

Prof. dr hab. inż. Grzegorz Dudek
Katedra Automatyki, Elektrotechniki i Optoelektroniki
Wydział Elektryczny
Politechnika Częstochowska
Al. Armii Krajowej 17
42-200 Częstochowa

Częstochowa, dn. 18 grudnia 2024 r.

Rada Naukowa Dyscypliny
INFORMATYKA TECHNICZNA
I TELEKOMUNIKACJA
Sekretariat
Data wpływu 03.01.2025
Numer.....

RECENZJA

rozprawy doktorskiej mgr. inż. Łukasza Lepaka pt. *Task automation using artificial intelligence methods*

Formalną podstawą opracowania recenzji jest pismo Rady Dyscypliny Naukowej Informatyka Techniczna i Telekomunikacja Politechniki Warszawskiej z dnia 25.10.2024 r. Oceny rozprawy doktorskiej dokonano według kryteriów określonych w ustawie z 20 lipca 2018 r. *Prawo o szkolnictwie wyższym i nauce*. Promotorem rozprawy doktorskiej jest dr hab. inż. Paweł Wawrzyński, prof. PW.

Charakterystyka rozprawy

Rozprawę doktorską stanowi zbiór pięciu artykułów naukowych opatrzonych częścią wstępną. Rozprawa napisana jest w języku angielskim. Część wstępna liczy 56 stron.

Zawartość części wstępnej:

1. Wprowadzenie. Rozdział ten wprowadza do tematyki pracy doktorskiej, czyli automatyzacji zadań z wykorzystaniem metod sztucznej inteligencji. Autor przedstawia główne tezy i cele swojej pracy, a także jej strukturę. Wskazuje również na pięć publikacji naukowych, które stanowią podstawę rozprawy.
2. Podstawy teoretyczne. W tym rozdziale Autor omawia kluczowe koncepcje i metody, które wykorzystuje w swoich badaniach: uczenie ze wzmocnieniem, architektury sieci neuronowych i algorytmy ich uczenia.
3. Automatyzacja handlu energią na rynku dnia następnego. Rozdział omawia artykuł [1] – przedstawia autorską strategię tworzenia ofert kupna i sprzedaży energii elektrycznej wykorzystującą uczenie ze wzmocnieniem. W oparciu o dostępne informacje, takie jak ceny energii, dane pogodowe i prognozy zużycia energii elektrycznej, strategia optymalizuje decyzje agenta handlowego w celu maksymalizacji zysków w rzeczywistych warunkach rynkowych.
4. Wykrywanie słów kluczowych w języku polskim w nagraniach audio. Rozdział omawia artykuł [2] – dotyczy problemu wykrywania słów kluczowych w nagraniach audio, ze szczególnym uwzględnieniem języka polskiego. Autor omawia dwa główne podejścia do tego zadania: konwersję sygnału audio na tekst i tekstowe wyszukiwanie słów kluczowych oraz wyszukiwanie oparte na podobieństwach sekwencji audio. Z uwagi na ograniczoną dostępność danych dla języka polskiego, Autor koncentruje się na drugim podejściu. Badania obejmują testy różnych modeli dopasowywania podobieństwa na różnych zbiorach danych w języku angielskim i polskim.
5. Strojenie hiperparametrów w trakcie uczenia sieci neuronowych. Rozdział omawia artykuł [3] – poświęcony jest optymalizacji hiperparametrów w trakcie procesu uczenia sieci neuronowych. Autor proponuje nowy algorytm, który dynamicznie optymalizuje wartości hiperparametrów podczas procesu uczenia się, mierząc ich krótko- i długoterminowy wpływ na wartość funkcji straty.

Wykazano, że metoda ta przewyższa inne popularne algorytmy pod względem stabilności i efektywności.

