



Fundusze Europejskie
dla Rozwoju Społecznego



Rzeczpospolita
Polska

Dofinansowane przez
Unię Europejską



Politechnika Świętokrzyska
Wydział Zarządzania i Modelowania Komputerowego

Kierunek studiów:
Inżynieria danych

Dariusz Dobrowolski

Materiały dydaktyczne do przedmiotu

INŻYNIERIA LINGWISTYCZNA

opracowane w ramach realizacji Projektu
**„Dostosowanie kształcenia
w Politechnice Świętokrzyskiej do potrzeb
współczesnej gospodarki”**
FERS.01.05-IP.08-0234/23

Kielce, 2025



Politechnika Świętokrzyska
Kielce University of Technology

*Projekt „Dostosowanie kształcenia w Politechnice Świętokrzyskiej do potrzeb współczesnej gospodarki”
nr FERS.01.05-IP.08-0234/23*



Spis treści

1. Wprowadzenie do inżynierii lingwistycznej	4
2. Podstawy teoretyczne inżynierii lingwistycznej	5
2.1 Lingwistyka teoretyczna a lingwistyka stosowana.....	5
2.2 Modele języka naturalnego	7
2.3 Reprezentacje wiedzy językowej	11
3. Historia i rozwój inżynierii lingwistycznej.....	14
3.1 Początki badań nad językiem naturalnym	14
3.2 Rozwój technologii językowych.....	16
3.3 Współczesne nurty badawcze	20
4. Technologie przetwarzania języka naturalnego	22
4.1 Analiza składniowa i semantyczna.....	22
4.2 Przetwarzanie morfologiczne	26
4.3 Generowanie języka naturalnego	28
5. Uczenie maszynowe w inżynierii lingwistycznej.....	33
5.1 Metody klasyczne	33
5.2 Uczenie głębokie.....	37
6. Zastosowania inżynierii lingwistycznej	41
6.1 Systemy tłumaczenia maszynowego	41
6.2 Analiza sentymentu.....	43
6.3 Wyszukiwanie informacji.....	45
6.4 Rozpoznawanie mowy	50
7. Etyka i wyzwania w inżynierii lingwistycznej	54
7.1 Ochrona prywatności użytkowników	54
7.2 Bias i sprawiedliwość w modelach językowych.....	56
8. Przyszłość inżynierii lingwistycznej.....	58
8.1 Integracja z multimodalnymi systemami AI	58
8.2 Rozwój interfejsów człowiek-maszyna.....	60



Fundusze Europejskie
dla Rozwoju Społecznego



Rzeczpospolita
Polska

Dofinansowane przez
Unię Europejską



8.3 Potencjalne kierunki badań	62
9 Zakończenie	64
10. Bibliografia	66



Materiały dydaktyczne objęte licencją Creative Commons BY 4.0.

Licencja dostępna pod adresem: <https://creativecommons.org/licenses/by/4.0/>



1. Wprowadzenie do inżynierii lingwistycznej

Inżynieria lingwistyczna zajmuje się konstruowaniem systemów i narzędzi umożliwiających komputerom analizę, interpretację oraz generowanie języka naturalnego. Kluczowym aspektem tej dziedziny jest łączenie metod informatyki z wiedzą z zakresu językoznawstwa, tak aby możliwe było odwzorowanie procesów komunikacji międzyludzkiej w formie cyfrowej. W praktyce oznacza to projektowanie mechanizmów potrafiących rozpoznawać strukturę wypowiedzi, ich znaczenie oraz kontekst użycia.

Takie podejście wymaga pracy na dużych zbiorach danych tekstowych, tzw. korpusach językowych, które stanowią punkt wyjścia do trenowania modeli statystycznych i sieci neuronowych wykorzystywanych w przetwarzaniu języka. Proces tworzenia aplikacji w tej dziedzinie często rozpoczyna się od reprezentacji języka w postaci sekwencji znaków lub słów, które następnie przekształca się w formy umożliwiające analizę komputerową. Przykładowo, tekst może być reprezentowany jako lista tokenów, co pozwala na operacje takie jak wyszukiwanie określonych jednostek leksykalnych czy kalkulacja częstotliwości występowania słów. Już proste statystyki, takie jak np. rozkład długości zdań czy częstość poszczególnych elementów gramatycznych, dostarczają informacji istotnych dla zadań takich jak klasyfikacja dokumentów czy ekstrakcja informacji.

Języki naturalne posiadają wiele niejednoznaczności, które utrudniają ich automatyczne przetwarzanie. Jednym z kluczowych wyzwań jest polisemiczność, czyli sytuacja, gdy jedno słowo może mieć różne znaczenia zależnie od kontekstu. Rozwiązywanie takich problemów wymaga technik analizy składniowej i semantycznej oraz uwzględnienia wiedzy o świecie, w którym funkcjonuje język. W tym celu korzysta się z zasobów takich jak słowniki elektroniczne, ontologie czy oznaczone korpusy.

Rozpoznawanie mowy stanowi kolejną ważną gałąź inżynierii lingwistycznej. Tutaj zamiast tekstu źródłem informacji jest sygnał akustyczny reprezentujący wypowiedź ludzką. Techniki takie jak automatyczne rozpoznawanie mowy (ASR) przekształcają te sygnały na tekst poprzez analizę widmową dźwięku oraz dopasowywanie wzorców fonetycznych. Po uzyskaniu formy tekstowej możliwe jest dalsze przetwarzanie, tłumaczenie na inne języki czy interpretacja poleceń dla systemu komputerowego.

Współczesne rozwiązania w inżynierii lingwistycznej coraz częściej opierają się na modelach uczenia maszynowego. Systemy te potrafią samodzielnie uczyć się reguł z danych treningowych bez manualnego definiowania każdego aspektu przetwarzania. Przykładem są algorytmy uczenia głębokiego stosowane zarówno do analizy treści pisanej, jak i mowy. Takie modele wykorzystują wielowarstwowe sieci neuronowe zdolne do reprezentacji bardzo złożonych zależności pomiędzy elementami języka.



Szczególnie dużą rolę odgrywają tu architektury transformatorowe, które okazały się niezwykle skuteczne w uczeniu kontekstowego znaczenia słów dzięki mechanizmowi tzw. uwagi.

Rozwój dużych modeli językowych (LLM) rozszerzył możliwości inżynierii lingwistycznej o generowanie treści spełniającej kryteria spójności i poprawności stylistycznej niemal na poziomie ludzkim. Modele takie jak GPT-3 uczone były na bilionach słów pozyskanych z książek, artykułów czy stron internetowych i potrafią wykonywać różnorodne zadania: od streszczania tekstu, przez tłumaczenie i odpowiadanie na pytania, po tworzenie nowych fragmentów prozy. W porównaniu do dawnych regułowych systemów NLP oferują one większą elastyczność i lepsze dopasowanie do nieprzewidywalnych scenariuszy komunikacyjnych.

Historia inżynierii lingwistycznej obejmuje również etap programowania prostych chatbotów bazujących na regułach lub prostych algorytmach wzorcowego dopasowania tekstu. Choć ich możliwości były ograniczone i najczęściej sprowadzały się do odpowiadania według ustalonego schematu dialogowego, stanowiły fundament pod bardziej zaawansowane systemy interakcji człowiek–komputer. Obecnie chatboty mogą korzystać z pełnej infrastruktury LLM, co pozwala im lepiej radzić sobie z utrzymaniem spójnej rozmowy i reagowaniem na nietypowe pytania użytkownika.

Nie można jednak zapominać o aspektach praktycznych tej dyscypliny. Tworzenie systemów inżynierii lingwistycznej wymaga znajomości nie tylko algorytmiki, ale także umiejętności integracji różnych komponentów, od warstwy akwizycji danych (np. mikrofony lub moduły OCR), przez warstwę analityczną (modele językowe), aż po prezentację wyników odbiorcy (syntezy mowy czy interfejsy wizualne). Korzystanie z bibliotek programistycznych takich jak NLTK czy specjalistycznych zasobów jak platforma Hugging Face Hub upraszcza dostęp do gotowych modeli oraz ułatwia eksperymentowanie z różnymi typami zadań NLP. Analiza jakości odpowiedzi powinna uwzględniać zarówno czynniki lingwistyczne, jak i etyczne problemy związane ze stronniczością obecną w danych uczących. Tylko całościowe spojrzenie umożliwia ocenę przydatności danego narzędzia przy konkretnym zastosowaniu, czy to będzie pomoc w tłumaczeniach, asystent głosowy obsługujący inteligentny dom, czy też moduł klasyfikacji treści dla redakcji internetowej.

2 Podstawy teoretyczne inżynierii lingwistycznej

2.1 Lingwistyka teoretyczna a lingwistyka stosowana

Lingwistyka teoretyczna zajmuje się formułowaniem modeli, które opisują strukturę i funkcjonowanie języka w sposób możliwie abstrakcyjny. Jej celem jest uchwycenie uniwersalnych reguł rządzących składnią, morfologią, semantyką czy fonologią.



Takie podejście odwołuje się często do formalizmów matematycznych oraz logiki, które pozwalają na precyzyjne formułowanie hipotez dotyczących organizacji języka. Przykładem może być tworzenie gramatyk opisujących wszystkie możliwe zdania poprawne w danym języku lub badań nad strukturą ontologii zawierających relacje między pojęciami.

W lingwistyce teoretycznej istotne jest unikanie nadmiernego uzależniania opisów od konkretnych technologii – chodzi o modele, które mogą funkcjonować niezależnie od implementacji komputerowych. Z kolei lingwistyka stosowana koncentruje się na praktycznym wykorzystaniu ustaleń teoretycznych w rozwiązywaniu realnych problemów komunikacyjnych. Tu przykładem są systemy tłumaczeń maszynowych, oparte zarówno na klasycznych metodach analizy składniowej, jak i nowoczesnych architekturach głębokich sieci neuronowych. Implementacje tego typu muszą uwzględniać ograniczenia danych wejściowych, wymagania wydajnościowe oraz aspekty użytkowe. Dlatego proces projektowania narzędzi często obejmuje etap dostosowania modeli teoretycznych do warunków rzeczywistych, z redukcją złożoności czy uproszczeniem reprezentacji. Oba nurty wchodzą ze sobą w dialog. Modele teoretyczne są testowane poprzez ich użycie w aplikacjach NLP, takich jak rozpoznawanie mowy czy analiza sentymentu tekstu. Wyniki wdrożeń mogą ujawniać braki i niedoskonałości teorii, prowokując jej dalsze rozwijanie. Na przykład mechanizmy przetwarzania języka oparte na korpusach pokazują, że część intuicyjnie przyjętych reguł gramatycznych bywa statystycznie mało reprezentatywna wobec rzeczywistego użycia języka. Taka informacja ma wartość zwrotną dla badaczy rozwijających ramy teoretyczne.

Warto zauważyć pewne napięcie między czysto naukowym podejściem a wymogami praktyki. Lingwistyka stosowana niekiedy odchodzi od rygorystycznego opisu struktury języka na rzecz osiągnięcia wysokiej wydajności lub akceptowalnego poziomu błędu w konkretnej aplikacji. W przypadku dużych modeli językowych LLM wykorzystujących transformery kluczową rolę gra dostosowanie parametrów modelu do zadania – co może odbiegać od pierwotnej wizji modelu jako uniwersalnego narzędzia opisu języka.

Lingwistyka teoretyczna obejmuje także analizę sposobów uczenia się języka przez człowieka – co inspiruje konstrukcję algorytmów uczących się ze zbiorów danych w sposób zbliżony do mechanizmów percepcji i akwizycji mowy u ludzi.

Tymczasem lingwistyka stosowana przenosi te idee np. do systemów automatycznej transkrypcji nagrań audio czy interfejsów głosowych wykorzystywanych w asystentach AI. Powstaje tu przestrzeń eksperymentalna: teoria podpowiada, jakie mechanizmy mogą działać skuteczniej, zaś praktyczne wdrożenie pokazuje granice tych możliwości. Równie interesującym polem styku jest przetwarzanie danych wielojęzycznych. Zasoby takie jak WordNet czy wielojęzyczne korpusy są



opracowywane z myślą o językoznawczej poprawności definicji i relacji między jednostkami leksykalnymi, co odpowiada preferencjom nurtu teoretycznego. Jednak dopiero narzędzia stosowane, wyszukiwarki semantyczne, translatory kontekstowe, weryfikują użyteczność tych danych w codziennych zadaniach komunikacyjnych.

Modelowanie języka w obszarze teorii często sprowadza się do zagadnień czysto abstrakcyjnych: jakie cechy są konieczne dla uzyskania poprawnej interpretacji czy reprezentacji znaczenia? Lingwistyka stosowana uzupełnia tę perspektywę o elementy inżynieryjne: jak zapisać słowo jako wektor (embedding), aby algorytm mógł operować na jego reprezentacji numerycznej bez utraty sensu semantycznego? Różnica polega więc na tym, że pierwsza dyscyplina skupia się na konceptualnym „dlaczego”, druga zaś na operacyjnym „jak”. Nie można pomijać faktu, że coraz częściej granica między tymi obszarami się zaciera. Automatyczne optymalizowanie zapytań przez modele LLM korzysta z wiedzy teoretycznej o tym, jakie cechy mają pytania lub instrukcje sprzyjające trafnym odpowiedziom, lecz jednocześnie jest czysto zastosowaniową techniką podnoszącą efektywność pracy użytkownika. Ten trend pokazuje większą integrację obu podejść.

W ujęciu pragmatycznym lingwistyka stosowana musi także zmagać się z kwestiami jakości danych wejściowych oraz skalowalności systemów. Teoria może oferować modele odporne na szum informacyjny czy niejednoznaczność leksykalną, ale dopiero wdrożenie ujawnia koszty obliczeniowe takich rozwiązań i ich wpływ na czas odpowiedzi systemu. Kryteria takie jak klarowność interfejsu użytkownika czy możliwość integracji z innymi aplikacjami informatycznymi stają się równie ważne jak zgodność z modelem gramatycznym. Analiza relacji między lingwistyką teoretyczną a stosowaną wskazuje więc na dynamiczne sprzężenie zwrotne: odkrycia jednej inspirują drugą, podczas gdy ograniczenia praktyki prowokują rewizję teorii. Można powiedzieć, że skuteczne projekty NLP powstają właśnie tam, gdzie oba nurty spotykają się we wspólnym punkcie, testując hipotezy na rzeczywistych danych i modyfikując je tak, aby służyły zarówno nauce, jak i technologii użytkowej.

2.2 Modele języka naturalnego

Gramatyki formalne

Gramatyki formalne stanowią jeden z fundamentalnych sposobów opisu struktury języka, pozwalający na jego modelowanie w sposób ścisły i jednoznaczny. W praktyce są to zestawy reguł określających, jakie ciągi symboli (najczęściej słów lub tokenów) można uznać za poprawne pod względem składniowym. Reguły te operują na abstrakcyjnych kategoriach gramatycznych, takich jak rzeczownik czy czasownik, które można dalej rozbudowywać o cechy morfo syntaktyczne. Taki formalizm jest przydatny zarówno w analizie języka, jak i jego generowaniu.



Można wyróżnić kilka typów gramatyk formalnych zgodnie z klasyfikacją Chomsky'ego, od najbardziej restrykcyjnych gramatyk regularnych po nieograniczone gramatyki typu zero. Każdy poziom tej hierarchii ma swoje zastosowania w przetwarzaniu języka naturalnego. Gramatyki bezkontekstowe (CFG) odgrywają szczególną rolę, ponieważ oferują kompromis między mocą wyrazu a prostotą obliczeniową. Składają się z reguł produkcji, które przekształcają symbole nieterminalne w inne symbole nieterminalne lub terminalne, co umożliwia konstruowanie drzewa składniowego dla danego zdania. Takie drzewa składniowe ujawniają hierarchiczną strukturę wypowiedzi, ilustrują relacje pomiędzy frazami nominalnymi, czasownikowymi czy przyimkowymi. Analiza za pomocą CFG pozwala też na wykrywanie niejednoznaczności syntaktycznych, kiedy dane zdanie może mieć więcej niż jedną poprawną reprezentację drzewa. Rozwiązaniem bywa rozszerzenie modelu o mechanizmy ważenia reguł lub dodatkowe ograniczenia wynikające z informacji kontekstowych. Jednakże sama składnia, nawet opisana rygorystycznie przez gramatykę formalną, bywa niewystarczająca do dokładnej interpretacji znaczenia zdań.

Dlatego powstały podejścia oparte na gramatykach z cechami, które wzbogacają symbol nieterminalny o zestaw atrybutów takich jak liczba, rodzaj czy przypadek. Przetwarzanie struktur cech wymaga bardziej zaawansowanych algorytmów parsowania, ale umożliwia lepsze odwzorowanie rzeczywistych właściwości języka. Te mechanizmy są istotne zwłaszcza w językach fleksyjnych, gdzie forma wyrazu dostarcza krytycznych informacji gramatycznych. Kontekstowe podejścia do gramatyk integrują także zależności między elementami zdania w ujęciu tzw. gramatyki dependencji. Tutaj relacje są przedstawiane jako zbiór powiązań między słowami, zamiast grupowania ich w większe frazy, co bywa bardziej naturalne dla pewnych zadań analitycznych. Instrukcja parsowania oparta na zależnościach często korzysta z metod statystycznych lub uczenia maszynowego dla ustalenia najbardziej prawdopodobnego grafu powiązań pomiędzy tokenami zdania.

Pod kątem implementacji komputerowej wybór odpowiedniej gramatyki musi uwzględniać złożoność obliczeniową i jakość dostępnych danych treningowych. Parsery oparte na CFG są relatywnie proste do wdrożenia i sprawdzają się przy danych czystych syntaktycznie. Natomiast w przypadku tekstów pochodzących ze źródeł nieformalnych, np. mediów społecznościowych, konieczne staje się połączenie regułowych opisów języka ze statystycznym modelowaniem błędów i odstępstw od normy.

Gramatyki formalne stanowią również podstawę systemów tłumaczy maszynowych działających według modelu transferowego: najpierw analizują zdanie wejściowe w języku źródłowym przy użyciu parsera zgodnego z jego gramatyką, następnie dokonują transformacji struktury syntaktycznej na odpowiadającą jej formę docelową zgodną z gramatyką języka wyjściowego. To pozwala zachować spójność struktur



oraz poprawność syntaktyczną tłumaczeń nawet wtedy, gdy znaczenia są przenoszone idiomatycznie.

W kontekście dużych modeli językowych pojawia się interesujące pojęcie: architektury transformatorowe mogą generować składniowo poprawne zdania bez jawnego kodowania jakiegokolwiek gramatyki formalnej, „uczą się” jej statystycznie na podstawie biliona przykładów tekstu. Mimo to formalne opisy pozostają wartościowe jako narzędzie diagnostyczne i kontrolne: pozwalają ocenić stopień zgodności generowanego tekstu z normą składniową oraz identyfikować odchylenia wynikające np. z halucynacji modelu. Na etapie rozwoju aplikacji NLP integracja parsera obsługującego formalną gramatykę z modułem semantycznym daje możliwość bardziej precyzyjnej interpretacji poleceń użytkownika w systemach sterowania głosem czy interfejsach konwersacyjnych. W przeciwieństwie do czysto statystycznych rozwiązań gwarantuje to przewidywalność reakcji systemu na rzadkie lub nietypowe konstrukcje składniowe.

Stosowanie gramatyk formalnych może obejmować także proces ich automatycznego generowania z danych korpusowych. Metody indukcji gramatycznej próbują znaleźć najmniejszy zestaw reguł produkcji mogący wygenerować wszystkie obserwowane zdania w zbiorze uczącym. Choć takie podejście jest atrakcyjne dla badań nad ewolucją języka albo adaptacją narzędzi NLP do nowych domen tematycznych, wymaga starannej kontroli jakości wynikowej gramatyki ze względu na ryzyko przeuczenia na specyficznych danych i utratę ogólności modelu. W tej perspektywie można zauważyć pewien dualizm: z jednej strony modele sieci neuronowych potrafią niejako „wewnętrznie” reprezentować zasady składni i semantyki bez ich wyraźnego zapisu; z drugiej strony zapis formalny daje możliwość przejrzystej dokumentacji i modyfikacji tych reguł przez człowieka, co jest niezwykle istotne tam, gdzie priorytetem jest wyjaśnialność działania systemu oraz jego zgodność ze standardami lingwistycznymi zasugerowanymi przez wcześniejsze badania teoretyczne.

Semantyka formalna

Semantyka formalna zajmuje się opisem znaczenia wypowiedzi w językach naturalnych za pomocą precyzyjnych, dobrze zdefiniowanych reguł i struktur matematycznych. W przeciwieństwie do gramatyk formalnych, które koncentrują się na składni, tutaj dąży się do uchwycenia relacji znaczeniowych pomiędzy wyrażeniami oraz ich powiązań z rzeczywistością. W praktyce oznacza to ustanowienie aparatu logicznego lub modelowego, w którym każdy element języka otrzymuje dokładną interpretację, często przy wykorzystaniu logiki zdań i logiki predykatów. Podstawowym krokiem jest określenie sposobu mapowania zdania na formułę logiczną reprezentującą jego sens. Na przykład proste zdanie, typu "Każdy student czyta książkę", można odwzorować jako formułę pierwszego rzędu:



$$\forall x(\text{"Student"}(x) \rightarrow \exists y(\text{"Książka"}(y) \wedge \text{"Czyta"}(x,y)))$$

która jednoznacznie definiuje zbiór modeli interpretujących prawdziwość tej wypowiedzi. Tego rodzaju konstrukcje umożliwiają analizę formalną, weryfikację spójności wypowiedzi oraz sprawdzanie logicznych konsekwencji zbioru zdań.

Z punktu widzenia zastosowań NLP semantyka formalna wymaga połączenia analizy syntaktycznej z mechanizmami interpretacji znaczenia. Parser składniowy dostarcza strukturę drzewa lub grafu zależności, a komponent semantyczny przekształca je na abstrakcyjne reprezentacje logiczne. Często konieczne jest uwzględnienie kontekstu dyskursu, wcześniejszych zdań lub wiedzy o świecie, aby poprawnie ustalić odniesienia i rozstrzygnąć niejednoznaczności. Języki naturalne pełne są idiomów czy metafor, które wymagają od systemu znajomości realiów kulturowych, co stanowi szczególne wyzwanie dla czysto formalnych metod. Systemy semantyki formalnej od dawna próbują zaimplementować zasady koherencji i spójności dialogowej w postaci reguł łączenia interpretacji kolejnych zdań. W ramach tzw. semantyki dialogowej każde nowe zdanie modyfikuje model sytuacyjny budowany przez system, dodaje nowe obiekty, aktualizuje relacje między nimi lub zmienia wartości predykatów. Takie podejście pozwala na obsługę anafor i zaimków (np. "on", "tam"), które bez odniesienia do wcześniejszej części tekstu pozostawałyby nieokreślone.

Wyzwaniem jest zderzenie teorii z praktyką implementacji w systemach NLP działających na dużą skalę. Tradycyjne metody formalne były związane z ręcznym definiowaniem reguł interpretacji oraz enumeracją wyjątków, co przy mnogości konstrukcji językowych staje się trudne do utrzymania. Rozwiązaniem okazało się połączenie semantyki formalnej z podejściem opartym na uczeniu maszynowym, przykładowo duże modele językowe LLM operują na wektorowych reprezentacjach zdań i potrafią generować parafrazy czy streszczenia o zgodnym sensie bez jawnego wykonywania kroków transformacji logicznych. Niemniej jednak zachowanie kontroli nad interpretacją wymaga wciąż możliwości translacji takich reprezentacji na postać formalną dla celów weryfikacyjnych czy automatycznego dowodzenia twierdzeń. Ciekawym obszarem zastosowań semantyki formalnej jest tłumaczenie maszynowe w ujęciu interlingua. Tutaj zdanie źródłowe jest najpierw przekształcane do abstrakcyjnej reprezentacji znaczeniowej niezależnej od języka, często właśnie poprzez zapis logiczny lub strukturę ramową, a dopiero potem generowane jest zdanie docelowe zgodne ze składnią danego języka. Takie rozwiązanie minimalizuje ryzyko utraty kluczowych relacji między obiektami i działaniami podczas tłumaczenia idiomatycznego. Przykład może dotyczyć rozumienia zdań warunkowych: system bazujący tylko na prostych asercjach może potraktować "Jeśli pada deszcz, impreza będzie odwołana" jako dwie luźne informacje o deszczu i odwołaniu imprezy. Semantyka formalna, za pomocą logiki zdań, odda zależność warunkową jako:

$$\text{"Deszcz"} \rightarrow \text{"OdwołanieImprezy"}$$



co pozwala przeprowadzać standardowe operacje inferencyjne i odpowiadać na pytania typu "Co stanie się, jeśli będzie słonecznie?". Współczesne badania integrują też modele probabilistyczne z semantyką formalną, co umożliwia ujęcie niepewności w interpretacjach, np. gdy dane wejściowe są szumowe lub część informacji jest pominięta. Modele takie mogą korzystać z logiki rozmytej albo rozkładów prawdopodobieństwa nad możliwymi wartościami predykatów, co zwiększa odporność systemu na niedokładności danych wejściowych.

