

**RADA NAUKOWA DYSCYPLINY
INFORMATYKA TECHNICZNA I TELEKOMUNIKACJA POLITECHNIKI WARSZAWSKIEJ**

zaprasza na
OBRONĘ ROZPRAWY DOKTORSKIEJ

mgr. inż. Łukasza Eugeniusza Lepaka

która odbędzie się w dniu **7 marca 2025 roku**, o godzinie **11:00** w trybie łączonym (stacjonarnie równolegle ze zdalnie)

Temat rozprawy:

„Task Automation Using Artificial Intelligence Methods”

Promotor: dr hab. inż. Paweł Wawrzyński

Recenzenci: prof. dr hab. inż. Grzegorz Dudek – Politechnika Częstochowska

prof. dr hab. inż. Jacek Rumiński – Politechnika Gdańska

dr hab. inż. Wojciech Kotłowski – Politechnika Poznańska

Obrona odbędzie się na **Wydziale Elektroniki i Technik Informatycznych Politechniki Warszawskiej, ul. Nowowiejska 15/19, 00-665 Warszawa w sali AC** oraz jednocześnie zdalnie na platformie MS Teams. Osoby zainteresowane uczestnictwem w obronie w formie zdalnej proszone są o zgłoszenie chęci uczestnictwa w formie elektronicznej na adres sekretarza komisji: dr. hab. inż. Rafał Biedrzycki, email : rafal.biedrzycki@pw.edu.pl, do dnia 4.03.2025 r., godz. 17:00.

Z rozprawą doktorską i recenzjami można zapoznać się w Czytelnicy Biblioteki Głównej Politechniki Warszawskiej, Warszawa, Plac Politechniki 1.

Streszczenie rozprawy doktorskiej i recenzje są zamieszczone na stronie internetowej: <https://www.bip.pw.edu.pl/Postepowania-w-sprawie-nadania-stopnia-naukowego/Doktoraty/Wszczete-po-30-kwietnia-2019-r/Rada-Naukowa-Dyscypliny-Informatyka-Tehniczna-i-Telekomunikacja/mgr-inz.-Lukasz-Eugeniusz-Lepak>

Przewodniczący Rady Naukowej Dyscypliny
Informatyka Techniczna i Telekomunikacja
Politechniki Warszawskiej
prof. dr hab. inż. Jarosław Arabas

Automatyzacja zadań z wykorzystaniem metod sztucznej inteligencji

Niniejsza rozprawa doktorska przedstawia serię pięciu publikacji dotyczących automatyzacji zadań z wykorzystaniem metod sztucznej inteligencji.

Przedstawiamy sposoby automatyzacji czterech różnych zadań oparte na metodach sztucznej inteligencji. W handlu energią elektryczną na rynku dnia następnego wykorzystujemy uczenie ze wzmocnieniem do stworzenia automatycznej strategii handlu, która generuje zbiór ofert w oparciu o dostępne informacje istotne z punktu widzenia handlu. Pokazujemy, że strategia ta pozwala agentowi maksymalizować zyski. Ponadto, może być ona stosowana w rzeczywistych zastosowaniach. W przypadku wykrywania słów kluczowych w nagraniach audio testujemy kilka architektur dopasowywania podobnych obiektów na różnych zbiorach danych i pokazujemy, że możliwe jest stworzenie detektora słów kluczowych dla języków z małą ilością dostępnych danych. W przypadku optymalizacji hiperparametrów on-line w metodach uczenia sieci neuronowych, proponujemy nowy algorytm, który optymalizuje wartości hiperparametrów podczas procesu uczenia się; eksperymenty prezentują jego skuteczność w porównaniu z innymi popularnymi algorytmami uczenia. W przypadku symultanicznego tłumaczenia maszynowego proponujemy architekturę uczoną ze wzmocnieniem, która automatycznie kontroluje opóźnienie tłumaczenia przy zachowaniu jego wysokiej jakości. Porównujemy ją z najnowocześniejszymi architekturami neuronowego tłumaczenia maszynowego, osiągając podobne wyniki przy ograniczonym kontekście.

Bazując na obserwacjach z zadania tłumaczenia maszynowego, prezentujemy także nową komórkę rekurencyjną, która pozwala na głęboką transformację stanu i jest odporna na problemy związane z propagacją gradientu. Przeprowadzone eksperymenty pokazują, że proponowana komórka osiąga porównywalne, a często lepsze wyniki niż popularne komórki rekurencyjne przy użyciu mniejszej liczby parametrów.

Podsumowując, przedstawiamy cztery automatyzacje oparte na sztucznej inteligencji, co potwierdza ich ogólność i użyteczność w różnych zadaniach i dziedzinach, a także nową jednostkę rekurencyjną, która może być używana jako część takich automatyzacji.

Słowa kluczowe: uczenie ze wzmocnieniem, sieci neuronowe, automatyzacja zadań, zastosowania sztucznej inteligencji

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.

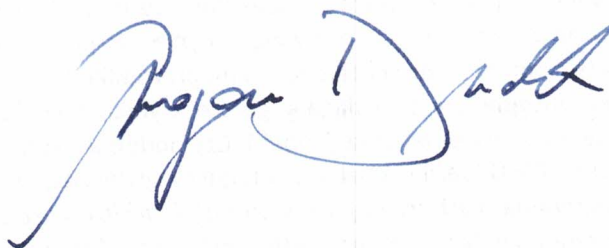
Prace [1] oraz [3]-[5] zostały zaprezentowane na prestiżowych konferencjach międzynarodowych, a artykuł [2] opublikowano w czasopiśmie o przyzwoitym poziomie naukowym ($IF=3,4$). Udział Autora w tych pracach jest znaczący, od 20% do 75%. Jego wkład obejmuje opracowanie głównych rozwiązań w pracach [1] 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.



Prof. dr hab. inż. Jacek Rumiński,
Katedra Inżynierii Biomedycznej
Wydział Elektroniki, Telekomunikacji i Informatyki
Politechnika Gdańska

Gdańsk, 10.12.2024 r.

