

dr hab. Adrian Horzyk, prof. AGH

Kraków, 07 stycznia 2026

Katedra Biocybernetyki i Inżynierii Biomedycznej
Wydział Elektrotechniki, Automatyki, Informatyki i Inżynierii Biomedycznej
Akademia Górniczo-Hutnicza w Krakowie
30-059 Kraków, al. Mickiewicza 30

Recenzja rozprawy doktorskiej mgr Piotra Tempczyka

Formalną podstawą sporządzenia niniejszej recenzji jest pismo Przewodniczącego Rady Naukowej Dyscypliny Informatyka Uniwersytetu Warszawskiego z dnia 30.10.2025 r., wydane na podstawie uchwały Rady Naukowej Dyscypliny Informatyka, w którym zwrócono się do mnie o sporządzenie recenzji rozprawy doktorskiej mgr Piotra Tempczyka, w postępowaniu w sprawie nadania stopnia doktora w dziedzinie nauk ścisłych i przyrodniczych, w dyscyplinie informatyka, zatytułowanej: „*Estimating local intrinsic dimension via density estimation*” („*Estymacja lokalnej wymiarowości rozmaitości danych za pomocą modeli estymacji gęstości*”), przygotowanej pod kierunkiem dr. hab. Marka Cygana, prof. ucz., w Instytucie Informatyki Wydziału Matematyki, Informatyki i Mechaniki Uniwersytetu Warszawskiego. Postępowanie prowadzone jest zgodnie z art. 186–190 ustawy z dnia 20 lipca 2018 r. Prawo o szkolnictwie wyższym i nauce oraz z załącznikiem nr 1 do uchwały nr 157 Senatu Uniwersytetu Warszawskiego z dnia 29 czerwca 2022 r. w sprawie określenia sposobu postępowania w sprawie nadania stopnia doktora, wraz z późniejszymi zmianami. Autorem ocenianej rozprawy doktorskiej jest mgr Piotr Tempczyk. Rozprawa liczy 122 strony głównego tekstu, poprzedzonego abstraktem i streszczeniem w języku polskim oraz zakończonych dodatkami technicznymi. Strukturalnie obejmuje siedem rozdziałów merytorycznych, notę redakcyjną, podziękowania oraz obszerny aneks. Praca została napisana w języku angielskim, z krótkim polskim streszczeniem i utrzymana jest w typowej dla informatyki teoretyczno-stosowanej formie monografii naukowej.

1. Ogólna charakterystyka rozprawy

Oceniana rozprawa podejmuje problem estymacji lokalnego wymiaru wewnętrznego (Local Intrinsic Dimension – LID) danych wysokowymiarowych, rozumianego jako liczba efektywnych stopni swobody w małym otoczeniu punktu w przestrzeni cech. Autor wychodzi od tzw. Hipotezy rozmaitości (Manifold Hypothesis), zgodnie z którą obserwacje w wielu rzeczywistych zadaniach uczenia maszynowego koncentrują się w pobliżu niżejwymiarowych rozmaitości

osadzonych w przestrzeniach o bardzo dużej liczbie wymiarów. W tym kontekście „wymiar” nie jest traktowany jako jedna globalna liczba, lecz jako własność lokalnych sąsiedztw, zależna zarówno od położenia, jak i od przyjętej skali obserwacji.

Doktorant stawia sobie za cel zaprojektowanie i zbadanie metod estymacji LID, które byłyby z jednej strony teoretycznie ugruntowane, a z drugiej – skalowalne do współczesnych zbiorów danych, w których liczba wymiarów i złożoność struktur znacznie przekracza możliwości klasycznych estymatorów opartych na najbliższych sąsiadach czy prostych skalowaniach fraktalnych. Centralnym osiągnięciem rozprawy jest zaproponowanie nowego estymatora LID – LIDL (Local Intrinsic Dimension via Likelihood) – wykorzystującego nowoczesne modele estymacji gęstości (normalizing flows) oraz analizę zachowania log-gęstości punktu przy kontrolowanym dodawaniu szumu Gaussa.