Warto zaznaczyć napięcie między precyzją a elastycznością modeli semantycznych. Formalny zapis zapewnia jednoznaczność interpretacji i możliwość automatycznego dowodzenia własności tekstu; jednak wymaga dużej mocy obliczeniowej oraz pełnych i poprawnie oznaczonych danych wejściowych. Z kolei modele neuronowe potrafią odwzorować znaczenie w sposób wystarczający dla wielu zastosowań praktycznych (np. chatbota), ale ich mechanizmy są trudniejsze do wyjaśnienia człowiekowi oraz mniej odporne na subtelne błędy znaczeniowe wynikające z halucynacji modelu.

Semantyka formalna znajduje zastosowanie także w zadaniach klasyfikacji tekstów według kryteriów tematycznych lub intencyjnych poprzez definiowanie kategorii jako określonych predykatów spełnianych przez dokumenty należące do klasy. Możliwość dedukcyjnego przypisania dokumentu do kategorii daje większą przejrzystość decyzji niż klasyczne modele statystyczne bez ścisłego aparatu logicznego. Patrząc szerzej, integracja semantyki formalnej z inżynierią lingwistyczną może przybrać formę hybrydową: moduły formalne przejmują odpowiedzialność za krytyczne fragmenty analizy znaczeniowej (np. polecenia sterujące systemem lub zapytania bazy danych), natomiast głębokie sieci neuronowe dostarczają wsparcia tam, gdzie wieloznaczność czy brak pełnego kontekstu wymagają uczenia z dużych zbiorów przykładów. Takie połączenie wydaje się być obecnie jednym z najbardziej obiecujących kierunków dla tworzenia aplikacji NLP odpornych zarówno na niespodziewane konstrukcje językowe, jak i typowe pułapki semantyczne wynikające z bogactwa naturalnych języków.

2.3 Reprezentacje wiedzy językowej

Reprezentacje wiedzy językowej stanowią fundament dla wielu zastosowań w przetwarzaniu języka naturalnego, ponieważ umożliwiają odwzorowanie informacji lingwistycznych w formie, którą można efektywnie analizować i przekształcać algorytmicznie. Wiedza o języku obejmuje zarówno zasady morfosyntaktyczne, semantyczne jak i pragmatyczne, a także elementy związane z użyciem języka w określonych kontekstach. Jej zapis może przybierać postać struktur formalnych, baz danych, sieci semantycznych czy wektorowych reprezentacji zdań i słów. Już proste słowniki elektroniczne lub bazy pojęciowe, takie jak ontologie, pozwalają na



zorganizowanie informacji o relacjach pomiędzy jednostkami leksykalnymi w sposób hierarchiczny. Ontologie zawierają definicje pojęć oraz ich powiązania typu nadrzędność–podrzędność czy relacje skojarzeniowe.

W kontekście NLP często stosuje się struktury takie jak WordNet, gdzie rzeczowniki, czasowniki, przymiotniki i przysłówki są grupowane w synsety reprezentujące konkretne znaczenie. Takie organizowanie wiedzy umożliwia systemom komputerowym nie tylko rozpoznanie formy leksykalnej, ale także inferencję znaczenia zależnie od kontekstu.

Reprezentacja może być także oparta na modelach statystycznych i uczeniu maszynowym. W tym podejściu wiedza językowa jest zakodowana w parametrach sieci neuronowych uczonych na dużych zbiorach tekstu. Architektury RNN czy LSTM potrafią przechowywać informacje o poprzednich tokenach w sekwencji, co jest kluczowe dla zachowania spójności interpretacji. Sieci te tworzą ukryte reprezentacje, które zawierają zarówno lokalne zależności syntaktyczne jak i globalne cechy semantyczne zdania. Tego rodzaju modelowanie jest przykładem tzw. uczenia reprezentacji, inspirowanego sposobem przyswajania wzorców językowych przez człowieka. Rozbudowane wersje tego podejścia prowadzą do powstania głębokich modeli takich jak transformery, które dzięki mechanizmowi uwagi mogą dynamicznie wybierać istotne fragmenty sekwencji wejściowej do obliczeń reprezentacji. Zamiast jednego stanu ukrytego dla całego zdania stosuje się tu macierze wag wagowych wskazujących powiązania pomiędzy wszystkimi parami tokenów. Powstałe w ten sposób wektorowe reprezentacje, często typu osadzanie, są gęstymi opisami semantyczno-syntaktycznymi każdego elementu zdania. W modelach takich jak GPT-3 czy BERT parametry odpowiadające za te reprezentacje są wynikiem trenowania na bilionach przykładów tekstu, co umożliwia im uchwycenie subtelnych różnic znaczeniowych.

Niezależnie od warstwy technologicznej ważnym zagadnieniem pozostaje separacja tzw. wiedzy parametrycznej od nieparametrycznej. Wiedza parametryczna jest bezpośrednio zakodowana w parametrach modelu podczas treningu. Natomiast wiedza nieparametryczna to informacje możliwe do pobrania lub odczytania z zewnętrznych źródeł w trakcie działania systemu, przykładowo dane osadzone w wektorowych bazach danych dostępnych przez komponenty takie jak łańcuch odzyskiwania konwersacyjnego. Takie podejście umożliwia aktualizację zasobu informacyjnego bez konieczności ponownego trenowania całego modelu. Implementacyjnie często korzysta się z bibliotek takich jak NLTK, które oferują gotowe struktury danych do zapisu słowników, gramatyk czy korpusów anotowanych. Pozwalają one na łatwe łączenie różnych poziomów opisów języka, od fonetyki po składnię i integrację z modułami analitycznymi oraz generującymi tekst.



W praktycznych zastosowaniach chmurowych platformy integrują te zasoby z mechanizmami przetwarzania równoległego oraz z usługami takimi jak IBM Watson czy Hugging Face Hub, co zapewnia skalowalność i dostęp do aktualnych modeli. Istotną kategorią są również formalizmy logiczne wykorzystywane do modelowania wiedzy językowej w formie nadającej się do dedukcji automatycznej (jak pokazano przy wzorze [eq:student_book]). Zapisy takie umożliwiają ścisłe odwzorowanie zależności między obiektami oraz definiowanie reguł wyprowadzania nowych faktów z istniejącej bazy wiedzy. Mogą one współistnieć z probabilistycznymi modelami semantyki rozmytej, które dodają informację o stopniu pewności danego twierdzenia.

Współczesne duże modele językowe (LLM) są swego rodzaju hybrydą powyższych metod: przechowują abstrakcyjne wzorce syntaktyczne i semantyczne wyuczone statystycznie oraz potrafią korzystać z kontekstowych zasobów w trybie operacyjnym. Łączenie reprezentacji wektorowych ze źródłami formalnej wiedzy umożliwia zwiększenie trafności odpowiedzi przy jednoczesnym zachowaniu kontroli nad interpretacją sensu wypowiedzi. Takie systemy mogą mieć np. dostęp do specjalistycznych dokumentacji technicznych zapisanych w formatach strukturalnych XML lub RDF; wewnętrzne osadzenia obsługują szybkie kojarzenie fraz użytkownika z odpowiednimi fragmentami tych dokumentów.

Nie mniej ważna jest transparentność użytych form reprezentacji, szczególnie gdy system ma być wykorzystywany w krytycznych zastosowaniach wymagających wyjaśnialności działania. Tutaj zapisy formalne oferują przewagę: pozwalają lingwiście lub inżynierowi ocenić poprawność relacji i definicji zawartych w bazie wiedzy oraz modyfikować ją ręcznie zgodnie ze standardem językowym lub branżowym. Z kolei gęste osadzenia sprawdzają się tam, gdzie najważniejsza jest szybkość wyszukiwania podobieństw kontekstowych.

Jednym z nowych kierunków badań jest także adaptacja modeli poprzez wzbogacenie ich parametrów o cechy specyficzne dla danej domeny tematycznej przy minimalnym dostrajaniu. Pozwala to zachować ogólną zdolność rozumienia języka przy jednoczesnym dodaniu specjalistycznej terminologii do wewnętrznych reprezentacji. Praca nad takimi adaptacjami pokazuje wpływ jakości źródłowych danych tekstowych, im lepsze anotacje i bardziej spójny korpus dziedzinowy, tym wierniejsze odwzorowanie sensu pojęć i aktów mowy przez system. Modele hybrydowe starają się tu znaleźć równowagę: część operacji interpretacyjnych przenosi się na poziom formalnych operatorów logiki predykatów lub grafowych baz wiedzy; reszta pozostaje zakodowana statystycznie jako osadzanie optymalizowane pod kątem szybkiego dopasowania treści wejściowej i odpowiedzi wyjściowej.



3 Historia i rozwój inżynierii lingwistycznej

3.1 Początki badań nad językiem naturalnym

Wczesne próby automatycznego przetwarzania języka naturalnego były silnie związane z rozwojem ogólnej sztucznej inteligencji. Już w połowie XX wieku pojawiły się próby stworzenia programów zdolnych do samodzielnego uczenia się wzorców językowych, inspirowane zarówno przez praktyczne potrzeby, jak i przez pytania o naturę inteligencji. Arthur Samuel, uznawany za jednego z pionierów uczenia maszynowego, sugerował możliwość konstruowania adaptacyjnych programów, które nie tylko wykonują zaprogramowane instrukcje, ale potrafią doskonalić swoje działanie na podstawie doświadczenia. Ten sposób myślenia wskazał drogę ku systemom NLP, które mogłyby wykraczać poza sztywne reguły gramatyczne. Kamieniem milowym było stworzenie przez Josepha Weizenbauma programu ELIZA w 1965 roku. ELIZA była jednym z pierwszych chatbotów i miała symulować rozmowę w języku naturalnym poprzez dopasowywanie określonych wzorców tekstowych i stosowanie prostych transformacji składniowych. Choć system był ograniczony – nie posiadał żadnej „prawdziwej” wiedzy czy rozumienia wypowiedzi – jego efekty często były dla użytkowników na tyle przekonujące, że przypisywali mu cechy ludzkiej inteligencji. To zjawisko antropomorfizacji podkreśliło psychologiczne aspekty interakcji człowieka z maszyną i odśloniło potencjalny wpływ konstrukcji dialogu na odbiór technologii. W tamtym okresie badania nad językiem skupiały się także na próbach formalizowania struktur lingwistycznych w kontekście implementacji komputerowych. Rozwijano algorytmy dopasowywania wzorców oraz podstawowe parsery składniowe wykorzystujące reguły gramatyczne. Był to czas dominacji podejść symbolicznych – reguły eksplicytne tworzone były ręcznie przez lingwistów i przenoszone na formę kodu możliwego do interpretacji przez programy. Nacisk kładziono na budowanie słowników komputerowych oraz baz wiedzy zawierających informacje gramatyczne i semantyczne dla poszczególnych słów.

Równolegle trwała dyskusja teoretyczna nad tym, jak oceniać, czy system naprawdę „rozumie” język. Alan Turing zaproponował koncepcję testu rozpoznawania inteligencji – testu Turinga – który opierał się m.in. na kompetencji językowej. Warunkiem pozytywnego wyniku było to, aby komputer potrafił prowadzić rozmowę z człowiekiem tak, by ten nie potrafił odróżnić odpowiadającej maszyny od innego człowieka. Wymagało to połączenia analizy składniowej, generowania sensownych odpowiedzi oraz posiadania szerokiego zestawu informacji o świecie. Pierwsze eksperymenty obejmowały też proste mechanizmy tłumaczeń maszynowych oparte na transferze struktur syntaktycznych między językami. Ze względu na ograniczoną moc obliczeniową oraz niewielkie zasoby cyfrowe tworzone niewielkie modele radzące sobie z określonym zakresem tematycznym i wąskimi domenami leksykalnymi. Wykorzystywano przy tym gramatyki bezkontekstowe oraz reguły



zamiany fraz jako podstawę generowania tłumaczeń. Podstawą tych badań było gromadzenie danych językowych – korpusów tekstów literackich, artykułów naukowych czy zapisów dialogów – które mogły służyć jako materiał referencyjny do opracowywania modeli. Ze względu na brak rozbudowanych bibliotek przetwarzania języka badacze tworzyli własne zestawy narzędzi i formaty danych.

Z czasem zaczęto standaryzować interfejsy dostępu do korpusów, co znalazło wyraz chociażby w późniejszym powstaniu modułów takich jak `nltk.corpus`, umożliwiających operacje wyszukiwania i filtrowania dokumentów według określonych kryteriów. Nie bez znaczenia był również rozwój metod tokenizacji i segmentacji tekstu. W początkowej fazie korzystano ze schematów dzielenia tekstu według prostych separatorów (spacje, znaki interpunkcyjne), jednak szybko okazało się to niewystarczające dla języków o skomplikowanej morfologii czy nieregularnych odstępach między znakami (np. w zapisie ręcznym lub OCR). Podjęto więc próby tworzenia bardziej zaawansowanych algorytmów segmentacyjnych uwzględniających statystyczne występowanie struktur w tekście. Rozwój analizy morfologicznej wymagał od badaczy ustalenia sposobów lematyzacji (sprowadzania wyrazów do formy podstawowej) oraz stemmingu (redukcji form fleksyjnych poprzez ucinanie końcówek). Choć metody te były dalekie od doskonałości – często generowały błędy dla słów nietypowych lub zapożyczeń – stanowiły istotny krok ku automatycznej ekstrakcji informacji ze zbiorów tekstowych.

Warto zauważyć, że początkowe sukcesy modeli takich jak ELIZA przewyższały oczekiwania głównie w warstwie percepcyjnej użytkowników. Faktyczna funkcjonalność była mocno ograniczona, a brak głębokiego rozumienia treści stawał się widoczny przy bardziej wymagających dialogach. Niemniej jednak te projekty stały się inspiracją dla kolejnych pokoleń badaczy pracujących nad realistyczną symulacją komunikacji naturalnej.

Podejmowano też pierwsze próby mierzenia jakości działania takich systemów w różnych zadaniach lingwistycznych, co zapoczątkowało późniejsze tworzenie standardowych zestawów testowych oceniających różne aspekty przetwarzania języka. Choć początkowo oceny miały charakter subiektywny lub były sporządzane ad hoc przez projektującego system badacza, idea ustandaryzowanych benchmarków umocniła się wraz z pojawieniem większej liczby konkurencyjnych narzędzi. Można stwierdzić, że pierwsze dekady badań koncentrowały się wokół podejścia regułowego wspieranego analizą lingwistyczną prowadzoną ręcznie przez ekspertów. Dopiero zwiększona dostępność danych oraz rosnąca moc obliczeniowa otworzyła perspektywę wykorzystania statystycznych modeli uczenia maszynowego dla NLP, perspektywę zapowiedzianą już we wczesnych pomysłach Samuela i zilustrowaną popularnością ELIZY jako eksperymentalnej platformy interakcji człowiek–maszyna.



3.2 Rozwój technologii językowych

Systemy eksperckie w językoznawstwie

Systemy eksperckie w kontekście językoznawstwa powstały jako próba przeniesienia wiedzy specjalistów na grunt komputerowy tak, aby procesy analizy czy interpretacji języka można było prowadzić automatycznie. Konstrukcja takich narzędzi wymagała uprzedniego sformalizowania reguł postępowania ekspertów, co w przypadku obszarów lingwistycznych oznaczało zapis zasad gramatycznych, morfologicznych oraz semantycznych w formie możliwej do kodowania. W praktyce oznaczało to, że system musiał być wyposażony w bazę wiedzy zawierającą komplet informacji potrzebnych do wykonania danego zadania, na przykład zestaw symptomów dla diagnozy medycznej lub wzorców składniowych dla tłumaczenia maszynowego. Brak mechanizmów samouczenia się uniemożliwiał adaptację do nowych sytuacji bez ręcznej aktualizacji bazy reguł.

W latach 70. i 80., wraz z rosnącą mocą obliczeniową i dostępnością komputerów osobistych, systemy eksperckie stały się istotnym narzędziem również w badaniach nad językiem. Dominowały wtedy modele oparte na schematach decyzyjnych przypominających drzewa decyzji, które krok po kroku replikowały rozumowanie specjalisty. W przypadku analiz językowych odpowiadało to procesom typu: identyfikacja kategorii gramatycznej słowa, ustalenie relacji syntaktycznych między elementami zdania, a następnie przypisanie interpretacji semantycznej. Mechanizm walki z niejednoznacznością opierał się na kolejności i priorytetach reguł, jeśli dana konstrukcja pasowała do kilku wzorców, używano dodatkowych heurystyk ustalonych przez projektantów. Taka statyczna natura powodowała jednak trudność z obsługą konstrukcji rzadkich lub nowych. Przykład INTERNIST-I pokazuje dobrze ograniczenia tej technologii. Choć w swoim pierwotnym zastosowaniu wspierał lekarzy w diagnostyce poprzez zbiór reguł opracowanych na podstawie wiedzy jednego eksperta Johna D. Myersa, podobny model mógł zostać teoretycznie przystosowany do dziedziny językoznawstwa: reguły określające poprawne frazy czy sposób rozpoznawania błędów gramatycznych mogły być zapisane na bazie wskazań lingwisty. Jednak brak integracji szerokiej wiedzy kontekstowej skutkowało problemami, system mógł operować tylko tymi informacjami, które zostały mu przekazane przy tworzeniu bazy. W tekście naturalnym pojawiały się często konstrukcje wykraczające poza jego kompetencje.

Rozkwit komercyjnego zastosowania systemów eksperckich obejmował także branżę lingwistyczną i tłumaczeniową. Firmy zaczęły inwestować w narzędzia oparte na kodowanych regułach, które mogły automatyzować fragmenty pracy redaktorów czy tłumaczy. Powstawały wyspecjalizowane rozwiązania integrujące parsery składniowe z modułami oceny poprawności stylistycznej oraz systemami sugestii zmian tekstu zgodnie z ustaloną normą językową. Tego rodzaju implementacje wymagały



precyzyjnych opisów struktur językowych i ich powiązań, często korzystano tu z tworzonych przez lingwistów ontologii terminologicznych czy baz wzorcowych zdań. Mimo swoich ograniczeń systemy eksperckie miały też zalety: oferowały przewidywalność wyników oraz zgodność z określonym zestawem standardów lingwistycznych. Tam, gdzie ważna była zgodność formalna, np. w korekcie tekstów akademickich, eliminowało to ryzyko związane ze stosowaniem algorytmów probabilistycznych mogących generować treści niezgodne z normą. Ze względu na swoją architekturę takie systemy dobrze nadawały się też do implementacji specjalistycznych słowników i glosariuszy branżowych; reguły mogły precyzować sposób traktowania terminologii fachowej odmiennie niż słownictwa potocznego. W pewnym momencie próbowano połączyć ten typ technologii z bardziej dynamicznymi metodami przetwarzania tekstu opartymi na statystyce i uczeniu maszynowym. Idea polegała na tym, aby rdzeń decyzyjny zachował charakter ekspercki (rule-based), ale wspomagać go modułem zdolnym do analizowania dużych zbiorów tekstów pod kątem nowych wzorców użycia języka. W obszarze NLP oznaczało to integrację komponentu uczącego się ze strumienia danych z istniejącym zestawem reguł składniowych lub semantycznych. Takie hybrydowe podejście mogło rozwiązać część problemów związanych z brakiem adaptowalności klasycznych systemów eksperckich.

Z perspektywy historycznej należy zauważyć sprzężenie zwrotne między tworzeniem takich narzędzi a rozwojem opisu teoretycznego języka. Kodowanie reguł wymuszało bardzo szczegółowe jednoznaczne definicje pojęć lingwistycznych; luki w tych definicjach ujawniały się szybko w trakcie implementacji. W rezultacie kolejne wersje systemów stawały się nie tylko dokładniejsze technicznie, ale też lepiej osadzone w teorii języka, wiele projektów konsultowano bezpośrednio z akademickimi lingwistami.

Obecnie część zasad opracowanych dla dawnych systemów eksperckich przenika do współczesnych dużych modeli językowych poprzez mechanizmy tzw. prompt engineering czy fine-tuningu. Choć modele neuronowe uczone są przede wszystkim statystycznie, można je ukierunkowywać do przestrzegania zbioru ustalonych „reguł” sformułowanych przez człowieka. W ten sposób zachowuje się korzyść płynąca ze stabilności regułowego podejścia przy jednoczesnym wykorzystaniu elastyczności uczenia maszynowego. Ostatecznie systemy eksperckie stanowią element historii technologii językowych, który wydaje się dobrze ilustrować przejście od podejść symbolicznych ku hybrydowym formom integrującym modele uczenia głębokiego. Choć ich czasy świetności minęły wraz z nadejściem epoki big data i sieci neuronowych zdolnych uczyć się zależności bez zapisanego kompletu instrukcji, nadal znajdują one miejsce tam, gdzie najważniejsza jest spójność formalna oraz możliwość audytu decyzji algorytmu przez człowieka.



Rozwój korpusów językowych

Rozwój korpusów językowych stanowił jeden z kluczowych etapów w dojrzewaniu technologii językowych, wpływając zarówno na badania teoretyczne, jak i na implementacje praktyczne. Początkowo tworzone niewielkie zbiory tekstów, opracowane ręcznie i dostosowane do potrzeb konkretnych projektów. Ich ograniczona wielkość oraz brak standaryzacji struktury utrudniały porównywanie wyników między zespołami badawczymi. Z czasem zaczęto dążyć do tworzenia dużych, reprezentatywnych zbiorów obejmujących różne gatunki tekstów, style wypowiedzi oraz odmiany języka. Powstały inicjatywy budowania korpusów narodowych, takich jak British National Corpus czy polski NKJP, które stanowią przekrój współczesnego użycia języka i obejmują miliony słów. Na wczesnym etapie korzystano głównie z tekstów pisanych – literatura piękna, artykuły prasowe, dokumenty urzędowe – jednak szybko zauważono potrzebę uwzględnienia także mowy spontanicznej i dialogów. Rozszerzenie korpusów o dane fonetyczne i transkrypcje nagrań umożliwiło trenowanie systemów rozpoznawania mowy oraz badanie cech prozodycznych. Wraz z tym rozwojem pojawiła się potrzeba anotacji wielopoziomowej: od prostego podziału na zdania i tokeny po przypisywanie kategorii gramatycznych, oznaczanie jednostek nazewniczych oraz struktur składniowych.

Anotacja semantyczna dodaje kolejną warstwę – relacje pomiędzy wyrażeniami oraz odwzorowanie ich odniesień w świecie rzeczywistym. Prace nad korpusami poszerzały się o aspekty techniczne: należało stworzyć formaty danych umożliwiające sprawne łączenie różnych typów informacji lingwistycznych. XML czy JSON stały się popularnymi sposobami reprezentacji anotacji, pozwalającymi jednocześnie zachować kompatybilność między narzędziami analitycznymi. Biblioteki takie jak NLTK oferują moduły pozwalające na dostęp do licznych korpusów językowych wraz z funkcjami wyszukiwania, filtrowania i wizualizacji statystyk – np. poprzez konstrukcje typu Warunkowy Rozkład Częstotliwości, które mogą służyć analizie trendów w użyciu określonych słów lub fraz.