Recenzja

pracy doktorskiej mgr. inż. Łukasza Lepaka

pt. „Task Automation Using Artificial Intelligence Methods”.

Promotor: dr hab. inż. Paweł Wawrzyński

Recenzja została przygotowana w odpowiedzi na pismo Rady Dyscypliny Informatyka Techniczna i Telekomunikacja, Politechniki Warszawskiej (z dnia 25 października 2024 r.) z informacją, że na podstawie uchwały Rady nr 739/2024 (podpisanej przez Przewodniczącego Rady, prof. dr. hab. inż. Jarosława Arabasa) zostałem wyznaczony jako recenzent rozprawy doktorskiej Pana mgr. inż. Łukasza Lepaka.

1. Tematyka, cele i tezy rozprawy

Tematyka pracy dotyczy szeroko rozumianej automatyzacji zadań z wykorzystaniem sztucznej inteligencji. W szczególności związana jest z rozwiązywaniem bardzo różnych problemów w obszarze AI, które mogą być zastosowane: w handlu energią elektryczną, detekcji słów kluczowych w nagraniach audio, symultanicznym tłumaczeniu maszynowym oraz dostrajaniu wartości hiperparametrów w trakcie procesu uczenia.

Częścią rozprawy jest seria pięciu artykułów. Całość jest napisana w języku angielskim, dlatego cytowane fragmenty w niniejszej recenzji będą przytaczane w oryginale.

Doktorant sformułował pięć hipotez:

H1: „*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.*”

H2: „*Finding keywords in call center recordings in a low-resource language is possible using audio similarity matching models.*”

H3: „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.”

H4: „The translation delay in a simultaneous machine translation can be automatically controlled while maintaining high translation quality for sequences of arbitrary length.”

H5: „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.”

Hipotezy są interesujące i wskazują ogólne problemy badawcze. Szkoda, że nie zostały uszczegółowione (m.in. poprzez precyzyjnie określone założenia), ponieważ tak sformułowane mogą być oczywiste w świetle aktualnego stanu wiedzy. Przykładowo postawione w H2 stwierdzenie, że coś jest możliwe, jest hipotetycznie zawsze realizowalne w świetle stanu wiedzy i w przedstawionym kontekście (nie jest uszczegółowiona np. czułość czy precyzja).

2. Zawartość rozprawy

Rozprawa napisana jest w języku angielskim i składa się z dziewięciu rozdziałów, spisu literatury oraz załączników, spośród których kluczowym jest zbiór pięciu publikacji.

W rozdziale pierwszym Autor wprowadza do tematyki rozprawy oraz przedstawia i krótko komentuje hipotezy badawcze.

Rozdział drugi prezentuje wybrane i podstawowe aspekty teoretyczne dotyczące uczenia ze wzmocnieniem i sztucznych sieci neuronowych.

Każdy z rozdziałów 3-7 zawiera krótki opis aspektów badawczych i wyników związanych z powiązanymi pięcioma artykułami, związanymi z tematyką: 1) automatyzacją handlu energią z jednodniową prognozą wyprzedzającą, 2) detekcji polskich słów w zapisach audio, 3) strojenia hiperparametrów w trakcie uczenia, 4) symultanicznego tłumaczenia maszynowego oraz 5) głębokiej transformacji stanów w sieciach rekurencyjnych.

Rozdział ósmy przedstawia podsumowanie rozprawy, a ostatni rozdział dziewiąty zawiera opis innych osiągnięć Doktoranta (wystąpienia konferencyjne, projekty).

Autor cytuje 79 pozycji literatury. Większość cytowanych prac, to recenzowane publikacje w czasopismach lub materiałach konferencyjnych. Wykaz zawiera jedynie 7 pozycji, w tym jeden raport techniczny z ostatnich trzech lat (2022-2024).

Prezentując własne osiągnięcia Autor wskazuje m.in. na publikacje w materiałach konferencyjnych 2023 International Joint Conference on Neural Networks (IJCNN), przedstawiając punktację Ministerstwa jako 140 punktów. Taka punktacja widnieje w wykazie publikacji w Politechnice Warszawskiej. Zgodnie z komunikatem Ministra Edukacji i Nauki z dnia 03 listopada 2023 r. „o zmianie i sprostowaniu komunikatu w sprawie wykazu czasopism naukowych i recenzowanych materiałów z konferencji międzynarodowych”, w załączonym wykazie dla materiałów konferencji IEEE International Joint Conference on Neural Networks [IJCNN] przypisano 70 punktów. Ta sama punktacja była w wcześniejszym wykazie z 2023 roku. Wymaga to wyjaśnienia (może były jeszcze inne korekty, o których nie jestem świadom). Nie jest to istotne dla oceny treści merytorycznych rozprawy, ale stanowi uwagę porządkową.

3. Oryginalne osiągnięcia

W pracy podjęto różnorodne problemy naukowe z zakresu sztucznej inteligencji. Do szczególnych, oryginalnych osiągnięć zaliczam:

- O1: Opracowanie i weryfikacja zautomatyzowanej strategii dla metody uczenia ze wzmocnieniem możliwej (w pewnym zakresie) do wykorzystania w handlu energią na Rynku Dnia Następnego (z jednodniową prognozą wyprzedzającą), w której informacje dostępne dla agenta są przetwarzane w krzywe popytu i podaży z wykorzystaniem zbiorów ofert o zmiennym rozmiarze, umożliwiając w ten sposób agentowi racjonalne zachowanie.
- O2: Wykazanie, że detekcja słów kluczowych w języku polskim w niskiej jakości zapisach audio jest częściowo możliwa z zastosowaniem reprezentacji generowanych na bazie modeli przetrenowanych na zbiorze w języku angielskim, szczególnie w przypadku zastosowania zaproponowanej metody post-processingu. Warto podkreślić są bardzo ciekawe eksperymenty oraz udział Doktoranta w przygotowaniu zbioru danych.
- O3: Udział w opracowaniu algorytmu „Autonomous Stochastic Descent with Momentum v2 (ASDM2)” w wersjach bazujących na z normalizacją (ASDM2/n, na bazie metody ADAM) i bez normalizacji gradientu (ASDM2/b, na bazie techniki momentum), a także udział w stosownej weryfikacji algorytmów.

