

FDITT-npi 28.08 2023

dr hab. Jan Bazan, prof. UR
Kolegium Nauk Przyrodniczych
Uniwersytet Rzeszowski,
ul. Pigonia 1, 35-310 Rzeszów
jbazan@ur.edu.pl

22 sierpnia 2023 r.

Recenzja rozprawy doktorskiej

mgr inż. Joanny Badury

zatytułowanej:

Solutions for selected problems of demand forecasting based on machine learning methods and domain knowledge

Promotor: *dr hab. inż. Marek Sikora, prof. w Pol. Śl.*

Dziedzina: nauki techniczne. Dyscyplina: informatyka techniczna i telekomunikacja

1. Cel i zakres rozprawy

Rozprawa doktorska wpisuje się w nurt badań związany z metodami eksploracji danych, który już od kilkudziesięciu lat znajduje się w centrum badań wielu silnych ośrodków badawczych na świecie. Tak duże zainteresowanie tym nurtem wynika z dużej dostępności coraz większych zbiorów danych oraz pilnej potrzeby eksploracji tych danych. Postęp w tym zakresie jest niezwykle istotny m.in. dla rozwoju systemów inteligentnych bazujących na wspomaganiu podejmowania decyzji, w oparciu o wyniki analizy dostępnych zbiorów danych. Dlatego opracowywanie metod eksploracji danych i wykazanie skuteczności tych metod w różnych zastosowaniach ma wciąż kluczowe znaczenie dla dalszego rozwoju systemów inteligentnych w wielu dziedzinach, jak choćby w medycynie, finansach, przemyśle, transporcie, telekomunikacji i innych, gdzie pojawia się wiele problemów których rozwiązanie stoi na drodze do konstrukcji systemów inteligentnych. Jednym z takich problemów związanych z efektywnym funkcjonowaniem współczesnego przedsiębiorstwa jest problem prognozowanie popytu, który najczęściej wiąże się z szacowaniem wielkości sprzedaży towarów lub usług oferowanych przez przedsiębiorstwo w konkretnym okresie. Celem prognozowania popytu jest bardziej efektywne zaplanowanie zasobów produkcyjnych, lepsza ocena potrzebnych zapasów oraz opracowanie korzystniejszego budżetu. Problem prognozowania popytu występuje praktycznie w każdym przedsiębiorstwie branży przemysłowej, handlowej lub usługowej. Prognozowanie popytu, jako rodzaj prognozowania, jest zadaniem przewidywania nieznanych wartości na podstawie dostępnej wiedzy. Przewidywane wartości dotyczą przyszłości, a do ich przewidywania potrzebne są dane historyczne dotyczące przewidywanych zjawisk oraz wiedza dziedzinowa ekspertów. W najprostszym przypadku dla prognozowania popytu, system predykcyjny można oprzeć tylko na danych zebranych z przeszłości, np. wykorzystać dane o historycznej sprzedaży konkretnego produktu. Jest to zatem zadanie prognozowania szeregów czasowych. Recenzowana rozprawa doktorska nie koncentruje się jednak na typowej problematyce prognozowania szeregów czasowych w tym sensie, że mając dane historyczne o popycie, prognoza zostanie wykonana dla kolejnych punktów w czasie. Celem rozprawy nie było również porównanie wszystkich znanych algorytmów prognozowania szeregów czasowych. Zamiast tego skupiono się na innych, często mniej znanych problemach związanych z prognozowaniem popytu, które można rozwiązać za pomocą uczenia maszynowego. Wszystkie badania przeprowadzono na rzeczywistych danych. Większość z nich pochodziła z dużych firm z wieloma oddziałami. Dane obejmowały często kilka lat, pochodziły z wielu lokalizacji i były ściśle powiązane z projektami badawczo-rozwojowymi współfinansowanymi z Funduszy Europejskich i realizowanymi w ostatnim okresie

we współpracy Politechniki Śląskiej z różnymi firmami. W projektach tych aktywnie uczestniczyła Autorka rozprawy. Są to np. następujące projekty.