6. Symultaniczne tłumaczenie maszynowe. Rozdział omawia artykuł [4] – proponuje system tłumaczenia maszynowego oparty na uczeniu ze wzmocnieniem, który automatycznie kontroluje opóźnienie tłumaczenia przy zachowaniu jego wysokiej jakości. System porównano z innymi rozwiązaniami tłumaczenia maszynowego opartymi na sieciach neuronowych, osiągając podobne wyniki przy ograniczonym kontekście.
7. Głęboka transformacja stanu w rekurencyjnych sieciach neuronowych. Rozdział omawia artykuł [5] – proponuje nowy moduł bazowy sieci rekurencyjnej, który rozwiązuje problemy z propagacją gradientu występujące w klasycznych sieciach rekurencyjnych i osiąga lepsze wyniki od modułów standardowych.
8. Uwagi końcowe. Podsumowanie wyników badań i ich znaczenia dla automatyzacji zadań przy użyciu sztucznej inteligencji.
9. Inne osiągnięcia. Lista wystąpień konferencyjnych Autora i projektów, w których brał udział.
 - Bibliografia zawierająca ok. 80 pozycji.
 - Załącznik A – zestawienie skrótów
 - Załącznik B – artykuły wchodzące w skład dysertacji:

- [1] Łukasz Lepak, Paweł Wawrzyński. "Reinforcement learning meets microeconomics: Learning to designate price-dependent supply and demand for automated trading", *European Conference on Machine Learning and Principles and Practice of Knowledge Discovery in Databases 2024*.
Contribution: The PhD candidate gathered the required data, designed, developed, tested, and analyzed the proposed trading strategy, and prepared the manuscript.
Ministerial score: 140
Percentage of contribution: 75%
- [2] Łukasz Lepak, Kacper Radzikowski, Robert Nowak, Karol J. Piczak. "Generalisation gap of keyword spotters in a cross-speaker low-resource scenario", *Sensors*, MDPI, 2021.
Contribution: The PhD candidate developed, tested and analyzed created solutions, gathered the required data, and prepared the manuscript.
Ministerial score: 100
Impact factor: 3.4
Percentage of contribution: 40%
- [3] Paweł Wawrzyński, Paweł Zawistowski, Łukasz Lepak. "Automatic hyperparameter tuning in on-line learning: Classic Momentum and ADAM", *2020 International Joint Conference on Neural Networks (IJCNN)*, IEEE, 2020.
Contribution: The PhD candidate developed, tested, and analyzed different versions of the proposed algorithm.
Ministerial score: 140
Percentage of contribution: 40%
- [4] Grzegorz Rypeś, Łukasz Lepak, Paweł Wawrzyński. "Reinforcement Learning for on-line Sequence Transformation", *2022 17th Conference on Computer Science and Intelligence Systems (FedCSIS)*, IEEE, 2022.

Contribution: The PhD candidate assisted in developing, testing, and analyzing the proposed system, reviewed the literature, prepared the manuscript, and presented the method at the conference.

Ministerial score: 70

Percentage of contribution: 50%

- [5] Łukasz Neumann, Łukasz Lepak, Paweł Wawrzyński. "Least Redundant Gated Recurrent Neural Network", 2023 International Joint Conference on Neural Networks (IJCNN), IEEE, 2023.

Contribution: The PhD candidate developed models to which the proposed solution was compared, and conducted part of the experiments.

Ministerial score: 140

Percentage of contribution: 20%

Tezy rozprawy:

Hypothesis 1. Utilizing information relevant to the day-ahead energy market trading process to automate bid creation increases the trading agent's profits and allows the strategy to be used in real-life scenarios.

Hypothesis 2. Finding keywords in call center recordings in a low-resource language is possible using audio similarity matching models.

Hypothesis 3. The short- and long-term influence of the learning algorithm's hyperparameters on the training process can be used to optimize these hyperparameters on the fly without prior knowledge of the solved problem, thus improving the learning performance.

Hypothesis 4. The translation delay in a simultaneous machine translation can be automatically controlled while maintaining high translation quality for sequences of arbitrary length.

Hypothesis 5. Deep, arbitrary transformation of states in a recurrent neural network allows to achieve results comparable with state-of-the-art methods while mitigating training problems and being memory-efficient.

Opinia na temat rozprawy

Tematyka

Tematyka rozprawy – automatyzacja zadań za pomocą metod sztucznej inteligencji (SI) – jest niezwykle szeroka i zróżnicowana. W ostatnich latach SI znacząco przyspieszyła rozwój technologii automatyzacji w wielu dziedzinach, takich jak przemysł, medycyna, logistyka czy sektor finansowy. Algorytmy SI umożliwiają nie tylko automatyzację powtarzalnych zadań, ale także realizację bardziej skomplikowanych procesów decyzyjnych, co prowadzi do wzrostu efektywności i redukcji kosztów. Inteligentne systemy mogą przewidywać i optymalizować procesy w skali niedostępnej dla człowieka, przyczyniając się do dynamicznego rozwoju tej dziedziny oraz jej zastosowań w rzeczywistych scenariuszach. Autor w swojej rozprawie prezentuje cztery takie scenariusze zastosowania: handel energią elektryczną, analizę danych audio, optymalizację hiperparametrów w uczeniu maszynowym oraz tłumaczenie maszynowe.