Ponadto pozytywnie oceniam inne formy osiągnięć związanych z rozwiązywanymi problemami badawczymi:

- I1: Rozwój i testowanie agenta RLST (Reinforcement Learning for on-line Sequence Transformation), zaprojektowanego do transformacji sekwencji o dowolnej długości. W szczególności uzyskanie lepszych wyników dla proponowanego rozwiązania względem metod referencyjnych dla dłuższych sekwencji (> 22 słów, Tatoeba), co opisano w publikacji P4.
- I2: Eksperymentalna weryfikacja modułu DMU (Deep Memory Update), poprzez opracowanie modeli referencyjnych i przeprowadzenie testów opisanych w publikacji P5.

Wyżej wymienione osiągnięcia (I1 i I2) nie zostały przeze mnie włączone do grupy szczególnych i oryginalnych osiągnięć doktoranta, ponieważ nie znalazłem wystarczających dowodów w przedstawionej dokumentacji w zakresie oryginalnego wkładu Doktoranta w autorstwo proponowanych rozwiązań. W przedstawionej dokumentacji widnieje jedynie:

P4: „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.”

P5: „The PhD candidate developed models to which the proposed solution was compared, and conducted part of the experiments.”

Powyższe osiągnięcia mieszczą się w obszarze dyscypliny informatyka techniczna i telekomunikacja i wskazują umiejętność samodzielnego prowadzenia pracy naukowej przez Doktoranta. Ponadto, opisane osiągnięcia na tle stanu wiedzy demonstrują, że rozprawa doktorska stanowi oryginalne rozwiązanie problemu naukowego.

4. Ocena szczegółowa i pytania

Przedstawiona do oceny rozprawa doktorska bazuje na pięciu publikacjach. Są one związane z dyscypliną informatyka techniczna i telekomunikacja i dotyczą aspektów sztucznej inteligencji. Niemniej są bardzo różnorodne tematycznie, o niejednorodnym wkładzie Doktoranta, wpływając na spójność rozprawy. Każdy z przedstawionych problemów badawczych jest ciekawy, aczkolwiek różnorodny pod względem stosowanych metod i zastosowań.

Przedstawiona dokumentacja w rozprawie doktorskiej pozwala mi przede wszystkim skupić się na hipotezach H1, H2 i H3 jako tych, których podjęcie bezpośrednio związane jest z wykazaniem umiejętności Doktoranta w zakresie samodzielnego prowadzenia pracy naukowej.

Praca badawcza związana z H1 jest bardzo szeroko udokumentowana zarówno w artykule jak i obszernym załączniku. Doktorant jest głównym autorem pracy, a opis wkładu merytorycznego nie pozostawia wątpliwości o jego zasadniczym udziale w prowadzeniu tej pracy naukowej. Autor ciekawie rozważył funkcję opisującą działanie agenta (wzór (10), publikacja P1), uwzględniającą reprezentacje sieci neuronowej (g^1 i g^2) dla określonych stanów analizując zmienne kontrolowane i niekontrolowane. W eksperymentach wykorzystał dostępne repozytorium Stable Baselines3. Zaproponowaną strategię zbadał w odniesieniu do innych, potencjalnie prostszych: strategii arbitralnej, bazującej wiedzy dotyczącej historycznie najlepszych godzin w zakresie zakupu i sprzedaży energii oraz strategii godzinowej, w ramach której w każdej godzinie składane są oferty zakupu i sprzedaży.

Kluczowym wkładem autora jest opracowana strategia i zaprojektowanie oraz przeprowadzenie stosownych eksperymentów. Mam jednak kilka uwag i pytań związanych zarówno ze strategią jak i eksperymentami.

Po pierwsze, Doktorant porównał wyniki zaproponowanej przez siebie strategii głównie do dwóch innych strategii, również zdefiniowanych przez Autora rozprawy. Szkoda, że nie odniósł się do innych metod proponowanych z literatury. Szkoda również, że Doktorant nie przedstawił wyników weryfikacji strategii najpierw na takich danych, na których można byłoby porównać znaczenie opracowanej metody z innymi, znanymi ze stanu wiedzy (a dopiero później pokazać dla danych krajowych, jako przypadku szczególnego).

Sformułowana strategia bazuje na prostej polityce zakładającej, że działanie w chwili t jest określane bezpośrednio przez wyjście sieci neuronowej (wzór 10) dla określonego stanu s . W związku z tym mam pytania:

P1. Jak dobór parametrów sieci neuronowej wpływał na uzyskiwane wyniki (np., dlaczego 300 neuronów)?

P2. Jaki wpływ na uzyskiwane wyniki miało dodanie szumu? Dlaczego jest to szum gaussowski?

Zakładam, że propozycja określonej strategii ma mieć charakter uniwersalny (przy określonych założeniach) i sprawdzać się nie tylko na danych historycznych (wówczas jest to jedynie wartość poznawcza), ale również w przyszłości. Dlatego, mam kolejne pytanie:

P3. Czy doktorant sprawdzał, jakie wyniki uzyska z zastosowaniem zaproponowanego podejścia dla nowszych danych, tj. po 2019 roku? Warto zauważyć, że od 2020 roku średnio, wartość maksymalna ceny energii w Polsce była notowana w okolicach godziny 19.00, nie 10.00.

Ponadto, nie jest dla mnie jasne, czy Doktorant rozważał wpływ innych agentów, biorących udział w aukcjach. Wydaje się to bardzo ważne.

Wyniki umieszczone w tabeli 7 załącznika do publikacji P1 wskazują, że kluczowy jest dobór algorytmów RL na uzyskiwane wyniki. Stąd kolejne pytanie:

P4: Czy Doktorant przeprowadzał eksperymenty, wskazujące, czy ważniejszy jest dobór strategii czy algorytmu RL?

Doktorant przyjmuje na wejściu 117 obserwacji, z czego 72 wartości związane są ze stanem pogody. Stąd pojawia się pytanie:

P5. Czy Doktorant zbadał wpływ parametrów pogody na uzyskiwane wyniki?

