Legislative Studies Quarterly” 1981, nr 1, s. 87–113.

a które odbiegają od wcześniejszych założeń (np. zwiększenie poziomu atrybutu obniża efektywność działania).

Ujęcie zaproponowane przez R. D. Hedlund i P. K. Freeman jest bardzo empiryczne i skupione na kwestiach organizacyjnych, a tym samym oderwane od tak istotnych narzędzi, jak np. ocena skutków regulacji prowadzona zarówno *ex ante*, jak i *ex post*. Bardziej kompleksowym narzędziem jest test efektywności legislacji zaproponowany przez Marię Mousmouti⁴⁰. W tym ujęciu efektywność jest wypadkową celu powstania aktu normatywnego, stosowanych ku temu środków (w tym m.in. jasności i spójności tekstu, proporcjonalności rozwiązań, ale i wskazówek prawodawcy, jak oceniać wpływ aktu na rzeczywistość) oraz skutków jego wejścia w życie (ocenianych według wskazań ustawodawcy, zawartych w raportach niezależnych podmiotów oraz w rozstrzygnięciach sądowych). Test efektywności legislacji kładzie mocny nacisk na zestawienie ze sobą ocen prowadzonych *ex ante* i *ex post*. Traktuje „cykl życia” aktu normatywnego jako jedną całość, gdzie poszczególne etapy procesu (począwszy od projektowania, a na ewaluacji wprowadzonych regulacji kończąc) są powiązane ze sobą na zasadzie skutku i przyczyny.

Podejścia zaproponowane przez R. D. Hedlund i P. K. Freeman bądź M. Mousmouti powinny stać się inspiracją do opracowania metody oceny efektywności polskiego procesu legislacyjnego, zgodnie ze wskazanym wyżej węższym ujęciem tematu. Podstawą do przeprowadzenia oceny byłyby dane pochodzące z logów projektowanego systemu dla obsługi e-legislacji. Pozyskane wartości liczbowe stanowiłyby dane wejściowe dla przeprowadzania analiz statystycznych (np. analizy regresji czy analizy wariancji) i w ten sposób dostarczałyby kryterium ilościowego dla oceny ja-

40 M. Mousmouti, *Operationalising Quality...*, *op. cit.*

kości legislacji. Jeżeli przyjąć, że oprogramowanie analityczne byłoby kolejną warstwą systemu informatycznego, pozyskującą dane poprzez tzw. interfejs programowania aplikacji (API), działania w dużej mierze mogłyby być podejmowane automatycznie. Niemalże w czasie rzeczywistym możliwe byłoby dostarczanie informacji o statystykach procesu legislacyjnego, prognozach dotyczących przebiegu prac oraz (w razie potrzeby) działaniach podejmowanych przez konkretnych użytkowników. Co więcej, system stwarzałby możliwość zestawienia ze sobą danych, które wcześniej były rozproszone między odseparowanymi systemami funkcjonującymi w różnych podmiotach publicznych. W wyniku zachodzenia zjawiska synergii informacji możliwe byłoby uzyskanie zupełnie nowego, być może zaskakującego, wglądu w działanie procesu legislacyjnego.

Jednakże należy mieć świadomość, że powstanie opisywanego systemu wymaga doprecyzowania sygnalizowanych wyżej potrzeb i wyraźnego określenia, które wartości powinny być przedmiotem analiz. Znacząca praca dotyczyć będzie także przekonania osób zaangażowanych w konkretne działania, że system nie ma stanowić narzędzia opresji, a raczej wsparcia przy podejmowaniu decyzji. Ostatecznie, celem warstwy analitycznej systemu obsługującego e-legislację powinna być sytuacja, gdy doświadczeni pracownicy podejmują decyzje w oparciu o wysokiej jakości dane i wiarygodne analizy⁴¹. Ma to szczególne znaczenie w sytuacji stałego niedoboru zarówno zasobów osobowych, jak i finansowych w administracji publicznej⁴². Tworzenie nieefektywnego prawa, o niskiej jakości, jest po prostu zbyt kosztowne.

41 Zob. E. Walters, *Introduction: Data Analytics for Law Firms: Using Data for Smarter Legal Services* [w:] *Data-Driven Law*, red. E. Walters, New York 2018.

42 Zob. L. Tafani, *Enhancing the quality of legislation: the Italian experience*, „The Theory and Practice of Legislation” 2022, nr 1, s. 5 – 21.

8. Wyszukiwanie, klasyfikacja i prezentacja informacji

W pracy legislatora nadrzędnym zadaniem o charakterze przygotowawczym jest wyszukanie wszystkich istotnych dokumentów – a w ramach tych dokumentów ich fragmentów – mających znaczenie z punktu widzenia przygotowania projektu aktu normatywnego. Dotyczy to w pierwszym rzędzie tekstów ogłoszonych aktów normatywnych, ale także aktów znajdujących się jeszcze na wcześniejszych etapach procesu tworzenia prawa, w tym ich projektów.

Istotnymi prerekwizytami poprawy rezultatów wyszukiwania tych informacji są: stosowanie odpowiednich standardów identyfikacji dokumentów (ELI, ELI-DL) i zbiorów metadanych, a także odpowiednie ustrukturyzowanie samych dokumentów. Wdrażanie aktualnie dostępnych metod i technik AI stwarza możliwość tworzenia i testowania rozwiązań mających na celu poprawę rezultatów wyszukiwania właściwych informacji również w odniesieniu do dokumentów niespełniających wyżej wymienionych standardów.