1. System wspomagania decyzji i zarządzania wiedzą dla branży handlu detalicznego realizowany z firmą 3Soft S.A.,
2. System automatyzujący i optymalizujący planowanie łańcucha dostaw, wykorzystujący metody inteligencji obliczeniowej do analizy online bardzo dużych zbiorów danych pozwalający na ograniczenie strat i zwiększenie efektywności działania firm branży handlu detalicznego realizowany z firmą 3Soft S.A.,
3. System wspomagania decyzji i zarządzania wiedzą operacyjną i procesową dla rynku dystrybutorów gazu LPG realizowany z firmą AIUT Sp z o.o.

Wyniki badań przedstawione w rozprawie mają zatem duże znaczenie dla biznesu i problematyka tych badań jest niebanalna. Dlatego uważam, że podjęta tematyka może z powodzeniem stanowić przedmiot rozprawy doktorskiej.

2. Zawartość rozprawy

Rozprawa doktorska została napisana w języku angielskim i składa się z 7 rozdziałów. Rozdział 1 to wprowadzenie obejmujące ogólne przedstawienie motywacji, problemu badawczego, celu i tezy rozprawy oraz opis struktury rozprawy. Rozdział 2 to wprowadzenie do tematu prognozowania popytu z wykorzystaniem uczenia maszynowego i wiedzy dziedzinowej. Przedstawiono tutaj autorską taksonomię zagadnień związanych z prognozowaniem popytu, opis zagadnień teoretycznych związanych z prognozowaniem oraz krótki przegląd metod uczenia maszynowego, które można wykorzystać do prognozowania popytu. Rozdział 3 przedstawia zagadnienie związane z prognozowaniem popytu, koncentrując się na temacie prognoz skuteczności promocji. W rozdziale 4 skupiono się na problematyce prognozowania popytu w przypadku prognozowania top-down, czyli rozbicia prognozy wyższego poziomu (np. dla grupy produktów) na prognozę niższego poziomu (np. prognozy dla konkretnego produktu). W rozdziale 5 przedstawiono zagadnienie wykorzystania dodatkowych szeregów czasowych do poprawy prognozowania badanego zjawiska. Chodzi tutaj o dodawanie danych z innych lokalizacji opisującym podobny problem lub dodawanie innych danych z tej samej lokalizacji. Rozdział 6 koncentruje się na zagadnieniu prognozowania prawdopodobieństwa zwrotu określonego produktu. Na koniec, rozdział 7 stanowi podsumowanie rozprawy i przedstawia kierunki dalszych prac. Rozprawę kończy bibliografia zawierająca 190 pozycji.

Warto dodać, że cechą charakterystyczną rozprawy jest to, że wyniki eksperymentów nie są przedstawione w osobnym rozdziale, ale umieszczono je w rozdziałach dotyczących poszczególnych problemów badawczych, których te eksperymenty dotyczyły. Takie postępowanie jest jak najbardziej właściwe.

3. Poprawność i oryginalność postawionej tezy (wkład Autorki)

Na podstawie doświadczeń Autorki rozprawy w zakresie prognozowania popytu, do realizacji wspomnianego wyżej celu rozprawy postanowiono wykorzystać metody uczenia maszynowego, prowadząc do następującej tezy rozprawy:

Rozbijając zadanie prognozowania na konkretne zagadnienia i biorąc pod uwagę wiedzę dziedzinową, można uzyskać użyteczne wartości prognoz z wykorzystaniem metod uczenia

maszynowego. Wiedzę dziedzinową można uwzględnić, wprowadzając do modelu nowe atrybuty lub nowe rekordy danych z podobnych problemów, co skutkuje poprawą ostatecznych wyników prognozy.

Należy stwierdzić, że postawienie takiej tezy było uzasadnione licznymi aktualnymi zastosowaniami związanymi z funkcjonowaniem znanych Autorce przedsiębiorstw oraz niewielką liczbą publikacji w literaturze na ten temat, szczególnie w przypadku specyficznych problemów rozpatrywanych w rozprawie.