Tematykę rozprawy oceniam jako niezwykle ważną i aktualną. Automatyzacja zadań za pomocą metod SI nie tylko odpowiada na bieżące potrzeby gospodarcze i społeczne, ale także kształtuje przyszłość wielu sektorów. Jej znaczenie będzie stale rosło, a badania w tej dziedzinie pozostaną kluczowe dla rozwoju technologicznego.

Tytuł

Tytuł rozprawy jest odpowiednio zwarty, komunikatywny i precyzyjny. Trafnie oddaje najistotniejsze elementy treściowe pracy. Wyraźnie wskazuje na problem badawczy (automatyzacja zadań) oraz główną metodologię (metody SI).

Spójność tematyczna artykułów

Artykuły [1]-[4] są spójne tematycznie, ponieważ wszystkie koncentrują się na automatyzacji zadań w różnych dziedzinach z wykorzystaniem metod SI. Artykuł [5] ma nieco inny charakter – opisuje nowy moduł bazowy sieci rekurencyjnych. Autor uzasadnia włączenie tego artykułu do rozprawy tym, że proponowana sieć może być wykorzystywana w różnych zastosowaniach opisanych w rozprawie, takich jak klasyfikacja sekwencji, tłumaczenie maszynowe lub modelowanie języka, przyczyniając się do poprawy efektywności metod SI stosowanych w automatyzacji zadań. Mimo tej argumentacji uważam, że włączenie artykułu [5] do rozprawy jest dyskusyjne. Należy jednak zauważyć, że nawet w przypadku jego pominięcia pozostałe artykuły tworzą spójną i wartościową rozprawę doktorską.

Problem badawczy i tezy

Problem badawczy został jasno sformułowany w podrozdz. 1.1 jako automatyzacja zadań przy użyciu metod SI w czterech obszarach: handlu energią elektryczną na rynku dnia następnego, symultanicznym tłumaczeniu maszynowym, wykrywaniu polskich słów kluczowych w nagraniach audio i dostrajaniu hiperparametrów w trakcie treningu sieci neuronowych.

Autor zdefiniował pięć tez badawczych powiązanych z poszczególnymi artykułami. Wszystkie tezy zostały zasadniczo poprawnie sformułowane, jednak teza 1 i teza 2 mają charakter zbyt ogólny i zyskałyby na czytelności, gdyby je doprecyzować. Tezy zostały przedstawione w sposób, który umożliwia ich weryfikację za pomocą odpowiednich eksperymentów, symulacji oraz standardowych miar oceny. Weryfikacja opiera się na zastosowaniu metod empirycznych i ilościowych, polegając głównie na porównaniu wyników osiągniętych przy użyciu proponowanych metod automatyzacji z wynikami uzyskanymi za pomocą metod konkurencyjnych.

Część wstępna

Część wstępna pracy ma logiczną i przejrzystą strukturę. Rozprawa zaczyna się wprowadzeniem, które jasno definiuje cel i tezy pracy, a także zawiera listę artykułów objętych rozprawą wraz z opisem wkładu Autora. W tym rozdziale Autor szczegółowo opisuje również, w jaki sposób zweryfikował postawione tezy. Rozdział poświęcony podstawom teoretycznym zapewnia tło merytoryczne, omawiając kluczowe pojęcia oraz metody istotne dla tematu pracy. W kolejnych pięciu rozdziałach szczegółowo omówiono artykuły wchodzące w skład rozprawy. Uwagi końcowe (rozdział 8) podsumowują osiągnięcia i główne wnioski płynące z rozprawy. Rozdział 9, dotyczący innych osiągnięć Autora, choć mniej istotny w ocenie rozprawy, podkreśla jego wszechstronność oraz szeroki zakres działalności naukowej.