Szczególnie należy zwrócić uwagę na następujące kategorie zadań:

1. wyszukiwanie w tekstach aktów normatywnych przepisów określonego typu – istotne zwłaszcza w związku z:
 - a) koniecznością uniknięcia redundancji (np. wprowadzenia po raz kolejny definicji już zdefiniowanego w danej gałęzi prawa terminu lub istniejącej regulacji dotyczącej pewnego stosunku prawnego) lub sprzeczności;
 - b) potrzebą ustalenia, do jakich przepisów mają odsyłać przepisy w projektowanym akcie normatywnym oraz jakie przepisy odsyłają do nowelizowanego aktu normatywnego;

- c) wymogiem ustalenia zakresu obowiązywania przepisów, zwłaszcza dat ich wejścia w życie oraz możliwości ich stosowania do stanów rzeczy zaistniałych przed ich wejściem w życie;
- 2. wyszukiwanie przepisów odnoszących się do określonych zagadnień lub posługujących się określoną terminologią – szczególnie ważne w związku z:
 - a) ponownie – koniecznością uniknięcia redundancji lub sprzeczności;
 - b) zasadnością zapoznania się z konwencjami redagowania przepisów w danej gałęzi prawa lub danym obszarze regulacyjnym – m.in. w celu zachowania konsekwencji stylistycznej albo zaproponowania odpowiednich zmian w tym zakresie;
 - c) wymogiem ustalenia zbioru przepisów, które mają podlegać zmianie lub uchyleniu.

Tytułem przykładu można przywołać niemiecki projekt klasyfikacji przepisów kodeksu cywilnego⁴³. Autorzy wyróżnili dziewięć typów semantycznych przepisów i wykonali eksperymenty, których celem była ich automatyczna klasyfikacja. Eksperymenty prowadzone były równolegle z wykorzystaniem podejścia regułowego oraz przy użyciu pięciu automatycznych klasyfikatorów opartych na uczeniu maszynowym. Uzyskany najlepszy wynik miary F1⁴⁴ przez klasyfikator Support Vector Classifier wyniósł 0,83, co należy uznać za satysfakcjonujący rezultat, biorąc pod uwagę stosunkowo niewielki rozmiar próbki treningowej (601 zdań będących elementami przepisów). Eksperymenty te są dobrym przykładem zastosowania różnych metod automatycznej klasyfikacji przepisów

43 B. Walzl et al., *Semantic types of legal norms in German laws: classification and analysis using local linear explanations*, „Artificial Intelligence and Law” 2019, nr 1, s. 43–71.

44 Miara F1 jest jedną z najczęściej wykorzystywanych miar skuteczności predykcji; jest to średnia harmoniczna precyzji (*precision*) i czułości (*recall*) testu.

oraz uczenia nadzorowanego. Prezentują także metodykę przygotowania próbki danych treningowych obejmującą w szczególności manualne etykietowanie przepisów. Realizacja takiego zadania wymaga dysponowania przygotowanym i weryfikowanym przez eksperta prawnego zestawem etykiet.

Jednak w niektórych wypadkach wyszukiwanie przepisów podobnych do planowanych/projektowanych może opierać się na cechach, które z trudem podają się ścisłej, wstępnej konceptualizacji. W tej sytuacji zasadne może być podejście oparte na podobieństwach w zakresie leksyki przepisów, w tym współwystępowania określonych wyrazów. Tego typu eksperyment został przeprowadzony w Estonii w związku z projektem CoReO (Copyright Reform Observatory)⁴⁵. Zastosowana metryka podobieństwa obejmowała występowanie tych samych czasowników, występowanie tych samych rzeczowników w połączeniu z określonymi czasownikami oraz częstość występowania takich par. W rezultacie została stworzona tabela porównująca stopień podobieństwa leksykalnego 386 aktów prawa estońskiego, a w ramach tych aktów – stopień podobieństwa poszczególnych przepisów. Wykryte znaczące podobieństwa pomiędzy przepisami (w tym przepisy powtarzające się w różnych aktach) zostały poddane analizie przez ekspertów prawnych. Omawiany eksperyment stanowi dobry przykład zastosowania rozwiązania obliczeniowego o charakterze eksploracyjnym, mającym przynieść nową wiedzę, trudną do zidentyfikowania bez wielkoskalowej analizy materiału normatywnego.

W niektórych wypadkach problem wyszukiwania może polegać na tym, że wiemy, jakich rodzajów przepisów szukamy, ale przygotowanie zbioru danych w celu zastosowania modeli ucze-

45 E. Täks *et al.*, *Creating CoReO, the Computer Assisted Copyright Reform Observatory* [w:] *Logic in the Theory and Practice of Lawmaking*, red. M. Araszkiewicz, K. Pleszka, Cham 2015, s. 253–297.