Dla wykazania tej tezy wykonano następujące prace, których wyniki są wkładem Autorki rozprawy w rozwój dyscypliny naukowej.

1. Zaproponowano taksonomię zagadnień związanych z prognozowaniem popytu, pozwalającą uzyskać szerszą perspektywę niż tylko prognozowanie szeregów czasowych.
2. Zaproponowano podejście do prognozowania efektu promocji za pomocą sześciu wskaźników i modeli uczenia maszynowego. Podejście to zostało następnie wykorzystane w symulatorze wpływu promocji na sprzedaż i zysk oraz w systemie rekomendacji dla promocji.
3. Przedstawiono autorski algorytm generowania reguł akcji dla przeżywalności promocji w celu wskazania zmian, które wpłyną na efektywność promocji.
4. Przedstawiono wykorzystanie metod uczenia maszynowego do prognozowania top-down czyli rozbicia prognozy wyższego poziomu (np. dla grupy produktów) na prognozę niższego poziomu (np. prognozy dla konkretnego produktu).
5. Przeprowadzono badania wskazujące na znaczenie wykorzystania dodatkowych szeregów czasowych do prognozowania popytu dla rozważanych scenariuszy, uwzględniających szeregi czasowe z podobnych lokalizacji (np. prognozując ilość sprzedanego paliwa przez sieć stacji benzynowych można wykorzystać dane z wielu stacji) i z tej samej lokalizacji (np. prognozując w domu lub w firmie zużycie energii elektrycznej można wykorzystać dane o zużyciu energii przez różnego typu urządzenia).
6. Zaproponowano algorytm łączenia transakcji zwrotu z transakcją zakupu. Podjęto również temat prognozowania zwrotu określonego produktu z konkretnej transakcji zakupu, osiągając dobre wyniki prognozowania.

Poniżej omawiam bardziej szczegółowo wyniki wyżej wymienionych prac.

Ad 1. W rozprawie zaproponowano autorską taksonomię zagadnień związanych z prognozowaniem popytu. W taksonomii zaproponowano podział przypadków prognozowania popytu na trzy kategorie: przypadki związane z prognozowaniem krótkookresowym, prognozowaniem średniookresowym i prognozowaniem długoterminowym. Taki podział podyktowany jest różnymi przypadkami biznesowymi, które pozwalają na wprowadzenie takiego podziału. Przedziały czasowe omawianych horyzontów mogą różnić się w zależności od przedsiębiorstwa, którego dotyczy prognozowanie. W pracy przypadki krótkoterminowe oznaczono jako prognozowanie na tydzień do przodu, średniookresowe jako prognozowanie na okres od 2 tygodni do 3 miesięcy, a długoterminowe jako prognozowanie na ponad 3 miesiące do przodu. Taksonomia zawiera szereg przypadków prognozowania, z których niektóre zostały uwzględnione w dalszej części pracy. Są to

np. następujące przypadki prognozowania, które często definiowano w taksonomii dla wymienionych wyżej przedziałów czasowych.

- Prognozowanie popytu na dany produkt z daną historią sprzedaży.
- Prognozowanie efektu planowanych promocji.
- Prognozowanie wielkości zwrotów produktów.
- Prognozowanie zapotrzebowania na nowe produkty.
- Prognozowanie popytu przy rozpoczynaniu nowej działalności.
- Prognozowanie zapotrzebowania na nowy produkt, ale w nowej lokalizacji.
- Prognozowanie top-down (np. na grupie produktów a nie na pojedynczym produkcie).

Zaproponowana taksonomia pokazuje szerszy obraz zagadnień związanych z prognozowaniem popytu. Definiuje także szereg problemów, które są rozwiązywane w dalszej części rozprawy.