Artykuł [1]

Artykuł przedstawia zastosowanie uczenia przez wzmocnienie (*reinforcement learning*) do zautomatyzowania handlu energią na rynku dnia następnego (*day-ahead*). Zakłada się, że agent handlowy może zużywać energię, wytwarzać ją ze źródeł odnawialnych oraz magazynować. Głównym celem jest wyznaczenie strategii składania ofert kupna i sprzedaży energii, która maksymalizuje zyski lub minimalizuje koszty. Problem ten jest kluczowy dla uczestników rynku w kontekście rosnącego udziału odnawialnych źródeł energii oraz magazynów energii w miksie energetycznym.

Nowatorskie podejście zaproponowane przez autorów polega na generowaniu kolekcji ofert bazujących na parametrycznych krzywych podaży i popytu. W procesie decyzyjnym wykorzystuje się dostępne informacje, takie jak prognozy pogody, poziom naładowania magazynu energii oraz dane kalendarzowe. Strategia jest optymalizowana za pomocą uczenia przez wzmocnienie, w którym nagrodą agenta jest suma przyszłych zdyskontowanych zysków. Autorzy proponują analityczne przedstawienie parametryzowanych krzywych podaży i popytu, które przekształcane są na funkcje schodkowe, co uwzględnia specyfikę składania ofert rynkowych. Parametry krzywych są uzależnione od akcji agenta, które zawierają komponent deterministyczny i stochastyczny – obydwa są produktem sieci neuronowej. Dzięki temu możliwe jest eksplorowanie różnych działań w trakcie uczenia, co jest istotne w algorytmach uczenia przez wzmocnienie.

Do optymalizacji strategii wykorzystano algorytm A2C. Autorzy przedstawili również wyniki uzyskane za pomocą innych popularnych algorytmów (PPO, SAC, TD3), co pozwala na szerszą ocenę efektywności proponowanego podejścia. Eksperymenty zostały przeprowadzone w specjalnie przygotowanym środowisku symulacyjnym. Nagroda uwzględniała zysk finansowy, zysk referencyjny oraz karę regularizacyjną (szczegółowo omówioną w dodatkach), która zapobiegała saturacji parametrów w modelu. W badaniach eksperymentalnych proponowana strategia osiągnęła najwyższe zyski, przewyższając strategię referencyjną we wszystkich analizowanych scenariuszach. Autorzy wskazują na kluczowe zalety strategii: racjonalne zarządzanie magazynem, efektywne składanie ofert oraz zdolność do wykorzystywania nieprzewidywalnych wahań cen na korzyść agenta (większy zakup, gdy ceny są niskie, oraz większa sprzedaż, gdy ceny są wysokie).

Wyniki badań zostały przedstawione w przejrzysty sposób, dokładnie przedyskutowane i wyciągnięto z nich poprawne wnioski. W artykule znajdują się również ilustracje prezentujące optymalne strategie oraz zarządzanie magazynem energii w różnych scenariuszach, co wzmacnia przekaz i umożliwia wizualną ocenę efektywności podejścia.

Uwagi i pytania:

- 1) W przyjętym podejściu zakłada się, że agent handlowy jest na tyle mały, iż nie wpływa na ceny rynkowe. Jak bardzo model by się skomplikował, gdyby uwzględniono wpływ agenta na ceny rynkowe? Czy Autor rozważał takie rozszerzenie?
- 2) Eksperymenty zostały przeprowadzone w środowisku symulacyjnym. Czy Autor może ocenić, w jakim stopniu wyniki w tym środowisku różnią się od wyników w rzeczywistych warunkach, biorąc pod uwagę ograniczenia modeli generacji energii z fotowoltaiki i wiatru, zapotrzebowania na energię oraz prognoz pogody opisanych w dodatkach C-E?

Artykuł [2]

Artykuł przedstawia autorski system rozpoznawania słów kluczowych w nagraniach audio, który ma na celu usprawnienie procesu przeszukiwania rozmów w języku polskim, takich jak nagrania z centrów obsługi klienta. Autorzy zwracają uwagę na problem niskiej dostępności zasobów językowych dla języków takich jak polski. Deficyt danych znacząco ogranicza skuteczność systemów automatycznego rozpoznawania mowy w porównaniu z językami o większych zasobach. Aby zredukować czasochłonne i kosztowne procesy adnotacji danych, autorzy proponują wykorzystanie metod dopasowywania wzorców akustycznych zamiast pełnej konwersji mowy na tekst. Argumenty za stosowaniem podejścia opartego na dopasowywaniu wzorców, szczegółowo omówione w podrozdziale 1.4, są przekonujące.