nia nadzorowanego jest zbyt skomplikowane lub czasochłonne. Jest tak w szczególności w przypadku konieczności identyfikacji różnego rodzaju odesłań w tekstach aktów normatywnych oraz dokonania ich klasyfikacji – zwłaszcza wtedy, gdy ustawodawca nie przestrzega wynikających z właściwych zasad techniki prawodawczej metod konstruowania odesłań. Przykładem eksperymentu nakierowanego na realizację tego celu z wykorzystaniem algorytmów grupowania (*clustering*) jest projekt automatyzacji wykrywania i klasyfikacji odesłań w legislacji amerykańskiej (*U.S. Code*)⁴⁶. W rezultacie eksperymentu autorzy dokonali automatycznej ekstrakcji dziewięciu różnych rodzajów odesłań w ustawodawstwie Stanów Zjednoczonych, przy czym rezultaty zostały zaprezentowane w postaci sieci cytowań (*citation network*) – techniki reprezentacji wiedzy ukazującej relacje pomiędzy badanymi obiektami w postaci grafu. Dysponowanie taką metodą wizualizacji odesłań pomiędzy przepisami pozwala legislatorowi efektywnie je identyfikować i analizować, w szczególności w związku z projektowaniem nowych odesłań.

Innym zaawansowanym eksperymentem nakierowanym na wyszukiwanie przepisów w obszernym materiale normatywnym jest badanie włoskich uczonych⁴⁷ z wykorzystaniem modelu LAMBERTA – dostrojonego na bazie włoskiego prawa wariantu modelu językowego BERT. Rezultaty eksperymentu okazały się bardzo obiecujące, ponieważ model ze znaczną dokładnością wykrywał właściwe przepisy włoskiego prawa cywilnego w reakcji na sześć różnych rodzajów zapytań (również z wykorzystaniem odniesień

46 A. Sadeghian *et al.*, *Automatic semantic edge labeling over legal citation graphs*, „Artificial Intelligence and Law” 2018, nr 2, s. 127–144.

47 A. Tagarelli, A. Simeri, *Unsupervised law article mining based on deep pre-trained language representation models with application to the Italian civil code*, „Artificial Intelligence and Law” 2022, nr 3, s. 417–473.

do takich źródeł jak orzecznictwo czy doktryna), korzystając z algorytmów uczenia nienadzorowanego (grupowania). Wynik ten jest znaczący także z tego względu, że model rozwiązywał też problemy wieloetykietowe (dany przepis mógł być zaliczony do kilku klas).

Istotnym zadaniem, którego realizacja może być wystarczająca dla celów legislatora, a która ponadto może być prerekwizytem dla zrealizowania bardziej zaawansowanych zadań analitycznych, jest segmentacja tekstów aktów normatywnych. Legislator może być zainteresowany w szczególności analizą określonych fragmentów niektórych aktów normatywnych (np. wyłącznie przepisów o wejściu w życie, załączników itp.). Segmentację tekstów aktów normatywnych można przeprowadzać np. z wykorzystaniem modeli LSTM oraz BERT, tak jak to miało miejsce w niedawnym brazylijskim studium⁴⁸. Autorzy uzyskali istotne wyniki w zakresie segmentacji brazylijskich tekstów aktów normatywnych, a następnie poprzez technikę uczenia transferowego zastosowali modele do segmentacji ustawodawstwa francuskiego, włoskiego, niemieckiego oraz amerykańskiego (Stanów Zjednoczonych).

9. Sprawdzanie jakości tworzonych tekstów

Bezsporne jest, że teksty aktów normatywnych powinny spełniać wiele kryteriów jakościowych. Teksty te powinny być, przede wszystkim, wolne od różnego rodzaju błędów. Po drugie, teksty pozbawione błędów, a przynajmniej błędów powodujących poważne konsekwencje, mogą być optymalizowane pod kilkoma istotnymi względami.

48 F. Siqueira *et al.*, *Segmenting Brazilian legislative text using weak supervision and active learning*, „Artificial Intelligence and Law” 2024, <https://doi.org/10.1007/s10506-024-09419-5> [dostęp: 23.09.2025 r.].

Możliwe jest konstruowanie różnych taksonomii błędów w tekstach aktów normatywnych. Z punktu widzenia celów tego opracowania dogodna jest taksonomia oparta na rodzaju wiedzy, potrzebnej dla stwierdzenia błędu lub co najmniej sformułowania hipotezy o jego występowaniu⁴⁹. Taksonomia nie aspiruje do kompletności i możliwe jest też tworzenie taksonomii znacznie dokładniejszych:

1. błędy literowe, interpunkcyjne i składniowe (w tym błędy w zakresie odmiany wyrazów, argumentowości funktorów, niejednoznacznych związków składniowych) – do ich stwierdzenia konieczna i wystarczająca jest wiedza z zakresu ortografii, interpunkcji i gramatyki języka polskiego; podkreślenia wymaga, że niektóre błędy tego rodzaju mogą mieć poważne konsekwencje normatywne;
2. błędy strukturalne rozumiane jako błędy dotyczące ustrukturyzowania poszczególnych części aktu normatywnego, w tym błędna numeracja jednostek redakcyjnych lub systematyzacyjnych, posłużenie się błędnymi jednostkami danego rodzaju lub ich wprowadzenie w nieuzasadniony sposób, pominięcie obligatoryjnych części tekstu aktu normatywnego – do ich stwierdzenia konieczna i wystarczająca jest znajomość zasad techniki prawodawczej oraz praktyki ich stosowania;
3. błędy polegające na naruszeniu postulatów racjonalności stawianych przed prawodawcą w obszarze cech systemu prawa, takich jak niesprzeczność tego systemu, jego kompletność (zrelatywizowana do pewnych przyjmowanych założeń), konsekwencja pojęciowa oraz brak występowania w systemie ele-