Ad 2. Zaproponowano podejście do prognozowania efektu promocji. W tym celu zaproponowano zestaw wskaźników, które można wykorzystać do oceny promocji. Są to np. takie wskaźniki jak: średnia liczba produktów sprzedanych każdego dnia, średnia liczba paragonów z promowanym produktem, średnia wartość koszyka zawierającego promowany produkt i inne. Wskaźniki te były podstawą do tworzenia atrybutów decyzyjnych w konstruowanych tablicach decyzyjnych. Jeśli zaś chodzi o atrybuty warunkowe, to były one związane: z ceną produktu, z czasem trwania promocji, nośnikami reklamy, sklepem i jego otoczeniem oraz wpływem innych promocji. Każdy wiersz konstruowanej tablicy decyzyjnej dotyczył wybranego okresu przed promocją (dane do treningu klasyfikatora) lub w trakcie promocji (dane dla testu klasyfikatora). Do tworzenia klasyfikatorów wykorzystano głównie algorytmy XGBoost, sztuczne sieci neuronowe głębokiego uczenia (biblioteki Keras i Hyperas) oraz sztuczne sieci neuronowe głębokiego uczenia z biblioteki H2O w języku R. Warto wspomnieć, że badano także inne modele konstrukcji klasyfikatorów, ale najlepsze wyniki uzyskiwał XGBoost. Ponieważ atrybuty decyzyjne miały wartości liczbowe, to do oceny jakości klasyfikatorów wykorzystano 4 miary specjalnie dobrane do tej sytuacji. W tym miejscu warto wspomnieć 2 spośród tych miar. Miarę MAE, która mierzy średnią różnicę wartości prognozowanych i rzeczywistych oraz miarę MAPE, która mierzy średnią różnicę wartości prognozowanych i rzeczywistych względem wartości rzeczywistej. Miara MAE jest użyteczna przy stwierdzeniu na ile prognoza różni się od wartości rzeczywistej wskaźnika, a miara MAPE mówi o tym na ile ta różnica jest duża w stosunku do rzeczywistej wartości wskaźnika. Wyniki wykonanych eksperymentów pokazały, że proponowane podejście jest użyteczne do prognozowania wyników promocji. Badano możliwość konstrukcji klasyfikatorów dla pojedynczych produktów i grup produktów wykazując, że klasyfikatory dla wielu produktów osiągają bardzo dobre wyniki. Na przykład dla przypadku promocji na owoce, miara MAPE dla wskaźnika *średnia liczba koszyków zakupowych* wynosiła tylko 0.16, co oznacza dużą zgodność wartości prognozowanej i rzeczywistej.

Następnie pokazano, w jaki sposób zaproponowane modele predykcyjne można wykorzystać w praktyce biznesowej. Pierwsze z zaproponowanych zastosowań to symulator wpływu promocji na sprzedaż i zysk. Przygotowano moduł GUI realizujący taką funkcjonalność. W module tym użytkownik podaje dla jakiego produktu chce zrobić symulację oraz inne parametry, które są konfigurowane przy planowaniu promocji (np. czas trwania promocji). Na wyjściu, użytkownik

otrzymuje informację jak zmienia się sprzedaż produktu przy różnych zmianach cen oraz jak zmienia się zysk. Drugie z zaproponowanych zastosowań to system rekomendacji promocji. Tutaj także przygotowano moduł GUI realizujący taką funkcjonalność. W module tym użytkownik podaje jakiego produktu ma dotyczyć promocja, jak jest cena regularna produktu i okres trwania promocji. Na wyjściu, użytkownik otrzymuje informację o efekcie danej promocji, co daje możliwość wybrania optymalnej promocji.