Ponadto mam następujące pytanie:

P6: Czy inne zmienne ekstremalne (np. wojna, katastrofy) mogą być uwzględnione w niepewności określonej przez zaproponowaną strategię lub co trzeba byłoby zrobić, aby to uwzględnić?

Hipoteza H1 jest zdefiniowana bardzo ogólnie i dla takiej jej postaci Doktorant dostarczył argumentów za jej potwierdzeniem.

Kolejna praca, związana z hipotezą H2, skupia się na wykorzystaniu dwóch modeli sieci neuronowych do uczenia reprezentacji danych w celu detekcji słów kluczowych w niskiej jakości zapisach audio wraz z ich weryfikacją na zbiorach danych w języku angielskim i polskim. Wytrenowany model dla zbioru języka angielskiego osiąga gorsze wyniki niż większość modeli znanych z literatury (tabela 3, publikacja P2). Autorzy w artykule stwierdzili, że „ (...) we deemed the differences not significant enough to warrant the trade-off these models introduce in terms of complexity and training time.”. Nie mogę się jednak zgodzić z tym wnioskiem, przy braku odpowiednio szczegółowych założeń w zakresie pracy badawczej związanej z P2 i hipotezą H2. Niemniej model ten jest wykorzystywany jako baza w szeregu kolejnych eksperymentów i dlatego mogę przyjąć, że sam w sobie stanowi założenie kolejnych eksperymentów.

W dalszych eksperymentach Doktorant przeprowadza ocenę detekcji słów kluczowych dla języka angielskiego i polskiego. Są to bardzo interesujące analizy. Szczególnie ciekawe są analizy przekrojowe, tj. w sytuacji, gdy model był trenowany na mieszanych zbiorach danych. W przypadku eksperymentów dla pojedynczego języka wyraźnie widać różnicę jakości detekcji słów kluczowych pomiędzy językiem polskim, a angielskim. Dla słów kluczowych w języku polskim, z wykorzystaniem syntetycznego zbioru ($KW_{PL,mix}$) wartość F-score jest ponad czterokrotnie niższa niż dla eksperymentów ze zbiorem z języka angielskiego (F-score= 0,22

vs. 0,94 w tabeli 4, P2). Dla rzeczywistego zbioru ewaluacyjnego ($KW_{PL,real}$) najlepsza wartość $F\text{-score}=0,02$. To oczywiście nieakceptowalne wyniki. W świetle oceny pracy Doktoranta pojawiają się następujące pytania:

P7. Proszę o wyjaśnienie, dlaczego w opisach eksperymentów i tabelach widnieje notacja wskazująca na 22 słowa kluczowe (np. tabela 5: „3: Evaluation is performed on Polish mixtures ($KW_{PL,mix}$) and real recordings ($KW_{PL,real}$) with 22 keywords.”), podczas gdy w tekście czytamy: „The dataset contains ten examples for each of the 22 target keywords, although only 19 keywords are actually present in the evaluation recordings.”?

P8. Jaką miarę $F\text{-score}$ zastosowano w pracy?

Klasyczna definicja to $F\text{-score} = 2 * (\text{precision} * \text{recall}) / (\text{precision} + \text{recall})$, podczas gdy w pracy P2 jawnie wskazano, że „The F-score is defined as $F = \text{precision} + \text{recall} / 2$, (...)”? Taka niestandardowa definicja powoduje, że niska wartość $F\text{-score}$ wskazuje na jeszcze mniejszą jakość, niż w przypadku klasycznej definicji. Uważam, że taki dobór miary $F\text{-score}$ nie jest właściwy. Co więcej, w tabeli 7 (P2) Autor wskazał wartości „micro” dla precyzji i czułości jako np. 0,4% i 18,5% i powiązana wartość $F\text{-score}$: 0,01. Nie wiem jak ta wartość (i inne z tabeli 7) zostały wyliczone. Klasyczna interpretacja trybu „micro” oznacza uzyskanie wartości globalnych precyzji i czułości. Zatem $F\text{-score}$ liczone według wzoru z P2 powinno wynosić: $(0,004+0,185)/2=0,189/2=0,0945$; natomiast liczone według klasycznej definicji $F\text{-score}=2*(0,004*0,185)/(0,004+0,185)=0,00148/0,189=0,00783$. Proszę o wyjaśnienie jak wyznaczono wartość $F\text{-score}$.

Przeprowadzone eksperymenty w trybie mieszanym wskazują, że np. model (Siamese) wytrenowany na zbiorze w języku angielskim daje lepsze wyniki ewaluacji na zbiorze w języku polskim, niż ten samo model wytrenowany albo na zbiorze z językiem polskim, albo na połączonych zbiorach w obu językach. Jest to bardzo zastanawiające. We wszystkich przypadkach wskazywane wartości $F\text{-score}$ są bardzo niskie. Najlepsze wyniki uzyskano, gdy zastosowano przetrenowany model Google (pozycja [6] literatury w P2) do generacji reprezentacji danych, wykorzystanych w detekcji polskich słów kluczowych ($KW_{PL,mix}$). Uzyskana wartość $F\text{-score}$ to 0,69. Stąd pojawia się kolejne pytanie do Doktoranta:

P9. Czy wyniki uzyskane dla przetrenowanego modelu Google (tabela 6, praca P2), wskazują, że wcześniej rozważane modele są bezużyteczne? Czy nie należało najpierw przeprowadzić eksperymentu z dostępnymi modelami ze stanu wiedzy, a dopiero potem podejmować decyzję o ewentualnym projektowaniu/wyborze modelu bazowego do generacji reprezentacji?

Podsumowując, wyniki przeprowadzonych eksperymentów pokazane w publikacji P2, pozwalają w pewnym sensie potwierdzić hipotezę H2: „*Finding keywords in call center recordings in a low-resource language is possible using audio similarity matching models*”. Jest to jednak możliwe tylko dlatego, że hipoteza została zdefiniowana zbyt nieprecyzyjnie i jest właściwie, w zakresie stanu wiedzy, oczywista.