Praca ma wyraźnie zarysowaną oś tematyczną: od formalnego wprowadzenia do pojęcia LID i przeglądu klasycznych metod, poprzez konstrukcję i analizę teoretyczną nowej metody LIDL, ujednolicenie współczesnych podejść neuronowych z użyciem aparatu procesu Wienera, aż po szeroko zakrojone eksperymenty porównawcze i zaproponowanie benchmarków, które lepiej oddają złożoność realnych danych. Rozprawa integruje aparat matematyczny (analizę rozkładów i operatorów różniczkowych), techniki uczenia głębokiego (normalizing flows, modele dyfuzyjne) oraz projektowanie zestawów testowych i transformacji danych specyficznych dla domeny obrazowej.

Kolejne rozdziały rozprawy nie są zbiorem luźno powiązanych publikacji, lecz logicznie zorganizowaną całością, w której nowa metoda, jej analiza i szerokie testowanie wzajemnie się uzupełniają. W tym sensie rozprawa odpowiada zarówno wymaganiu przedstawienia oryginalnego rozwiązania istotnego problemu naukowego, jak i wymaganiu wykazania przez autora umiejętności samodzielnego prowadzenia pracy naukowej – w rozumieniu art. 187 ustawy z dnia 20 lipca 2018 r. – Prawo o szkolnictwie wyższym i nauce (Dz. U. z późn. zm.).

2. Szczegółowa ocena zawartości pracy

Rozprawa składa się z części wprowadzającej, części metodologicznej, rozdziałów poświęconych konstrukcji i analizie nowej metody, porównań z istniejącymi algorytmami oraz bogatej części eksperymentalnej. Autor w przejrzysty sposób przedstawia problematykę LID, omawia dotychczasowe podejścia (m.in. metody geometryczne, oparte na odległościach oraz estymacji rozkładu), a następnie wprowadza własny aparat formalny oparty na analizie krótkoczasowej dyfuzji rozkładu. Szczególnie wartościowe jest powiązanie interpretacji probabilistycznych z intuicjami geometrycznymi – czytelnik otrzymuje jasny

obraz tego, w jaki sposób zachowanie gęstości w otoczeniach punktów ujawnia lokalną strukturę danych.

Rozdział 1 ma charakter rozbudowanego wprowadzenia. Autor najpierw precyzyjnie formułuje problem badawczy, osadzając go w kontekście Hipotezy różnorodności oraz potrzeby lokalnego, zależnego od skali ujęcia wymiaru. Następnie szeroko omawia zastosowania LID: od analizy dynamiki uczenia sieci neuronowych, przez dobór wymiaru przestrzeni ukrytej w modelach generatywnych, diagnostykę jakości reprezentacji, detekcję próbek out-of-distribution, po analizę generalizacji i zjawisk, takich jak zapamiętywanie przykładów w modelach generatywnych. Ta część stanowi wartościowy przegląd motywacji praktycznych i pokazuje, że Autor swobodnie porusza się w aktualnej literaturze dotyczącej uczenia głębokiego i analizy reprezentacji. Kolejna część rozdziału poświęcona jest historii badań nad wymiarem wewnętrznym – od klasycznych metod PCA i estymatorów opartych na sąsiadach oraz wymiarach fraktalnych, poprzez estymatory maksymalno-likelihoodowe, metody kątowe i oparte na prostych, aż po współczesne formalizacje LID oraz ich zastosowania w wyszukiwaniu podobnych obiektów, detekcji anomalii i badaniu odporności modeli na ataki.

W podrozdziale „My Contribution” Autor systematycznie wylicza własne osiągnięcia: zaproponowanie estymatora LIDL, sformułowanie perspektywy procesu Wienera dla współczesnych algorytmów LID oraz zaprojektowanie nowego zestawu benchmarków. Każdy z tych wkładów jest następnie rozwinięty w osobnym podrozdziale, z opisem głównych idei, zalet i ograniczeń. W części „Related Work” Doktorant dokonuje starannej kategoryzacji literatury (surveys, fundamenty, estymatory oparte na sąsiadach, kątach, prostych, formalizacje LID, inne neuronowe estymatory, zastosowania w uczeniu reprezentacji) i lokuje na tym tle własne prace. Rozdział zamyka „Roadmap”, która bardzo przejrzysto opisuje, które rozdziały bazują na jakich publikacjach i jak całość rozprawy jest zorganizowana.