W celu rozpoznawania wzorców w nagraniach audio zastosowano konwolucyjne sieci neuronowe typu bliźniaczego (*Siamese*) oraz sieci prototypowe, które trenowano na różnych zbiorach nagrań w językach angielskim i polskim. Główną zaletą tego podejścia jest ograniczone zapotrzebowanie na dane – zarówno pod względem ilości, jak i jakości adnotacji. Sieci konwolucyjne mają za zadanie mapowanie par przykładów na pary wektorów osadzeń (*embeddings*). Proces optymalizacji dąży do tego, aby osadzenia przykładów należących do tej samej klasy (pasuje/nie pasuje) były umieszczane blisko siebie w przestrzeni osadzeń, podczas gdy osadzenia przykładów z różnych klas były od siebie oddalane. W modelach prototypowych klasy są reprezentowane przez tzw. prototypy, a porównanie odbywa się między nowymi przykładami a prototypami. Dla obu modeli zdefiniowano odpowiednie funkcje straty.

Wyniki badań pokazują, że proponowane modele dobrze sprawdzają się w języku angielskim, ale mają trudności w przypadku języka polskiego. W odpowiedzi na ten problem autorzy wdrożyli detektor oparty na osadzeniach audio generowanych przez wstępnie wytrenowany model udostępniony przez Google. Model ten lepiej radzi sobie w zadaniach międzyjęzykowych, jednak jego skuteczność na danych spoza zbioru treningowego pozostaje ograniczona. To wskazuje na trudności w transferze międzyjęzykowym, nawet przy wykorzystaniu zaawansowanych metod osadzania mowy. Ważnym elementem badań jest analiza zdolności modeli do generalizacji na dane spoza rozkładu, która uwidacznia problemy z przenoszeniem wiedzy między językami.

Podobnie jak w [1], wyniki badań zostały przedstawione w przejrzysty sposób, przedyskutowane, a wnioski logicznie wyciągnięte. Autorzy wskazali również kierunki dalszych badań, wykazując orientację w aktualnych trendach badawczych związanych z tą tematyką. Mimo że nie osiągnięto wszystkich oczekiwanych rezultatów, analiza wyników dostarcza cennych wskazówek dla przyszłych prac w zakresie rozpoznawania mowy w językach o ograniczonych zasobach.

Uwagi i pytania:

- 1) Zaproponowane rozwiązania opierają się na sieciach konwolucyjnych. Czy zdaniem Autora zastosowanie innych typów sieci, takich jak sieci rekurencyjne lub transformery, które są szczególnie przystosowane do przetwarzania danych sekwencyjnych (np. nagrań audio) i zawierają wbudowane mechanizmy przetwarzania tego typu danych, mogłoby poprawić uzyskiwane wyniki? Czy Autor rozważał modele bazujące na takich architekturach?
- 2) W podrozdziale 2.2.1 opisano architekturę sieci konwolucyjnej, podając konkretne wartości jej hiperparametrów. Jak przebiegał proces doboru tych hiperparametrów? Dlaczego zrezygnowano z dostrajania hiperparametrów do każdego zadania badawczego?

Artykuł [3]

Artykuł przedstawia metodę automatycznego dostrajania hiperparametrów, takich jak współczynnik uczenia (*step-size*) i współczynnik zaniku *momentum* (*momentum decay*), w procesie uczenia sieci neuronowych. Autorzy proponują algorytm, który dynamicznie dostosowuje te hiperparametry na podstawie analizy ich krótkoterminowego i długoterminowego wpływu na wartość funkcji straty. Algorytm ocenia wpływ wcześniejszych wartości hiperparametrów na bieżącą wartość parametrów modelu oraz ich długoterminowe zmiany (wykorzystując model wygładzania wykładniczego). Aby utrzymać pożądany trend parametrów zarówno w krótkim, jak i długim okresie, Autorzy wprowadzają wskaźnik jakości, minimalizowany względem hiperparametrów, uwzględniający składnik kary za fluktuacje parametrów. W rozdziale V definiują algorytm ASDM2, który rozwiązuje dodatkowe problemy, takie jak inicjalizacja, reparametryzacja, regulacja hiperparametrów oraz zapobieganie ich rozbieganiu.