Ad 3. Zaproponowano autorski algorytm generowania reguł akcji dla przeżywalności promocji (ang. *survival action rules*) w celu wskazania zmian, które wpłyną na efektywność promocji. Algorytm ten jest algorytmem pokryciowym liczenia reguł akcji, ale działa w specyficznej sytuacji, gdy atrybut decyzyjny ma wartości będącymi krzywymi przeżycia. Krzywe te są estymatorami Kaplana-Meiera, a porównanie ich ze sobą (jako porównanie wartości decyzji) wykonywane jest w oparciu o porównanie wielkości pól pod tymi krzywymi. Im to pole jest większe, tym leprze rokowanie w stosunku do promocji z takim polem. Każdy punkt na krzywej odpowiada jednemu z dni promocji, a przeżywalność promocji w danym dniu oznacza, że promocja przynosi pozytywny efekt (np. sprzedaż produktu z promocji jest jeszcze odpowiednio wysoka). Przedstawiono także przykład zastosowania tego algorytmu do analizy zmian, jakie można wprowadzić, aby zmniejszyć prawdopodobieństwo, że wartość prognozowanego wskaźnika rozumianego jako efekt promocji, spadnie poniżej zadanego progu.

Ad 4. Kolejnym rozważanym w rozprawie poruszonym tematem była kwestia podziału prognoz wyższego rzędu na prognozy niższego rzędu. Prognoza wyższego poziomu może być prognozą popytu na całą grupę produktów zaproponowaną przez eksperta dziedzinowego. Prognoza niższego poziomu to prognoza popytu na określony produkt. Prognozowanie top-down może mieć miejsce wszędzie tam, gdzie można pogrupować produkty. Ma to znaczenie, gdy sprzedaż danego produktu jest bardzo niska lub jest rzadka, co uniemożliwia stworzenie modelu prognostycznego dla konkretnego produktu. W rozprawie skoncentrowano się na porównaniu wybranych algorytmów prognostycznych w kontekście prognozowania top-down. Pierwszy z algorytmów był oparty na kostce danych. Dla każdej kostki wyznaczanej na danym okresie (np. tygodniowym) wyznaczano istotne parametry dla prognozowania popytu na dany produkt. Dzięki wyznaczaniu kostek dla wielu okresów, można przewidywać wartości tych parametrów w innych okresach (np. w oparciu o kostki tygodniowe w danym roku, można przewidywać popyt na dany produkt w analogicznym tygodniu w kolejnym roku). Drugi z algorytmów wykorzystywał metodę najbliższego sąsiada (k-NN). Podobnie jak poprzedni, przy prognozowaniu popytu wybierał kostkę o podobnej sytuacji sprzedaży i wykorzystywał ją do przewidywania popytu na dany produkt w testowanym okresie. Jednak przy wyborze wykorzystywał odległość euklidesową pomiędzy wektorami parametrów kostek. Zamiast algorytmu k-NN badano jeszcze na tych samych atrybutach algorytmy regresji liniowej i XGBoost. Osiągnięciem tego punktu jest także możliwość wykorzystania tablicy odmienności typów produktów, która pozwala na wykorzystanie typu produktu w algorytmach opartych na metrykach odległości. Chodzi o to, że jeśli historycznie brak jest informacji o sprzedaży jakiegoś produktu, to można skorzystać o informacji o sprzedaży produktów podobnych. Na koniec warto dodać, że jeśli mamy prognozę dla grupy produktów (np. podaną przez ekspertów), to za pomocą opracowanej w pracy metodologii możemy z tej prognozy wyluskać prognozę dla produktów składowych poprzez

predykcję tego jak dużą część sprzedaży grupy produktów zajmie sprzedaż danego produktu składowego. Niestety, w pracy nie przedstawiono miary jakości MAPE dla wykonanych eksperymentów, a jedynie miarę MAE. Powoduje to, że, zdaniem recenzenta, czytelnik nie może ocenić rzeczywistej jakości prezentowanych metod w kontekście zastosowań praktycznych.