49 Por. M. Araszkiewicz, T. Żurek, *Identification of legislative errors through knowledge representation and interpretive argumentation* [w:] *AI Approaches to the Complexity of Legal Systems XI–XII*, red. V. Rodriguez-Doncel et. al., Cham 2021, s. 15–30.

mentów zbędnych; dla wykrycia błędów tego typu konieczna jest wiedza prawnicza, logiczna oraz – często – wiedza o przedmiocie danej regulacji⁵⁰. Przykładowymi błędami legislacyjnymi lokowanymi w tej grupie są:

- a) przepisy dające podstawę do rekonstrukcji norm niezgodnych (sprzecznych, przeciwnych lub niezgodnych prakseologicznie), przy czym niezgodność ta nie powinna być w prosty sposób usuwalna z wykorzystaniem reguł kolizyjnych lub dyrektyw wykładni prawa;
- b) regulacje dające podstawę do stwierdzenia wystąpienia luki rzeczywistej, tzn. braku elementu regulacji ocenianego negatywnie z uwagi na pewne zobiektywizowane kryteria; przykładami takiej luki są np. brak przepisów określających kompetencje organu administracji publicznej; brak przepisów formułujących istotne elementy procedury nieuregulowanej możliwymi do zastosowania przepisami ogólnymi; brak przepisów ustalających warunki nabycia uprawnienia itd.; dla oceny powagi tego błędu istotne jest również, aby uzupełnienie luki w drodze analogii lub jej usunięcie w drodze wykładni było uznawane za niedopuszczalne lub co najmniej sporne;
- c) przepisy powtarzające się albo zawierające wyrażenia zbędne w celu sformułowania treści normatywnych już zawartych w innych częściach przepisu.

Istotne jest, że trafne stwierdzenie występowania błędów kategorii 3) wymaga umiejętności biegłego posługiwania się dyrektywami interpretacji prawa. Co więcej, należy wyróżnić tutaj perspektywę wewnętrzną i zewnętrzną – w wypadku pierwszej

50 M. Araszkiewicz, K. Pleszka, *The Concept of Normative Consequence and Legislative Discourse* [w:] *Logic in the Theory and Practice of Lawmaking*, red. M. Araszkiewicz, K. Pleszka, Cham 2015, s. 253–297.

kontekstem dla stwierdzenia wystąpienia błędu jest tekst danego (projektu) aktu normatywnego, natomiast przyjmując perspektywę zewnętrzną należy wziąć pod uwagę kontekst tekstów innych aktów normatywnych.

Kwalifikowaną postacią błędów kategorii 3) są sytuacje, w których wadliwość tekstu aktu normatywnego może dać podstawę do stwierdzenia jego niezgodności z Konstytucją poprzez naruszenie zasad przyzwoitej legislacji wyprowadzanych z ogólnej zasady demokratycznego państwa prawa (art. 2 Konstytucji RP), w tym wymogu dostatecznej określoności przepisów prawa.

Obok sytuacji zdefiniowanych jako błędy w tekstach aktów normatywnych należy wyróżnić fragmenty tekstów, które ogólnie można uznać za nadające się do optymalizacji. Wymaga podkreślenia, że w tym miejscu nie podejmujemy dyskusji postulatów związanych z tzw. prostym językiem prawnym czy dotyczących zmiany aktualnego paradygmatu komunikowania norm prawnych w ogóle (np. odejścia od formy przepisów na rzecz formy narracyjnej lub wykorzystującej w szerokim zakresie infografiki czy materiały audiowizualne). Przyjmując zatem za standard komunikowanie norm prawnych z wykorzystaniem przepisów, należy dostrzec, że w wielu wypadkach formułowane są one w sposób suboptymalny, np. nadmiernie obszerny, zawiły, z wykorzystaniem niepotrzebnych terminów technicznych, zwrotów nadmiarowych, zbędnych konstrukcji wielokrotnie złożonych itd. Występowanie tych zjawisk niewątpliwie obniża komunikatywność i zrozumiałość tekstów aktów normatywnych, niekoniecznie z korzyścią dla jego precyzji.

Należy zatem podjąć próbę odpowiedzi na pytanie, w jaki sposób narzędzia AI mogą być stosowane przez legislatora (a w pewnym zakresie również przez redaktora tekstu aktu normatywnego), tak aby ów tekst, tworzony w sposób klasyczny (pisany

przez człowieka), miał jak najwyższą jakość, tzn. był pozbawiony określonych powyżej błędów typu 1–3 oraz w odpowiedni sposób zoptymalizowany pod względem komunikatywności oraz stopnia precyzji.

Trzeba zauważyć, że aktualnie dostępne systemy generatywnej AI oparte na dużych modelach językowych mogą być efektywnie wykorzystane w zakresie oceny oraz poprawiania tekstu pod względem interpunkcji i składni, a także wskazanych wyżej cech mających znaczenie dla komunikatywności tekstu. W szczególności, systemy oparte na LLM są w stanie sformułować kilka alternatywnych sposobów wyrażenia danej myśli w przepisie prawnym (ocena trafności tych parafraz należy jednak do użytkownika i może wymagać pogłębionej wiedzy prawnej). LLM są też w stanie wykrywać potencjalne usterki tekstu, takie jak niekonsekwencje terminologiczne, powtórzenia, posługiwanie się niezdefiniowanymi terminami technicznymi itp.