Metoda została przetestowana na płytkich sieciach neuronowych, klasycznych autoenkoderach oraz autoenkoderach konwolucyjnych. Rozważano dwie implementacje automatycznego dostrajania hiperparametrów: w klasycznym algorytmie uczenia z *momentum* (CM) oraz w algorytmie ADAM. Wyniki porównano z popularnymi algorytmami (CM, ADAM, AdaGrad i AdaDelta), przy optymalizacji hiperparametrów w zewnętrznej pętli (*grid search*). Algorytm ASDM2 w większości przypadków przewyższał lub dorównywał metodom bez wbudowanego mechanizmu optymalizacji hiperparametrów. Jednak w pewnych sytuacjach, szczególnie przy małych wartościach parametru pamięci (γ), obserwowano suboptymalne zachowanie.

Dodatkowy proces optymalizacji hiperparametrów włączony w trening sieci ma swoje koszty obliczeniowe: wsteczna propagacja gradientu oraz obliczanie Hesjanu wykonywane są dwukrotnie w każdym kroku pętli treningowej. Powoduje to co najmniej trzykrotne wydłużenie czasu treningu w porównaniu z algorytmami bez wbudowanego mechanizmu optymalizacji hiperparametrów.

Artykuł [3] prezentuje pomysłowość autorów w podejściu do integracji optymalizacji parametrów i hiperparametrów w jednym procesie. Wyniki zostały przedstawione w sposób przejrzysty. Dokonano ich analizy i wyciągnięto poprawne wnioski. Zaproponowano dalsze badania nad poprawą stabilności algorytmu w różnych warunkach. Pomimo ograniczeń związanych z większą złożonością obliczeniową oraz suboptymalnością, proponowana metoda stanowi istotny wkład w dziedzinę optymalizacji procesu uczenia sieci neuronowych.

Uwagi i pytania:

- 1) Optymalizacja hiperparametrów jest kluczowym wyzwaniem w modelach uczenia maszynowego. Standardowe podejścia, takie jak *grid search* czy optymalizacja bayesowska, umożliwiają optymalizację zarówno hiperparametrów ciągłych, jak i porządkowych oraz nominalnych. Czy Autor rozważał możliwość integracji optymalizacji tego typu hiperparametrów bezpośrednio z procesem treningu? Jakim wyzwaniom należałoby sprostać, aby zapewnić skuteczność takiej integracji?
- 2) Proponowana metoda automatycznie optymalizuje dwa kluczowe hiperparametry w procesie uczenia sieci neuronowych, jednocześnie wymagając określenia kilku dodatkowych hiperparametrów, które mają istotny wpływ na skuteczność tego procesu (siedem z nich wymieniono w rozdziale V F). Czy to nie prowadzi do paradoksu, gdzie eliminacja dwóch

hiperparametrów wprowadza konieczność zarządzania większą liczbą nowych hiperparametrów?

Artykuł [4]

Artykuł dotyczy zastosowania uczenia ze wzmocnieniem do problemu tłumaczenia maszynowego w trybie symultanicznym (*simultaneous machine translation*, SMT). W przeciwieństwie do tradycyjnych metod tłumaczenia neuronowego, które wymagają przetworzenia całego wejściowego ciągu danych przed wygenerowaniem wyniku, SMT generuje wyjściowe elementy sekwencji na bieżąco, co wymaga podejmowania decyzji w czasie rzeczywistym. Kluczowym wyzwaniem jest znalezienie kompromisu między opóźnieniem tłumaczenia a jego jakością.

Autorzy proponują architekturę opartą na uczeniu ze wzmocnieniem, nazwaną RLST (*reinforcement learning for on-line sequence transformation*), która iteracyjnie podejmuje decyzje o odczycie kolejnych tokenów wejściowych lub generowaniu tokenów wyjściowych. Polityka agenta określa, w jaki sposób wybiera on akcje i generuje tokeny wyjściowe na podstawie tokenów dotychczas odczytanych i wygenerowanych. Celem agenta jest maksymalizacja oczekiwanej wartości zdyskontowanych nagród w przyszłości, które zależą od jakości tłumaczenia oraz poprawności podejmowanych akcji. Agent ma dostęp jedynie do odczytanych i wygenerowanych tokenów, co ogranicza pełną wiedzę o wejściowej sekwencji. Jest to zgodne z rzeczywistymi scenariuszami, w których decyzje muszą być podejmowane na podstawie niepełnych danych. System wykorzystuje politykę opartą na rekurencyjnej sieci neuronowej (RNN), która uczy się podejmowania decyzji sekwencyjnych, minimalizując jednocześnie błąd tłumaczenia.