Rozdział 2 stanowi kluczowy trzon rozprawy i przedstawia autorską metodę szacowania lokalnego wymiaru wewnętrznego LIDL (Local Intrinsic Dimension using approximate Likelihood). Autor rozpoczyna od intuicyjnego wyjaśnienia idei stojącej za metodą: dodanie do danych izotropowego szumu Gaussa o wariancji δ^2 można traktować jako stopniowe „nadmuchiwanie” różnorodności, na której dane są skupione. Analizując, jak zmienia się logarytm gęstości w punkcie x w funkcji $\log \delta$, Autor argumentuje, że tempo tej zmiany odzwierciedla lokalny wymiar struktury danych.

Następnie przechodzi do formalnego ujęcia problemu. Wprowadza pojęcie gładkich miar na różnorodnościach oraz modeluje dane jako rozkład wsparty na gładkiej, zanurzonej różnorodności $S \subset \mathbb{R}^D$. Zaburzony rozkład uzyskuje poprzez

splot z rozkładem normalnym $N(0, \delta^2 I)$, a następnie bada zależność gęstości $\rho_\delta(x)$ od parametru δ . Rozbudowana część teoretyczna prowadzi do kluczowego wyniku: dla odpowiednio małych, lecz niezerowych wartości δ , nachylenie zależności $\log \rho_\delta(x)$ od $\log \delta$ jest w przybliżeniu równe $d - D$, gdzie d oznacza lokalny wymiar rozmaitości, zaś D – wymiar przestrzeni otaczającej. Autor rozwija przy tym aparat pozwalający oddzielić kierunki styczne i normalne, posługując się projekcjami oraz lokalnymi rozwinięciami Taylora.

W dalszej części rozdziału parametrowi δ nadana zostaje interpretacja parametru skali, co pozwala kontrolować, które struktury danych są widoczne, a które ulegają „wygładzeniu”. Rozważania zostają następnie rozszerzone na sytuacje bardziej złożone: rozmaitości niepołączone, przecinające się oraz przypadki o skomplikowanej geometrii, z wykorzystaniem pojęcia dobrych immersji i odpowiednich projekcji. Szczególnie wartościowe są liczne przykłady obliczeniowe, obejmujące m.in. rozkład normalny w $\mathbb{R}^D$, punkty na prostej, rozkłady jednorodnie na odcinku oraz na zakrzywionych rozmaitościach, które pokazują, jak zachowuje się estymator zarówno w warunkach „idealnych”, jak i w obecności brzegów lub krzywizny. Rozdział zamykają eksperymenty numeryczne w sytuacjach, w których ρ_δ można obliczyć dokładnie lub z dużą precyzją (np. dla odcinka, okręgu czy sąsiadujących rozmaitości), co pozwala jakościowo zweryfikować przewidywania teoretyczne i zidentyfikować źródła obciążenia estymatora.

Rozdział 3 stanowi próbę ujednolicenia współczesnych algorytmów estymacji lokalnego wymiaru wewnętrznego w jednym, spójnym aparacie analitycznym. Autor interpretuje dodawanie kontrolowanego szumu Gaussa jako krótkoczasową ewolucję rozkładu danych wzdłuż procesu Wienera i korzysta z równania dyfuzji (tzw. drugiego prawa Ficka), aby opisać, w jaki sposób zmienia się gęstość danych w czasie. W takiej perspektywie różne metody (w szczególności LIDL, FLIPD, ID-NF, ID-DM oraz NB) wykorzystują w istocie tę samą informację o zachowaniu gęstości ρ_t , choć czynią to w odmienny sposób: poprzez analizę nachylenia log-gęstości, widma Jacobianu odpowiednich odwzorowań do rozkładu odniesienia, bądź geometrii pola gradientu log-gęstości (score field). W dalszej części wprowadzona zostaje wielkość $\beta_t(x) = t \frac{\Delta \rho_t(x)}{\rho_t(x)}$, a następnie analizowana jest jej granica przy $t \rightarrow 0$. Autor interpretuje ją jako miarę lokalnej reakcji rozkładu na dyfuzję i wykazuje jej ścisły związek z lokalnym wymiarem rozmaitości. Konstrukcja ta pozwala spojrzeć na szereg istniejących algorytmów z jednej, wspólnej perspektywy.