Znacznie większym wyzwaniem jest identyfikacja błędów należących do kategorii 3), ponieważ wymaga to posłużenia się wiedzą prawniczą, a szczególnie dokonania rekonstrukcji logicznej i juredyczno-pojęciowej struktury przepisów (elementów norm prawnych i struktury stosunków prawnych wyznaczanych przez normy prawne). Jest tak z tego względu, że np. sprzeczności pomiędzy normami prawnymi nieczęsto polegają na tym, że wyrażenie reprezentujące jedną z nich jest prostą negacją drugiego, lecz chodzi raczej o możliwość wyprowadzenia z nich wniosku, iż dane zachowanie jest jednocześnie przedmiotem uprawnienia i zakazu, albo że brak jest korelatywnego obowiązku dla danego uprawnienia. Automatyczna identyfikacja błędów tego rodzaju wymaga więc połączenia elementów reprezentacji wiedzy oraz modelu maszynowego uczenia. Jeśli wziąć pod uwagę pierwszy aspekt, propozycja metodyki rekonstrukcji struktury norm praw-

nych została przedstawiona w badaniu Michała Araszkiewicza, Enrico Francesconiego i Tomasza Żurka⁵¹ i obejmuje:

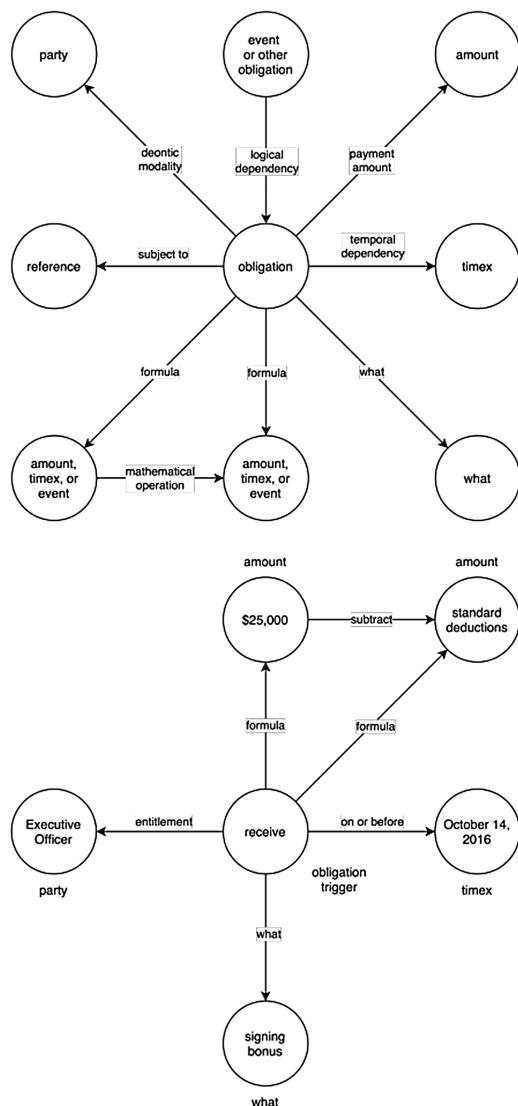
- rekonstrukcję składni przepisu prawnego;
- zmianę semantyki przepisu prawnego obejmującą elementy stosunku prawnego wyznaczanego przez przepis (np. podmiotu uprawnienia, podmiotu obowiązującego korelatywnie, treści uprawnienia, różnych sposobów korzystania z uprawnienia etc.);
- możliwość reprezentowania alternatywnych interpretacji przepisu prawnego, będących rezultatem wykładni rozszerzającej lub zwężającej interpretację pierwotnie przyjętą;
- zestaw reguł dotyczących się dopuszczalnych sposobów interpretacji prawa w odniesieniu do danego obszaru prawa.

Stosowanie takiego podejścia pozwala na reprezentowanie struktury oraz semantyki przepisów prawa z wykorzystaniem standardów Sieci Semantycznej, np. jako grafów. Taka reprezentacja umożliwia szybkie wychwycenie źródła wadliwości tekstu i dokonanie stosownej korekty przez legislatora. Pewną wadą powyższego podejścia – jako takiego – jest konieczność wykonania żmudnej pracy merytorycznej polegającej na przygotowaniu modelu regulacji. Jednakże obecnie dysponujemy modelami maszynowego uczenia pozwalającymi na dokonywanie automatycznej ekstrakcji grafów wiedzy z tekstu wyrażonego w języku naturalnym. Przykładem udanego eksperymentu z zakresu automatycznego generowania grafów z tekstów umów jest studium opublikowane przez zespół pod kierownictwem S. Servanteza⁵². Autorzy zastosowali modele BERT, SpERT i GPT-3 w celu ekstrakcji

51 M. Araszkiewicz, E. Francesconi, T. Żurek, *Identification of Legislative Errors* [w:] *ICAIL '23: Proceedings of the Nineteenth International Conference on Artificial Intelligence and Law*, red. M. Grabmair, New York 2023, s. 2–11.

52 S. Servantez et al., *Computable Contracts...*, *op. cit.*

grafów reprezentujących semantykę klauzul umownych. Schemat i konkretna instancja takiego grafu mogą zostać zaprezentowane w sposób pokazany na Rys. 1.