Dostępność korpusów zaczęła obejmować nie tylko projekty akademickie. Setki anotowanych zbiorów tekstu i mowy istnieją obecnie w dziesiątkach języków; część z nich objęta jest licencjami niekomercyjnymi przeznaczonymi do nauki i badań, inne mają licencje komercyjne wymagające opłat. Takie zróżnicowanie pozwala firmom wdrażać modele trenowane na wyspecjalizowanych zasobach przy zachowaniu kontroli nad prawami własności. Wraz ze wzrostem możliwości obliczeniowych zaczęto tworzyć korpusy o skali wcześniej niespotykanej – terabajty lub nawet petabajty danych tekstowych. W kontekście dużych modeli językowych (LLM) taka skala jest niezbędna, aby sieci neuronowe mogły uchwycić bogactwo wzorców językowych: od niskopoziomowej morfologii po zależności semantyczne obecne w kontekście całych akapitów czy dokumentów. Przetrenowanie modeli na tak ogromnych zbiorach surowego tekstu (książki, artykuły naukowe, Wikipedia) odbywa



się często w trybie uczenia samo, gdzie model sam generuje etykiety na podstawie wejściowego kontekstu. Warto zaznaczyć rolę różnorodności danych dla jakości działania systemu: modele uczone na jednorodnym korpusie będą miały zawężony zakres kompetencji i trudniej poradzą sobie w sytuacjach odbiegających od zapisanego wzorca. Dlatego współczesne projekty gromadzenia korpusów dbają o balans między danymi literackimi a technicznymi, formalnymi a potocznymi, standardowym językiem a dialektami czy odmianami regionalnymi.

Rozwój dużych zasobów pociągnął za sobą także kwestie etyczne i jakość danych wejściowych. Korpusy pozyskiwane z Internetu mogą zawierać treści toksyczne lub stroniczne; filtracja automatyczna jest tu trudna do stuprocentowego wykonania ze względu na subtelny charakter uprzedzeń czy ironii w tekście. Niektóre projekty wdrażają procedury ręcznej kontroli próbek albo wykorzystują systemy klasyfikacji treści do wykrywania potencjalnie szkodliwych fragmentów przed ich anotacją. W kontekście zastosowań przemysłowych powstaje potrzeba tworzenia korpusów specjalistycznych z danej dziedziny tematycznej, medycyny, prawa czy finansów, co umożliwia dostrajanie modeli ogólnych do pracy w wyspecjalizowanym środowisku przy minimalnym nakładzie treningowym.

Firmy coraz częściej inwestują we własne korpusy chronione przed dostępem stron trzecich ze względu na poufność danych, szczególnie gdy modele są uruchamiane bezpośrednio na urządzeniach użytkowników. Technologie przechowywania i udostępniania zasobów korpusowych ewoluowały razem z samymi zbiorami: od lokalnych baz plikowych po rozproszone systemy chmurowe wspierające równoległą analizę milionów dokumentów. Integracja z platformami takimi jak Hugging Face Hub zapewnia łatwy dostęp programistyczny do tysięcy datasetów wraz z ich metadanymi oraz narzędzia ułatwiające trenowanie własnych modeli na wybranym podzbiorze danych.

Nie można pominąć roli standaryzacji oznaczeń, schematy takie jak Universal Dependencies dla anotacji składniowej czy FrameNet dla opisu ram semantycznych pozwalają na porównywalność między projektami oraz ułatwiają transfer wiedzy między językami. Na bazie jednolitych standardów łatwiej jest też szkolić modele wielojęzyczne zdolne operować porównawczo na strukturach lingwistycznych. Rozbudowa zasobów sprzyja powstawaniu nowych metod badawczych: analizy diachroniczne wykorzystują historyczne edycje korpusów, aby śledzić zmiany znaczeniowe lub częstotliwości użycia określonych słów w czasie; eksperymenty psycholingwistyczne mogą korzystać z Internetu jako źródła realnych przykładów mowy potocznej do testowania hipotez o przetwarzaniu języka przez człowieka. Korpus przestaje być statycznym magazynem, staje się dynamicznym narzędziem synchronizującym badania teoretyczne z praktyką komputerową w NLP. Ostatecznie rozwój infrastruktury korpusowej stworzył fundament dla dzisiejszych systemów przetwarzania języka: parsery składniowe testowane są dziś masowo na



ustandaryzowanych podzbiorach danych; algorytmy generowania odpowiedzi mogą być oceniane pod kątem zgodności stylistycznej względem ustalonego benchmarku; a duże modele językowe osiągają swoje możliwości właśnie dzięki temu, że mają dostęp do szerokich i bogatych źródeł tekstu reprezentujących różnorodne konteksty komunikacyjne.

3.3 Współczesne nurty badawcze

Obecne kierunki badań w inżynierii lingwistycznej wyraźnie korzystają z dynamicznego postępu technologicznego i rosnącej dostępności mocy obliczeniowej, co otwiera możliwości dla architektur zdolnych operować na skalę wcześniej nieosiągalną. Wyraźnym przykładem jest rozwój dużych modeli językowych bazujących na mechanizmach uwagi (attention), które zastąpiły wcześniejsze podejścia sekwencyjne ze względu na możliwość równoległego przetwarzania długich kontekstów i selektywnego skupiania się na istotnych fragmentach danych wejściowych. Modele transformatorowe takie jak GPT czy BERT stały się głównym punktem odniesienia w zadaniach NLP – ich wektorowe reprezentacje oraz segmentacja kontekstu umożliwiają budowanie bardziej precyzyjnych odpowiedzi oraz interpretacji tekstu w rozmowie z użytkownikiem.

Jednym z nurtów, który zyskuje na znaczeniu, jest integracja tych modeli z wiedzą nieparametryczną oraz narzędziami uzupełniającymi. Architektury implementujące pamięć konwersacyjną potrafią śledzić historię dialogu, korzystając zarówno z parametrów zapisanych podczas treningu (wiedza parametryczna), jak i danych pobieranych na bieżąco z zewnętrznych baz lub systemów wyszukiwania. To pozwala uniknąć problemu utraty kontekstu w wieloetapowych interakcjach i zwiększa adaptowalność systemu do nowych obszarów tematycznych bez konieczności pełnego trenowania od podstaw. Widać także rosnące zainteresowanie osadzaniem modeli w środowiskach brzegowych, gdzie inferencja odbywa się bezpośrednio na urządzeniu końcowym – smartfonie, sensorze przemysłowym czy systemie IoT. Takie rozwiązania zmniejszają opóźnienia, ograniczają przesył danych do chmury i wzmacniają prywatność użytkownika.

W praktyce oznacza to możliwość stosowania analizy języka w czasie rzeczywistym nawet tam, gdzie połączenie sieciowe jest niestabilne lub niedostępne. Z perspektywy inżynierskiej przyzwala to też na projektowanie systemów odpornych na awarie infrastruktury centralnej. Kolejnym intensywnie badanym obszarem jest personalizacja modeli poprzez techniki dostrajania wymagające minimalnej liczby dodatkowych parametrów – adaptery czy p-tuning – które pozwalają dostosować bazowy model do konkretnego zadania lub domeny przy ograniczonym nakładzie obliczeń. Tego typu podejścia są szczególnie cenne dla aplikacji niszowych, gdzie zdobycie ogromnych zbiorów treningowych jest nieopłacalne lub niemożliwe. Przy



tym poprawa jakości predykcji wynika często z bliskości problemu docelowego do pierwotnego zadania treningowego modelu.

Rozwój generatywnej sztucznej inteligencji stanowi kolejny ważny aspekt obecnych badań. Głębokie sieci neuronowe – zaprojektowane do generowania nowych treści – znajdują zastosowanie w tworzeniu obrazu, muzyki czy tekstu na podstawie wzorców wyuczonych z istniejących danych. Szczególną rolę odgrywa tu przetwarzanie języka naturalnego i jego synteza: od realistycznych dialogów w chatbotach po kreatywne narracje pisarskie inspirowane wskazówkami użytkownika. Efekty widoczne są nie tylko w sektorze rozrywkowym, ale także w dziedzinach takich jak edukacja czy obsługa klienta. Integracja uczenia głębokiego z innymi technologiami – łańcuchami bloków, sieciami 5G czy obliczeniami kwantowymi – wskazuje dalsze możliwe ścieżki wzrostu możliwości NLP. Przykładowo, 5G może wspierać strumieniową transmisję dużych wolumenów danych między serwerami a modelami uruchomionymi u użytkownika, a łańcuch bloków zapewnić autentyczność i integralność rekordów wykorzystywanych jako dane treningowe lub źródła informacji nieparametrycznej.

Znaczącą uwagę zwraca się obecnie również na zagadnienia etyczne. Procedury takie jak Constitutional AI mają za zadanie kształtować modele tak, aby były zgodne z wartościami człowieka i unikały generowania treści toksycznych czy promujących działania nielegalne. Badacze eksplorują sposoby wdrażania filtrów treści i regulacji zachowań modelu zarówno na etapie jego trenowania, jak i interakcji operacyjnej. Wyzwaniem pozostaje balans pomiędzy eliminacją ryzykownych wypowiedzi a zachowaniem swobody formy oraz bogatego spektrum języka. Interesującym trendem jest wykorzystywanie otwarto źródłowych modeli takich jak Mistral-7B-instruct dla badań porównawczych oraz adaptacji w różnych środowiskach biznesowych. Możliwość pobierania pełnych wag modelu przez platformy dystrybucyjne (np. Hugging Face Hub) ułatwia eksperymentowanie zarówno naukowcom akademickim, jak i zespołom komercyjnym bez konieczności inwestowania w kosztowny proces trenowania od podstaw.

Nie można pominąć pracy nad architekturami agentowymi współpracującymi z LLM. Agenty wyposażone w zestawy narzędzi są zdolne autonomicznie wybierać działanie adekwatne do otrzymanego zadania w analizie języka: mogą pobrać zapytanie użytkownika, wyszukać powiązane treści w bazach wektorowych i wygenerować odpowiedź bazującą na najbardziej relewantnym kontekście. Taki mechanizm przypomina modułowy sposób pracy człowieka – selekcja informacji poprzedza etap formułowania wypowiedzi.

Obecne badania pokazują coraz silniejsze sprzężenie zwrotne między jakością korpusów a efektywnością modeli generatywnych i interpretacyjnych. Zdobycie odpowiednio różnorodnego zbioru danych pozostaje fundamentem dla każdej nowej architektury; jednak równie istotna staje się transparentna ocena jakości wynikowej



interakcji system–użytkownik pod kątem spójności merytorycznej i adekwatności stylistycznej. Nowoczesna inżynieria lingwistyczna zmierza więc ku architekturom hybrydowym: częściowo regułowym tam, gdzie wymagana jest formalna poprawność albo wyjaśnialność decyzji; częściowo statystycznym lub głębokim wszędzie tam, gdzie priorytetem jest płynność komunikacji i zdolność radzenia sobie z nieprzewidywalnymi konstrukcjami językowymi. To podejście wydaje się oferować kompromis między kontrolą a kreatywnością systemu – aspekt krytyczny wraz ze wzrostem roli AI w codziennym życiu społecznym i gospodarczym.

4 Technologie przetwarzania języka naturalnego

4.1 Analiza składniowa i semantyczna

Parsery składniowe

Parsery składniowe pełnią kluczową rolę w systemach przetwarzania języka naturalnego, umożliwiając komputerowej interpretacji struktury syntaktycznej zdań. Ich zadaniem jest przekształcenie sekwencji słów lub tokenów w reprezentację opisującą relacje między elementami wypowiedzi – najczęściej w postaci drzewa składniowego lub grafu zależności. Interpretacja taka pozwala uwzględnić kolejność i hierarchię fraz, co jest niezbędne w dalszych etapach analizy semantycznej czy tłumaczenia maszynowego. W przypadku modeli bazujących na gramatykach formalnych, parser stosuje reguły produkcji określone w danym formalizmie (np. gramatyki bezkontekstowej), tworząc strukturę zgodną z ustalonym zestawem dopuszczalnych konstrukcji. Istnieje kilka podejść do implementacji parserów. W tradycyjnych systemach regułowych stosuje się parsowanie odgórne (top-down), które zaczyna od symbolu startowego gramatyki i rozwija go do terminali zgodnie z dopasowanymi regułami, lub oddolne (bottom-up), które łączy tokeny wejściowe w większe frazy aż do osiągnięcia symbolu startowego. Oba warianty mają swoje wady – parsowanie odgórne może generować wiele niepotrzebnych gałęzi drzewa dla zdań niepoprawnych, podczas gdy podejście oddolne bywa podatne na eksplozję liczby możliwych kombinacji przy niejednoznaczności.

Współczesne parsery statystyczne zdecydowanie ograniczają te problemy dzięki mechanizmom probabilistycznym – każdej regule przypisana jest estymowana wartość prawdopodobieństwa jej zastosowania na podstawie danych korpusowych. Parsery oparte na gramatykach zależności reprezentują zdanie jako zbiór krawędzi łączących słowa w pary głowa–dziecko. Taka struktura bywa bardziej intuicyjna w wielu zastosowaniach, zwłaszcza tam, gdzie analizuje się bezpośrednio relacje między czasownikiem a jego argumentami. Budowa grafów zależności może być realizowana poprzez algorytmy przejść lub metody maksymalizujące dopasowanie globalne, często wspierane przez modele uczenia maszynowego przewidujące



kolejne kroki budowy struktury. W implementacjach opartych na narzędziach takich jak NLTK dostępne są parsery korzystające zarówno z jawnie zapisanych reguł, jak i z modeli statystycznych trenowanych na anotowanych korpusach. NLTK daje możliwość łatwego eksperymentowania z różnymi technikami parsowania: od prostych parserów rekurencyjno–zstępujących po zaawansowane algorytmy typu Earley parser czy chart parser, które minimalizują koszty obliczeń przez zapamiętywanie częściowych wyników analizy i ich ponowne użycie.

Naturalne języki niosą ze sobą problem niejednoznaczności syntaktycznej, który parser musi rozstrzygać. Zdanie „Widzę chłopca z lornetką” można interpretować co najmniej dwojako – jako sytuację, gdy mówiący dysponuje lornetką lub gdy chłopiec ją posiada. Parser probabilistyczny może wybrać najbardziej prawdopodobną interpretację na podstawie tendencji zauważonych w korpusie treningowym, lecz konieczne bywa uzupełnienie tego procesu wiedzą semantyczną zewnętrznymi zasobów.

Kolejnym aspektem jest obsługa języków fleksyjnych o bardziej skomplikowanych zależnościach między formą wyrazu a jego funkcją gramatyczną. Parser dla takich języków wymaga integracji z modułem analizy morfologicznej dostarczającym informacji o odmianie wyrazów i ich cechach morfo syntaktycznych (np. przypadek, liczba), co pozwala uniknąć błędnej segmentacji fraz i poprawnie odwzorować strukturę zdania. Rozszerzeniem klasycznych metod parsowania są podejścia hybrydowe integrujące reguły lingwistyczne z modelami głębokiego uczenia.

Architektury transformatorowe potrafią przewidywać relacje syntaktyczne pomiędzy tokenami bez jawnego zapisu gramatyki – uczą się wzorców składniowych w procesie trenowania na bardzo dużych zbiorach tekstu. Mimo to łączenie takiego predykcyjnego komponentu z modułem regułowym pozwala utrzymać kontrolę nad strukturą generowanego drzewa oraz zapewnić zgodność ze standardami anotacji syntaktycznej ustalonymi np. przez schemat zależności uniwersalnych. W zadaniach wymagających dużej wydajności, jak analiza strumienia treści czy przetwarzanie wiadomości w czasie rzeczywistym, istotną staje się optymalizacja algorytmu pod kątem szybkości działania przy zachowaniu poprawnego rozpoznania struktur składniowych. Przykładowo parser może pracować inkrementalnie – konstruując drzewo lub graf zależności „w locie” wraz z napływem kolejnych fragmentów tekstu – co redukuje opóźnienia w aplikacjach konwersacyjnych. W takich zastosowaniach istotna jest także możliwość adaptacji modelu do specyfiki komunikatów krótkich i nieformalnych, które często odbiegają od norm pisemnych.

Warto podkreślić także aspekt ewaluacji jakości pracy parsera: mierniki takie jak dokładność etykiet zależności czy procent poprawnie rozpoznanych relacji syntaktycznych stanowią standard oceny modeli na benchmarkach korpusowych. Regularna walidacja wyników pozwala identyfikować typowe błędy analizy i ulepszać



zarówno komponenty lingwistyczne, jak i te oparte na uczeniu maszynowym. Z perspektywy rozwoju technologii NLP rola parserów ewoluuje wraz z pojawianiem się nowych architektur przetwarzania tekstu – jednak niezależnie od metody implementacji pozostają one fundamentem interpretacji strukturalnej dokumentów i interakcji językowych. Dzięki nim możliwe jest powiązanie sekwencji leksykalnej z jej abstrakcyjnym opisem hierarchicznym lub sieciowym, co stanowi punkt wyjścia dla dalszych faz analitycznych ukierunkowanych na znaczenie wypowiedzi oraz generowanie odpowiedzi spójnych względem tej struktury.

Analiza semantyczna

Analiza semantyczna stanowi etap przetwarzania języka, w którym celem jest ustalenie znaczenia struktur uzyskanych w fazie analizy składniowej. Interpretacja treści obejmuje mapowanie elementów zdania na bardziej abstrakcyjną formę reprezentacji znaczeniowej i uwzględnia zarówno relacje logiczne, jak i kontekst sytuacyjny. Fundamentem takich metod może być formalny aparat logiki predykatów pierwszego rzędu, który pozwala opisać zbiór obiektów, ich atrybuty oraz związki między nimi. Model języka formalnego wskazuje, jakie warunki muszą być spełnione, aby dane zdanie było prawdziwe w określonej sytuacji.

W praktyce system przetwarzający treść wyrażen korzysta z parsera syntaktycznego jako źródła struktury wejściowej. Następnie komponent semantyczny wykonuje odwzorowanie na formuły logiczne lub inne uporządkowane schematy interpretacyjne, na przykład ramy semantyczne definiujące akt, uczestnika i okoliczności. Takie przejście wymaga rozpoznania typów pojęć oraz charakteru relacji i często wiąże się z rozwiązywaniem problemów odniesienia, kiedy dane wyrażenie wskazuje na konkretny obiekt lub zdarzenie. Dla wyrażen prostych proces ten bywa stosunkowo jednoznaczny: „Piotr czyta książkę” można zapisać jako predykat Czyta(Piotr,Książka) odpowiadający binarnemu związkowi wykonawcy czynności i obiektu. Złożoność rośnie jednak gwałtownie w przypadku idiomów czy konstrukcji metaforycznych. Tutaj klasyczny aparat formalny wymaga uzupełnienia o procedury interpretacyjne bazujące na wiedzy kontekstowej lub zasobach lingwistycznych takich jak WordNet.

Synsety dostarczają zestaw ekwiwalentnych pojęć wraz z informacją o relacjach hiponimii, antonimii czy meronimii. Dzięki temu możliwe jest rozszerzenie znaczenia słowa o jego warianty użycia i powiązane pojęcia. W implementacjach praktycznych analizę semantyczną wspiera się zależnościami zewnętrznymi źródeł wiedzy, od sieci semantycznych po ontologie dziedzinowe. Ontologie opisują hierarchię pojęć i reguły ich stosowania w danym obszarze; integracja analizy semantycznej z ontologią umożliwia nie tylko rozpoznanie znaczenia terminu, ale także sprawdzenie zgodności wypowiedzi ze znanym zestawem faktów. Takie podejście bywa



szczególnie istotne dla systemów odpowiadających w trybie eksperckim, gdzie poprawność merytoryczna jest priorytetem.

Rozbudowane architektury korzystające z uczenia maszynowego tworzą alternatywę dla jawnego konstruowania bazy reguł semantycznych. Modele uczenia nadzorowanego lub samonadzorowanego potrafią przekształcać wejściowe reprezentacje syntaktyczne w wektory osadzeń będące statystycznym odwzorowaniem sensu. Transformery, poprzez mechanizm uwagi, wychwytyują zależności między tokenami pozwalające odwzorować powiązanie znaczeniowe nawet dla oddalonych słów w sekwencji. Predykcja zależności semantycznych odbywa się tu bez jawnego budowania formuły logicznej; zamiast tego model operuje na przestrzeni wysokowymiarowych reprezentacji, które można porównywać pod kątem podobieństwa kontekstowego. Modele takie jak BERT czy GPT różnią się podejściem: BERT jako model dwukierunkowy analizuje jednocześnie lewy i prawy kontekst danego tokena, co sprzyja precyzyjnej interpretacji znaczeń słów wieloznacznych; GPT natomiast opiera się na kontekście lewostronnym stosowanym sekwencyjnie, co wpływa na sposób budowania odpowiedzi.

W zadaniach wymagających ujednoznacznienia sensu wyrazu zastosowanie BERT-a okazuje się efektywne dzięki równoczesnemu dostępowi do całego otoczenia tokena. Integracja semantyki formalnej z podejściem statystycznym tworzy tzw. modele hybrydowe. W takim rozwiązaniu część tekstu analizowana jest za pomocą ręcznie definiowanych reguł interpretacyjnych (np. polecenia sterujące systemem), a pozostałe fragmenty poddaje się uogólnionej inferencji wektorowej opartej na modelach przetworzonych dużymi zbiorami danych. Kluczowe jest tu utrzymanie spójności między tymi modułami: jeśli komponent neuronowy proponuje parafrazę polecenia użytkownika, musi ona zostać zweryfikowana przez warstwę formalną tak, aby zachować identyczny sens działania.

Wiele współczesnych prac badawczych skupia się także na probabilistycznym rozszerzeniu aparatu semantycznego. Logika rozmyta czy modele Bayesowskie pozwalają opisać niepewność co do prawdziwości predykatu lub wystąpienia relacji, szczególnie istotne przy analizie danych głośnych lub niepełnych. Taki opis daje możliwość wydedukowania najbardziej prawdopodobnego scenariusza nawet przy braku pełnej informacji. Nieodzownym elementem analizy semantycznej pozostaje obsługa anafor i zaimków. Wyrażenia typu „on”, „tam” wymagają od systemu śledzenia historii dyskursu oraz aktualnego stanu modelu sytuacyjnego. Moduły dyskursowe aktualizują tę reprezentację wraz z napływem nowych zdań tak, aby późniejsze odniesienia były interpretowane poprawnie. W praktyce oznacza to integrację rejestru encji ze strukturami skorelowanych predykatów, każda nowa encja dodana do kontekstu otrzymuje unikalną identyfikację wykorzystywaną przez kolejne etapy analizy.



Analiza semantyczna znajduje bezpośrednie zastosowanie w tłumaczeniu maszynowym metodą interlingua: zdanie źródłowe zamieniane jest na niezależną od języka strukturę reprezentującą jego sens (może to być zapis logiczny albo wektorowy profil ramowy), a następnie generowana jest wersja docelowa zgodna ze składnią języka odbiorcy. Takie rozwiązanie pomaga zachować pełnię relacji między obiektami zdania i uniknąć utraty znaczeń idiomatycznych. Końcowym aspektem wartym podkreślenia jest ewaluacja jakości analizy semantycznej. W projektach przemysłowych stosuje się testy porównawcze polegające na ocenie trafności przypisań encji i relacji względem ręcznie anotowanych zbiorów referencyjnych. Wyniki takich testów pozwalają identyfikować typowe luki w działaniu komponentu oraz ukierunkowywać kolejne etapy optymalizacji algorytmicznej. Podsumowując, analiza semantyczna łączy formalne metody ścisłego odwzorowania znaczeń z elastycznością modeli uczonych statystycznie i głęboko. Skuteczność jej wdrożenia zależy od jakości zasobów leksykalnych, bogactwa korpusu treningowego oraz umiejętności systemu do integrowania informacji parametrycznej z dynamicznie pozyskiwaną wiedzą kontekstową.

4.2 Przetwarzanie morfologiczne

Przetwarzanie morfologiczne koncentruje się na analizie budowy wyrazów oraz identyfikowaniu zależności pomiędzy ich formami fleksyjnymi a znaczeniem i funkcją w zdaniu. Jest to proces, który łączy informacje o rdzeniu wyrazu, afiksach oraz nieregularnych odmianach, tak aby uzyskać reprezentację możliwą do dalszej obróbki algorytmicznej. Typowe operacje obejmują lematyzację, czyli sprowadzenie form odmiennych do postaci podstawowej (lematu), oraz stemming, redukcję końcówek i afiksów bez pełnej rekonstrukcji zasady gramatycznej. Stemming bywa mniej precyzyjny, ale szybszy obliczeniowo; dla wielu zastosowań technicznych wystarcza usunięcie typowych sufiksów według prostych reguł regularnych. Przykładowo, stosowanie wyrażań regularnych do separowania rdzenia od końcówki jest niezależne od słownika, dzięki czemu może działać na językach o bogatej morfologii. Jednak algorytmy tego typu często napotykają problemy przy rozpoznawaniu złożonych form lub nietypowych odmian, co prowadzi do błędów w grupowaniu oraz interpretacji.