Następnie Autor przeprowadza systematyczną analizę szeregu kanonicznych przykładów: rozkładu jednorodnego w $\mathbb{R}^D$, rozkładu normalnego, rozkładów jednorodnych na przedziałach i hipersześcianach, sum rozkładów wspartych na rozłącznych lub przecinających się rozmaitościach, a także kombinacji wypukłych

rozkładów. Analiza ta pozwala prześledzić, jak zachowują się poszczególne estymatory w sytuacjach kontrolowanych, w których struktura geometryczna i probabilistyczna danych jest dobrze znana. Na tej podstawie Autor identyfikuje główne źródła obciążenia estymatorów LID (wynikające z niejednorodności gęstości, krzywizny rozmaitości, obecności brzegów oraz specyficznych interakcji między komponentami rozkładu) oraz proponuje poprawki w postaci wzorów zamkniętych, umożliwiających lepsze zrozumienie odchyleń obserwowanych w eksperymentach od wartości „idealnych”. Rozdział ten stanowi silny element teoretyczny rozprawy i pokazuje, że Autor nie ogranicza się do konstrukcji jednego algorytmu, lecz dąży do osadzenia całej klasy metod w przejrzystym i matematycznie spójnym formalizmie.

Rozdział 4 ma charakter porównawczy. Na początku Autor krótko przypomina ideę normalizing flows jako modeli gęstości (odwołując się do literatury, szczegóły implementacyjne przenosząc do dodatków), a następnie opisuje protokół porównawczy między LIDL a wybranymi klasycznymi estymatorami LID – w szczególności opartymi na odległościach i sąsiadach.

Koncentruje się na kilku aspektach: wpływie regresji liniowej na jakość estymacji w LIDL, skalowalności do wysokich wymiarów rzędu tysięcy, zachowaniu na rozmaitościach wielkoskalowych, gdzie lokalny wymiar zmienia się z poziomem obserwacji, a także na rozmaitościach zakrzywionych i sumach rozmaitości. Autor pokazuje, że w sytuacjach, w których klasyczne estymatory ulegają silnemu obciążeniu (biasowi) lub stają się niestabilne (efekt koncentracji miary, wysokie wymiary, rozmaitości wieloskalowe), LIDL zachowuje względnie niski błąd, o ile gęstość jest odpowiednio estymowana. Z drugiej strony Autor uczciwie wskazuje ograniczenia związane z kosztem uczenia modeli gęstości oraz doбором okna poziomu szumu.

Rozdział 5 przenosi analizę na dane rzeczywiste – zbiory obrazów MNIST, Fashion-MNIST i CelebA. Autor bada, jak rozkład LID odpowiada intuicyjnej „złożoności” obrazów oraz struktury klas, pokazując m.in., że obrazy prostsze (np. proste cyfry) mają niższy LID niż bardziej skomplikowane (np. obrazy twarzy w CelebA). Analizowany jest zakres operacyjny δ : na prostych danych syntetycznych obserwuje się oczekiwane „skokowe” zmiany wraz z przełączaniem między skalami, natomiast na obrazach zbyt duże wartości δ prowadzą do zlewania się cienkich rozmaitości, a w konsekwencji do zaniżania wymiaru. Autor bada również wpływ dekwantyzacji obrazów oraz zwiększania liczby modeli gęstości na jakość estymacji, pokazując, że większa liczba niezależnych modeli prowadzi do monotonicznego zmniejszania błędu. Ważnym elementem jest pokazanie związków LID z jakością rekonstrukcji w autoenkoderach wariacyjnych oraz z dokładnością klasyfikacji. Okazało się, iż najwyższe wartości LID korelują z większymi błędami rekonstrukcji i niższą pewnością predykcji, co sugeruje użyteczność LID jako wskaźnika trudności.