Ad 5. Przeprowadzono badania wskazujące na znaczenie wykorzystania dodatkowych szeregów czasowych w celu poprawy skuteczności prognozy rozpatrywanego zagadnienia. Najpierw pokazano możliwość wykorzystania analogicznych szeregów czasowych z innych lokalizacji. W tym kierunku analizowano dwa następujące zbiory danych. Pierwszy dotyczył prognozy wysokości sprzedaży paliwa przez sieć wielu stacji benzynowych. Drugi zaś dotyczył wysokości sprzedaży gazu płynnego (LPG) z różnych zbiorników. Dla obu badanych zbiorów wykazano, że tworzenie modeli predykcyjnych z wykorzystaniem wszystkich danych poprawia wyniki uzyskiwanych prognoz. Przedstawiono także różne możliwości grupowania i dodawania informacji z podobnych szeregów czasowych, ale nie sprawdzono tych metod eksperymentalnie. Następnie zbadano problem prognozowania, w którym dodatkowe dane szeregów czasowych z tej samej lokalizacji zostały uwzględnione w modelach predykcyjnych. W tym kierunku analizowano dane dotyczące zużycia energii elektrycznej z podziałem na różnego typu urządzenia gospodarstwa domowego. Zbadano 5 przypadków danych dotyczących 2 różnych mieszkań, 2 różnych domów i biura. Uzyskanie danych na temat zużycia energii elektrycznej przez poszczególne urządzenia gospodarstwa domowego i wprowadzenie ich do modelu prognostycznego umożliwiło poprawę prognozowania ogólnego zużycia energii. Wniosek z tych badań jest taki, że dodanie danych z innych lokalizacji lub z tej samej lokalizacji może znacznie poprawić jakość prognoz. Jednak sposób dodawania tych informacji może być różny w zależności od zestawu danych.

Ad 6. Ostatnim rozważanym w rozprawie problemem był problem prognozowania zwrotów. Przegląd literatury wykazał, że do tej pory była tu luka badawcza. Problem prognozowania zwrotów pojawiał się w literaturze, ale w ujęciu globalnym (np. liczby całkowitych zwrotów w danym okresie). Poprzednie badania nie dotyczyły przewidywania prawdopodobieństwa zwrotu określonego produktu. Dla przygotowania danych do eksperymentów, zaproponowano niestandardowy algorytm łączenia transakcji zwrotu z transakcją zakupu. Podany algorytm był niezbędny do stworzenia zbioru uczącego, ponieważ w analizowanych danych były osobne listy sprzedanych produktów oraz dokonanych zwrotów. Dlatego trzeba było dokonać połączenia zwrotów z konkretnymi produktami, aby utworzyć tablicę decyzyjną na potrzeby generowania i testowania klasyfikatorów. Do eksperymentów wykorzystano trzy następujące algorytmy tworzenia klasyfikatorów: regresja logistyczna, las losowy oraz XGBoost. Stosowano je dla kilku różnych zbiorów danych, w tym dla przypadku balansowania klas decyzyjnych za pomocą ogólnie znanej metody nadpróbkowania mniejszych klas (SMOTE). Optymalizowano także hiperparametry dla klasyfikatorów. Najlepszą jakość klasyfikacji $AUC = 0.806$ uzyskano z użyciem algorytmu XGBoost, gdzie AUC to pole pod krzywą ROC. Z praktycznego punktu widzenia wynik ten jest bardzo dobry.

Podsumowując, wszystkie rozpatrywane zadania prognozowania popytu zostały pomyślnie rozwiązane z wykorzystaniem metod uczenia maszynowego i zastosowania wiedzy dziedzinowej. Pewnym zaskoczeniem było to, że zastosowane podejścia Deep Learning nie przyniosły lepszych rezultatów niż typowe metody uczenia maszynowego. Jest to jednak zgodne z wcześniejszymi

doniesieniami na ten temat, że dla danych tablicowych, zaawansowane metody drzewiaste (takie jak XGBoost) przewyższają metody oparte na sztucznych sieciach neuronowych.