Lematyzacja wymaga dostępu do zasobów leksykalnych, słowników lub baz danych, które zawierają pełen zestaw dopuszczalnych form fleksyjnych wraz z ich właściwościami morfosyntaktycznymi. Implementacje mogą wykorzystywać NLTK jako bibliotekę pośredniczącą w obsłudze takich zasobów: oferuje ona struktury danych reprezentujące słowniki i reguły morfologiczne oraz funkcje pozwalające na integrację analizy morfologicznej z tokenizacją czy parserem składniowym. Wiedza zapisana w słowniku umożliwia poprawne przypisanie informacji o części mowy,



rodzaju gramatycznym czy liczbie do każdego słowa w korpusie, co ma kluczowe znaczenie dla późniejszych etapów analizy semantycznej i syntaktycznej. W przetwarzaniu korpusowym analiza morfologiczna pełni rolę filtra normalizującego dane wejściowe. Zanim model językowy zacznie uczyć się zależności statystycznych, dane tekstowe są oczyszczane ze zbędnych wariantów fleksyjnych poprzez sprowadzenie ich do ustalonej formy kanonicznej. Pozwala to zmniejszyć wymiar przestrzeni cech i zwiększa gęstość obserwacji dla danego lemu w zbiorze treningowym. Ma to szczególne znaczenie dla dużych modeli językowych, zmniejszenie liczby unikalnych tokenów skraca czas trenowania i redukuje zapotrzebowanie na moc obliczeniową. Przetrenowanie takich modeli obejmuje analizę tekstu na poziomie pojedynczych tokenów; rozpoznanie prefiksów i sufiksów już na tym etapie może pomóc lepiej uchwycić strukturę wewnętrzną słowa.

Zintegrowanie przetwarzania morfologicznego z pipeline'em NLP wymaga uwzględnienia różnorodności językowej. Języki izolujące (np. chiński) praktycznie nie korzystają z odmiany gramatycznej, więc analiza skupia się głównie na segmentacji znaków. Natomiast języki fleksyjne (np. polski) albo aglutynacyjne (np. turecki) stawiają przed algorytmami zadanie dekodowania informacji zawartej w afiksach i rdzeniu jednocześnie. Podejścia oparte na tzw. morfologii automatycznej implementują reguły tworzenia form odmienionych przy użyciu przetworników skończonych (FST), które pozwalają przeprowadzać transformacje między zapisem kanonicznym a formą powierzchniową zgodnie z zestawem opisanych reguł fonetycznych i ortograficznych.

W praktyce coraz częściej łączy się metody deterministyczne z uczeniem maszynowym. Modele sekwencyjne takie jak LSTM czy GRU potrafią przewidywać formę bazową słowa lub jego cechy morfologiczne na podstawie kontekstu zdania bez jawnego dostępu do słownika. Transformery rozszerzają te możliwości dzięki mechanizmowi uwagi, umożliwiającemu dynamiczne podkreślenie zależności między sąsiadującymi tokenami a potencjalnym wzorcem odmiany. W modelach generatywnych parametry opanowują statystyczną reprezentację relacji pomiędzy ciągami liter a kategoriami gramatycznymi, co pozwala generować poprawne słowa odmienne w odpowiednim kontekście bez ręcznego definiowania zasad.

Osobnym przypadkiem pracy nad morfologią jest analiza nazw własnych oraz terminologii specjalistycznej, często wymagają one innych reguł dekompozycji niż język ogólny. W systemach tłumaczeń maszynowych dbałość o zachowanie poprawnej formy morfologicznej ma wpływ na zrozumiałość i zgodność przekładu ze standardem języka docelowego. Transfer modeli między językami o różnym stopniu złożoności fleksyjnej wymaga optymalizacji komponentu morfologicznego tak, aby minimalizować utratę informacji zakodowanej w afiksach podczas transformacji. W pewnym sensie przetwarzanie morfologiczne jest też zadaniem segmentacji strumienia tekstowego, rozbijania ciągu znaków na jednostki minimalne niosące



informację lingwistyczną. Może to być realizowane przez proste procedury usuwania sufiksów działające według określonego wzorca wyrażenia regularnego, albo przez strukturalną analizę umożliwiającą przypisanie wszystkich cech gramatycznych konkretnej formie wyrazu. Choć pierwsze podejście bywa wystarczające w wyszukiwarkach tekstowych lub indeksowaniu dokumentów, drugie jest konieczne tam, gdzie zachowanie ścisłości interpretacji ma nadrzędne znaczenie.

Integracja wiedzy parametrycznej, zakodowanej w modelach podczas trenowania, z wiedzą nieparametryczną pozyskiwaną dynamicznie z baz leksykalnych pozwala systemom NLP radzić sobie zarówno z typowymi formami odmiennymi, jak i neologizmami czy zapożyczeniami. Dzięki temu możliwe jest reagowanie na zmiany użycia języka w czasie rzeczywistym bez konieczności ponownego trenowania całego modelu od podstaw. Analiza jakości wynikowej segmentacji i lematyzacji odbywa się zwykle przy użyciu zbiorów testowych zawierających ręcznie oznaczone pary forma–lemma wraz z etykietami części mowy. Uzyskane miary trafności wskazują zarówno mocne strony algorytmu (np. obsługa regularnych wzorców odmiany), jak i problemy wymagające korekty (np. wieloznaczne homonimy). W kontekście zastosowań przemysłowych utrzymanie wysokiej jakości analizy morfologicznej przekłada się bezpośrednio na trafność wyszukiwania informacji, jakość tłumaczeń czy spójność generowanych treści przez chatboty i systemy głosowe. Przetwarzanie morfologiczne pozostaje więc istotnym elementem łańcucha technologii językowych, który wpływa zarówno na przygotowanie danych wejściowych dla bardziej złożonych analiz syntaktycznych i semantycznych opisanych wcześniej, jak i na precyzję operacji generujących nowy tekst zgodny z regułami danego języka naturalnego.

4.3 Generowanie języka naturalnego

Systemy dialogowe

Systemy dialogowe stanowią klasę technologii językowych, której celem jest realizacja interakcji językowej między człowiekiem a maszyną w sposób możliwie naturalny. Ich architektura obejmuje komponenty odpowiedzialne za rozpoznanie intencji użytkownika, interpretację kontekstu oraz generowanie odpowiedzi. Kluczowym elementem jest moduł przetwarzania języka naturalnego (NLP), który łączy analizę składniową i semantyczną z mechanizmami zarządzania kontekstem konwersacji. W najprostszej formie system dialogowy może działać na zasadzie dopasowania wzorców do wcześniej zdefiniowanych szablonów pytań i odpowiedzi, jednak takie podejście szybko ujawnia ograniczenia przy obsłudze bardziej złożonych lub nieprzewidzianych zapytań. Modele współczesnych systemów dialogowych korzystają często z kombinacji reguł deterministycznych oraz metod statystycznych. Reguły zapewniają przewidywalność i zgodność z określonym



standardem lingwistycznym, natomiast algorytmy uczące się pozwalają na adaptację do nowych sposobów formułowania pytań przez użytkowników.

Przykład prostego scenariusza dialogowego to sytuacja, gdy użytkownik pyta o godzinę emisji filmu – system rozpoznaje strukturę pytania i poprzez logikę biznesową wywnioskowuje, że intencją użytkownika jest uzyskanie informacji o repertuarze. Aby uniknąć sytuacji, w której odpowiedź byłaby nieprzydatna (np. „Tak”), implementuje się reguły dostosowane do specyfiki aplikacji, każde pytanie zaczynające się od „Kiedy...” czy „O której...” automatycznie uruchamia moduł wyszukiwania dat i godzin. Historyczne komercyjne systemy dialogowe opierały się na silnym ograniczeniu domeny, przez co były zdolne realizować jedynie zbiór wcześniej przewidzianych funkcji. Mogły np. udzielać informacji o repertuarze kin, ale nie potrafiły odpowiedzieć na pytania dotyczące innych tematów.

Współcześnie dzięki dużym modelom językowym (LLM) możliwe staje się poszerzenie zakresu domen bez konieczności ręcznej enumeracji ogniw konwersacji. LLM-y uczone są na olbrzymich zbiorach danych tekstowych i potrafią generować odpowiedzi syntaktycznie oraz semantycznie poprawne dla szerokiego spektrum zapytań, choć ryzyko generowania treści niedokładnych lub halucynacyjnych nadal istnieje. Naturalność interakcji wymaga śledzenia stanu konwersacji, tzw. pamięci dialogowej, która przechowuje odniesienia do wcześniejszych wypowiedzi użytkownika.

Zaawansowane platformy integrują pamięć parametryczną (informacje zakodowane w wagach modelu podczas treningu) z pamięcią nieparametryczną pochodzącą ze źródeł zewnętrznych. Przykładowo, komponent taki jak Łańcuch pobierania konwersacyjnego umożliwia wyszukanie dodatkowych informacji w bazie wektorowej dokumentów i użycie ich do tworzenia spójnej kontynuacji rozmowy. W praktyce budowa systemu dialogowego obejmuje również elementy rozpoznawania mowy oraz syntezy mowy, szczególnie w aplikacjach głosowych. Umożliwia to interakcję bezpośrednio za pomocą wypowiedzi ustnych – wypowiedź wejściowa po rozpoznaniu treści jest przekazywana do modułu NLP, a następnie wygenerowana odpowiedź zostaje przekształcona w mowę. Wersje dla środowisk wielojęzycznych dodatkowo integrują mechanizmy tłumaczenia maszynowego tak, aby obsługiwać komunikację pomiędzy rozmówcami mówiącymi różnymi językami.

Istotnym wyzwaniem dla projektantów takich systemów jest obsługa niejednoznaczności, zarówno syntaktycznej, jak i semantycznej. Modele oparte na regułach mogą stosować hierarchię priorytetów lub prośby o doprecyzowanie polecenia („Czy chodziło Ci o...?”), natomiast modele neuronowe mogą próbować rozstrzygać znaczenie przy wykorzystaniu statystycznego podobieństwa osadzeń. Dobór metody zależy od wymagań aplikacji: w środowisku krytycznym preferuje się rozwiązania zapewniające pełną wyjaśnialność decyzji. Nowoczesne systemy



dialogowe zaczynają też działać jako autonomiczne agenty wykonujące zadania na rzecz użytkownika, planowanie podróży, analizowanie dokumentów czy wyszukiwanie danych, poprzez sekwencję zapytań do innych narzędzi i usług. Integracja takich agentów wymaga od systemu posiadania warstwy decyzyjnej zdolnej wybrać optymalny przebieg działań: kiedy pobrać dane z bazy, kiedy skonsultować je z modelem LLM, a kiedy poprosić użytkownika o dodatkowe informacje.

Projektując interfejs komunikacyjny należy uwzględnić również czynniki psycholingwistyczne. Użytkownicy skłonni są przypisywać systemom cechy ludzkie, jeśli reakcja maszyny jest płynna i adekwatna emocjonalnie, rośnie akceptacja jej działania. Stąd pojawia się trend implementowania wariantów tonalnych odpowiedzi dopasowanych do intencji rozmowy (np. poważny ton przy danych technicznych i swobodny przy small talk). Optymalizacja jakości pracy systemu odbywa się poprzez ewaluację opartą na benchmarkach konwersacyjnych, mierząc trafność odpowiedzi względem oczekiwań oraz spójność między kolejnymi turami rozmowy. W przypadku wielojęzycznych wersji testuje się także poprawność tłumaczeń oraz płynność przełączania kontekstu między językami. Szczególnym nurtem są rozwiązania wdrażane wirtualnych asystentów, które łączą funkcje wyszukiwarki i interpretatora poleceń z generatywnym modelem treści. Asystent może odpowiadać na pytania faktograficzne („Jaka jest pogoda?”), wykonywać proste czynności („Włącz muzykę”) albo generować oryginalny tekst („Stwórz opis produktu”). W odróżnieniu od historycznych chatbotów ograniczonych liczbą scenariuszy dialogowych, takie asystenty operują na wszechstronnym modelu będącym częścią większego ekosystemu AI.

Różnica między klasycznymi chatbotami a pełnymi narzędziami AI polega na zakresie funkcji: chatbot jest zwykle przeznaczony do pojedynczych zadań lub działający w jednej domenie tematycznej; narzędzie AI może integrować wiele modeli i zasobów jednocześnie. Trend hybrydyzacji pozwala łączyć szybkość reakcji chatbota z wiedzą gromadzoną przez model LLM wspierany zasobami nieparametrycznymi. Integracja powyższych elementów pokazuje ewolucję systemów dialogowych od prostych programów opartych na dopasowaniu wzorca po rozbudowane architektury agentowe operujące w czasie rzeczywistym. Konsekwentne rozwijanie komponentów NLP wraz ze skalowaniem modeli neuronowych oraz rozszerzeniem ich źródeł wiedzy wydaje się kierunkiem dominującym w kolejnych generacjach takich technologii, pozwalając osiągnąć większą elastyczność wypowiedzi bez utraty kontroli nad jej zgodnością informacyjną.



Tworzenie tekstów automatycznych

Tworzenie tekstów automatycznych obejmuje generowanie wypowiedzi w języku naturalnym przez systemy komputerowe na podstawie zadanej instrukcji, danych wejściowych lub kontekstu rozmowy. Systemy realizujące tę funkcję mogą korzystać z różnych podejść, od szablonów uzupełnianych danymi po modele generatywne uczone na dużych korpusach. Obecnie dominują algorytmy oparte na architekturze transformatorowej, które potrafią tworzyć spójne i wieloznaczne narracje, odpowiadające stylowi i formie żądanej przez użytkownika. Wykorzystanie mechanizmu uwagi umożliwia dynamiczne określanie, które fragmenty kontekstu są istotne w każdym kroku generacji, co przekłada się na zachowanie koherencji wypowiedzi nawet w dłuższych akapitach. Proces tworzenia treści zaczyna się od etapu interpretacji polecenia. Tutaj istotną rolę odgrywa projektowanie promptów, polegające na formułowaniu instrukcji w taki sposób, aby model miał możliwie precyzyjne wskazówki dotyczące celu i oczekiwanej struktury tekstu. Słowa kluczowe zawarte w promptach uruchamiają skojarzenia statystyczne zapisane w parametrach modelu, np. termin „prescribe” przywoła sieć powiązań z medycyną, lekarzami i chorobami. Odpowiednia konstrukcja promptu może ograniczać ryzyko powstawania tzw. halucynacji, czyli informacji nieprawdziwych generowanych przez model. W historii implementacji automatycznego pisania długich form pojawia się istotny problem: wiele modeli radzi sobie dobrze z krótkimi odpowiedziami czy podsumowaniami, ale utrzymanie wysokiej jakości treści rozbudowanej jest trudniejsze. Modele generatywne mają tendencję do powtarzania struktur zdaniowych lub tracenia spójności argumentacji w długich fragmentach. Rozwiązaniem bywa podział zadania na mniejsze segmenty, najpierw tworzenie szkieletu treści, a następnie rozwijanie poszczególnych punktów. Ten podział pozwala zachować logikę wyводу i uniknąć nadmiernego rozwlekania myśli. Tworzenie treści automatycznych znajduje zastosowanie od prostych opisów produktów po skomplikowane raporty analityczne.

W przypadku tekstów kreatywnych, opowiadań czy artykułów publicystycznych, istotne jest dostosowanie tonu i stylu wypowiedzi do odbiorcy. Modele mogą zostać dostrojone (fine-tuning) na korpusach reprezentujących konkretny gatunek literacki czy rodzaj publikacji, co skutkuje lepszym dopasowaniem stylistycznym odpowiedzi. Przy pracy nad naukowymi artykułami czy dokumentacją techniczną wymagane jest jednak większe wsparcie ze strony użytkownika: model potrafi przygotować logicznie spójną konstrukcję zdań, lecz wymaga skontrolowania poprawności merytorycznej.

Istotnym wyzwaniem pozostaje integracja wiedzy aktualnej z procesem generacji tekstu. Wiedza zapisana parametrycznie w modelu często pochodzi sprzed kilku lat i nie obejmuje bieżących wydarzeń czy zmian definicyjnych. W takich przypadkach stosuje się mechanizmy wyszukiwania kontekstowego, np. poddawanie generatora dodatkowym informacjom pobranym z wektorowych baz danych lub zewnętrznych



API przed rozpoczęciem procesu tworzenia tekstu. Podejście to pozwala zachować świeżość treści oraz jej zgodność z aktualnym stanem wiedzy.

Oprócz produkcji tekstów o charakterze informacyjnym coraz większą popularność zdobywa generowanie utworów artystycznych, poezji, scenariuszy czy piosenek, które wykorzystują zdolności modeli do kreatywnego łączenia motywów tematycznych. Narzędzia takie jak SciWriter specjalizują się w pisaniu naukowym, podczas gdy inne platformy skupiają się na aspektach artystycznych, jak kompozycja muzyki lub tworzenie dialogów fabularnych. Jakość automatycznie tworzonych tekstów ocenia się według kilku kryteriów: spójności logicznej (czy kolejne zdania wynikają jedno z drugiego), adekwatności stylistycznej (dopasowanie języka do odbiorcy), poprawności językowej (gramatyka i ortografia) oraz trafności merytorycznej względem źródła informacji. Ewaluacja często wymaga zarówno automatycznych miar podobieństwa semantycznego, jak i oceny ludzkiej, szczególnie tam, gdzie ważny jest ton czy subtelna perswazja.

Pewnym rozszerzeniem jest wykorzystanie hybrydowych architektur integrujących zasady regułowe z uczeniem statystycznym. Takie rozwiązanie pozwala narzucić sztywne ramy formatowe lub terminologiczne przy jednoczesnym zachowaniu elastyczności syntaktycznej generatu. Przykładem może być system tworzący raport medyczny: sekcje raportu są stałe (nagłówki, struktura wyników badań), natomiast opis obserwacji powstaje dynamicznie na podstawie danych pacjenta i wskazanych wytycznych stylistycznych. Realizacja procesu generacji musi też uwzględniać ograniczenia obliczeniowe. Modele dużej skali wymagają znacznych zasobów GPU/TPU; ich uruchamianie lokalne bywa możliwe jedynie dla skróconych wersji przystosowanych do specyficznych zadań. Wersje chmurowe zapewniają pełny dostęp do możliwości obliczeniowych kosztem zależności od infrastruktury dostawcy usług.

Z perspektywy praktycznej warto zaznaczyć korelację pomiędzy jakością promptu a jakością wynikowego tekstu, niewłaściwe sformułowanie instrukcji może prowadzić do powstania treści niezrozumiałych lub niezgodnych z intencją nadawcy. Edukacja użytkowników w zakresie efektywnego projektowania zapytań dla generatora staje się więc integralną częścią wdrożeń tych technologii. Automatyczne tworzenie tekstów ewoluuje obecnie w kierunku większej interaktywności: użytkownik nie tylko wydaje pojedyncze polecenie, ale może iteracyjnie poprawiać fragmenty wygenerowanego materiału, dostosowując je tak długo, aż osiągną wymagany poziom jakości. Ostateczny kształt systemu generującego zależy od rozstrzygnięcia kompromisu między szybkością działania a jakością merytoryczną i stylistyczną uzyskiwanego tekstu. W środowiskach krytycznych dopuszcza się wolniejsze tempo produkcji treści w zamian za rygorystyczną zgodność faktograficzną; w aplikacjach kreatywno-rozrywkowych przeciwnie, liczy się płynność dialogu czy narracji przy tolerancji dla drobnych odstępstw od rzeczywistych faktów.



5 Uczenie maszynowe w inżynierii lingwistycznej

5.1 Metody klasyczne

Modele n-gramowe

Modele n-gramowe należą do najprostszych, a jednocześnie historycznie istotnych metod modelowania języka w systemach NLP. Ich konstrukcja opiera się na analizie ciągów n kolejnych elementów, najczęściej słów lub znaków, występujących w tekście i estymacji prawdopodobieństw pojawienia się danego ciągu w określonym kontekście. Gdy $n=1$, mamy do czynienia z modelem unigramowym, który ignoruje zależności między sąsiadującymi tokenami, traktując każde słowo jako niezależne zdarzenie. W przypadku bigramów ($n=2$) uwzględnia się relację pomiędzy bieżącym a poprzednim słowem, natomiast trigramy ($n=3$) i wyższe wartości n pozwalają uchwycić dłuższy kontekst lokalny. Formalnie, model n-gramowy przyjmuje założenie Markowa, że prawdopodobieństwo wystąpienia słowa w zdaniu zależy wyłącznie od $n-1$ poprzednich słów:

$$P(W_1^K) = \prod_{k=1}^K P(w_k | W_{k-n}^{k-1}), \quad 1 < n \leq N$$

co upraszcza obliczenia względem pełnego modelu sekwencji, który musiałby uwzględniać wszystkie wcześniejsze tokeny. W praktyce wartość [eq:ngram_prob] estymuje się na podstawie częstości występowania odpowiednich sekwencji w korpusie treningowym. Im większy zbiór danych, tym lepsze oszacowanie tych prawdopodobieństw, jednak wraz ze wzrostem n szybko rośnie liczba możliwych kombinacji i problem rzadkich lub nieobserwowanych n-gramów staje się bardziej dotkliwy. Rozwiązaniem problemu zerowych prawdopodobieństw jest stosowanie technik wygładzania, takich jak Laplace'a, Good–Turing czy interpolacja Knesera–Neya. Zabiegi te przesuwają część masy prawdopodobieństwa z obserwowanych n-gramów na te niezaobserwowane, co poprawia ogólną zdolność modelu do generowania realistycznych sekwencji. Modele interpolowane łączą informacje z różnych rzędów n , np. bigramów i trigramów, ważonych współczynnikami odzwierciedlającymi ich trafność w przewidywaniu kolejnego tokena. Budowa modeli n-gramowych wiąże się ze wstępnym przetwarzaniem tekstu: tokenizacja dzieli dane wejściowe na jednostki (słowa lub znaki), a następnie tworzy listę wszystkich wystąpień ciągów długości n . Przed obliczeniem statystyk często stosuje się normalizację, sprowadzenie liter do jednego formatu czy usunięcie znaków interpunkcyjnych, aby uniknąć nadmiernej fragmentacji danych. W wielu przypadkach wykorzystuje się również listy tzw. stopwords, eliminując z analizy bardzo częste słowa o niskiej wartości informacyjnej; jednak jak wskazują badania praktyczne, niekiedy metoda ta pogarsza wyniki klasyfikacji tekstu, co oznacza potrzebę eksperymentowania z parametrami przetwarzania (np. zakresem n-gramów czy minimalną częstością występowania min_df).



Modele n-gramowe znajdują zastosowanie m.in. w systemach rozpoznawania mowy (ASR), gdzie przewidywanie najbardziej prawdopodobnego ciągu słów wspiera dekodery rozpoznające sygnały akustyczne; oraz w zadaniach sugestii kończenia tekstu czy autokorekty. W tłumaczeniach maszynowych starszego typu (statystycznych SMT) modele takie były kluczowym komponentem zapewniającym płynność składniową zdań docelowych poprzez ocenę prawdopodobieństwa całych fraz. Choć we współczesnych architekturach neuronowych zastąpiono je reprezentacjami wektorowymi głębokiego uczenia, to nadal stanowią one bazową technologię referencyjną, prostą do implementacji i interpretacji. W pracy nad klasyfikacją dokumentów modele n-gramowe mogą funkcjonować jako metoda ekstrakcji cech wejściowych dla algorytmów takich jak modele Bayesowskie czy regresja logistyczna. Na przykład, bigramy uchwytujące połączenia „nie polecam” czy „bardzo dobre” są silnymi zmiennymi sentymentu opinii. Z kolei trigramy mogą lepiej modelować frazy idiomatyczne lub typowe konstrukcje językowe istotne dla danej domeny. Implementacja tych modeli w bibliotekach takich jak Scikit-Learn w Pythonie oferuje elastyczność parametrów i łatwość integracji z dowolnymi algorytmami klasyfikacyjnymi. Równolegle pojawia się kwestia optymalnego wyboru wartości n . Dla języków o skomplikowanej morfologii duże n zapewnia lepsze odwzorowanie kontekstu fleksyjnego; jednak rosnąca liczba unikalnych sekwencji wymaga intensywnej filtracji i efektywnego zarządzania pamięcią. Stąd typowe implementacje ograniczają rozmiar słownika poprzez usuwanie rzadkich n-gramów lub stosowanie ważenia TF-IDF redukującego wpływ bardzo częstych fragmentów o małej treści informacyjnej.