Rozdział 6 jest rozbudowaną propozycją ram benchmarkowych dla nowoczesnych algorytmów estymacji LID. Autor rozpoczyna od krytycznej obserwacji, że dotychczasowe testy albo opierały się na bardzo prostych rozkładach syntetycznych (gdzie znany jest dokładny wymiar, ale brakuje realizmu), albo na danych złożonych, jak obrazy, gdzie brakuje prawdziwej wartości odniesienia (ground truth) umożliwiającej obiektywną ocenę metod. W odpowiedzi projektuje zestaw transformacji domeny, takich jak Inverse Domain Representation, Monotonic Embedding, Ambient Space Extension, Auxiliary Dimension Injection oraz Manifold Synthesis, które pozwalają przenosić rozmaitości o znanym LID do bardziej złożonych przestrzeni (np. obrazów) przy zachowaniu kontroli nad wymiarem. Następnie przedstawia szczegółowe eksperymenty porównujące reprezentatywne algorytmy: klasyczny ESS oraz neuronowe: LIDL, NB, FLIPD, w scenariuszach stresujących dla estymatorów: przy niejednorodnych gęstościach, krzywiźnie rozmaitości, obecności brzegów, cienkich struktur, bliskości komponentów czy zmiennej liczebności próby. Autor pokazuje, że każda z metod ma obszary, w których radzi sobie bardzo dobrze, oraz takie, w których ponosi porażkę; w szczególności, architektura sieci oraz dobór parametrów szumu mają istotny wpływ na wyniki. Rozdział kończy się zestawieniem wniosków dla poszczególnych algorytmów oraz wskazaniem, że proponowane benchmarki mogą stać się standardem testowania w tej dziedzinie.

Rozdział 7 syntetyzuje wyniki wcześniejszych części rozprawy, podkreślając trzy główne osie wkładu: konstrukcję i analizę LIDL, ujęcie perspektywy procesu Wienera jako wspólnego języka dla nowoczesnych metod LID, oraz zaprojektowanie benchmarków zdolnych lepiej odzwierciedlać trudności realnych danych. Autor uczciwie omawia ograniczenia swojej pracy, m.in. koszty obliczeniowe, specyfikę danych obrazowych, oraz zarysowuje kierunki dalszych badań – zarówno teoretycznych (np. rozszerzenie analizy na bardziej ogólne klasy rozkładów), jak i praktycznych (np. zastosowanie LID w kontroli treningu modeli generatywnych).

Załączony aneks zawiera szczegóły techniczne: alternatywne sformułowania LIDL, parametry eksperymentów, konfiguracje modeli ESS, NB, LIDL, FLIPD oraz dodatkowe wyniki numeryczne, których włączenie do głównego tekstu zaburzyłoby jego czytelność.

3. Osiągnięcia naukowe oraz wpływ na rozwój dyscypliny naukowej

Centralnym osiągnięciem Doktoranta rozprawy jest zaproponowanie i dogłębna analiza nowego estymatora lokalnego wymiaru wewnętrznego – LIDL (Local Intrinsic Dimension via Likelihood). Proponowana konstrukcja łączy klasyczne intuicje geometryczne z nowoczesnymi modelami estymacji gęstości (normalizing

flows) oraz z analizą zachowania log-gęstości przy kontrolowanym dodawaniu izotropowego szumu Gaussa. Autor nie ogranicza się do zaprojektowania algorytmu: przedstawia spójne uzasadnienie teoretyczne, bada własności estymatora w modelach wzorcowych, identyfikuje źródła obciążenia i proponuje sposoby jego ograniczania. W tym sensie opracowany estymator stanowi oryginalne rozwiązanie problemu naukowego w rozumieniu art. 187 ustawy – Prawo o szkolnictwie wyższym i nauce.

Po drugie, perspektywa procesu Wienera, zastosowana do ujednolicenia i analizy współczesnych estymatorów neuronowych (LIDL, FLIPD, ID-NF, ID-DM, NB), ma wyraźny potencjał porządkujący całe pole badawcze. Pokazuje ona, że różne metody „czytają” ten sam sygnał dyfuzyjny w odmiennych reprezentacjach (gęstości, score, transformacjach odwzorowań), co pozwala lepiej zrozumieć zarówno ich mocne strony, jak i ograniczenia.

Po trzecie, zaproponowany przez Autora rozbudowany framework benchmarkowy – oparty na kontrolowanych transformacjach domeny oraz zestawie danych syntetyczno-realistycznych – ustanawia nowe standardy testowania algorytmów LID. Pozwala on oceniać metody w warunkach znacznie bliższych zastosowaniom praktycznym niż tradycyjne, często silnie uproszczone scenariusze.