Trzecia hipoteza podjęta została w publikacji P3. Doktorant jest trzecim autorem tej pracy, ale w rozprawie wskazał swój udział jako: „developed, tested, and analyzed different versions of

the proposed algorithm". Na tej podstawie zakładam, że również poprzez adresowanie H3 w P3 Doktorant zdobywał kompetencje do formułowania problemów badawczych, stawiania hipotez, dobierania i stosowania właściwej metodologii a także analizowania wyników prowadzonych badań. Zaproponowany algorytm (i jego wersje) są rozwinięciem pracy promotora z 2017 roku (na co jawnie wskazano w rozprawie), w którym optymalizowane są wartości wag algorytmów bazowych typu ADAM i SGD z techniką momentum. Wprowadzone modyfikacje są interesujące i wykazano eksperymentalnie ich przydatność, co oceniam pozytywnie. Mam jednak dwa pytania:

P10: Czy początkowa wartość kroku β (krok 3 algorytmu) może być rzeczywiście dowolna i jak wpływa na przebieg optymalizacji?

P11: W algorytmie ADAM, momentum definiowane jest jako: $m_t = \beta_1 m_{t-1} + (1 - \beta_1) g_t$, lub po uwzględnieniu współczynnika szybkości uczenia (kroku): $\alpha \cdot m_t = (\alpha \beta_1) m_{t-1} + (\alpha(1 - \beta_1)) g_t$. Biorąc pod uwagę notację w równaniu 2 w publikacji P3 oraz przebieg proponowanego algorytmu, czy w przeprowadzonych eksperymentach obserwowana była relacja pomiędzy λ_t i β_t ?

Powyższe uwagi i pytania wskazują pewne problemy dotyczące wybranych aspektów rozprawy. Będę wdzięczny za ustosunkowanie się do nich w przypadku publicznej obrony.

5. Konkluzja recenzji

Podsumowując, stwierdzam, że przedstawione w rozprawie hipotezy zostały zweryfikowane w stopniu, aby stwierdzić, że powiązane problemy badawcze zostały rozwiązane, przynajmniej na podstawowym poziomie. Metody badawcze zostały w większości właściwie dobrane i zastosowane, a eksperymenty interesująco przeprowadzone.

Biorąc pod uwagę przedstawione wyżej analizy i uzasadnienia stwierdzam, że rozprawa doktorska prezentuje ogólną wiedzę teoretyczną Doktoranta w dyscyplinie informatyka techniczna i telekomunikacja. Doktorant wykazuje umiejętność samodzielnego prowadzenia pracy naukowej a rozprawa doktorska zawiera oryginalne rozwiązanie problemu naukowego.

Dlatego stwierdzam, że rozprawa mgr. inż. Łukasza Lepaka spełnia wymagania sformułowane w Ustawie z dnia 20 lipca 2018 roku – Prawo o szkolnictwie wyższym i nauce, Dz. U. 2018, poz. 1668, z późniejszymi zmianami. Wnioskuje do Rady Dyscypliny Informatyka Techniczna i Telekomunikacja Politechniki Warszawskiej o dopuszczenie mgr. inż. Łukasza Lepaka do dalszych etapów postępowania w sprawie nadania stopnia naukowego doktora.



prof. dr hab. inż. Jacek Rumiński, prof. PG
Katedra Inżynierii Biomedycznej, ETI, PG

Dr hab. inż. Wojciech Kotłowski, prof. PP
Instytut Informatyki Politechniki Poznańskiej
ul. Piotrowo 2, 60-965 Poznań
tel: (+48) 61 665 2936
wkotlowski@cs.put.poznan.pl



Poznań, 1 lutego 2025 r.

Recenzja rozprawy doktorskiej

mgr inż. Łukasza Lepaka

Task automation using artificial intelligence methods

Problem badawczy podejmowany w pracy

Rozprawa doktorska Łukasza Lepaka pt. *Task automation using artificial intelligence methods* podejmuje niezwykle aktualny i istotny temat w dziedzinie uczenia maszynowego – automatyzację zadań za pomocą metod sztucznej inteligencji. W obliczu dynamicznego rozwoju technologii informacyjnych oraz ich rosnącego wpływu na różne dziedziny życia i gospodarki, efektywne wykorzystanie sztucznej inteligencji do automatyzacji procesów staje się jednym z kluczowych wyzwań współczesnej inżynierii i technologii.

Praca obejmuje dość szeroki zakres zastosowań głębokiego uczenia maszynowego oraz uczenia ze wzmocnieniem, przedstawiając sposób automatyzacji czterech różnych zadań: handlu energią elektryczną na rynku dnia następnego, wykrywania słów kluczowych w nagraniach audio dla języka polskiego (jako przykładu języków niskozasobowych, tzn. z małą ilością dostępnych danych), przyrostowej optymalizacji hiperparametrów w metodach uczenia sieci neuronowych, oraz symultanicznego tłumaczenia maszynowego. W tym ostatnim przypadku Autor proponuje również nową komórkę rekurencyjną, która pozwala na głęboką transformację stanu i jest odporna na problemy związane z propagacją gradientu.

Przedstawione w pracy rozwiązania mają istotne znaczenie zarówno z punktu widzenia naukowego, ponieważ każde z rozwiązań jest równocześnie propozycją nowego, autorskiego algorytmu, jak i z praktycznego punktu widzenia, gdyż każde z rozwiązań weryfikowane jest na rzeczywistych zbiorach danych w obszernych eksperymentach obliczeniowych.