Rys. 1. Generowanie grafu z tekstów umów: schemat (u góry), instancja (na dole); źródło: S. Servantez et al., *Computable Contracts by Extracting Obligation Logic Graphs* [w:] *ICAIL '23: Proceedings of the Nineteenth International Conference on Artificial Intelligence and Law*, red. M. Grabmair, New York 2023, s. 267–276.

Technologie i metodyki generowania grafów reprezentujących semantykę tekstów prawnych i prawniczych są teraz na relatywnie wczesnym etapie rozwoju, ale stanowią właściwe podejście mające na celu identyfikację błędów tekstów aktów normatywnych kategorii 3). Trzeba podkreślić, że trafna ocena poprawności grafu wymaga wiedzy prawniczej użytkownika systemu.

Z kolei należy zauważyć, że systemy oparte na dużych modelach językowych stosunkowo słabo radzą sobie z rozwiązywaniem zadań bezpośrednio dotyczących struktury aktów normatywnych (w szczególności powoływania określonych jednostek redakcyjnych tekstu)⁵³. Nie jest to zaskakujące, zważywszy na sposób działania tych modeli, ale jest to cecha, którą należy wziąć pod uwagę, planując system mikrouslug dla e-legislacji.

10. Generowanie dokumentów spełniających określone kryteria

Należy zwięźle omówić ewentualność tworzenia tekstów aktów normatywnych od podstaw z wykorzystaniem systemów generatywnej AI opartych na modelach językowych – zarówno modelach dużych, trenowanych na wielkich zbiorach tekstów dostępnych w sieci, jak i modelach specjalistycznych, trenowanych na mniejszych zbiorach dokumentów prawnych.

Modele językowe bazujące na architekturze transformera i mechanizmie *self-attention* stanowią podstawę współczesnych systemów generatywnych AI⁵⁴. Architektura transformera opiera

53 A. Blair-Stanek, N. Holzenberger, B. Van Durme, *Can GPT-3 Perform Statutory Reasoning?* [w:] *ICAIL '23: Proceedings of the Nineteenth International Conference on Artificial Intelligence and Law*, red. M. Grabmair, New York 2023, s. 22–31.

54 U. Kamath *et al.*, *Large Language Models: A Deep Dive: Bridging Theory and Practice*, Cham 2024.

się na mechanizmie *self-attention*, który umożliwia modelowi analizować zależności pomiędzy wszystkimi elementami sekwencji wejściowej, niezależnie od ich długości. Dzięki temu model może skutecznie uwzględniać kontekst zarówno lokalny, jak i globalny, co jest niezwykle ważne dla generowania spójnych, akceptowalnych dla użytkownika tekstów. *Self-attention* działa poprzez przypisywanie wag poszczególnym elementom tekstu, co pozwala modelowi skupić się na najbardziej istotnych informacjach w danym kontekście.

Trzeba zaznaczyć, że ewaluacja wyników dostarczanych przez system generatywnej AI może tworzyć różnorodne wyzwania. W przeciwieństwie bowiem do rezultatów generowanych przez klasyczne systemy klasyfikacyjne, rekomendacyjne czy predykcyjne, wyniki działania modeli językowych (teksty) zazwyczaj nie są klasyfikowane wprost jako poprawne lub błędne, lecz oceniane z punktu widzenia pewnych kryteriów jakościowych, takich jak odpowiedni dobór słów, uporządkowanie wyводу, jasność i zrozumiałość tekstu, adekwatność odpowiedzi do postawionego pytania, oddanie stylu charakterystycznego dla danego rejestru języka itp. Uwaga ta nie odnosi się do kwestii generowania przez system fałszywych informacji (tzw. halucynacji, zob. poniżej).

Ocena jakości generowanych tekstów może oczywiście mieć miejsce w odniesieniu do rezultatów konkretnych eksperymentów, ale można sformułować kilka stwierdzeń związanych z korzystaniem z tej technologii. Należy zatem oczekiwać, że rezultaty działań modeli (jako takich) mogą być lepsze w przypadku projektowania całkowicie nowych regulacji niż np. nowelizacji regulacji istniejących, a to z uwagi na problemy związane z przetwarzaniem informacji dotyczących struktury tekstów prawnych odzwierciedlonej w jednostkach redakcyjnych i systematyzacyjnych.

Mimo tych zaawansowanych możliwości modele LLM mają istotne wady, jak np. zjawisko halucynacji. Polega ono na generowaniu treści niemających pokrycia w danych wejściowych ani w rzeczywistości, co może prowadzić do poważnych błędów w kontekście prawnym, takich jak odniesienia do nieistniejących przepisów lub niewłaściwe cytowanie aktów normatywnych. Ponadto modele często generują treści ogólnikowe, które mogą być nieprzydatne w specyficznych zastosowaniach, oraz bywają podatne na mylące zapytania (*prompts*), co może skutkować nieadekwatnymi lub błędnymi odpowiedziami. Z tego powodu istotne znaczenie ma *prompt engineering*, czyli umiejętność projektowania zapytań kierowanych do modelu. Dobrze skonstruowany *prompt* powinien być precyzyjny, kontekstowy i zawierać niezbędne informacje, które ukierunkują model na generowanie trafnych odpowiedzi. Oczywiście z uwagi istotę działania modeli językowych żaden system zapytań nie daje gwarancji uzyskania prawidłowych odpowiedzi, gdyż modele językowe jako takie nie są opracowywane po to, by dostarczać informacji faktograficznej.