W eksperymentach przeprowadzonych na siedmiu zadaniach tłumaczenia maszynowego metoda RLST osiągnęła wyniki porównywalne z nowoczesnymi metodami opartymi na transformerach i modelach typu enkoder-dekoder z mechanizmem uwagi. Szczególnie dobrze sprawdziła się w tłumaczeniu długich zdań, przewyższając inne podejścia. Autorzy zauważyli, że stan pamięci agenta tłumaczącego skuteczniej zachowuje kontekst generowanych słów niż mechanizmy uwagi w architekturach referencyjnych. Z drugiej strony, RLST nie jest najlepszą metodą do tłumaczenia zdań o umiarkowanej długości, które nie mają dodatkowego kontekstu. Artykuł podkreśla potencjał RLST w aplikacjach wymagających tłumaczenia sekwencji o dowolnej długości.

Głównym osiągnięciem Autorów jest opracowanie nowej architektury tłumaczenia symultanicznego, która umożliwia przekształcanie sekwencji online bez konieczności wcześniejszego definiowania kompromisu między opóźnieniem a jakością. Ta innowacyjna architektura, wyposażona w starannie zaprojektowaną funkcję straty (14), uczy się podejmowania decyzji za pomocą uczenia ze wzmocnieniem. Artykuł został napisany w sposób przejrzysty, zarówno w części teoretycznej, jak i eksperymentalnej. Wnioski zostały poparte wynikami eksperymentów.

Uwagi i pytania:

- 1) W rozdziale V wspomniano o manualnej optymalizacji hiperparametrów modeli. Na czym dokładnie polegał ten proces? Czy RLST wymaga dobrania większej liczby hiperparametrów niż modele referencyjne? Które z hiperparametrów RLST mają największy wpływ na jego wydajność i jak ich wybór wpływa na wyniki tłumaczenia?
- 2) Czy istnieją scenariusze, w których RLST osiąga wyraźnie gorsze wyniki niż modele referencyjne, np. w przypadku tłumaczenia idiomów, metafor lub innych konstrukcji

językowych wymagających globalnego kontekstu? Jeśli tak, to jakie są przyczyny takich ograniczeń i czy można je przezwyciężyć?

Artykuł [5]

Artykuł wprowadza nową architekturę sieci rekurencyjnej opartą na module bazowym o nazwie *Deep Memory Update* (DMU). Celem tej architektury jest rozwiązanie problemów eksplozji i zaniku gradientów, które znacząco utrudniają trenowanie standardowych sieci rekurencyjnych. W odróżnieniu od istniejących rozwiązań, takich jak *Long Short-Term Memory* (LSTM) czy *Gated Recurrent Unit* (GRU), DMU umożliwia elastyczną nieliniową transformację stanu pamięci oraz wektora wejściowego, przy jednoczesnym zachowaniu stabilności i zwiększeniu szybkości procesu uczenia.

Standardowe rozwiązania wykorzystują kilka rodzajów bramek, które przekształcają stan pamięci i dane wejściowe za pomocą pojedynczych warstw neuronów. Autorzy zamiast tego wprowadzają wielowarstwową sieć typu FNN (*feedforward neural network*), która generuje dwa wektory służące do modyfikacji stanu pamięci poprzez jego redukcję i wzbogacenie. Dzięki głębokim, nieliniowym transformacjom FNN oferuje większe możliwości w zakresie kształtowania wektorów modyfikujących w porównaniu z klasycznymi bramkami. W artykule skrupulatnie opisano architekturę DMU, proces inicjalizacji oraz sposób uczenia proponowanej sieci rekurencyjnej. Autorzy, analizując propagację gradientu w module DMU, wykazali, że przy spełnieniu określonych warunków gradient pozostaje stabilny i nie rozbiega się.

W badaniach eksperymentalnych DMU zostało porównane z innymi architekturami, takimi jak RNN, LSTM, GRU i *Recurrent Highway Networks* (RHN), na zadaniach syntetycznych (np. dodawanie liczb, modelowanie sekwencji) oraz rzeczywistych (modelowanie muzyki polifonicznej, modelowanie języka naturalnego, tłumaczenie maszynowe). Wyniki wskazują, że DMU często przewyższa inne metody, zwłaszcza w zadaniach wymagających głębokich transformacji stanu sieci. Stabilność treningu w DMU została zapewniona dzięki odpowiedniemu doborowi współczynnika uczenia.