Z rozdziału 1.5 oraz listy publikacji wynika, że istotna część wyników rozprawy była opublikowana, upubliczniona i pozytywnie zweryfikowana w procesach recenzyjnych: LIDL został przedstawiony na konferencji ICML, perspektywa procesu Wienera na AAAI, zaś prace benchmarkowe są w toku recenzji. Wskazuje to na uznanie środowiska międzynarodowego i potwierdza dojrzałość prowadzonych badań.

Za najważniejsze osiągnięcia naukowe Autora należy zatem uznać:

1. opracowanie nowego estymatora lokalnego wymiaru wewnętrznego LIDL (Local Intrinsic Dimension via Likelihood), opartego na analizie log-gęstości i nowoczesnych modelach estymacji gęstości,
2. zaproponowanie spójnego formalizmu obejmującego kilka klas metod LID,
3. skonstruowanie zaawansowanego środowiska eksperymentalnego do systematycznej oceny metod,
4. przeprowadzenie krytycznej analizy wyników, pozwalającej określić zakres stosowalności poszczególnych podejść.

Z punktu widzenia rozwoju dyscypliny informatyka rozprawa wnosi zarówno nowe narzędzia obliczeniowe, jak i aparat teoretyczny służący do analizy

struktury danych wysokowymiarowych. Może mieć znaczenie dla projektowania modeli głębokich, doboru wymiarów przestrzeni ukrytych, diagnostyki procesu uczenia oraz analizy zachowania modeli generatywnych. W mojej ocenie całość stanowi oryginalny i znaczący wkład w rozwój metod analizy struktur danych w uczeniu maszynowym oraz wskazanej dyscypliny naukowej.

4. Uwagi dyskusyjne

Moja ocena merytoryczna recenzowanej rozprawy jest jednoznacznie pozytywna, niemniej uważam za celowe wskazanie kilku uwag dyskusyjnych, które nie podważają końcowego wniosku, ale mogą być pomocne przy dalszym rozwijaniu tego kierunku badań.

Uwaga 1. Choć Autor w wielu miejscach podkreśla wrażliwość LIDL na wybór modelu gęstości (normalizing flow) oraz na jakość jego trenowania, analiza eksperymentalna w tym zakresie jest ograniczona. Większość przykładów wykorzystuje jedną lub kilka blisko spokrewnionych architektur, a wpływ architektury, głębokości czy klasy zastosowanych flowów na jakość estymacji LID nie jest systematycznie zbadany. W praktyce użytkownik metody może stanąć przed problemem wyboru konkretnej architektury i zestawu hiperparametrów, a rozprawa nie daje pełnego, empirycznie ugruntowanego przewodnika po tym wymiarze decyzji, co może być istotnym ograniczeniem w popularyzacji przedstawionego podejścia i metody. Być może wynika to z ograniczeń objętości pracy, ale dla kompletności wartościowe byłoby choć częściowe studium ablacyjne, analogiczne do tego, jakie Autor przeprowadza dla parametrów związanych z poziomem szumu.

Uwaga 2. Chociaż metody LIDL i pokrewne mają charakter ogólny i mogłyby być stosowane do różnych typów danych (teksty, dane sekwencyjne, tablicowe), zasadnicza część badań na danych rzeczywistych koncentruje się na obrazach (MNIST, FMNIST, CelebA). Jest to zrozumiałe ze względu na bogactwo i popularność tych benchmarków, ale jednocześnie ogranicza zakres empirycznej walidacji. Brakuje choćby pojedynczych studiów przypadku na innej klasie danych, które pokazałyby, na ile wnioski z obrazów przenoszą się na tekst czy sygnały czasowe. W przyszłych pracach warto byłoby rozszerzyć eksperymenty o takie scenariusze, zwłaszcza że modele gęstości i dyfuzyjne są dziś rozwijane również w tych domenach.