Retrieval-Augmented Generation (RAG) to technika łącząca generatywne możliwości modeli językowych z dostępem do zewnętrznych baz wiedzy. Mechanizm działania RAG polega na tym, że model przed wygenerowaniem odpowiedzi wyszukuje informacje w zbiorach danych, takich jak dokumenty prawne, bazy danych czy specjalistyczne publikacje. Następnie wyniki tego wyszukiwania są integrowane w procesie generacji tekstu, co pozwala na zwiększenie dokładności i wiarygodności wyników, przy czym podkreślenia wymaga: iż nie daje jednak gwarancji ich poprawności.

W kontekście aktów normatywnych RAG umożliwia dynamiczne odwoływanie się do aktualnych przepisów, powiązanych

regulacji lub też innych relewantnych dokumentów (na przykład orzeczeń sądowych dotyczących interpretacji podobnie brzmiących przepisów), co jest szczególnie istotne w przypadku nowelizacji. Dzięki stosowaniu RAG możliwe jest ograniczenie ryzyka halucynacji, ponieważ model generuje treści na podstawie zweryfikowanych informacji, zamiast polegać wyłącznie na danych wewnętrznych używanych podczas treningu. Podkreślenia wymaga, że RAG, potencjalnie ograniczając ryzyko halucynacji lub nadmiernej ogólnikowości generowanego rezultatu, nie gwarantuje eliminacji tych zjawisk. Skuteczność stosowania tej techniki zależy między innymi od charakterystyki zewnętrznej bazy danych: zakresu gromadzonych danych, połączeń między nimi oraz ich merytorycznej poprawności (por. rozważania powyżej na temat danych i ich wykorzystania).

Integracja modeli LLM, podejścia RAG oraz technik reprezentacji wiedzy stanowi kolejny krok w kierunku zapewnienia wysokiej dokładności i merytorycznej poprawności generowanych tekstów aktów normatywnych. Reprezentacja wiedzy daje możliwość formalizacji złożonych struktur prawnych oraz semantyki aktów normatywnych, a w szczególności odwzorowywania hierarchii przepisów, sieci odesłań czy relacji systemowych. Dzięki temu można odwzorować strukturę logiczną i semantyczną aktu prawnego, co umożliwia automatyczne wykrywanie potencjalnych problemów natury jurydycznej, takich jak sprzeczności czy luki w projektowanej regulacji. Trzeba zaznaczyć, że podejście to znajduje się obecnie we wstępnej fazie rozwoju⁵⁵.

Na koniec należy jeszcze odnieść się do zagadnienia, czy dla celów rozwiązywania problemów w procesach prawotwórczych

55 N.H. Thanh, K. Satoh, *KRAG Framework for Enhancing LLMs in the Legal Domain*, 2024, (online) <https://arxiv.org/abs/2410.07551>

bardziej dogodne jest korzystanie z dużych modeli językowych, dodatkowo trenowanych na materiale normatywnym lub z wykorzystaniem technik RAG, czy też lepszą strategią jest tworzenie i stosowanie specjalistycznych, prawniczych modeli językowych, trenowanych przede wszystkim na dokumentach specjalistycznych. Jest to oczywiście pytanie empiryczne i w ciągu najbliższych lat należy oczekiwać postępu efektywności działania modeli obydwu typów⁵⁶.

Bibliografia

- Araszkiewicz M., *Algorytmizacja myślenia prawniczego. Modele, możliwości, ograniczenia* [w:] *Legal tech. Czyli jak bezpiecznie korzystać z narzędzi informatycznych w organizacji, w tym w kancelarii oraz dziale prawnym*, red. D. Szostek, Warszawa 2021.
- Araszkiewicz M., *HOLAI: A Methodology for Holistic Legal AI, integrating interdisciplinary, multilayered, multifunctional, cross-domain, and adaptive legal intelligence* [w:] *Human-Centred AI for Good – AI, Ethics and Law, Proceedings of the 28th International Legal Informatics Symposium IRIS 2025*, red. E. Schweighofer et al., Bern 2025.
- Araszkiewicz M., Francesconi E., Żurek T., *Identification of Legislative Errors* [w:] *ICAIL '23: Proceedings of the Nineteenth International Conference on Artificial Intelligence and Law*, red. M. Grabmair, New York 2023.
- Araszkiewicz M., Pleszka K., *The Concept of Normative Consequence and Legislative Discourse* [w:] *Logic in the Theory and Practice of Lawmaking*, red. M. Araszkiewicz, K. Pleszka, Cham 2015.