Mimo swoich ograniczeń, ignorowania dalekich zależności i semantycznego znaczenia słów poza lokalnym kontekstem, podejście n-gramowe posiada zalety diagnostyczne. Statystyka częstości fraz ujawnia zmiany trendów językowych w czasie czy charakterystyczne kolokacje dla wybranego gatunku wypowiedzi. Analiza taka może stanowić punkt wyjścia do projektowania bardziej złożonych modeli, łącząc proste cechy ilościowe z bogatszym opisem morfologiczno-syntaktycznym. Współczesne badania pokazują także możliwość integracji tradycyjnych modeli n-gramowych z sieciami neuronowymi jako dodatkowej warstwy wejściowej, parametry statystyczne fraz lokalnych wspierają proces uczenia tam, gdzie dostępny korpus jest stosunkowo niewielki. Pozwala to na wzbogacenie reprezentacji i stabilizację procesu trenowania poprzez dostarczenie modelowi informacji o częstotliwościach użycia konkretnych sekwencji tokenów. Podsumowując, modele n-gramowe zachowują swoją rolę jako prosta technika modelowania języka oraz punkty odniesienia dla oceny jakości bardziej złożonych systemów NLP. Ich walorem jest przejrzystość mechanizmu działania oraz możliwość szybkiego wdrożenia prototypowego rozwiązania analitycznego bez potrzeby inwestowania dużych zasobów obliczeniowych czy skomplikowanych procedur uczenia głębokiego.



Drzewa decyzyjne w analizie języka

Drzewa decyzyjne należą do klasycznych metod uczenia maszynowego, które w kontekście analizy języka naturalnego pozwalają w przejrzysty sposób odwzorować proces decyzyjny oparty na cechach wyekstrahowanych z danych tekstowych. Struktura drzewa składa się z węzłów reprezentujących wybór cechy, gałęzi określających warunki dla wartości tej cechy oraz liści, które odpowiadają ostatecznym decyzjom, w zadaniach językowych mogą to być etykiety klas przy klasyfikacji lub wartości predykcji przy regresji. Każdy węzeł analizuje pojedynczą cechę (np. obecność konkretnego tokena, wynik funkcji podobieństwa semantycznego czy wartość wskaźnika TF-IDF), a przejście do kolejnego poziomu struktury odbywa się według reguły optymalizującej kryterium jakości podziału.

Proces budowy drzewa decyzyjnego jest iteracyjny i rozpoczyna się od korzenia zawierającego całość danych treningowych. Na każdym etapie wybierana jest cecha oraz próg jej wartości dzielący zbiór na dwa lub więcej podzbiorów tak, aby najlepiej odseparować klasy lub zmniejszyć wariancję wartości docelowej. Dobór cechy odbywa się według miary nieczystości dla klasyfikacji, np. wskaźnika Giniego czy entropii bądź miary błędu (sumy kwadratów odstępstw lub błędu bezwzględnego) dla regresji. W analizie języka entropia może uwzględniać rozkład kategorii dokumentów względem wartości danej cechy; im bardziej jednorodny podzbiór powstaje po podziale, tym lepsza jakość węzła. Zastosowanie progów liczbowych (np. „częstość słowa > 3”) pozwala mierzyć relewantność cechy w automatycznym klasyfikowaniu tekstu.

Dla danych tekstowych wyzwaniem pozostaje odpowiednie przygotowanie cech wejściowych, surowe tokeny muszą zostać przekształcone do form numerycznych przydatnych dla algorytmu drzewa. Popularne rozwiązania obejmują reprezentacje worka słów, n-gramy opisane wcześniej lub osadzenia wektorowe, które przechowują relacje semantyczne pomiędzy jednostkami leksykalnymi. W każdym przypadku decyzja co do rodzaju i zakresu cech wpływa na przyszłą strukturę drzewa oraz jego zdolność generalizacji.

Projektując drzewo należy kontrolować głębokość struktury oraz minimalną liczbę próbek w liściu. Zbyt głębokie drzewa dopasowują się nadmiernie do danych treningowych, co skutkuje przeuczeniem modelu i spadkiem skuteczności na danych nowych. Ograniczenie głębokości oraz liczby próbek w liściu jest typowym warunkiem stopu, pozwalającym uniknąć wydobywania wzorców ze szumu charakterystycznego dla danych językowych o dużej zmienności. Można też ustawić minimalną liczebność zbioru wymaganą do wykonania podziału, model przerwie dalsze rozgałęzienia, gdy liczba przykładów spadnie poniżej ustalonego progu (np. 50 lub 200 elementów). W praktyce drzewo decyzyjne może wykonywać zarówno zadania klasyfikacyjne jak i regresyjne w obszarze NLP.



Przy klasyfikacji sentymentu opinia użytkownika intencji jest oznaczana etykietą „pozytywna”, „neutralna” albo „negatywna”; w regresji można przewidywać wartość ciągłą, np. ocenę recenzji w skali 1–5 punktów na podstawie treści komentarza. Różnica tkwi w rodzaju wartości umieszczanej w liściach: kategorie klas vs. średnia wartość predykcji. Budowa drzewa może być wspierana przez narzędzia takie jak DecisionTreeClassifier czy DecisionTreeRegressor z biblioteki Scikit-Learn. Implementacja wykorzystuje algorytm CART (Classification and Regression Trees) budujący binarne drzewo poprzez kolejne podziały i umożliwia wybór kryterium optymalizacji (np. „gini” lub „entropy”). Parametry takie jak max_depth, min_samples_split czy min_samples_leaf pozwalają sterować procesem wzrostu drzewa oraz redukować ryzyko przeuczenia. Efekt końcowy można wizualizować, ukazując które cechy są kluczowe w procesie decyzyjnym i jakie progi zostały zastosowane na poszczególnych poziomach struktury.

Entropia sprawdzana dla kolumn danych językowych działa tu jako miara jednorodności pojęć obecnych po dokonaniu podziału. Jeśli dystrybucja jest bliska równomiernej (wartość entropii wysoka), podział nie przynosi istotnej poprawy czystości zbioru; jeśli entropia maleje po rozgałęzieniu, oznacza to wyraźniejszą separację kategorii. Dla zmiennych ciągłych wyznacza się próg optymalny przez testowanie różnych punktów cięcia i wybór tego, który daje największą redukcję nieczystości lub błędu. Choć pojedyncze drzewo bywa łatwe do interpretacji, jego stabilność względem małych zmian danych może być ograniczona, niewielkie modyfikacje zbioru uczącego mogą prowadzić do znacznego przekształcenia struktury drzewa. Problem ten łagodzą podejścia zespołowe takie jak lasy losowe, gdzie wiele drzew o różnych losowo wybranych podzbiorach danych i cech pracuje równolegle, a wynik decyzji opiera się na agregacji przewidywań wszystkich modeli cząstkowych.

W kontekście NLP lasy losowe mogą wykorzystywać różne kombinacje n-gramów czy osadzeń słów jako atrybuty wejściowe dla poszczególnych drzew. Z perspektywy analitycznej drzewa wspierają interpretowalność dzięki możliwości wskazania ścieżki decyzyjnej, analityk czy lingwista może zobaczyć, które słowo kluczowe lub wynik funkcji podobieństwa był podstawą przypisania dokumentowi danej etykiety. Tego typu transparentność jest cenna tam, gdzie konieczne jest uzasadnienie wyników modelu przed użytkownikiem końcowym lub instytucją kontrolną. Integracja wiedzy lingwistycznej z procesem budowy drzewa bywa praktykowana poprzez selekcję cech uwzględniających specyficzne własności języka badanego, np. dla polskiego dobór informacji o formach fleksyjnych słów pozwala lepiej uchwycić różnice między konstrukcjami zdaniowymi. Można również obliczać specjalistyczne wskaźniki semantyczne (np. częstość użycia określonego sensu wyrazu) i traktować je jako parametry decyzyjne. Podsumowując, drzewa decyzyjne stanowią uniwersalne narzędzie analizy języka naturalnego zarówno dla zadań klasyfikacyjnych jak i



regresyjnych, utrzymując równowagę między prostotą implementacyjną a możliwością tworzenia złożonych wielopoziomowych modeli zależności między cechami tekstowymi a etykietami wynikowymi. Ich efektywność zależy od świadomego wyboru miernika jakości podziałów, kontroli parametrów wzrostu struktury oraz przygotowania reprezentatywnych i informacyjnych cech na wejściu systemu NLP.

5.2 Uczenie głębokie

Sieci neuronowe rekurencyjne

Sieci neuronowe rekurencyjne (RNN) stanowią istotny typ architektury w uczeniu głębokim, zaprojektowany specjalnie do analizy sekwencji danych, w tym języka naturalnego. Podstawową cechą odróżniającą RNN od jednokierunkowej sieci neuronowej jest możliwość przechowywania stanu wewnętrznego (tzw. wektor ukryty), który jest aktualizowany na każdym kroku przetwarzania sekwencji i przenosi informację z wcześniejszych elementów do kolejnych etapów analizy. Dzięki temu RNN mogą modelować zależności w czasie i pozycji tokenów w zdaniu, co jest kluczowe dla zachowania kontekstu językowego. Każdy krok obliczeniowy RNN przyjmuje bieżący element wejściowy oraz stan ukryty z poprzedniego kroku jako argumenty funkcji aktywacji, generując nowy stan ukryty oraz ewentualnie wartość wyjściową. Ten mechanizm sprawia, że sieć działa iteracyjnie na kolejnych tokenach. Jednakże klasyczne RNN mają ograniczoną zdolność do przechowywania informacji o bardzo odległych fragmentach sekwencji. Zjawisko znane jako problem zanikającego lub eksplodującego gradientu powoduje trudności w nauce długoterminowych zależności – gradient obliczany w trakcie propagacji wstecznej przez wiele kroków czasowych może maleć niemal do zera lub rosnąć wykładniczo, utrudniając skuteczne dopasowanie wag. Próby rozwiązania tego problemu doprowadziły do powstania modyfikacji architektury, takich jak długa pamięć krótkotrwała (LSTM), które wprowadzają dodatkowe mechanizmy sterujące przepływem informacji za pomocą bramek wejścia, wyjścia i zapomnienia. Bramki te wykorzystują operacje mnożenia punktowego, aby decydować, jakie informacje przejść z wejścia lub poprzedniego stanu ukrytego, a jakie usunąć. Dzięki takiemu rozwiązaniu LSTM mogą utrzymywać istotne informacje przez dziesiątki lub setki kroków czasowych.

Architektura LSTM składa się z komórki pamięci oraz zestawu wspomnianych bramek kontrolujących jej zawartość. Komórka może przechowywać wartości niezależnie od tego, co dzieje się w krótkim okresie – pozwala to np. utrzymać podmiot zdania w pamięci podczas przetwarzania wielu fraz przypisanych mu semantycznie. W zastosowaniach translacyjnych analogiczne działanie pozwala zachować spójność referencji między językami nawet wtedy, gdy struktura



syntaktyczna wymusza znaczną odległość między powiązаныmi słowami. LSTM stały się dominującym rozwiązaniem problemu modelowania długich sekwencji zanim pojawiły się architektury transformatorowe. Innym wariantem RNN są sieci GRU, które upraszczają konstrukcję LSTM poprzez redukcję liczby bramek i uproszczony model komórki pamięci. GRU łączą bramkę zapominania i wejścia w jedną strukturę aktualizacji, co zmniejsza liczbę parametrów przy zachowaniu zdolności do zarządzania przepływem informacji.

W kontekście inżynierii lingwistycznej takie uproszczenie pozwala osiągnąć podobną jakość predykcji przy mniejszym koszcie obliczeniowym. RNN były naturalnym wyborem dla pierwszych udanych modeli tłumaczeń maszynowych opartych na architekturze enkoder–dekoder. W takich systemach moduł enkodera sekwencyjnie czytał zdanie źródłowe, aktualizując swój stan ukryty tak, aby ostatecznie zgromadzić skondensowaną reprezentację całej wypowiedzi w końcowym wektorze. Następnie moduł dekodera generował tłumaczenie słowo po słowie jako sekwencję predykcji zależną od swojego aktualnego stanu i wygenerowanych wcześniej tokenów. Podejście to działało dobrze dla umiarkowanie krótkich zdań – dla dłuższych treści następował spadek jakości wynikający z utraty części kontekstu oryginalnego tekstu. Problem utraty kontekstu szczególnie uwydatniał się przy tłumaczeniu zdań ze skomplikowanymi relacjami między słowami lub tam, gdzie znaczenie zależało od fragmentów oddzielonych dużymi odstępami.

Choć LSTM częściowo łagodziły ten efekt dzięki mechanizmowi utrzymania istotnych informacji przez wiele kroków czasowych, ich natura sekwencyjna powodowała wolniejsze działanie – każde nowe słowo wymagało pobrania wyników poprzednich kroków zanim mogło zostać przetworzone. Obróbka równoległa była właściwie niemożliwa w klasycznych implementacjach RNN ze względu na inherentny charakter przetwarzania krok-po-kroku.

W analizie sentymentu RNN sprawdzają się dzięki zdolności kumulowania wpływu kolejnych słów na wynik emocjonalny wypowiedzi, rozpoznając np. że wyrażenie „wcale nie taki dobry” ma inny wydźwięk niż „bardzo dobry”, mimo występowania wspólnych pozytywnych przymiotników. Zastosowanie warstw dwukierunkowych umożliwia dodatkowo analizę kontekstu zarówno od początku do końca wypowiedzi, jak i w odwrotnej kolejności; zwiększa to trafność interpretacji wieloznacznych fragmentów dzięki korzystaniu z pełnej perspektywy otoczenia danego słowa. Sieci rekurencyjne można również łączyć z innymi rodzajami warstw – np. konwolucyjnymi – tworząc hybrydowe architektury eksploatujące jednocześnie możliwości lokalnego wykrywania wzorców i globalnego śledzenia interakcji między elementami sekwencji. Takie kombinacje znalazły zastosowanie m.in. w transkrypcji mowy na tekst czy detekcji nazw własnych w dokumentach.



Współczesne podejście często traktuje RNN jako komponent większych systemów bądź stosuje je tam, gdzie zależy nam na wyjątkowej zdolności uchwycenia relacji czasowych, np. w analizie strumieni danych mowy lub czata online wymagających ciągłego przetwarzania nowych wiadomości bez restartu modelu. Modele te potrafią utrzymywać wewnętrzny stan pomiędzy kolejnymi porcjami danych nawet jeśli dostęp do pełnego kontekstu jest ograniczony. W praktyce oznacza to możliwość implementowania systemów o długiej „pamięci rozmowy” działających na bazie jednolitej sesji użytkownika. W porównaniu z transformatorami RNN charakteryzuje prostota konstrukcji i mniejsze zapotrzebowanie na moc obliczeniową przy krótkich kontekstach, co czyni je nadal atrakcyjnym wyborem dla urządzeń o ograniczonych zasobach sprzętowych. Jednak ich definicyjna sekwencyjność pozostaje przeszkodą przy trenowaniu oraz inferencji dla bardzo dużych zbiorów danych tekstowych, tu transformery akcelerały działanie dzięki równoległemu przetwarzaniu całego kontekstu naraz.

Historia rozwoju RNN pokazuje ewolucję od prostych modeli posługujących się jedynie podstawowym mechanizmem powtarzalnego stanu ku bardziej wyrafinowanym komórkom pamięci zarządzanym bramkami oraz integracji z mechanizmami uwagi ustawiającymi nacisk na najbardziej relewantne fragmenty wejścia. Choć obecne trendy faworyzują architektury nierzadko pozbawione komponentu rekurencyjnego, jego idea nadal znajduje zastosowanie wszędzie tam, gdzie struktura danych ma naturalny charakter chronologiczny lub liniowy, co obejmuje sporą część przypadków użycia inżynierii lingwistycznej.

Transformery i modele językowe

Architektura transformera stała się fundamentem współczesnych dużych modeli językowych, wypierając w wielu zastosowaniach wcześniejsze sieci rekurencyjne. Jej kluczowym elementem jest mechanizm uwagi, a w praktyce wariant wielogłowy, umożliwiający jednoczesne śledzenie różnych typów powiązań pomiędzy tokenami wejściowymi. Transformery przetwarzają całą sekwencję równolegle, co eliminuje opóźnienia charakterystyczne dla architektur krok-po-kroku i sprzyja efektywności trenowania na bardzo dużych korpusach. W typowym bloku transformera znajdują się moduły maskowanej uwagi (w modelach generatywnych) oraz warstwy przekazywania w przód, które działają na każdej pozycji niezależnie, przekształcając wektory reprezentujące tokeny przy zachowaniu wymiarowości.

Mechanizm uwagi identyfikuje zależności semantyczne i syntaktyczne między elementami zdania, natomiast warstwy przekazywania w przód pozwalają nadać każdemu elementowi zindywidualizowaną transformację. Bezpośrednio z architektury transformera wywodzi się wiele nowoczesnych LLM-ów opracowanych przez różne firmy. Przykłady obejmują GPT-3 Davinci (175B) od OpenAI, Jurassic-1-Jumbo (178B) od AI21 Labs czy Megatron-Turing NLG (530B) będący wspólnym projektem



NVIDIA i Microsoft. Model PaLM Google'a z 540 miliardami parametrów ilustruje skalę, do jakiej doszło przetrenowanie sieci na danych tekstowych, setkach miliardów słów pochodzących z literatury, Internetu czy źródeł specjalistycznych. Wielkość tych modeli określana jest liczbą parametrów zapisanych podczas treningu; zasób ten stanowi wiedzę parametryczną systemu. Mechanizm tokenizowania dzieli dane wejściowe na słowa lub sub-słowa, które następnie są mapowane na wektory w przestrzeni cech.

Analiza rozkładu tokenów umożliwia ustalenie kontekstu wypowiedzi na podobnej zasadzie jak człowiek czytający zdanie zaczyna zwracać większą uwagę słowom kluczowym. Transformery są wyjątkowo elastyczne, można je stosować zarówno w trybie enkodera (np. BERT, nastawiony na analizę treści), jak i dekodera (np. GPT, nastawiony na generowanie) czy jako pełny stos enkoder-dekoder (np. T5). W każdym przypadku fundamentem pozostaje równoległa obróbka sekwencji i globalny dostęp do kontekstu. Modele uczone w konfiguracji dekodera stosują maskowanie uwagi tak, aby każde przewidywanie kolejnego tokenu opierało się tylko na wcześniejszych elementach sekwencji, co odpowiada naturalnemu procesowi pisania tekstu lub tłumaczenia. Metody trenowania oparte na ogromnych korpusach prowadzą do zakodowania w modelu szerokiego spektrum wiedzy ogólnej, co pozwala mu odpowiadać sensownie nawet w tematach nieobjętych nauką dedykowaną konkretnemu zadaniu. Po etapie wstępnego treningu modele często poddaje się dostrajaniu z wykorzystaniem danych zadaniowych lub instrukcji. Może to być instrukcyjny tuning do generowania odpowiedzi dialogowych zgodnych ze stylem oczekiwanym przez użytkownika albo adaptacja do języka specjalistycznego poprzez dodanie dokumentacji technicznej domeny.

Ciekawym aspektem jest oddzielanie wiedzy parametrycznej od nieparametrycznej poprzez integrację z komponentami wyszukującymi treść w zewnętrznych bazach; takie rozwiązanie pozwala modelowi odwoływać się do aktualnej informacji bez ponownego trenowania. Mechanizm prompt engineering nabiera szczególnego znaczenia przy pracy z LLM, właściwe sformułowanie polecenia decyduje o jakości odpowiedzi. Prompt działa jako inicjalny kontekst przekazywany modelowi; jego konstrukcja wpływa na to, które fragmenty zapisanej wiedzy parametrycznej zostaną aktywowane podczas generacji treści. Przykładowo, do tego samego pytania można uzyskać odpowiedzi różniące się szczegółowością lub stylem w zależności od zawartej w promptcie specyfikacji formatu czy punktu widzenia. Wersje komercyjne transformatorów bywają własnościowe i niedostępne w kodzie źródłowym, ale rośnie również ekosystem projektów otwarto źródłowych oferujących pełne modele gotowe do pobrania przez platformy takie jak Hugging Face Hub czy Azure AI Studio. Takie rozwiązania pozwalają badaczom eksperymentować z architekturą oraz wdrażać ją lokalnie w aplikacjach – co ma znaczenie dla środowisk wymagających prywatności danych lub działania offline.



Transformery obsługujące kontekst długi ułatwiają też integrację procesów utrzymywania pamięci dialogowej. Modele mogą rozszerzać pole widzenia o kilkanaście tysięcy tokenów wejściowych tak, aby zachować sens interakcji wieloetapowej bez gubienia odniesień do wcześniejszej wymiany zdań. To przydatne np. w systemach agentowych wykonujących ciąg operacji bazującym na LLM, agent decyduje, kiedy pobrać dane z bazy wektorowej, a kiedy kontynuować generację odpowiedzi. Historia architektury transformera sięga propozycji Google'a z 2017 r., gdzie opisano ją jako „nowatorską sieć neuronową opartą na mechanizmie samouwagi” będącą usprawnieniem technologii Seq2Seq wykorzystywanej wcześniej m.in. w tłumaczeniach maszynowych. Od tego czasu architektura ta stała się standardem dla większości zaawansowanych modeli językowych i inspirowała kolejne iteracje takich narzędzi, od LaMDA przez PaLM aż po LLaMA Meta AI.

Integracja ChatGPT z wyszukiwarką Bing i pojawienie się Barda pokazują też komercyjny wymiar ekspansji transformerów: urządzenia codziennego użytku stają się interfejsami dla tych architektur. W porównaniu do klasycznych podejść sekwencyjnych opisanych wcześniej, transformery dzięki zdolności równoległego przetwarzania wszystkich tokenów zwiększają skalowalność procesu treningowego oraz umożliwiają budowę modeli liczących setki miliardów parametrów bez strat związanych z długością kontekstu. Otwiera to przestrzeń dla nowych zastosowań inżynierii lingwistycznej: od precyzyjnej analizy tekstu po generację długich dokumentów technicznych zachowujących spójność koncepcyjną. Jednocześnie wyzwaniem pozostaje kontrola jakości i faktograficzna poprawność takich wygenerowanych treści, stąd rosnące znaczenie metod oceny wydajności modelu według benchmarków tematycznych takich jak MMLU czy testy akademickie prezentujące dokładność odpowiedzi LLM.

6 Zastosowania inżynierii lingwistycznej

6.1 Systemy tłumaczenia maszynowego

Systemy tłumaczenia maszynowego odgrywają istotną rolę w inżynierii lingwistycznej, przekształcając wypowiedzi pomiędzy językami w sposób automatyczny. Współczesne podejścia korzystają z dorobku zarówno klasycznych metod regułowych, jak i nowoczesnych architektur głębokich sieci neuronowych. Tradycyjny model oparty na transferze syntaktycznym wymaga wcześniejszej analizy struktury zdania źródłowego przy użyciu parserów zgodnych z gramatyką danego języka, a następnie przekształcenia tej struktury do formy zgodnej z gramatyką języka docelowego. Ten schemat pozwala zachować poprawność składniową, jednak w przypadku idiomów czy konstrukcji metaforycznych potrzebne są dodatkowe mechanizmy interpretacyjne, aby przenieść znaczenie wiernie.



W ramach rozwoju technologii tłumaczeń szczególną uwagę zwrócono na uczenie maszynowe. Modele n-gramowe, stosowane we wczesnych systemach statystycznych (SMT), oceniano pod kątem prawdopodobieństwa sekwencji słów w języku docelowym, co skutecznie poprawiało płynność przekładu. Ograniczeniem było jednak nieuwzględnianie dalekiego kontekstu i skomplikowanych relacji semantycznych. Postęp nastąpił dzięki zastosowaniu sieci rekurencyjnych i ich wariantów LSTM czy GRU. W konfiguracji enkoder–dekoder moduł kodujący sekwencyjnie odczytywał zdanie źródłowe, budując jego zwarty wektor reprezentacyjny, a dekodery generował tłumaczenie słowo po słowie. Mechanizm ten pozwalał modelować zależności między odległymi tokenami w długim zdaniu dzięki bramkom kontrolującym przepływ informacji.