Uwaga 3. W części teoretycznej opartej na procesie Wienera przyjmowane są stosunkowo silne założenia regularności gęstości oraz geometrii rozmaitości (gładkość, brak osobliwości), co jest typowe dla tej klasy wyników. W rozprawie brakuje jednak nieco bardziej szczegółowej dyskusji, na ile naruszenia tych założeń, np. ostre brzegi, „szpiczaste” rozmaitości, rozkłady z ciężkimi ogonami –

przekładają się na praktyczne zachowanie estymatorów. Autor pośrednio ilustruje część tych zjawisk w rozdziałach eksperymentalnych, ale połączenie wyników teoretycznych z obserwacjami empirycznymi mogłoby być mocniejsze, np. w postaci sekcji wyraźnie zestawiającej teorię z konkretnymi „patologiami” danych.

Uwaga 4. Proponowany framework benchmarkowy, choć bardzo cenny, opiera się na konkretnych wyborach parametrów transformacji domeny (IDR, ME, ASE, ADI, MS). Jest to nieuniknione, jednak wybory te mają charakter częściowo arbitralny. Autor słusznie wskazuje, że celem jest „stresowanie” algorytmów pod kątem określonych cech różnorodności, ale w przyszłości warto byłoby uzupełnić ten program o bardziej systematyczne poszukiwanie najtrudniejszych przypadków, np. w duchu adversarial benchmark design, tak aby można było powiedzieć, że dany estymator wytrzymuje nie tylko zaprojektowane scenariusze, ale także „najgorsze” konfiguracje w ramach zadanych klas danych.

Uwaga 5. Choć rozprawa przedstawia interesujące korelacje między LID a jakością rekonstrukcji w VAE i wynikami klasyfikacji, zastosowania te mają charakter głównie eksploracyjny. Brakuje przynajmniej jednego bądź dwóch pogłębionych studiów przypadku, w których LID byłby faktycznie użyty jako narzędzie decyzyjne, np. do projektowania architektury, wyboru wymiaru latent space czy dynamicznego włączania/wyłączania augmentacji, jak i pokazania, że taka interwencja prowadzi do mierzalnej poprawy działania modelu. Byłoby to bardzo mocnym argumentem na rzecz praktycznej użyteczności koncepcji, którą Autor tak przekonująco uzasadnia teoretycznie.

Powyższe uwagi nie podważają jednak głównego wniosku co do wysokiej wartości naukowej i oryginalności rozprawy, a stanowią bardziej konstruktywną krytykę i wsparcie dla dalszego rozwoju i uogólnień przedstawionego podejścia i skonstruowanej metody.

Uwagi natury językowej i edytorskiej:

Rozprawa jest napisana bardzo poprawnym językiem angielskim, typowym dla współczesnej literatury z zakresu uczenia maszynowego. Styl jest klarowny, precyzyjny, dobrze dopasowany do treści. Struktura rozdziałów, podrozdziałów i oznaczeń (twierdzeń, definicji, przykładów) jest logiczna. Ilustracje, wykresy i schematy, szczególnie te przedstawiające zachowanie estymatorów LID na różnorodnościach syntetycznych oraz wyniki na zbiorach obrazów, są czytelne i dobrze opisane, choć w kilku miejscach można by rozważyć zwiększenie rozdzielczości lub doprecyzowanie podpisów, aby ułatwić odczyt wartości numerycznych. Drobne literówki, powtórzenia czy błędy interpunkcyjne są bardzo rzadkie i nie wpływają zasadniczo na przekaz treści merytorycznej pracy ani jej ocenę, np.:

Str. 46: "By this point, the notion of *nice* density **is shall be** more clear."

Str. 47: "In Fig. 3.1a**[brak przecinka]** we can observe"

Str. 51: "In the following discussion**[brak przecinka]** it turns out that"

Str. 59: „Tables. 4.1, 4.2, A.2).” bez kropki po Tables.

Polskie streszczenie jest poprawne językowo, choć zawiera pojedyncze drobne literówki (np. sporadyczne braki polskich znaków diakrycznych), nieliczne niezgrabności składniowe, nie wpływające jednak na zrozumiałość tekstu, np.:

Str. 6: „estymator LID **uzywajacy** estymatorów gęstości”

Str. 6: „o modele gęstości **uzywajacye** sieci neuronowych”

Str. 6: „LID dla **różnych** gęstości prawdopodobieństwa”

Str. 6: „w pobliżu **niżejwymiarowych** struktur geometrycznych” brzmi mało naturalnie po polsku, może warto byłoby stosować bardziej rozbudowaną wersję „w pobliżu **struktur o niższym wymiarze geometrycznym**” lub „w pobliżu **struktury o mniejszej wymiarowości**”.