56 Por. J. Cui et al., *ChatLaw: Open-Source Legal Large Language Model with Integrated External Knowledge Bases*, 2023, (online) <https://arxiv.org/abs/2306.16092>

- Araszkiewicz M., Żurek T., *Identification of legislative errors through knowledge representation and interpretive argumentation* [w:] *AI Approaches to the Complexity of Legal Systems XI-XII*, red. V. Rodriguez-Doncel et. al., Cham 2021.
- Arrieta A.B. et al., *Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI*, „Information Fusion” 2020, nr 58.
- Ashley K.D., *Artificial Intelligence and Legal Analytics. New Tools for Legal Practice in the Digital Age*, Cambridge 2017.
- Blair-Stanek A., Holzenberger N., Van Durme B., *Can GPT-3 Perform Statutory Reasoning?* [w:] *ICAIL '23: Proceedings of the Nineteenth International Conference on Artificial Intelligence and Law*, red. M. Grabmair, New York 2023.
- Branting L.K. et al., *Scalable and explainable legal prediction*, „Artificial Intelligence and Law” 2021, nr 2.
- Casanovas P. et al., *Special Issue on the Semantic Web for the Legal Domain Guest Editors' Editorial: The Next Step*, „Semantic Web” 2016, nr 3.
- Ciurak P., *The relationship between the number of statutes coming into force and the number of regulations coming into force: Findings from empirical studies*, „Ruch Prawniczy, Ekonomiczny i Socjologiczny,” 2024, nr 4.
- Ciurak P., *Wpływ czasowego utrzymania w mocy aktów wykonawczych na wypełnianie dyspozycji ustawowych przez gminy – wnioski z badań*, „Samorząd Terytorialny” 2018, nr 11.
- Ciurak P., Głowacka A., *Analiza empiryczna długości vacatio legis ustaw i rozporządzeń w praktyce legislacyjnej w latach 1992–2018*, „Państwo i Prawo” 2021, nr 10.
- Cui J. et al., *ChatLaw: Open-Source Legal Large Language Model with Integrated External Knowledge Bases*, 2023, <https://arxiv.org/abs/2306.16092> [dostęp: 23.09.2025 r.].

- Dunbar K., *Problem Solving* [w:] *A Companion to Cognitive Science*, red. W. Bechtel, G. Graham, Oxford 1998.
- Dunn O.J., *Multiple Comparisons Among Means*, „Journal of the American Statistical Association” 1961, nr 293.
- Dunn O.J., *Multiple Comparisons Using Rank Sums*, „Technometrics” 1964.
- Finkelstein M.O., Levin B., *Statistics for Lawyers*, New York 2015.
- Flasiński M., *Introduction to Artificial Intelligence*, Cham 2016.
- Hedlund R. D., Freeman P. K., *A Strategy for Measuring the Performance of Legislatures in Processing Decisions*, „Legislative Studies Quarterly” 1981, nr 1.
- Kamath U. et al., *Large Language Models: A Deep Dive: Bridging Theory and Practice*, Cham 2024.
- Kubat M., *An Introduction to Machine Learning*, Cham 2017.
- Lynch C. et al., *Concepts, structures, and goals: Redefining ill-definedness*, „International Journal of Artificial Intelligence in Education” 2009, nr 3.
- Łomnicki A., *Wprowadzenie do statystyki dla przyrodników*, Warszawa 2007.
- Mousmouti M., *Operationalising Quality of Legislation through the Effectiveness Test*, „Legisprudence” 2012, nr 2.
- OECD, *Explanatory memorandum on the updated OECD definition of an AI system*, „OECD Artificial Intelligence Papers” 2024, nr 8, <https://doi.org/10.1787/623da898-en> [dostęp: 23.09.2025 r.].
- Russell S., Norvig P., *Artificial Intelligence. A Modern Approach*, 3rd ed., Upper Saddle River 2016.
- Sadeghian A. et al., *Automatic semantic edge labeling over legal citation graphs*, „Artificial Intelligence and Law” 2018, nr 2.
- Servantez S. et al., *Computable Contracts by Extracting Obligation Logic Graphs* [w:] *ICAAIL '23: Proceedings of the Nineteenth*

- International Conference on Artificial Intelligence and Law*, red. M. Grabmair, New York 2023.
- Siqueira F. *et al.*, *Segmenting Brazilian legislative text using weak supervision and active learning*, „Artificial Intelligence and Law” 2024, <https://doi.org/10.1007/s10506-024-09419-5> [dostęp: 23.09.2025 r.].
- Tafari L., *Enhancing the quality of legislation: the Italian experience*, „The Theory and Practice of Legislation” 2022, nr 1.
- Tagarelli A., Simeri A., *Unsupervised law article mining based on deep pre-trained language representation models with application to the Italian civil code*, „Artificial Intelligence and Law” 2022, nr 3.
- Täks E. *et al.*, *Creating CoReO, the Computer Assisted Copyright Reform Observatory* [w:] *Logic in the Theory and Practice of Lawmaking*, red. M. Araszkiewicz, K. Pleszka, Cham 2015.
- Thanh N.H., Satoh K., *KRAG Framework for Enhancing LLMs in the Legal Domain*, 2024, <https://arxiv.org/abs/2410.07551> [dostęp: 23.09.2025 r.].
- Voss J. F., *Toulmin’s Model and the Solving of Ill-Structured Problems* [w:] *Arguing on the Toulmin Model. New Essays in Argument Analysis and Evaluation*, red. D. Hitchcock, B. Verheij, Dordrecht 2006.
- Walters E., *Introduction: Data Analytics for Law Firms: Using Data for Smarter Legal Services* [w:] *Data-Driven Law*, red. E. Walters, New York 2018.
- Waltl B. *et al.*, *Semantic types of legal norms in German laws: classification and analysis using local linear explanations*, „Artificial Intelligence and Law” 2019, nr 1.
- Wierczyński G., *Profesjonalne źródła informacji o prawie*, Warszawa 2024.
- Yu L., *Introduction to the Semantic Web and Semantic Web Services*, Boca Raton 2019.