Znaczący przełom przyniosło wprowadzenie architektury transformera do tłumaczenia maszynowego. Równoległe przetwarzanie sekwencji oraz mechanizm wielogłowej uwagi umożliwiły uchwycenie skomplikowanych relacji syntaktycznych i semantycznych między słowami niezależnie od ich pozycji. Model może „patrzeć” na całe zdanie naraz, identyfikując powiązania nawet między bardzo oddalonymi elementami. Wersje enkoder–dekoder wykorzystują uwagę tak, aby dekodery mogły koncentrować się na fragmentach wejścia najbardziej istotnych dla bieżącego kroku generacji. Taki mechanizm redukuje błędy wynikające z utraty kontekstu, typowe dla sekwencyjnego przetwarzania RNN. Wstępnie wytrenowane modele neuronowe, takie jak GPT czy specjalistyczne warianty T5 i MarianMT, uczone są na gigantycznych korpusach wielojęzycznych obejmujących miliardy zdań. Dzięki temu mogą działać w trybie zero-shot lub few-shot dla par języków nieobjętych bezpośrednim treningiem, poprzez generowanie tłumaczeń bazujących na abstrakcyjnych wzorcach językowych wyuczonych wcześniej.

Proces dostrajania na dedykowanych korpusach dziedzinowych pozwala uzyskać wysoką zgodność terminologiczną w tłumaczeniach specjalistycznych (np. medycyna czy prawo). Przy tym rosnącą rolę odgrywa selekcja danych treningowych pod kątem jakości i równowagi językowej; niewłaściwy dobór korpusu może prowadzić do wprowadzenia biasu lub akwizycji błędnej składni. Interesującym nurtem jest wykorzystanie podejścia interlingua, gdzie tekst źródłowy transformowany jest najpierw do reprezentacji znaczeniowej niezależnej od języka (np. struktury logicznej lub ram semantycznych), a dopiero potem generowany jest tekst docelowy zgodny ze składnią odbiorcy. Takie rozwiązanie minimalizuje ryzyko utraty kluczowych relacji podczas przełożenia idiomów czy nietypowych struktur składniowych.

Praktyczne systemy coraz częściej integrują moduły dodatkowe: pamięć konwersacyjną pozwalającą śledzić historię tłumaczonych dialogów oraz komponenty wyszukiwania nieparametrycznej wiedzy dla terminologii rzadkiej. Daje to możliwość dynamicznego uzupełniania braków informacyjnych podczas tłumaczenia – np. przez pobranie definicji technicznego terminu spoza parametrów



modelu. Podobne operacje stają się łatwe dzięki platformom dystrybucji modeli; dostęp do otwarto źródłowych rozwiązań takich jak MarianMT czy M2M100 ułatwia badaczom eksperymentowanie z parametrami architektury i źródłem danych. Rozwiązania komercyjne często łączą tłumaczenie maszynowe z innymi funkcjami NLP, przykład stanowią aplikacje umożliwiające rozmowę w czasie rzeczywistym między osobami mówiącymi różnymi językami poprzez automatyczne rozpoznawanie mowy i natychmiastową translację. Wysoka jakość komunikacji (np. 98% trafności) bywa osiągana dzięki integracji dedykowanych modeli akustycznych i językowych z warstwą kontekstowej analizy treści. To pokazuje przesunięcie akcentu od czysto tekstowego przekładu ku pełnej multimodalności komunikacji.

Istotne pozostają zagadnienia ewentualnych błędów, system może generować treść poprawną syntaktycznie, ale niezgodną faktograficznie bądź kulturowo. Dlatego implementuje się filtry jakości lub mechanizmy oceny wyników według wyników benchmarków i miar BLEU czy COMET odpowiadających jakości odwzorowania znaczeń względem referencji ludzkiej. Własne procedury adaptacyjne pozwalają uwzględniać lokalne normy językowe oraz warianty regionalne, szczególnie tam, gdzie standaryzacja jest mniej istotna niż naturalność wypowiedzi dla rodzimych użytkowników danego języka.

Analiza trendów wskazuje też na rosnącą świadomość etyczną w projektowaniu systemów tego typu: kontrola nad treściami generowanymi przez tłumacze obejmuje unikanie uprzedzeń oraz szkodliwych sformułowań obecnych czasem w surowych danych treningowych. Transparentny zapis procesu translacji – np. poprzez wyjaśnialne modele integrujące formalną analizę składni z powiązaniem wyuczonymi neuronowo – daje zarówno użytkownikowi końcowemu, jak i badaczowi możliwość audytu zachowania systemu. W rezultacie rozwój systemów tłumaczenia maszynowego przesuwają się od restrykcyjnych reguł transferowych ku skalowalnym architekturom głębokim zdolnym obsługiwać setki par językowych jednocześnie. Jednak niezależnie od postępu technologicznego rdzeniem efektywnego przekładu pozostaje precyzyjne odwzorowanie intencji nadawcy, zadanie wymagające ścisłego połączenia wiedzy lingwistycznej z algorytmiczną adaptacją modelu do zmieniających się warunków komunikacyjnych.

6.2 Analiza sentymentu

Analiza sentymentu jest procesem klasyfikowania treści tekstowej pod kątem emocjonalnego wydźwięku, zwykle w kategoriach pozytywny, neutralny lub negatywny. W praktyce może ona przybierać zarówno tryb binarny (np. 0 dla negatywnego i 1 dla pozytywnego nastawienia), jak i bardziej szczegółowy, w którym ocena jest wyrażana jako wartość ciągła pomiędzy 0.0 a 1.0, oznaczająca prawdopodobieństwo pozytywności wypowiedzi. Przykładowo recenzja o treści



„świetny produkt w przyzwoitej cenie” uzyska wysoką punktację, podczas gdy „przepełniony gadżet, który prawie nie działa” zostanie sklasyfikowana blisko dolnej granicy skali.

Proces zaczyna się od przygotowania danych wejściowych, co obejmuje oczyszczenie tekstu z elementów mogących utrudniać analizę. Typowe kroki to sprowadzenie liter do formy małej (tak aby „Super” i „super” były traktowane jednakowo), usunięcie znaków interpunkcyjnych oraz ewentualne wyeliminowanie słów często występujących o niewielkiej wartości semantycznej. Następnie tekst musi zostać przekształcony w reprezentację numeryczną, najczęściej przez wektoryzację opartą na liczbie wystąpień słów lub ich wartościach TF-IDF. W prostym wariantcie model zlicza w każdym dokumencie, ile razy pojawia się dany token, tworząc macierz cech, gdzie każda kolumna odpowiada konkretnemu słowu lub n-gramowi.

Dla uczenia modeli niezbędny jest zbiór danych opatrzony etykietami sentymentu. Publiczny serwis recenzji filmowych z IMDb zawiera po 25 000 przykładów opinii pozytywnych i negatywnych, co stanowi dobry materiał treningowy do budowy binarnego klasyfikatora sentymentu. W przypadku implementacji klasycznej regresji logistycznej można użyć metody `predict_proba()`, aby otrzymać probabilistyczną ocenę danego tekstu, informującą o pewności modelu względem przypisania go do klasy pozytywnej czy negatywnej. Choć mechanizmy analizy sentymentu mogą być projektowane algorytmicznie (reguły językowe oraz słowniki nacechowania), skuteczniejsze okazują się podejścia oparte na uczeniu maszynowym. Pozwalają one uchwycić subtelne zmiany nacechowania wynikające z kontekstu; przykładowo fraza „nie taki zły” jest formalnie zaprzeczeniem „zły”, ale semantycznie wyraża raczej umiarkowaną aprobatę niż krytykę.

Modele statystyczne uczone na dużych korpusach rozpoznają tego rodzaju konstrukcje poprzez współwystępowanie tokenów w różnorodnych kontekstach. W zastosowaniach komercyjnych analiza sentymentu wspiera decyzje biznesowe: firmy badają percepcję marki w mediach społecznościowych, e-commerce automatycznie sortuje produkty według opinii klientów, a działy obsługi monitorują reakcje użytkowników na kampanie marketingowe. Integracja z systemami klasyfikacji pozwala połączyć ocenę emocjonalną z analizą tematyczną wypowiedzi, uzyskując pełniejszy obraz sytuacji. Z perspektywy technologicznej istnieje również miejsce na integrację wiedzy parametrycznej zapisanej w wagach modelu podczas treningu z wiedzą pozyskiwaną dynamicznie ze źródeł nieparametrycznych. Taki mechanizm umożliwia wzbogacenie procesu analizy sentymentu o aktualne informacje kontekstowe, np. nastroje wobec nowych wydarzeń lub produktów, bez konieczności ponownego trenowania całego modelu.



W praktycznych implementacjach narzędzi takich jak łańcuch pobierania konwersacyjnego można pobrać najnowsze posty lub recenzje i wykorzystać je jako dodatkowy kontekst interpretacyjny. Warto wspomnieć o ograniczeniach: modele uczone na jednym typie dokumentów mogą mieć problem z generalizacją do innych domen czy stylów języka. Recenzje filmowe różnią się od komentarzy politycznych nie tylko tematyką, ale też sposobem ekspresji emocji; slang, ironia czy frazy memiczne są trudne do uchwycenia metodami opartymi na prostym liczeniu słów. Stąd rośnie znaczenie transformatorowych architektur uczenia głębokiego zdolnych analizować relacje semantyczne między tokenami na różnych dystansach w tekście.

Łączenie analizy sentymentu z rozpoznawaniem mowy otwiera drogę do przetwarzania wypowiedzi ustnych, od nagrań rozmów call center po materiały audiowizualne publikowane online. Po transkrypcji wypowiedź podlega takim samym procedurom jak tekst pisany, choć należy uwzględnić specyfikę języka mówionego: większą liczbę przerw, powtórzeń czy niegramatycznych konstrukcji.

Opracowanie wskaźników jakości dla modeli sentymentu obejmuje miary takie jak dokładność klasyfikacji, precyzja i przypomnienia dla poszczególnych klas oraz analiza krzywej ROC ilustrującej kompromis między czułością a swoistością predykcji. W środowiskach produkcyjnych monitorowanie tych wskaźników pozwala ocenić stabilność działania systemu przy zmianach strumienia danych wejściowych. Rozszerzone wersje systemów analizy sentymentu próbują także oceniać intensywność odczuć (np. lekko pozytywny vs. bardzo pozytywny) lub identyfikować konkretne emocje: radość, smutek, gniew itd., co wymaga bardziej szczegółowych etykiet oraz odpowiednio przygotowanego korpusu treningowego. Taki kierunek badań wskazuje potencjał adaptacji technologii do wielowymiarowej oceny emocjonalnej tekstów, korzystając zarówno z możliwości statystycznych modeli binarnych, jak i elastyczności architektur głębokiego uczenia zdolnych interpretować kontekst wielopłaszczyznowo.

6.3 Wyszukiwanie informacji

Indeksowanie semantyczne

Indeksowanie semantyczne odnosi się do technik organizacji i wyszukiwania danych tekstowych w taki sposób, aby uwzględniać znaczenie słów i zależności między nimi, a nie jedynie ich literalną postać. W odróżnieniu od klasycznego podejścia opierającego się na wyszukiwaniu fraz wprost w treści dokumentu, tutaj dąży się do odwzorowania relacji pojęciowych tak, aby system rozumiał, że zapytanie „motyl” może być powiązane z dokumentami zawierającymi termin „owad”, a nawet bardziej szczegółowe nazwy gatunków. Implementacja takich funkcji wymaga reprezentacji



języka w formie umożliwiającej obliczeniowe porównywanie znaczeń, najczęściej poprzez osadzanie słów lub zdań w przestrzeni wektorowej.

Podstawę indeksowania semantycznego może stanowić model typu osadzanie słów, trenowany na dużych korpusach tekstowych tak, aby umieścić wyrazy o podobnym kontekście blisko siebie w przestrzeni cech. Takie rozmieszczenie pozwala realizować wyszukiwanie „na znaczenie”, system szuka wektorów maksymalnie zbliżonych do wektora zapytania użytkownika, co przekłada się na odnajdywanie dokumentów związanych tematycznie, nawet jeśli słowa nie pokrywają się dokładnie. Zaletą tej metody jest uchwycenie synonimii i innych relacji leksykalnych trudnych do odwzorowania przez tradycyjne mechanizmy indeksowania. Indeks semantyczny często korzysta z zasobów leksykalnych takich jak WordNet, które grupują słowa w synsety oraz zawierają informacje o hiponimach, meronimach czy antonimach. System może rozszerzyć zapytanie użytkownika o te powiązane pojęcia, np. wyszukując „samochód” znajdzie również materiały o „pojazdach”, „autach” czy konkretnych markach, dzięki czemu zwiększa trafność wyników. Proces takiego wzbogacania zapytań może działać dynamicznie; po wejściu hasła przez użytkownika następuje analiza semantyczna ciągu tokenów z użyciem sieci semantycznej i generowanie listy alternatywnych lub uzupełniających terminów.

W implementacjach przemysłowych duże modele językowe (LLM) coraz częściej są wykorzystywane jako komponent kreatorów indeksów semantycznych. Modele takie jak BERT czy GPT potrafią wytwarzać gęste osadzenia dla całych zdań lub akapitów, co skraca dystans pomiędzy wyszukiwaniem semantycznym a pełno tekstowym. Takie osadzenia przechowywane są w specjalistycznych bazach wektorowych (np. FAISS), które umożliwiają szybkie porównywanie milionów rekordów pod kątem podobieństwa kosinusowego czy euklidesowego. Istotne znaczenie ma wybór metody tworzenia tego rodzaju osadzeń. Technikę można dostosować poprzez tzw. fine-tuning na korpusach domenowych tak, aby model interpretował zależności zgodnie z potrzebami danej branży. Przykładowo baza wiedzy medycznej będzie wymagała wyróżnienia relacji pomiędzy nazwami chorób, objawami i lekami; model przeszkolony na takim korpusie będzie lepiej kojarzył wpisane przez lekarza symptomy z odpowiednimi kartotekami przypadków.

Indeksowanie semantyczne wymaga również integracji wiedzy parametrycznej zakodowanej w modelu podczas treningu z wiedzą nieparametryczną pobieraną dynamicznie z zewnętrznych źródeł, chociażby aktualnych baz danych czy dokumentacji technicznej. Dzięki temu system jest w stanie odpowiadać adekwatnie nawet na pytania dotyczące nowych pojęć lub danych zmieniających się w czasie, bez konieczności ponownego uczenia całego modelu od podstaw. Warto zwrócić uwagę na etap tokenizacji i normalizacji danych przed ich osadzeniem. Poprawna segmentacja słów, usunięcie zbędnych znaków oraz rozpoznanie form podstawowych (lematyzacja) zapewnia spójność reprezentacji i minimalizuje ryzyko



sztucznego oddalania wektorów z powodu różnic fleksyjnych czy wariantów zapisowych. Istotne jest także radzenie sobie z wieloznacznością; ten sam token może mieć różne znaczenia w zależności od kontekstu i dlatego system powinien generować osadzenia kontekstowe zamiast statycznych, rozwiązanie to stało się standardem dzięki modelom dwukierunkowym jak BERT. Optymalizacja działania indeksu obejmuje strategię przechowywania i aktualizacji osadzeń.

W środowiskach wysokiej dynamiki danych stosuje się reindeksację inkrementalną, nowe dokumenty otrzymują swoje wektory natychmiast i są dodawane do bazy bez konieczności przebudowy całej struktury. Równocześnie bardzo ważna pozostaje filtracja jakościowa źródeł: dane skażone błędami lub stroniczne mogą powodować fałszywe dopasowania, które obniżą wartość praktyczną wyszukiwarki.

W kontekście integracji z systemami wyszukiwawczymi istotna jest możliwość łączenia wyników indeksowania semantycznego z klasycznym rankingiem opartym o słowa kluczowe. Takie hybrydowe podejście pozwala zachować precyzję dla zapytań dosłownych oraz zwiększyć pokrycie wynikowe dla zapytań bardziej ogólnych czy metaforycznych. Mechanizm ten bywa realizowany poprzez agregację wyników obu metod wraz z nadaniem im odpowiednich wag decydujących o kolejności prezentowania dokumentów użytkownikowi. Pomiar jakości indeksowania semantycznego opiera się najczęściej na analizie współczynników trafności i kompletności. Dla aplikacji krytycznych wdraża się również testy A/B sprawdzające wpływ zmian modelu osadzającego i algorytmu porównania na doświadczenie użytkownika końcowego oraz szybkość działania systemu. Monitorowanie tych parametrów jest zasadne zarówno przy dużych wdrożeniach chmurowych, jak i lokalnych implementacjach.

Zastosowanie indeksowania semantycznego obejmuje m.in. wyszukiwarki akademickie agregujące publikacje według obszaru badań niezależnie od użytej terminologii w tytułach oraz repozytoria kodu źródłowego pozwalające znaleźć implementacje wzorcowo opisujące rozwiązanie danego problemu mimo innego nazewnictwa funkcji czy klas. W sektorze biznesowym techniki te wspierają obsługę klienta poprzez szybkie znajdowanie dokumentacji pasującej do zgłaszanego problemu albo rekomendowanie produktów odpowiadających opisanym potrzebom, nawet jeśli klient nie używa terminologii katalogowej firmy. Rozwój architektury agentowych wzmacnia rolę indeksowania semantycznego jako modułu dostarczającego kontekst LLM przed wygenerowaniem odpowiedzi. Agent pobiera intencję użytkownika, identyfikuje potrzebne dane w bazie wektorowej opartej na indeksie semantycznym i przekazuje je dalej do modelu językowego jako ustrukturyzowany kontekst, co poprawia trafność oraz spójność informacji w tworzonej komunikacji. Taka synergia mechanizmów pokazuje, jak łączenie metod NLP opisanych wcześniej z zaawansowaną organizacją danych pod kątem



znaczenia stanowi dziś kluczową strategię budowy nowoczesnych systemów wyszukiwania informacji.

Systemy pytanie-odpowiedź

Systemy pytanie–odpowiedź stanowią szczególną kategorię narzędzi wyszukiwania informacji, których celem jest udzielenie użytkownikowi odpowiedzi w formie bezpośredniej, możliwie zwięzłej i adekwatnej do zapytania, zamiast zwracania pełnych dokumentów. Działają one na zasadzie przetwarzania języka naturalnego, interpretując ciąg znaków jako pytanie, rozpoznając jego typ oraz kluczowe elementy semantyczne, a następnie wyszukując w zbiorze danych (tekstowym lub mieszanym) fragment mogący stanowić odpowiedź. Proces ich pracy łączy analizę składniową i semantyczną zapytania z mechanizmami przeszukiwania treści według kontekstu znaczeniowego. W odróżnieniu od klasycznych wyszukiwarek tekstowych, systemy te muszą wykrywać, czy pytanie dotyczy faktu (np. „Kto napisał Hamleta?”), opinii („Czy technologia X jest skuteczna?”), czy też wymaga obliczenia lub syntezy informacji z wielu źródeł. Architektura typowego rozwiązania obejmuje moduł interpretacji pytania, ekstrakcję encji, relacji oraz określenie kategorii odpowiadającej spodziewanej odpowiedzi.

Następnie stosowany jest komponent wyszukiwania dokumentów lub fragmentów najbardziej podobnych semantycznie do zapytania. Tutaj używa się osadzeń wektorowych generowanych przez modele językowe, co pozwala mierzyć podobieństwo znaczeniowe nawet wtedy, gdy forma zapytania nie pokrywa się literalnie z treścią dokumentu. Przykładowo „Autor Makbeta” i „Kto stworzył dramat Makbet?” otrzymają bardzo bliskie wektory reprezentacji i będą prowadzić do tych samych danych źródłowych. Ważnym elementem systemów pytanie–odpowiedź jest proces selekcji odpowiedzi w granicach jednego dokumentu. Po wybraniu kandydatów następuje etap ekstrakcji, np. model BERT przetrenowany metodą prognozy rozpiętości przewiduje początkowy i końcowy indeks frazy w tekście stanowiącej odpowiedź na zadane pytanie.

Rozwiązanie takie bywa rozszerzane o filtrację wyników pod kątem wiarygodności: każda odpowiedź otrzymuje ocenę pewności, na podstawie której system może zdecydować o przedstawieniu wyniku lub przekazaniu użytkownikowi komunikatu o braku dostatecznej informacji. Integracja z wiedzą nieparametryczną ma tu istotne znaczenie. Modele bazujące wyłącznie na wiedzy zapisanej parametrycznie mogą mieć problem z aktualnymi danymi (np. datami wydarzeń). Dlatego coraz częściej stosuje się mechanizmy pobierające kontekst ze źródeł zewnętrznych, baz wektorowych, repozytoriów dokumentów czy API serwisów informacyjnych, a następnie włączające go do procesu generowania odpowiedzi. Tak działają architektury typu łańcuch pobierania konwersacyjnego, które umożliwiają utrzymanie spójności rozmowy przy wielokrotnym zadawaniu pytań powiązanych tematycznie.



Ocena jakości systemów pytanie–odpowieź wymaga testów bardziej wyspecjalizowanych niż standardowe miary wyszukiwania informacji. Stosowane są benchmarki takie jak TruthfulQA, zaprojektowane tak, aby sprawdzać rzetelność odpowiedzi poprzez pytania wymagające wiedzy ogólnej, zdrowego rozsądku oraz interpretacji lingwistycznej. Wyniki takich testów pokazują podatność modeli na „halucynacje”, generowanie poprawnych językowo, lecz fałszywych merytorycznie treści. W celu ograniczenia tego efektu projektuje się procedury filtrowania lub wskazane podpowiedzią strategię formułowania pytań. Konstrukcja promptu może nakazać modelowi najpierw przedstawić krok po kroku tok rozumowania lub podać źródło informacji. Systemy te dzielą się także pod względem dynamiki interakcji: jednowariantowe Q&A odpowiadają na pojedyncze pytanie, podczas gdy systemy konwersacyjne utrzymują historię dialogu i uwzględniają poprzednie wypowiedzi użytkownika przy generowaniu kolejnej odpowiedzi. W takim przypadku pamięć dialogowa przeplata warstwę parametryczną modelu z odniesieniami pobranymi dynamicznie, co pozwala zachować konsekwentny kontekst nawet po dwóch czy trzech wymianach pytań.

W kontekście zastosowań biznesowych systemy pytanie–odpowieź są integrowane z platformami obsługi klienta i chatbotami. Dzięki zdolności interpretacji zapytań otwartych mogą odpowiadać nie tylko na proste kwestie faktograficzne („Jaki jest status mojego zamówienia?”), ale także udzielać porad („Jak mogę rozwiązać problem X?”), korzystając przy tym ze strukturalnych baz wiedzy firmy. Z kolei w sektorach takich jak medycyna czy prawo integracja analiz Q&A z dedykowanymi korpusami umożliwia szybkie pozyskiwanie konkretnych zapisów akt prawnych lub wyników badań naukowych zgodnych z kryteriami użytkownika.

Istotnym trendem jest też łączenie tego typu systemów z agentami wykonującymi zadania pomocnicze, np. wyszukiwanie plików w repozytorium lub wykonywanie obliczeń, i zwracającymi wynik w formie uzyskanej odpowiedzi tekstowej. W tej architekturze sam proces odpowiadania staje się jednym z modułów większego ekosystemu działań. Powiązane rozwiązania często służą jako podsystem dostarczający zoptymalizowane dane wejściowe, czyli fragmenty najbardziej istotne semantycznie względem pytania. Praktyka wdrożeniowa pokazuje również konieczność dopasowania modelu do specyfiki języka wejściowego: dla języków fleksyjnych integruje się analizę morfologiczną tak, aby lematyzacja wspierała proces rozpoznawania intencji i struktury zapytania. Ułatwia to prawidłowe kojarzenie fraz niezależnie od ich form odmiany.

Podobnie ważne okazuje się rozpoznawanie konstrukcji idiomatycznych; bez tego model może udzielić literalnej odpowiedzi pozbawionej sensu w kontekście kulturowym. Podsumowując, systemy pytanie–odpowieź łączą zaawansowane metody lingwistyczne z algorytmiką wyszukiwania informacji tak, aby dostarczyć użytkownikowi precyzyjną odpowiedź w minimalnym czasie. Postęp architektur LLM



oraz integracja wiedzy nieparametrycznej poszerza zakres ich zastosowań, od prostych asystentów FAQ po wielomodułowe agenty zdolne interaktywnie prowadzić rozmowę i realizować działania na rzecz użytkownika przy rygorystycznej kontroli jakości merytorycznej wygenerowanych treści.

6.4 Rozpoznawanie mowy

Przekształcanie mowy na tekst

Przekształcanie mowy na tekst jest kluczowym etapem w wielu aplikacjach inżynierii lingwistycznej, gdzie celem jest odwzorowanie wypowiedzi ustnej w postaci ciągu znaków możliwych do dalszej obróbki komputerowej. Proces ten obejmuje kilka komponentów technologicznych, które współdziałają w celu analizy sygnału akustycznego i interpretacji go jako reprezentacji lingwistycznej. W punkcie startowym znajduje się akwizycja danych, rejestracja sygnału mowy za pomocą mikrofonu w odpowiednich warunkach akustycznych. Już na tym etapie kontrola jakości nagrania decyduje o zwiększeniu skuteczności kolejnych etapów, szumy tła, pogłos czy zakłócenia mogą wpływać na zdolność systemu do poprawnego odczytu fonemów.