Struktura pracy i cele badawcze

Rozprawa jest napisana bardzo dobrym językiem angielskim, liczy 151 stron i składa się z 9 rozdziałów oraz szeregu załączników. Ma charakter syntezy treści pięciu powiązanych ze sobą tematycznie artykułów naukowych, które opisane są w rozdziałach 3-7, oraz załączone do pracy w oryginalnej formie w załącznikach B.1-B.5. Rozdział pierwszy stanowi wprowadzenie do tematyki badań, oraz zbiorcze przedstawienie hipotez badawczych wraz z opisaniem wyników, a także przedstawia listę publikacji wraz ze wskaźnikami bibliograficznymi (punkty MNiSW, współczynnik *Impact Factor*, procentowy wkład Autora rozprawy). W rozdziale drugim zawarto podstawy teoretyczne, dotyczące uczenia ze wzmocnieniem (przedstawienie problemu, definicje podstawowych pojęć, lista najbardziej popularnych algorytmów z podziałem na typy), architektur sieci neuronowych (sieci gęste, autoenkodery, sieci rekurencyjne – w tym LSTM i GRU, metodykę Seq2Seq – w tym architektury encoder/decoder i transformer, architektury uczenia się metryk podobieństwa – sieci syjamskie i sieci prototypowe), metod uczenia sieci neuronowych (w tym klasycznego SGD, metod z członem „pędu” (*momentum*) takich jak CM, NAG, oraz metod adaptacyjnych - AdaGrad, RMSProp, Adadelta, a także – zapewne najpopularniejszej z nich – Adam). Rozdział 8 to podsumowanie, a rozdział 9 opisuje inne osiągnięcia doktoranta, nie stanowiące wkładu w rozprawę doktorską.

Podsumowując, Doktorant starannie opisał i dobrze umotywowował dziedzinę badawczą, zamieścił obszerny opis podstaw teoretycznych, streścił każdą z załączonych prac i zawarł również wszystkie potrzebne dane bibliograficzne dotyczące publikacji. Ponieważ struktura pracy jest w zasadzie podporządkowana „publikacyjnemu” charakterowi rozprawy, trudno mieć do niej zastrzeżenia.

Głównym celem pracy była automatyzacja zadań. Pod tym bardzo obszernym tematem kryje się kilka konkretnych hipotez badawczych, wymienionych przez Autora we wprowadzeniu:

1. Automatyzacja tworzenia ofert w handlu energią na rynku dnia następnego prowadzi do zwiększenia zysków handlującego i prowadzi do strategii możliwej do zastosowania w rzeczywistych scenariuszach.
2. Wyszukiwanie słów kluczowych w nagraniach z centrów obsługi w języku niskozasobowym może zostać przeprowadzone przy użyciu modeli dopasowujących nagrania przy użyciu miar podobieństwa.
3. Krótko- i długoterminowy wpływ hiperparametrów algorytmu uczącego się na proces treningowy można wykorzystać do ich dynamicznej optymalizacji bez wcześniejszej znajomości rozwiązywanego problemu.
4. Opóźnienie tłumaczenia w symultanicznym tłumaczeniu maszynowym może być automatycznie kontrolowane przy jednoczesnym utrzymaniu wysokiej jakości tłumaczenia dla sekwencji o dowolnej długości.
5. Głęboka, ogólna transformacja stanu w rekurencyjnej sieci neuronowej pozwala osiągnąć wyniki porównywalne z metodami *state-of-the-art*, jednocześnie ułatwiając proces uczenia i zapewniając efektywne wykorzystanie pamięci.

Do każdej z hipotez przyporządkowana jest jedna z pięciu załączonych publikacji.

Trudno w tym momencie nie zwrócić uwagi na stosunkowo niewielkie powiązanie tematyczne samych prac. Autorowi udało się znaleźć „wspólny mianownik” w postaci problematyki automatyzacji zadań, ale robi on jednak wrażenie dodanego trochę sztucznie, w celu spięcia prac w jedną rozprawę. Same prace dotyczą dziedzin tak różnych jak handel algorytmiczny, przetwarzanie języka naturalnego czy optymalizacja hiperparametrów. Z mojej perspektywy nie stanowi to jednak istotnego problemu. Proces prowadzenia badań naukowych jest trudny do odgórnego zaplanowania, a jego kierunek zmienia się często w sposób ciężki do przewidzenia na etapie konstruowania planu badawczego. Dodatkowo, dorobek naukowy obejmujący tak wiele różnych dziedzin świadczy na korzyść wiedzy i umiejętności Doktoranta. Stąd luźnego powiązania tematycznego wyników nie traktuję jako rzeczy niekorzystnej w ocenie pracy.

Ocena wkładu oryginalnego

Rozprawa jest oparta na pięciu artykułach, których współautorem jest Doktorant:

- [P1] 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 (ECML PKDD) 2024 [Punkty MNiSW: 140]
- [P2] Łukasz Lepak, Kacper Radzikowski, Robert Nowak, Karol J. Piczak: *Generalisation gap of keyword spotters in a cross-speaker low-resource scenario*. Sensors, MDPI, 2021 [Punkty MNiSW: 100]
- [P3] Paweł Wawrzyński, Paweł Zawistowski, Łukasz Lepak: *Automatic hyperparameter tuning in on-line learning: Classic Momentum and ADAM*. International Joint Conference on Neural Networks (IJCNN), IEEE, 2020 [Punkty MNiSW: 140]
- [P4] Grzegorz Rypeś, Łukasz Lepak, Paweł Wawrzyński: *Reinforcement Learning for on-line Sequence Transformation*. 17th Conference on Computer Science and Intelligence Systems (FedCSIS), IEEE, 2022 [Punkty MNiSW: 70]
- [P5] Łukasz Neumann, Łukasz Lepak, Paweł Wawrzyński: *Least Redundant Gated Recurrent Neural Network*. International Joint Conference on Neural Networks (IJCNN), IEEE, 2023 [Punkty MNiSW: 140]

Wszystkie prace zostały opublikowane w punktowanych materiałach konferencyjnych lub czasopismach. W szczególności, robią wrażenie trzy prace na prestiżowych konferencjach uczenia maszynowego: ECML PKDD oraz IJCNN (wszystkie po 140 punktów MNiSW) – jak wiemy, większość kluczowych wyników w uczeniu maszynowym publikowana jest właśnie na głównych konferencjach dziedziny, stąd aktywny w nich udział traktuję zawsze bardzo pozytywnie. Sumaryczne 590 punktów jest wynikiem bardzo dobrym jak na dorobek w trakcie doktoratu. Nie będę odnosił się tutaj do współczynnika *Impact Factor*, ponieważ materiały konferencyjne najprawdopodobniej go nie posiadają. Wkład w każdą z prac został przez Doktoranta precyzyjnie

określony i we wszystkich pracach z wyjątkiem ostatniej przewyższa średni wkład przypadający na pojedynczego współautora. Jedynym mankamentem jest dla mnie niewielka liczba cytowań (13 wg. *Google Scholar*, odczytane na przełomie styczeń/luty 2025), co jest dość zaskakujące z uwagi na bardzo dobre wyniki eksperymentalne zaproponowanych przez Doktoranta metod. Nie zmienia to jednak mojej pozytywnej opinii o całości dorobku naukowego.