Przedstawiona w artykule propozycja nowej architektury sieci rekurencyjnej stanowi istotny wkład w rozwój sieci neuronowych. Artykuł został napisany rzetelnie – model został poprawnie zdefiniowany, a badania eksperymentalne potwierdzają jego wysoką efektywność w szerokim zakresie zastosowań.

Uwagi i pytania:

- 1) W rozdz. V autorzy wspominają o tym, że zwiększanie głębokości modelu nie przyniosło spodziewanej poprawy wyników i spekulują, że może to wynikać z braku zaawansowanych mechanizmów regularyzacji (używają jedynie *weight decay*). Jakie metody regularyzacji mogłyby być potencjalnie skuteczne w tym przypadku?
- 2) Sieci rekurencyjne często osiągają gorsze wyniki, gdy dane są ograniczone (np. krótkie sekwencje lub niewielka liczba próbek). Czy podobne ograniczenia obserwowane są również w przypadku DMU? Jeśli tak, jakie strategie mogłyby złagodzić ten problem?
- 3) Teza 5 wspomina o arbitralnej transformacji („*arbitrary transformation*”) oraz o efektywnym wykorzystaniu pamięci („*memory-efficient*”). Jak należy rozumieć arbitralność transformacji? W jaki sposób wykazano efektywne wykorzystanie pamięci w DMU?

Podsumowanie

Artykuły [1]-[5] cechują się poprawną strukturą, typową dla prac naukowych. Każda sekcja jest logicznie uporządkowana, co ułatwia czytanie i zrozumienie treści. Problematyka badawcza oraz cele są jasno przedstawione, a metodyka szczegółowo opisana, co umożliwia odtworzenie przeprowadzonych badań. Autorzy wykazali się solidnym warsztatem naukowym, wykorzystując zaawansowane narzędzia oparte na sieciach neuronowych i proponując innowacyjne rozwiązania postawionych problemów. Wnioski zostały rzetelnie poparte wynikami eksperymentów, a źródła literaturowe dobrano adekwatnie do tematyki. Jedynym zauważonym mankamentem w analizie rezultatów jest pominięcie odpowiednich testów statystycznych, które mogłyby obiektywnie porównać wyniki uzyskane przez różne modele.

Prace [1] oraz [3]-[5] zostały zaprezentowane na prestiżowych konferencjach międzynarodowych, a artykuł [2] opublikowano w czasopiśmie o przyzwoitym poziomie naukowym ($IF=3,4$). Udział Autora w tych pracach jest znaczący, od 20% do 75%. Jego wkład obejmuje opracowanie głównych rozwiązań w pracach [1] i [3], testowanie i analizę proponowanych rozwiązań w [2] i [4], a także opracowanie modeli referencyjnych w pracy [5]. Uznaję ten wkład za wystarczający, by ubiegać się o stopień doktora.

Nie mam wątpliwości, że Autor skutecznie rozwiązał postawione problemy, osiągnął założone cele i obronił tezy, stosując odpowiednie metody badawcze. Oryginalność rozprawy polega na opracowaniu nowych metod automatyzacji rozwiązywania różnych typów zadań z wykorzystaniem sztucznej inteligencji, w szczególności sieci neuronowych. Praca czerpie z najnowszych osiągnięć naukowych opisanych w światowej literaturze. Autor wykazał się zdolnością do poprawnego i przekonującego przedstawienia uzyskanych wyników. Oceniam wysoko znaczenie rezultatów pracy dla rozwoju dyscypliny informatyka techniczna i telekomunikacja, szczególnie w kontekście ich potencjalnych zastosowań.

Wniosek końcowy

Rozprawa doktorska mgr. inż. Łukasza Lepaka stanowi oryginalne rozwiązanie problemu naukowego i wskazuje na wysoki poziom wiedzy teoretycznej Autora z dyscypliny informatyka techniczna i telekomunikacja, a także na umiejętność samodzielnego prowadzenia przez niego pracy naukowej.

Stwierdzam, że opiniowana rozprawa doktorska spełnia wymogi ustawy z 20 lipca 2018 r. *Prawo o szkolnictwie wyższym i nauce*. Oceniam ją pozytywnie. Wnoszę o dopuszczenie mgr. inż. Łukasza Lepaka do publicznej obrony pracy doktorskiej.