Wiedza i umiejętności Autorki do poprawnego i przekonującego przedstawienia uzyskanych przez siebie wyników

W rozprawie Autorka zamieściła przegląd aktualnego stanu wiedzy w zakresie metod potrzebnych do przewidywania popytu, w tym metod prognozowania szeregów czasowych i metod uczenia maszynowego stosowanych obecnie do eksploracji danych (rozdział 2). Opisy te zostały wykonane z wysoką starannością i z odwołaniami do literatury, co pozwoliło umiejscowić w literaturze przedmiotu prezentowane badania. W kolejnych rozdziałach także zamieszczono odwołania do literatury związane z używanymi pojęciami, metodami lub używanymi do eksperymentów bibliotekami oprogramowania. Bez wątplenia świadczy to o dużej wiedzy Kandydatki. Ponadto, należy mocno podkreślić, że napisanie rozprawy wymagało wcześniejszego skonstruowania, w tym zaprogramowania, środowiska eksperymentalnego. Autorka aktywnie uczestniczyła w wykonaniu tego środowiska, które zapewne powstało w powiązaniu z pracami prowadzonymi we wspomnianych wyżej projektach badawczych realizowanych we współpracy Politechniki Śląskiej z przedsiębiorstwami.

Jak już wspominałem wyżej, pozycja rozprawy w stosunku do obecnego stanu wiedzy polega na tym, że rozprawa wypełnia lukę w literaturze światowej dotyczącą specyficznych problemów związanych z prognozowaniem popytu. Chodzi tutaj głównie o przewidywanie efektu promocji, prognozowania popytu dla grup produktów, wykorzystanie razem różnych szeregów czasowych do prognozowania oraz przewidywania zwrotów produktów.

Zaproponowane nowe metody z pewnością będą użyteczne w prognozowaniu popytu, gdyż algorytmy uczenia maszynowego mogą być niewątpliwie skuteczną alternatywą lub uzupełnieniem wykorzystywanych do tej pory metod prognozowania szeregów czasowych.

Odnosząc się do umiejętności Autorki rozprawy w zakresie poprawnego i przekonującego przedstawiania uzyskanych wyników należy stwierdzić, że w tym zakresie Autorka rozprawy wykazała się dobrymi umiejętnościami. Rozprawa napisana jest dość przejrzystie i bardzo dobrze zaprojektowano jej strukturę, która została uwidoczniła w spisie treści.

4. Uwagi na temat rozprawy

Z formalnego i matematycznego punktu widzenia rozprawa nie budzi zastrzeżeń recenzenta.

Podczas lektury rozprawy nasuwają się jednak pewne uwagi, które można traktować jako sformułowanie pewnych mankamentów lub sugestii co do dalszych prac. Dobrze by było, aby Kandydatka odniosła się do nich podczas obrony rozprawy.

1. Jednym z najważniejszych wyników rozprawy jest algorytm generowania reguł akcji dla przeżywalności promocji. Niestety, zdaniem recenzenta, nie został on opisany wyczerpująco. W opisie algorytmu brak jest istotnych informacji, jak choćby sposobu znalezienia najlepszego warunku do reguły (funkcja `BestElementaryCondition`), opisu miary log-rank oceniającej jakość reguły podczas jej rozszerzania i przycinania, algorytmu sposobu przycinania reguły (funkcja

Prune), a przede wszystkim szczegółów na temat tego jak decyzja wyrażona za pomocą krzywej KM wpływa na wyliczone reguły. Te szczegóły powinny zostać przedstawione podczas obrony rozprawy.