Rozprawa zawiera szereg oryginalnych i nowatorskich wyników, które zostały już pozytywnie zweryfikowane w procesie recenzyjnym na etapie przyjmowania prac na konferencje i do czasopism:

- W pracy [P1] opracowano strategię składania ofert kupna i sprzedaży energii elektrycznej na rynku dnia następnego. Strategia oparta jest na metodzie uczenia ze wzmocnieniem. Tradycyjne podejścia do tego typu problemu opierały się na prostych strategiach parametrycznych, optymalizacji liniowej, lub metodach uczenia się proponujących jedynie jedną parę ofert kupna i sprzedaży w ciągu całego dnia (prowadzących do bardzo elementarnych krzywych popytu i podaży). Autorzy pracy zaproponowali strategię, w której generowane są kolekcje ofert kupna i sprzedaży energii dla każdej godziny z osobna, oparte na parametryzacji ogólnych krzywych popytu i podaży za pomocą 100-wymiarowej przestrzeni akcji agenta.

Eksperymenty przeprowadzone na symulacji polskiego rynku energii pokazały, że proponowana strategia osiąga znacznie wyższe zyski niż metody „klasyczne” bądź wcześniejsze metody oparte na uczeniu ze wzmocnieniem, nawet przy uwzględnieniu ograniczeń związanych z magazynowaniem energii i zmiennością podaży z odnawialnych źródeł. Istotnym aspektem tej pracy jest możliwość jej praktycznego zastosowania, co zostało potwierdzone zainteresowaniem ze strony przemysłu energetycznego.

- W artykule [P2] opracowany został system do wykrywania słów kluczowych w nagraniach rozmów telefonicznych w języku polskim. Ze względu na ograniczoną dostępność danych dla języka polskiego, zadanie to stanowiło duże wyzwanie badawcze. Autorzy zastosowali metody dopasowywania nagrań oparte na metryce podobieństwa, używając do tego różnych narzędzi: sieci syjamskich, sieci opartych na prototypach, oraz wektorach zanurzeń pochodzących z pre-trenowanego modelu firmy Google, z dodatkowym przetwarzaniem wyników.

Największym wyzwaniem była praca z ograniczoną liczbą danych w języku polskim i konieczność posiłkowania się danymi w języku angielskim. Na uwagę zasługuje wysiłek autorów w pozyskaniu licznych zbiorów danych, w tym samodzielnym utworzeniu znacznej części danych w języku polskim. Pomimo tych trudności autorowi udało się stworzyć system, który osiągnął dość przyzwoite wyniki detekcji słów kluczowych, w najlepszym wypadku uzyskując wynik miary F na poziomie 0.69.

- W pracy [P3] zaproponowano algorytm *Autonomous Stochastic Descent with Momentum 2 (ASDM2)*, który automatycznie dostosowuje hiperparametry metod używających członu *momentum* (takich jak CM lub Adam) podczas procesu uczenia sieci neuronowych. Jest to więc rozwinięcie istniejących metod optymalizacji sieci, wyposażające je w mechanizm,

który zamiast odgórnego określania hiperparametrów pod problem (zwykle metodą prób i błędów) pozwala na swobodną ewolucję ich wartości podczas procesu uczenia celem przyspieszenia minimalizacji błędu na zbiorze uczącym. Sposób adaptacji hiperparametrów to błyskotliwa metoda „spadku wzdłuż gradientu wewnątrz metody spadku wzdłuż gradientu”, opierająca się na wygładzaniu wykładniczym parametrów modelu i propagacji wpływu, jaki mają na nie hiperparametry (tzn. pochodnych), wykorzystując do tego odpowiednie rekurencje.

Wyniki eksperymentów przeprowadzonych na różnych architekturach sieci neuronowych, w tym płytkich sieci gęstych (dla danych tabelarycznych), gęstych, głębokich autoenkoderach i autoenkoderach konwolucyjnych, potwierdziły skuteczność proponowanego algorytmu: ASDM2 osiągnął lepsze wyniki niż popularne algorytmy uczenia w ogromnej większości przypadków, w szczególności przebił algorytm Adam z *optymalnie* strojonymi parametrami! Czyni go to obiecującym narzędziem do zastosowań w praktycznych problemach uczenia głębokiego.

- W pracy [P4] zaproponowano system *Reinforcement Learning for on-line Sequence Transformation (RLST)*, który automatycznie kontroluje opóźnienie tłumaczenia symultanicznego przy zachowaniu wysokiej jakości, używając metod uczenia ze wzmocnieniem – agent dokonuje nie tylko wyboru tokenu sekwencji wyjściowej, ale przede wszystkim wybiera jedną z dwóch akcji: odczytanie tokenu na wyjściu bądź generację tokenu wyjściowego, co pozwala na kontrolę przetargu między opóźnieniem, a jakością tłumaczenia. W przeciwieństwie do tradycyjnych metod tłumaczenia maszynowego, RLST pozwala na generowanie tłumaczeń w czasie rzeczywistym. Eksperymenty obliczeniowego na rzeczywistych danych pokazały, że RLST osiąga wyniki porównywalne z nowoczesnymi architekturami tłumaczenia maszynowego, szczególnie dla dłuższych sekwencji, gdzie utrzymanie kontekstu jest kluczowe dla wysokiej jakości tłumaczeń.
- W pracy [P5] opracowano nową komórkę rekurencyjną nazwaną *Deep Memory Update (DMU)*. Głównym celem tego rozwiązania było umożliwienie głębokiej transformacji stanu przy jednoczesnym uniknięciu problemów związanych z propagacją gradientu. W szczególności, autorzy zaproponowali użycie sieci neuronowej *feed-forward* jako nauczalnej transformacji stanu, wyposażonej w odpowiednio dobrane funkcje aktywacji. Z jednej strony transformacja taka jest uniwersalnym aproksymatorem, z drugiej strony – używając odpowiedniej struktury sieci można zapewnić (co autorzy udowadniają), że sieć taka nie prowadzi do eksplozji lub zanikania gradientów. Dodatkowo, zaproponowano odpowiedni dobór szybkości uczenia takiej transformacji celem utrzymania stabilności procesu trenowania całościowej sieci.