W implementacjach praktycznych stosuje się najpierw przetwarzanie wstępne sygnału, obejmujące normalizację amplitudy oraz filtrację pasmową eliminującą częstotliwości spoza zakresu ludzkiej mowy. Kolejnym krokiem jest segmentacja strumienia audio na krótkie ramki (typowo 20 ms z nakładaniem 10 ms), co umożliwia analizę zmian spektrum w czasie. Z każdej ramki ekstraktowane są cechy akustyczne, najczęściej w postaci współczynników MFCC (Mel-Frequency Cepstral Coefficients), które odwzorowują percepcję częstotliwości przez człowieka. Cechy te stanowią wejście dla modułu rozpoznawania wzorców.

Klasyczne systemy korzystały z modeli statystycznych takich jak ukryte modele Markowa (HMM), które opisywały prawdopodobieństwa przejścia między stanami reprezentującymi kolejne fonemy oraz emisji obserwowanych cech akustycznych z tych stanów. Integracja HMM z modelami n-gramowymi języka pozwalała łączyć interpretację fonetyczną z przewidywaniem najbardziej prawdopodobnej sekwencji słów. Dla języków o bogatej morfologii komponenty te wymagały dopasowania słowników wymowy tak, aby odzwierciedlały reguły fleksji i odmiany. Rozwój uczenia głębokiego zmienił architekturę ASR (Automatic Speech Recognition).

Obecnie stosuje się sieci neuronowe wielowarstwowe, zdolne do bezpośredniego mapowania cech akustycznych na etykiety fonemowe lub całe tokeny tekstowe. Sieci rekurencyjne LSTM czy GRU pozwalają uchwycić zależności czasowe pomiędzy kolejnymi ramkami dźwięku i radzą sobie lepiej niż HMM przy rozpoznawaniu kontekstu długo–terminowego. Ich warianty dwukierunkowe analizują sygnał



zarówno od początku, jak i od końca sekwencji, co redukuje błędy wynikające z niepełnego kontekstu.

Istotnym elementem nowoczesnych systemów staje się zastosowanie architektury transformatorowej także dla rozpoznawania mowy. Mechanizm uwagi jest tu wykorzystywany do skupienia się na kluczowych fragmentach sygnału w odniesieniu do całej sekwencji. Modele end-to-end takie jak Speech-to-Text BERT czy Whisper uczone są na multimodalnych korpusach audio–tekst i potrafią generować transkrypcję bez pośredniej reprezentacji fonemowej. Daje to większą odporność na różnorodność akcentów i szybkości mówienia.

Ważnym aspektem jest integracja wiedzy parametrycznej zapisanej podczas trenowania modelu z wiedzą nieparametryczną dostarczaną dynamicznie. Przykładowo komponent typu łańcuch pobierania konwersacyjnego może pobrać aktualną terminologię branżową z bazy wektorowej i uzupełnić nią proces transkrypcji w czasie rzeczywistym, tak aby poprawnie odwzorować nazwy własne lub skróty specyficzne dla danego środowiska. Rozwiązanie takie pozwala obniżyć ryzyko błędnego zapisu rzadkich słów. Proces zamiany mowy na tekst kończy się fazą dekodowania, tutaj stosuje się algorytmy przeszukiwania przestrzeni hipotez, takie jak wyszukiwanie wiązki, aby wybrać najbardziej prawdopodobną sekwencję słów po uwzględnieniu zarówno wyniku sieci akustycznej, jak i modelu językowego. Balans pomiędzy tymi dwoma źródłami informacji wpływa bezpośrednio na trafność wynikowej transkrypcji, nadmierna dominacja modelu językowego może powodować „podciąganie” wypowiedzi pod częściej spotykane frazy kosztem rzadkich konstrukcji.

W zastosowaniach przemysłowych rośnie znaczenie automatycznej personalizacji modeli ASR poprzez adaptację parametrów do głosu konkretnego użytkownika. Techniki takie jak speaker adaptation lub fine-tuning na próbkach mowy pozwalają poprawić dokładność rozpoznawania nawet przy obecności silnego akcentu lub nietypowej wymowy. Podobny efekt osiąga się przez dodanie warstw adaptacyjnych minimalizujących konieczność pełnego przetrenowania sieci. Zagadnienia jakości obejmują też obsługę aspektów paralingwistycznych: pauzy, intonacja czy natężenie głosu mogą wносить dodatkowe informacje interpretacyjne wymagające translacji do formy pisanej, np. poprzez uwzględnienie znaków interpunkcyjnych lub opisów emocji w transkrypcji. Integracja systemów przekształcania mowy na tekst z platformami wielojęzycznymi umożliwia rozszerzenie funkcjonalności o tłumaczenie maszynowe w czasie rzeczywistym.

Po transkrypcji następuje segmentacja zdań i ich przekład zgodny ze składnią języka docelowego, co pozwala prowadzić rozmowę między użytkownikami różnych języków niemal bez opóźnień. W kontekście komercyjnych zastosowań przykłady takie jak integracje IBM Watson pokazują możliwość łączenia rozpoznawania mowy



z innymi usługami AI – od analizy sentymentu po wyszukiwanie informacji powiązanych z treścią wypowiedzi. Postęp technologii wskazuje także kierunek fuzji rozpoznawania mowy z multimodalnością, GenAI jest zdolna nie tylko konwertować mowę na tekst, ale też synchronizować ją z treścią wizualną lub generować audiodeskrypcję dla materiałów filmowych przy zachowaniu naturalności głosu. Takie podejście otwiera możliwość stworzenia systemów dostępnych dla osób o różnych potrzebach komunikacyjnych, integrując funkcje zamiany mowy na tekst z automatycznym tworzeniem napisów czy streszczeń treści nagrań.

Identyfikacja mówcy

Identyfikacja mówcy jest zadaniem polegającym na przypisaniu próbki sygnału mowy do konkretnej osoby lub grupy osób na podstawie cech charakterystycznych jej głosu. Może być realizowana w trybie weryfikacji (sprawdzenie, czy dana próbka należy do zgłoszonego mówcy) lub identyfikacji otwartej (wybór najbardziej prawdopodobnego kandydata spośród znanego zbioru głosów). W odróżnieniu od procesu przekształcania mowy na tekst, tutaj kluczowe jest uchwycenie indywidualnych właściwości akustycznych i artykulacyjnych, a nie semantycznej treści wypowiedzi.

Pierwszym etapem typowego systemu jest akwizycja sygnału mowy oraz jego wstępne przetwarzanie: normalizacja amplitudy, filtracja szumów tła i segmentacja na ramki czasowe o stałej długości (20 ms–30 ms). Następnie wykonywana jest ekstrakcja cech opisujących głos. Najczęściej stosowane są współczynniki oparte na skali Mel (MFCC), które odwzorowują subiektywną percepcję częstotliwości przez ucho ludzkie. Oblicza się również parametry związane z barwą głosu, takie jak współczynniki LPC (Linear Predictive Coding) czy widmowe momenty statystyczne. W przypadku bardziej zaawansowanych systemów cechy te są rozszerzane o wektory x-vectors lub i-vectors – zwarte reprezentacje wysokopoziomowych właściwości sygnału akustycznego uzyskiwane poprzez modelowanie gęstości gaussowskich (GMM) połączonych z analizą głębokich warstw sieci neuronowych.

W klasycznych podejściach rozpoznawanie tożsamości mówcy opierało się na porównywaniu wektorów cech między próbką testową a próbkami wzorcowymi zapisanymi w bazie. Miara podobieństwa mogła być obliczana jako odległość euklidesowa, kosinusowa bądź wynik dopasowania przez modele probabilistyczne, takie jak GMM-UBM (Universal Background Model). Modele tego typu oceniają prawdopodobieństwo wygenerowania danych cechowych przez model odniesienia dla danego mówcy względem modelu tła reprezentującego ogólną populację. Rozwój uczenia głębokiego umożliwił zastąpienie ręcznie projektowanych wskaźników pełnymi reprezentacjami uczonymi end-to-end.

Sieci konwolucyjne przetwarzające spektrogramy próbki mowy traktują je jak obrazy – filtrując pasma i wydobywając lokalne wzorce energetyczne – natomiast sieci



rekurencyjne LSTM albo GRU pozwalają zachować kolejność temporalną sygnału. Mechanizm uwagi, szeroko stosowany w transformatorkach, zaczął być adaptowany także do identyfikacji mówcy, umożliwiając podkreślenie fragmentów nagrania niosących najwięcej informacji biometrycznej. Istotnym zagadnieniem jest odporność systemu na zmienność warunków nagrania: różne mikrofony, odległości od źródła dźwięku czy obecność pogłosu mogą istotnie wpłynąć na wartości wyekstrahowanych cech.

Modele uczone na zróżnicowanych korpusach – obejmujących dane zbierane w warunkach studyjnych i terenowych – lepiej generalizują do nowych scenariuszy. W praktyce oznacza to konieczność stosowania procedur augmentacji danych akustycznych takich jak dodawanie szumu, modulacja wysokości tonu czy zmiana prędkości wypowiedzi w celu zwiększenia wariacji próbek treningowych. Integracja wiedzy parametrycznej zapisanej w wagach modelu podczas trenowania z wiedzą nieparametryczną pozyskiwaną dynamicznie pozwala poprawić wyniki identyfikacji w środowiskach zmiennych lub wymagających uwzględnienia dodatkowych metadanych. Przykładem może być pobranie profilu głosowego użytkownika z zabezpieczonej bazy przed sesją rozpoznania i dopasowanie modelu do tego profilu „w locie” bez konieczności ponownego uczenia wszystkich parametrów sieci.

Systemy komercyjne często łączą identyfikację mówcy z innymi funkcjami bezpieczeństwa – autoryzacją dostępu czy podpisem biometrycznym dla dokumentów elektronicznych. Połączenie analizy biometrów głosu z autentykacją wieloskładnikową (np. hasło + zapis głosu) zwiększa poziom ochrony zasobów. Jednocześnie pojawia się problem podatności takich rozwiązań na ataki podszywania się przy użyciu syntezy głosu (voice spoofing). W odpowiedzi badacze opracowują techniki detekcji sygnałów syntetycznych bazujące na subtelnych artefaktach widmowych generowanych przez obecne systemy TTS – zarówno w fazie analiz parametrów częstotliwości, jak i poprzez klasyfikatory uczone rozróżniać nagrania naturalne i sztucznie wygenerowane.

Kolejnym wyzwaniem pozostaje wydajność działania systemu przy rozpoznawaniu dużej liczby użytkowników lub pracy w czasie rzeczywistym. Zastosowanie wyspecjalizowanych struktur danych dla przechowywania reprezentacji wektorowych (np. bazy wektorowe FAISS) umożliwia szybkie porównania cech między nową próbką a bazą wzorcową. Wersje wdrożeniowe stosują często algorytmy rezydujące bezpośrednio na urządzeniach brzegowych, minimalizując potrzebę przesyłania surowych nagrań do chmury, co sprzyja zachowaniu prywatności i skróceniu opóźnień inferencji. Ocena efektywności systemu prowadzona jest przy użyciu wskaźników takich jak EER (Equal Error Rate) – wartość punktu przecięcia krzywych fałszywych akceptacji i fałszywych odrzuceń – oraz ROC (Receiver Operating Characteristic), prezentujących kompromis pomiędzy wykrywaniem poprawnych tożsamości a odpornością na błędne dopasowania. Testy wykonuje się zarówno na



zamkniętych zestawach referencyjnych obejmujących znanych mówców, jak i w scenariuszach otwartej populacji. Identyfikacja mówcy ma rosnące znaczenie poza klasycznymi zastosowaniami biometrycznymi, w personalizacji usług, gdzie profil głosowy użytkownika może aktywować dostosowane odpowiedzi asystenta AI czy ustawienia interfejsu konwersacyjnego. Możliwe jest również łączenie tego modułu z rozpoznawaniem emocji ze ścieżki akustycznej; wspólnie dają one obraz nie tylko tego, kto mówi, ale też w jakim stanie emocjonalnym się znajduje, co może wpływać np. na decyzje systemu obsługi klienta. W tak szerokim krajobrazie integracje wymagają jednak precyzyjnego balansowania między skutecznością techniczną a ochroną prywatności danych głosowych użytkowników.

7 Etyka i wyzwania w inżynierii lingwistycznej

7.1 Ochrona prywatności użytkowników

W kontekście systemów inżynierii lingwistycznej prywatność użytkowników staje się jednym z najbardziej newralgicznych aspektów, który wymaga zarówno technicznych, jak i proceduralnych mechanizmów ochrony. Nowoczesne narzędzia NLP, rozpoznawania mowy czy generatywnej sztucznej inteligencji operują na dużych wolumenach danych językowych; często są to dane o wysokiej wrażliwości: korespondencje, transkrypcje rozmów lub zapytania skierowane do asystentów głosowych, które mogą ujawniać szczegóły życia prywatnego. Już na etapie akwizycji konieczne jest ustalenie zasad gromadzenia danych – pojawia się tu wymóg jasno zdefiniowanej zgody, najlepiej w modelu opt-in, gdzie przetwarzanie rozpoczyna się dopiero po jednoznacznym potwierdzeniu przez użytkownika. Zakazane powinny być tzw. ciemne wzorce projektowania, które domyślnie oznaczają zgodę lub ukrywają jej faktyczny zakres.

W praktyce oznacza to również obowiązek informowania użytkownika o czasie przechowywania jego danych oraz o tym, czy będą przekazywane podmiotom trzecim do analizy lub magazynowania. Jeśli takie przekazanie występuje, należy zapewnić pełną anonimizację lub pseudonimizację informacji przed wysłaniem ich poza pierwotny system. Stosuje się w tym celu strategie tokenizacji, szyfrowania na poziomie rekordów czy zabezpieczenia dostępu na poziomie wiersza.

Kontraktowana współpraca z zewnętrznymi dostawcami powinna obejmować obowiązek regularnego audytu oprogramowania i procedur weryfikujących zgodność działań z zapisami umowy. Z perspektywy projektowania architektury systemu istotne jest dążenie do minimalizacji ilości gromadzonych danych – tzw. zasada „privacy by design” milczy o zbieraniu wszystkiego „na zapas”, wskazując raczej na celowość i adekwatność informacji względem planowanego przetwarzania. Coraz częściej wdraża się rozwiązania pozwalające przenosić wybrane operacje analizujące język



na urządzenie końcowe, co redukuje konieczność przesyłu surowych danych do chmury i ogranicza ryzyko ich przechwycenia w tranzycie.

Pomocne okazują się też narzędzia zapewniające odseparowane środowiska działania modelu, w których logi i dane wejściowe są widoczne tylko dla konkretnego użytkownika, przykładem może być opracowany przez Harvard sandbox dla dużych modeli językowych GPT-4, gwarantujący brak wykorzystania wpisów do dalszego trenowania modelu. Nie bez znaczenia pozostaje fakt, że wiele systemów NLP korzysta z usług chmurowych oferowanych przez dostawców globalnych. W takim przypadku kluczowym warunkiem ochrony prywatności jest znajomość jurysdykcji prawnych dotyczących lokalizacji serwerów oraz regulacji odnoszących się do transferu danych transgranicznych. Regulacje takie jak europejskie RODO wymagają jasnego określenia celu i podstawy prawnej przetwarzania oraz umożliwiają użytkownikowi m.in. prawo do bycia zapomnianym. Bezpośrednia integracja restrykcji regulacyjnych z procesem implementacji, np. poprzez automatyczne kasowanie danych po określonym terminie, zapobiega późniejszym naruszeniom.

Na płaszczyźnie praktycznej przydatne są metody generowania syntetycznych danych zamiast korzystania ze zbiorów rzeczywistych zawierających dane osobowe. Dzięki temu można przeprowadzać trenowanie modeli NLP bez realnego ryzyka ekspozycji poufnych treści; jednocześnie syntetyczne bazy powinny zachowywać strukturę statystyczną oryginałów tak, aby wyniki były reprezentatywne. Integracja takich zbiorów wspiera zachowanie zgodności z wymogami etycznymi i prawnymi. Kolejnym aspektem jest bezpieczeństwo przeciwko potencjalnym nadużyciom ze strony cyberprzestępców. Systemy GenAI są podatne na wykorzystanie do phishingu czy socjotechniki, możliwe jest generowanie wiarygodnych wiadomości podszywających się pod instytucje poprzez wyspecjalizowane prompty.

Prywatność użytkowników można chronić poprzez blokowanie treści jawnie proszących o podanie informacji identyfikujących, a także filtrowanie odpowiedzi modelu pod kątem prób obejścia reguł bezpieczeństwa. Ochrona prywatności wymaga także budowy świadomości społecznej, programy edukacyjne w zakresie działania systemów AI pomagają użytkownikom rozumieć jakie dane są udostępniane i jakie mogą być tego konsekwencje. Publiczne fora dotyczące wpływu AI pozwalają na otwartą dyskusję nad granicami pozyskiwania informacji w relacjach człowiek–maszyna. Dobrym uzupełnieniem wymogów prawnych jest stosowanie wewnętrznych kodeksów etycznych organizacji tworzących technologie AI. Wskazują one m.in. na konieczność konsultacji między zespołami inżynieryjnymi a ekspertami prawa czy ochrony danych osobowych oraz określenie procedur reakcji na incydenty naruszenia prywatności. Wielodyscyplinarny charakter takich grup roboczych zwiększa szanse wykrycia słabych punktów projektu jeszcze przed jego wdrożeniem. Wreszcie należy uwzględnić specyfikę danych dźwiękowych i wizualnych wykorzystywanych np. w rozpoznawaniu mowy czy analizie obrazu; sygnały te



często zawierają elementy umożliwiające dodatkową identyfikację (barwa głosu, szczegóły otoczenia). Ich przetwarzanie powinno być objęte równymi rygorami ochrony co tekst pisany. Połączenie opisanych środków – od technik anonimizacji po regulacje formalne – stanowi podstawę realnej ochrony prywatności osób korzystających z technologii opartych na analizie języka naturalnego. Niezależnie od postępu technologicznego najważniejsze pozostaje zapewnienie użytkownikowi kontroli nad własnymi danymi oraz transparentności sposobu ich wykorzystania.

7.2 Bias i sprawiedliwość w modelach językowych

Zagadnienie biasu w modelach językowych odnosi się do zjawiska, w którym systemy AI uczą się i odtwarzają istniejące w zbiorach treningowych uprzedzenia, utrwalając nierówności obecne w danych źródłowych. W praktyce polega to na tym, że statystyczne wzorce pozyskane przez model mogą faworyzować lub dyskryminować określone grupy osób, często niezamierzenie. Klasycznym przykładem jest trenowanie algorytmów rekrutacyjnych na danych historycznych, gdzie nadreprezentacja mężczyzn w branży technicznej skutkuje preferowaniem ich aplikacji przez model, mimo braku formalnych reguł różnicujących płeć. Podobny mechanizm ujawnia się przy rozpoznawaniu twarzy: jeśli proces uczenia odbywał się głównie na obrazach osób o jasnej karnacji, to efektywność systemu względem innych grup demograficznych będzie niższa. Dlatego analiza biasu wymaga równoczesnego rozpatrywania źródeł danych i procesu budowania modelu. Zbiory treningowe powinny być reprezentatywne względem cech populacji, którą mają opisywać lub której dotyczą decyzje modelu.

Nie zrównoważenie danych prowadzi do stronniczych predykcji, w kontekście języka naturalnego może oznaczać częstsze kojarzenie określonych zawodów z jedną płcią czy danym pochodzeniem etnicznym. Problem ten jest szczególnie znaczący dla modeli generatywnych, które operują na miliardach słów i mogą nieświadomie reprodukować stereotypy zakodowane w treściach publicznych. W sytuacjach krytycznych konieczne jest stosowanie metod wyrównywania udziału danych dla różnych grup, zarówno w fazie przygotowania zbiorów danych, jak i w trakcie trenowania. Tworzenie syntetycznych danych jest jednym z narzędzi minimalizujących efekt biasu. Dane takie mogą być generowane tak, aby odwzorowywać strukturę wypowiedzi pozbawionych uprzedzeń lub aby dopełniać brakujące przykłady dla grup niedoreprezentowanych. Dzięki temu model uczy się bardziej zbalansowanego obrazu rzeczywistości bez konieczności przetwarzania dużych ilości prawdziwych danych osobowych. Istotnym aspektem sprawiedliwości jest ocena czy decyzje modelu są transparentne i możliwe do wyjaśnienia użytkownikowi lub organowi kontrolnemu.



Modele językowe bywają traktowane jako „czarne skrzynki”, ich wewnętrzne mechanizmy generowania odpowiedzi i wyboru tokenów są trudne do interpretacji bez specjalistycznych narzędzi. Wyjaśnialność ma tutaj podwójny sens: umożliwia wykrywanie źródeł uprzedzeń oraz wzmacnia zaufanie do systemu poprzez przedstawienie procesu decyzyjnego. Podczas projektowania interakcji z modelem warto kontrolować także warstwę wejść, prompty używane do inicjowania odpowiedzi mogą same w sobie nosić cechy stronnicze. Źle sformułowane pytanie, zawierające sugestie lub stereotypowe założenia, może spowodować odpowiedź wzmacniającą uprzedzenia. Dlatego zaleca się regularną rewizję szablonów rozmów czy instrukcji wykorzystywanych w środowiskach produkcyjnych.

Potrzebny jest również etap testowania pod kątem sprawiedliwości: modele można badać przy użyciu benchmarków zaprojektowanych specjalnie do ujawniania nierówności w wynikach między grupami demograficznymi czy socjoekonomicznymi. Analiza wyników powinna wskazać różnice jakości odpowiedzi, np. procent poprawnych odpowiedzi dla każdej badanej grupy i umożliwić wdrażanie działań naprawczych poprzez retrening albo korektę wag. Nie można pomijać faktora kulturowego: modele trenowane na materiałach z jednego kręgu kulturowego będą miały ograniczoną zdolność adekwatnego reagowania na dane spoza swojego zakresu referencyjnego. Może to skutkować błędną interpretacją metafor czy idiomów oraz niepoprawnym przypisywaniem konotacji emocjonalnych. Odpowiedzią na to jest wzbogacenie korpusów o wielojęzyczne zasoby harmonizowane semantycznie oraz opracowanie procedur mapowania relacji znaczeniowych między kulturami.

W przypadku zastosowań komercyjnych sprawiedliwość modeli językowych ma także wymiar regulacyjny: przepisy prawa nakładają obowiązek unikania dyskryminacji osób ze względu na cechy chronione (płeć, rasa, religia itd.). Niespełnienie tych wymogów może prowadzić nie tylko do szkód społecznych, ale też konsekwencji finansowo-prawnych dla organizacji wdrażającej system AI. Zwraca się również uwagę na istniejące ryzyko wtórnego przenoszenia biasu poprzez adaptację modelu. Jeżeli dane użyte do dostrojenia parametrycznego reprezentują zawężony światopogląd lub specyficzny profil użytkownika docelowego, wszystkie wcześniejsze wysiłki redukujące uprzedzenia mogą zostać osłabione. Stąd wynika potrzeba dokumentowania procesu trenowania i utrzymywania spójności standardów jakości przy każdej iteracji prac nad modelem.

Równocześnie należy pamiętać o aspekcie ochrony przed tzw. manipulacją promptami, celowym podsuwaniem modelowi danych lub form pytań prowadzących go ku stronniczym odpowiedziom. Wdrożenie filtrów treści zapobiegających reagowaniu na tego rodzaju próby pozwala utrzymać zgodność działania z założeniami etyki cyfrowej. Wszystkie te elementy składają się na szeroki proces zapewniania sprawiedliwości modeli językowych: począwszy od tworzenia



zbalansowanych zbiorów treningowych, przez obiektywne projektowanie interfejsu użytkownik–model, aż po mechanizmy kontroli wyjść pod kątem rzetelności i neutralności treści. Ostatecznym celem pozostaje stworzenie systemów NLP zdolnych operować w sposób rzetelny wobec wszystkich użytkowników, niezależnie od tego, jakie dane je kształtowały w fazie edukacyjnej.