2. W rozprawie nie przedstawiono dyskusji związanej ze złożonością obliczeniową algorytmu generowania reguł akcji przeżywalności promocji. Taka dyskusja jest konieczna, ponieważ dane dla których takie reguły można wyznaczać często są bardzo duże.
3. Wobec pojawiających się obecnie możliwości technicznych masowego zrównoleglenia obliczeń w chmurach obliczeniowych, pojawia się pytanie czy zaproponowany w rozprawie algorytm generowania reguł akcji może zostać zrównoleglony, a jeśli tak, to jakie przeniosłoby to efekty w zakresie jego przyspieszenia.
4. Ponieważ atrybuty decyzyjne w problemach dotyczących prognozowania promocji miały wartości liczbowe, to do oceny jakości klasyfikatorów w rozprawie wykorzystano 4 miary specjalnie dobrane do tej sytuacji. Były to miary MAE, RMSE, MAPE i WMAPE. Zdaniem recenzenta, szczególnie użyteczna jest miara MAPE, która mierzy średnią różnicę wartości prognozowanych i rzeczywistych względem wartości rzeczywistej dla ustalonej listy produktów. Niestety, przy prezentacji wartości tych miar nie podano odchylenia standardowego dla tych miar. Skutkiem tego nie ma w rozprawie informacji o tym czy jakość prognozowania promocji jest zawsze na podobnym poziomie, czy też dla jednych produktów może być znacząco lepsza, a dla innych znacząco gorsza.
5. Jak już pisałem, do przeprowadzenia eksperymentów związanych z rozprawą, wykorzystano wiele zbiorów danych. Jednak, z powodów poufności, tylko dwa spośród nich mogły być udostępnione. Dlatego dla innych zbiorów danych spośród wykorzystanych zbiorów, czytelnik nie może wykonać eksperymentów porównawczych. Stąd nasuwa się pytanie: czy Autorka rozprawy nie mogła udostępnić tych danych przynajmniej w wersji skróconej w formie zakodowanej? Chodzi o sytuację, gdy np. w danych, gdzie występuje nazwa produktu, zamiast tej nazwy piszemy „ProduktX” itp. Być może jednak nawet na to nie mogło być zgody partnerów biznesowych.

5. Działalność publikacyjna

Autorka rozprawy opublikowała łącznie 13 publikacji, w tym 6 w czasopiśmie naukowych (zwykle z wysokim współczynnikiem IF (ang. *impact factor*) oraz 7 w recenzowanych materiałach pokonferencyjnych. Zdaniem recenzenta, jest to znaczący dorobek jak na osobę, która ukończyła studia magisterskie w 2019 roku.

Aktualne główne wskaźniki bibliometryczne Autorki rozprawy to:

- według źródła Scopus: index Hirscha=3, liczba cytowań=26,
- według źródła Google Scholar: index Hirscha=4, liczba cytowań=48.

Są to zatem wskaźniki na dobrym poziomie, zważając na młody wiek Autorki rozprawy.

6. Podsumowanie

Uzyskane wyniki są interesujące zarówno z teoretycznego, jak i praktycznego punktu widzenia. Dlatego niezależnie od wspomnianych wyżej mankamentów, uważam pracę za wartościową. Autorka wykazała dobre opanowanie wielu różnorodnych technik matematycznych i

informatycznych. Sposób wykorzystania tych technik wskazuje na rzetelne opanowanie przez Nią warsztatu naukowego.

Biorąc pod uwagę opinie zaprezentowane w poprzednich punktach, moja ocena rozprawy pod względem trzech zwyczajowo rozpatrywanych kryteriów jest następująca:

1. rozprawa doktorska zawiera oryginalne rozwiązanie problemu naukowego
2. Kandydatka posiada ogólną wiedzę teoretyczną w dyscyplinie,
3. Kandydatka ma umiejętność samodzielnego prowadzenia pracy naukowej.

W związku z powyższym, **wnioskuję o dopuszczenie rozprawy doktorskiej do publicznej obrony.**

Ponadto, mając na uwadze osiągnięcia publikacyjne kandydatki uzyskane w okresie niespełna 4 lat (13 publikacji, z których 6 w czasopiśmie z listy A z wysokim IF, a 7 w recenzowanych materiałach konferencji międzynarodowej), znaczący index Hirscha i liczba cytowań, wysoka wartość techniczna przedstawionej rozprawy, której napisanie wymagało wykorzystania złożonego środowiska sprzętowo-programowego (konieczne było opanowanie wielu metod, języków i technologii) oraz realne zastosowanie zaproponowanych w rozprawie metod w analizie rzeczywistych danych pochodzących z wielu firm, **rekomenduję wyróżnienie rozprawy doktorskiej.**