Eksperymenty wykazały, że DMU osiąga wyniki porównywalne, a często lepsze niż popularne sieci rekurencyjne, takie jak LSTM, GRU czy RHN. Proponowana komórka rekurencyjna ma więc potencjał do szerokiego zastosowania w zadaniach przetwarzania sekwencyjnego.

Uważam, że każda z wymienionych publikacji daje istotny wkład do dziedziny uczenia maszynowego i uznaję stąd, że wymienione na wstępie cele badawcze udało się Doktorantowi w

pełni osiągnąć.

Uwagi dyskusyjne

Nie kwestionując wartości całościowych wyników zawartych w rozprawie, chciałbym zgłosić poniżej kilka uwag, głównie w formie pytań bądź dyskusji.

- W pracy [P1] pojawia się kilka niejasnych dla mnie elementów. Po pierwsze, do czego potrzebna jest randomizacja strategii w równaniu (10)? Oczywiście, losowe zachowanie strategii w procesie uczenia (np. ϵ -greedy) jest kluczowe dla eksploracji, natomiast – o ile nie zakładamy wpływu strategii na rynek energii – losowość nie wydaje się potrzebna na etapie stosowania strategii na rynku energii. Po drugie, czy parametryzacja zbioru ofert kupna/sprzedaży przez akcje a jest wystarczająco elastyczna do modelowania ogólnych krzywych popytu i podaży zgodnie z dyskusją pod Rysunkiem 1? Innymi słowy: czy można pomocą równań 4-7 utworzyć dowolną kolekcję ofert kupna/sprzedaży (a jeśli nie – to jakie są ograniczenia?). Po trzecie, zastanawia mnie, do jakiego stopnia zadanie to jest problemem uczenia ze wzmocnieniem, a nie problemem uczenia nadzorowanego (ze złożonym wyjściem)? Kluczową trudnością w uczeniu się ze wzmocnieniem jest wpływ wykonywanych akcji na stan agenta, przez co lokalnie opłacalna akcja (dająca wysoką nagrodę) może nie być opłacalna w dłuższym okresie (niewielka wartość funkcji Q). W rozważanym problemie jedynym chyba czynnikiem propagującym się na kolejne dni, na który wpływ mają akcje agenta danego dnia (rozpatrywane całościowo, jako zespół ofert kupna i sprzedaży), jest stan baterii pod koniec dnia, czyli możliwość sprzedaży nadwyżki energii w dni kolejne. Ale z danych eksperymentalnych wynika, że pod koniec dnia bateria jest zwykle bliska rozładowaniu. Stąd być może można by problem zapisać jako problem predykcji złożonego wyjścia (akcji) ocenianego przez złożoną funkcję nagrody będącą całkowitym zyskiem w ciągu dnia?
- Również odnośnie pracy [P1]: kluczowymi dla określenia zysku jest produkcja i zużycie energii przez agenta. Wartości te mogą być – zgodnie z przyjętym modelem – wyznaczone w oparciu o pogodę. Czy Autor rozważał dostarczenie agentowi na wejściu prognoz produkcji i zużycia energii w ciągu dnia wyznaczonych na podstawie prognoz pogody?
- W pracy [P2] wspomniano, że z powodu potencjalnych problemów zrezygnowano z użycia metod rozpoznawania mowy (*speech-to-text*). Czy Doktorant próbował użyć takich metod do rozwiązania problemu i porównywał je z zaproponowanymi metodami opartymi na podobnych wzorcach?
- W pracy [P3] zaproponowany algorytm cechuje się znakomitą jakością mierzoną zdolnością minimalizacji błędu uczącego, bijąc nawet (prawie) optymalnie dostrojoną metodę Adam. Autorzy pracy nie wspominają natomiast o wartości błędu na zbiorze testowym, która nie zawsze musi być dobrze skorelowana z błędem treningowym, a także nie podają czasów obliczeń algorytmu ASDM2 względem algorytmów bazowych. Byłby wdzięczny za komentarz, z ewentualnym przybliżonym podaniem tych – w praktyce bardzo istotnych – wartości, jeśli były w eksperymentach mierzone.

Konkluzja końcowa

Recenzowaną rozprawę oceniam bardzo dobrze. Załączone prace potwierdzają wiedzę Doktora w zakresie uczenia maszynowego i sztucznej inteligencji, w szczególności zastosowania ich do automatyzacji zadań, oraz umiejętność prowadzenia prac badawczych. Problemy, z którymi zmierzył się mgr inż. Łukasz Lepak, są ambitne i istotne dla postępu w dziedzinie. Wyniki eksperymentów obliczeniowych na rzeczywistych danych, zawarte w rozprawie, potwierdzają wysoką jakość zaproponowanych rozwiązań. Stwierdzam, że *przedstawiona do oceny Rozprawa Doktorska spełnia warunki określone w Art. 187 Ustawy z dnia 20 lipca 2018 r./Prawo o szkolnictwie wyższym i nauce / (Dz.U. z 2022 r. poz. 574 z późn. zm.), oraz uzasadnia nadanie stopnia naukowego doktora w dyscyplinie informatyka.*

W związku z tym **rozprawę oceniam jako spełniającą wymogi stawiane pracom doktorskim i wnoszę o dopuszczenie mgr. inż. Łukasza Lepaka do dalszych etapów postępowania w sprawie nadania stopnia doktora, równocześnie sugerując wyróżnienie pracy.**

Dr hab. inż. Wojciech Kotłowski