8 Przyszłość inżynierii lingwistycznej

8.1 Integracja z multimodalnymi systemami AI

Integracja inżynierii lingwistycznej z multimodalnymi systemami AI nabiera coraz większego znaczenia, ponieważ nowe architektury operują już nie tylko na tekście, ale równocześnie na obrazie, dźwięku czy nawet danych środowiskowych. W odróżnieniu od rozwiązań, które koncentrują się głównie na pracy z językiem pisanym, tutaj wyzwaniem jest projektowanie mechanizmów zdolnych rozumieć i generować treści w oparciu o kombinację różnych modalności. Multimodalny model może analizować transkrypcję wypowiedzi i równolegle interpretować gesty osoby mówiącej lub obraz przedstawiający obiekt rozmowy. Powiązanie ścieżek przetwarzania tak odmiennych typów danych wymaga nie tylko większej mocy obliczeniowej, ale również spójnego mechanizmu reprezentacji, potrzeba mapowania słów, pikseli i próbek dźwięku do kompatybilnych przestrzeni cech.

Obecne projekty zakładają użycie modeli typu enkoder dla każdej modalności, po czym następuje etap fuzji reprezentacji, polegający na łączeniu wektorów opisujących poszczególne źródła informacji. W implementacjach wykorzystujących transformery stosuje się tzw. multi-head cross-attention, gdzie uwaga obejmuje tokeny różnych modalności jednocześnie. Taki układ pozwala np. powiązać fragment obrazu z odpowiadającym mu opisem w tekście oraz odpowiednim momentem w zapisie audio. Ważnym aspektem jest synchronizacja czasowa, szczególnie przy danych wideo, gdzie ruch i kontekst wizualny zmieniają się w sposób ciągły. System musi wtedy utrzymywać zsynchronizowany model czasowy zarówno dla danych audio/mowy, jak i dla zmian klatki obrazu. Postęp w tej dziedzinie został zilustrowany pojawieniem się narzędzi takich jak ChatGPT 4o, zdolnego przyjmować wejścia tekstowe, graficzne oraz foniczne i generować wielokanałowe odpowiedzi. Model taki może otrzymać zdjęcie schematu technicznego, pytanie opisowe oraz próbkę wypowiedzi ustnej użytkownika i – uwzględniając pełny kontekst – wygenerować odpowiedź obejmującą zarówno interpretację obrazu, jak i wyjaśnienie szczegółów pytania.

Z punktu widzenia inżynierii lingwistycznej oznacza to konieczność rozszerzenia tradycyjnych pipeline'ów NLP tak, aby integracja ze ścieżkami wizualnymi czy akustycznymi była natywna. Rozwiązania takie jak Copilot AI wykorzystują te



możliwości w kontekście produktywności biurowej, proponują obrazy na podstawie opisu tekstowego lub podsumowują zawartość dokumentów uzupełnionych materiałem graficznym. Multimodalność ma także swój wymiar generatywny: modele zdolne tworzyć spójne narracje łączące warstwę tekstu z generowanymi obrazami (np. DALL-E) czy materiałem wideo budują nowe formy komunikacji cyfrowej. Przykładem zastosowania jest tworzenie filmów szkoleniowych łączących wystąpienie lektora (syntetyzowanego na bazie wejścia tekstowego) z ilustracjami wygenerowanymi przez algorytmy dyfuzyjne.

Z perspektywy badawczej wspólna przestrzeń reprezentacji umożliwia także bardziej zaawansowane zadania – wyszukiwanie semantyczne pomiędzy różnymi typami zasobów: użytkownik może zapytać o „produkt przypominający ten ze zdjęcia” poprzez wpis lub krótką wypowiedź głosową. Wyzwania technologiczne obejmują konieczność tworzenia korpusów multimodalnych, zestawów danych zawierających skorelowane pary lub tria informacji tekstowych, obrazowych i dźwiękowych. Brak wysokiej jakości anotowanych korpusów obejmujących wszystkie potrzebne modalności hamuje rozwój niektórych zastosowań; ponadto procedury anotacji muszą zapewniać zgodność znaczeniową między modalnościami. W praktyce oznacza to, że opis słowny powinien wiernie odwzorowywać zawartość obrazu czy przebieg nagrania audio – nawet drobne rozbieżności mogą dezorientować model podczas treningu.

Kolejnym aspektem jest integracja wiedzy parametrycznej zapisanej w dużych modelach językowych z wiedzą nieparametryczną pozyskiwaną dynamicznie z zasobów multimodalnych. Jeśli system ma udzielić odpowiedzi bazującej na aktualnym obrazie produktu dostępnego w katalogu firmy, powinien umieć pobrać odpowiednie dane wizualne z bazy wektorowej i powiązać je z opisem technicznym przed wygenerowaniem multimodalnej odpowiedzi. Architektury agentowe stanowią tutaj naturalny wybór: agent analizuje intencję użytkownika, decyduje o pobraniu danych wizualnych bądź audio i następnie uruchamia moduł NLP do przygotowania pełnej wypowiedzi tekstowej.

Coraz częściej takie rozwiązania stosuje się w środowiskach przemysłowych oraz badawczych – przykładem są systemy asystentów osobistych, które poznają preferencje użytkownika zarówno poprzez dialogi tekstowe, jak i obserwację jego zachowań w aplikacjach multimedialnych. Integracja modalności pozwala im antycypować potrzeby odbiorcy: mogą zauważyć wzór kolorystyczny używany przez grafikę w materiałach promocyjnych i proponować podobne palety barw przy kolejnych projektach. W kontekście edukacyjnym multimodalna inżynieria lingwistyczna otwiera możliwość tworzenia interaktywnych materiałów dydaktycznych reagujących zarówno na pytania ucznia zadane głosowo, jak i na przesłane szkice czy wykresy. System może skorygować błędy w rysunku chemicznym wskazane



wizualnie oraz dostarczyć słowną instrukcję korekty – tym samym spajając analizę języka naturalnego z interpretacją grafiki.

Ważnym zagadnieniem pozostaje standaryzacja formatów przesyłu i przechowywania danych multimodalnych. Potrzeba definicji jednolitych protokołów API pozwalających modelom językowym konsultować się z komponentami przetwarzania obrazu lub dźwięku niezależnie od platformy sprzętowej. Rozproszone środowiska chmurowe sprzyjają tu modularyzacji, poszczególne usługi mogą działać niezależnie jako mikroserwisy odpowiadające za obsługę konkretnej modalności. Patrząc perspektywicznie integracja inżynierii lingwistycznej z systemami multimodalnymi będzie sprzyjała powstawaniu bardziej inteligentnych interfejsów człowiek–maszyna zdolnych rozpoznawać intencję nie tylko po treści słów, lecz także po ich tonie czy ilustrującym je materiale wizualnym. Efektywność takich systemów zależy jednak od jakości wszystkich komponentów składających się na proces fuzji modalności oraz od umiejętnego zarządzania kontekstem pochodzącym jednocześnie z wielu źródeł danych, co stawia przed twórcami wymagania wykraczające poza klasyczne ramy projektowania narzędzi NLP.

8.2 Rozwój interfejsów człowiek-maszyna

Projektowanie interfejsów człowiek–maszyna ewoluuje w kierunku coraz bardziej naturalnej i kontekstowej komunikacji, gdzie język staje się podstawowym medium interakcji. Rozwój dużych modeli językowych (LLM) oraz technik multimodalnych tworzy nowe możliwości dla systemów, które rozumieją nie tylko treść wypowiedzi, ale również intencje i emocje rozmówcy. Interfejs przestaje być jedynie kanałem wymiany komend – staje się środowiskiem konwersacyjnym zdolnym adaptować się do stylu komunikacji użytkownika. Kluczowym aspektem tej transformacji jest zastosowanie architektur agentowych, które łączą komponenty rozumienia języka naturalnego z narzędziami wykonawczymi. Agenty mogą samodzielnie wybierać działania – od pobrania danych z wyszukiwarki semantycznej po generowanie spersonalizowanej odpowiedzi – na podstawie ciągłego modelowania kontekstu rozmowy.

W praktyce oznacza to, że interfejs potrafi nie tylko reagować na pojedyncze polecenia, lecz utrzymywać dialog dopasowany do historii interakcji. Mechanizm pamięci konwersacyjnej integruje wiedzę parametryczną modelu z danymi pobieranymi dynamicznie z baz wektorowych, co eliminuje problem „resetowania” kontekstu przy zmianie tematu. Rozwój takich interfejsów napędza również rosnąca rola personalizacji. Systemy mogą analizować wcześniejsze zachowania użytkownika – preferowany format odpowiedzi, częstotliwość zapytań czy nawet słownictwo charakterystyczne – aby dostosować ton i strukturę komunikatów. Personalizacja może obejmować także modalność przekazu: użytkownik preferujący



komunikację głosową otrzyma odpowiedź zsynchronizowaną z synteizatorem mowy, natomiast dla odbiorcy tekstowego interfejs wygeneruje rozwiniętą wersję pisemną wraz z odsyłaczami do dodatkowych źródeł. Istotnym trendem jest przechodzenie od prostych chatbotów opartych na dopasowaniu wzorców do pełnoprawnych narzędzi AI integrujących wiele domen wiedzy i funkcji systemowych. Takie narzędzia potrafią np. przeanalizować plik graficzny przesłany przez użytkownika, powiązać go z opisem tekstowym i wygenerować rekomendację lub raport łączący oba źródła informacji.

Warunkiem skuteczności jest płynne przenikanie się warstw przetwarzania – rozpoznawania obrazu, analizy języka i logiki biznesowej – bez widocznych granic dla użytkownika końcowego. Z punktu widzenia ergonomii komunikacji pożądane są mechanizmy interpretujące mowę ciała czy modulację głosu w środowiskach wyposażonych w czujniki audio-wizualne. Interfejs może w ten sposób modyfikować odpowiedź w czasie rzeczywistym – podając więcej szczegółów przy wykryciu niepewności odbiorcy lub skracając komunikat, gdy sygnały niewerbalne sugerują pośpiech. Włączenie analizy kontekstów pozajęzykowych zwiększa poczucie „rozumienia” przez system i redukuje ryzyko frustracji wynikającej z niedopasowania przekazu.

Coraz większe znaczenie przypisuje się strategiom projektowania promptów (prompt engineering), które pozwalają kontrolować zachowanie interfejsu na poziomie mikrointerakcji. Starannie skonstruowane instrukcje sterują sposobem interpretacji danych wejściowych przez model językowy – mogą wymusić krokową prezentację rozwiązania problemu lub uzasadnienie rekomendacji produktowej. Wdrażanie takich strategii staje się elementem optymalizacji doświadczenia użytkownika. Ważnym nurtem pozostaje też integracja elementów etyki i transparentności działania w projekcie interfejsu. Użytkownik powinien mieć możliwość śledzenia źródeł informacji użytych do generacji odpowiedzi oraz dostępu do mechanizmów kontroli nad zakresem przetwarzanych danych osobowych. Funkcje wyjaśniające pomagają budować zaufanie – zwłaszcza gdy interfejs dokonuje wyboru pomiędzy sprzecznymi opcjami lub sugeruje działania o konsekwencjach finansowych czy prawnych. Integracja technologii rozpoznawania mowy oraz syntezy głosu pozwala projektować interfejsy obsługiwane całkowicie głosowo.

Takie rozwiązania wykraczają poza same translatory poleceń – umożliwiają prowadzenie wielowątkowej rozmowy, przy czym system sam zarządza strukturą dialogu i priorytetami odpowiedzi. Wersje wielojęzyczne wspierają natychmiastowe tłumaczenie wypowiedzi, co ułatwia współpracę międzynarodowych zespołów bez potrzeby korzystania z oddzielnych tłumaczy. Postęp mocy obliczeniowej urządzeń brzegowych otwiera możliwość uruchamiania modeli bezpośrednio na sprzęcie użytkownika, co zmniejsza opóźnienia i zwiększa prywatność dzięki ograniczeniu przesyłu danych do chmury. Lokalne inferencje przyspieszają reakcję systemu przy



zachowaniu wysokiego poziomu jakości generowanych treści. Nowoczesne interfejsy człowiek–maszyna podążają więc w kierunku synergii wielu technologii: multimodalnego wejścia i wyjścia, agentowego zarządzania zadaniami, dostosowania języka komunikatu do odbiorcy oraz dbania o transparentność procesu decyzyjnego. Ich dalszy rozwój będzie zależał od równowagi między kreatywnością systemu a kontrolą merytoryczną, co staje się kluczowe wraz ze wzrostem roli AI w codziennych procesach pracy i życia prywatnego.

8.3 Potencjalne kierunki badań

Kierunki badań w inżynierii lingwistycznej będą prawdopodobnie podążały w stronę coraz głębszej integracji modeli językowych z technologiami multimodalnymi, co wynika z obserwowanego trendu poszerzania funkcjonalności systemów o zdolność interpretowania obrazu, mowy, tekstu i innych sygnałów w sposób skoordynowany. Modele te mogą działać w oparciu o architekturę transformera z mechanizmem uwagi krzyżowej, umożliwiającym powiązanie kontekstu wizualnego z treścią wypowiedzi oraz synchronizację tych informacji w czasie. Wraz ze wzrostem zastosowań agentów AI, pojawia się potencjał badań nad rozwiązaniami zdolnymi do autonomicznego wyboru modalności dla realizacji zadania czy odpowiedź powinna być wyrażona opisem tekstowym, syntezą mowy, czy też wizualizacją danych.

Jednym z ciekawych pól rozwoju jest personalizacja generowanych treści i interfejsów konwersacyjnych poprzez dostosowanie modeli zarówno na poziomie parametrycznym, jak i przez dynamiczne połączenie z bazami wiedzy nieparametrycznej. Może to oznaczać opracowanie algorytmów dostrajania modeli przy minimalnym zużyciu zasobów obliczeniowych, np. poprzez adaptory lub metody p-tuning, tak, aby obok ogólnej kompetencji językowej model dysponował specjalistyczną terminologią branżową. Potencjalnym nurtem są także badania nad adaptacją modeli generatywnych do środowisk brzegowych, co opisywano w kontekście rozwoju interfejsów człowiek–maszyna. Byłoby to istotne dla urządzeń mobilnych czy systemów IoT, które mogłyby interpretować polecenia głosowe lub tekstowe lokalnie, bez konieczności wysyłania danych do chmury. Umożliwiłoby to zmniejszenie opóźnień oraz poprawę prywatności użytkowników.

Inny kierunek obejmuje dalsze rozwijanie mechanizmów pamięci konwersacyjnej pozwalającej na prowadzenie długotrwałych dialogów tematycznych bez utraty spójności kontekstu. Badania w tym obszarze mogłyby skupić się na hybrydach łączących wewnętrzną pamięć modelu (poznaną podczas treningu) z modułami przeszukiwania baz wektorowych lub repozytoriów dokumentów. Taka integracja sprzyja zwiększeniu aktualności odpowiedzi, model mógłby „doutczać” się bieżących wydarzeń z zewnętrznych źródeł za pomocą API lub indeksowania semantycznego. Obserwuje się tendencję do utraty spójności przy generowaniu długich wypowiedzi;



projektowanie procesów hierarchicznej generacji (najpierw struktura, potem szczegółowa treść) może ograniczyć ten problem. Planowanie wieloetapowych promptów i analiza wpływu konstrukcji zapytań na jakość odpowiedzi pozostają tu istotnym polem eksperymentów.

Kolejne wyzwanie badawcze dotyczy nawet subtelnych form biasu ujawniającego się podczas działania dużych modeli językowych. Można spodziewać się rosnącej liczby prac nad metodami wykrywania uprzedzeń oraz ich neutralizowania już w procesie szkolenia i dostrajania modeli. W tym kontekście szczególnie istotne staje się zapewnienie wielokulturowej reprezentacji danych i harmonizacja znaczeń między językami. Tworzenie benchmarków specjalnie nastawionych na badanie rzetelności i sprawiedliwości działania modeli otwiera drogę do lepszych narzędzi oceny jakości odpowiedzi pod kątem neutralności kulturowej i społecznej. Badania mogą też objąć dalszy rozwój narzędzi do automatycznej kontroli faktograficznej odpowiedzi generowanych przez modele – łączenie analizy lingwistycznej z porównywaniem informacji względem zweryfikowanych źródeł.

Integracja modułów sprawdzania faktów z procesem generacji treści wymaga utrzymania równowagi między szybkością tworzenia wypowiedzi a kompleksowym sprawdzaniem danych wejściowych. Szczególną rolę mogą tu odegrać architektury agentowe korzystające z planowania działań: przed wygenerowaniem końcowej odpowiedzi system wykonuje kroki pobrania tła faktograficznego, analizy kontekstowej i walidacji spójności semantycznej. Inspiracją dla tego podejścia są scenariusze zastosowań biznesowych, gdzie błędna informacja może mieć bezpośrednie konsekwencje finansowe czy prawne.

Nie można pominąć potencjału badań nad rozszerzeniem zdolności LLM-ów o precyzyjną analizę modalności pozajęzykowych dla interpretacji emocji użytkownika podczas interakcji. Implementacja takich funkcji wymaga synergii technologii przetwarzania mowy, analizy obrazu i NLP; celem jest stworzenie systemów reagujących nie tylko na treść słowną, ale też na ton głosu czy wyraz twarzy odbiorcy. W sektorach takich jak obsługa klienta mogłoby to skutkować adaptacją strategii komunikacyjnej interfejsu w czasie rzeczywistym.

Na horyzoncie pozostaje też rozwój korpusów multimodalnych wysokiej jakości – zestawów danych ściśle skorelowanych pod względem znaczeniowym między różnymi modalnościami. Bez takich zasobów projektowanie uniwersalnych systemów potrafiących płynnie przełączać tryb pracy między analizą obrazu a interpretacją tekstu będzie ograniczone. Badania w tym obszarze mogą skupić się na metodach automatycznej anotacji multimodalnej oraz translacji znaczenia między różnymi typami sygnałów. Dalsze prace będą także obejmować rozwój architektur hybrydowych, łączących reguły lingwistyczne zapewniające formalną poprawność ze statystycznymi modelami neuronowymi gwarantującymi elastyczność reakcji.



Podejście takie wydaje się kluczowe wszędzie tam, gdzie wymagana jest zarówno kreatywność w komunikacji jak i zgodność ze standardem językowym lub branżowym.

Podsumowując przewidywane kierunki badań obejmują jednocześnie pogłębianie integracji modalności danych wejściowych i wyjściowych, personalizację działania modeli przy zachowaniu efektywności obliczeniowej, ograniczanie biasu oraz kontrolę jakości generowanej treści. Wprowadzenie agentowego planowania działań wraz z rozbudowaną pamięcią konwersacyjną może stać się fundamentem dla kolejnej generacji systemów językowych zdolnych do samodzielnego zarządzania procesem komunikacyjnym od zebrania kontekstu aż po przygotowanie wyjaśnialnych i adekwatnych odpowiedzi.

9. Zakończenie

Analiza przedstawionych zagadnień ukazuje, jak szerokie i złożone jest pole inżynierii lingwistycznej, łączące elementy lingwistyki teoretycznej, stosowanej oraz zaawansowanych technologii informatycznych. Ewolucja od prostych systemów opartych na regułach do nowoczesnych modeli głębokiego uczenia, takich jak transformery i duże modele językowe, umożliwiła osiągnięcie poziomu przetwarzania języka naturalnego zbliżonego do ludzkiego w wielu aspektach. Jednocześnie integracja wiedzy formalnej z reprezentacjami statystycznymi pozwala na tworzenie systemów bardziej elastycznych, a zarazem kontrolowalnych i wyjaśnialnych.

Rozwój korpusów językowych oraz narzędzi do ich analizy stanowi fundament dla trenowania i oceny modeli, a także dla ich adaptacji do specyficznych dziedzin i języków. Współczesne technologie umożliwiają realizację zadań takich jak tłumaczenie maszynowe, rozpoznawanie mowy, analiza sentymentu czy systemy pytanie–odpowiedź z wysoką skutecznością, choć wciąż wymagają dalszych usprawnień, zwłaszcza w zakresie radzenia sobie z niejednoznacznościami, błędami interpretacyjnymi oraz generowaniem treści zgodnych z faktami.

Ważnym aspektem pozostają kwestie etyczne, w tym ochrona prywatności użytkowników oraz eliminacja biasu i zapewnienie sprawiedliwości w działaniu modeli językowych. Konieczne jest stosowanie odpowiednich procedur, transparentność oraz ciągłe monitorowanie i doskonalenie systemów, aby minimalizować ryzyko utrwalania uprzedzeń i naruszania praw osób korzystających z technologii. Wdrażanie mechanizmów kontroli oraz edukacja użytkowników stanowią integralną część odpowiedzialnego rozwoju tej dziedziny.

Przyszłość inżynierii lingwistycznej wiąże się z dalszą integracją modeli językowych z multimodalnymi systemami AI, które łączą analizę tekstu, obrazu, dźwięku i innych sygnałów, co pozwala na tworzenie bardziej naturalnych i kontekstowych interfejsów



człowiek–maszyna. Rozwój agentów AI zdolnych do autonomicznego podejmowania decyzji oraz zarządzania pamięcią konwersacyjną otwiera nowe możliwości w zakresie personalizacji i efektywności komunikacji. Wyzwania techniczne, takie jak synchronizacja modalności, standaryzacja formatów danych czy optymalizacja działania na urządzeniach brzegowych, będą wymagały dalszych badań i innowacji.

W perspektywie badawczej istotne pozostaje również doskonalenie metod kontroli jakości generowanych treści, w tym ograniczanie halucynacji modeli oraz rozwijanie mechanizmów sprawdzania faktów. Współpraca interdyscyplinarna, łącząca lingwistykę, informatykę, etykę i prawo, będzie kluczowa dla tworzenia systemów NLP, które będą nie tylko zaawansowane technologicznie, ale także bezpieczne, sprawiedliwe i użyteczne w codziennych zastosowaniach. Ostatecznie rozwój tej dziedziny przyczyni się do coraz bardziej efektywnej i naturalnej komunikacji między ludźmi a maszynami, wspierając różnorodne obszary życia społecznego, gospodarczego i naukowego.



10. Bibliografia

1. Alto V. (2024) *Generatywna sztuczna inteligencja z ChatGPT i modelami OpenAI: podnieś swoją produktywność i innowacyjność za pomocą GPT3 i GPT4*. Gliwice, Wyd. Helion.
2. Alto V. (2024) *LLM w projektowaniu oprogramowania. Tworzenie inteligentnych aplikacji i agentów z wykorzystaniem dużych modeli językowych*. Gliwice, Wyd. Helion
3. Jain A., Fernandez M. (2024) *AI-assisted programming for web and machine learning*. Packt Publishing
4. Microsoft (2024) *Microsoft 365 copilot at work*. <https://www.microsoft.com/en-us/microsoft-365-copilot>
5. Baker P. (2024) *Generative AI*. Wyd. For Dummies ISBN 978-1394270743
6. Bergeret O., Abbasi A. and Farvault J. (2025) *GenAI on AWS a practical approach to building generative AI applications on AWS*. Wyd. Wiley ISBN 978-1394281282
7. Bird S., Klein E., Loper E. (2009) *Natural language processing with Python*. Wyd. O'Reilly Media, ISBN 978-05-965-5571-9
8. Chockalingam, A. et al. (2023) *A beginner's guide to large language models part 1*. <https://resources.nvidia.com/en-us-large-language-model-ebooks/llm-ebook-part1>
9. Coveyduc J.L., Anderson J.L. (2020) *Artificial intelligence for business: A roadmap for getting started with AI*. Wyd. Wiley, ISBN 978-1119651734
10. Deitel P., Deitel H. (2020) *Python dla programistów. Big Data i AI. Studia przypadków*. Wyd. Helion, Gliwice, ISBN 978-83-283-6395-3
11. Dodhia R. (2024) *AI for social good*. Wyd. Wiley, ISBN 978-1394205783
12. Lindsell-Roberts S. (2024) *Business writing with AI For Dummies*. Wyd. For Dummies, ISBN 978-1394261734
13. Marr B. (2024) *Generative AI in Practice*. Wyd. John Wiley & Sons, ISBN: 978-1-394-24556-7
14. Nagy Z. (2018) *Artificial Intelligence and Machine Learning Fundamentals. Develop real-world applications powered by the latest AI advances*. Wyd. Packt Publishing ISBN 978-17-898-0920-6
15. Prosise J. (2022) *Applied Machine Learning and AI for Engineers*. Wyd. O'Reilly Media ISBN 978-10-981-3629-1
16. Raschka S. (2024) *Build a Large Language Model (from Scratch)*. Wyd. Manning Publications, ISBN 978-1633437166