Redakcja matematyczna – zapisy wzorów, numeracja równań i odwołania – jest konsekwentna; zauważyłem jedynie sporadyczne drobiazgi, jak niejednolitość w zapisie nazw własnych (np. w bibliografii część tytułów jest dużymi a część tylko rozpoczyna się od dużej, a reszta małymi literami, czy też brak kropek po skrótach imion, np. „Michael **E** Houle” czy też „Jakob **H** Macke”), które nie mają znaczenia merytorycznego. Całość sprawia wrażenie pracy starannie zredagowanej.

5. Ocena dorobku naukowego

Z informacji zawartych w rozprawie, w szczególności w części „Roadmap” oraz w bibliografii, wynika, że zasadnicze wyniki pracy zostały opublikowane bądź zgłoszone do publikacji w renomowanych czasopismach i na konferencjach z zakresu uczenia maszynowego: praca o LIDL ukazała się na ICML, praca o perspektywie procesu Wienera na AAAI, a manuskrypt dotyczący benchmarków LID jest w trakcie recenzji. Autor rozprawy jest pierwszym i głównym autorem tych dwóch publikacji konferencyjnych, a są to konferencje o bardzo wysokich standardach recenzyjnych, co potwierdza jakość prowadzonych badań oraz uzyskanych wyników badawczych Autora rozprawy.

Choć pełny, usystematyzowany wykaz publikacji kandydata nie był udostępniony w materiałach recenzyjnych, z samej rozprawy wynika, że Doktorant jest aktywnym współautorem kilku prac z obszaru analizy wymiarowości i uczenia

reprezentacji. Na podstawie prezentacji wyników i stylu ich omówienia można wnioskować, że odgrywał w nich wiodącą rolę, zarówno koncepcyjną, jak i implementacyjną. Dorobek ten uznaję za adekwatny, a nawet wysoki w odniesieniu do standardów stawianych przewodom doktorskim w dziedzinie nauk ścisłych i przyrodniczych, dyscyplinie informatyka.

6. Podsumowanie i wniosek końcowy

Podsumowując, stwierdzam, że rozprawa doktorska mgr Piotra Tempczyka pt. „*Estimating local intrinsic dimension via density estimation*” stanowi **oryginalne i wartościowe osiągnięcie naukowe w dziedzinie nauk ścisłych i przyrodniczych w dyscyplinie informatyka**. Praca rozwiązuje istotny problem naukowy, estymacji lokalnego wymiaru wewnętrznego danych wysokowymiarowych, proponując nowy estymator LIDL, ugruntowując klasę współczesnych metod w perspektywie procesu Wienera oraz wprowadzając nowe, ambitne benchmarki dla tej klasy algorytmów. Rozprawa świadczy o bardzo dobrej znajomości literatury przedmiotu, wysokich kompetencjach matematycznych i algorytmicznych Autora rozprawy oraz jego umiejętności samodzielnego prowadzenia badań, formułowania wniosków i krytycznej analizy własnych wyników.

W świetle przeprowadzonej analizy stwierdzam, że oceniana **rozprawa spełnia wymagania** określone w art. 187 ustawy z dnia 20 lipca 2018 r., Prawo o szkolnictwie wyższym i nauce, stawiane rozprawom w postępowaniu o nadanie stopnia doktora w dziedzinie nauk ścisłych i przyrodniczych, w dyscyplinie informatyka, a także wymagania określone w załączniku nr 1 do uchwały nr 157 Senatu Uniwersytetu Warszawskiego z dnia 29 czerwca 2022 r. w sprawie określenia sposobu postępowania w sprawie nadania stopnia doktora.

Wobec powyższego wnioskuję do Rady Naukowej Dyscypliny Informatyka Uniwersytetu Warszawskiego **o przyjęcie rozprawy jako rozprawy doktorskiej oraz o dopuszczenie mgr Piotra Tempczyka do dalszego postępowania, w tym publicznej obrony**, zgodnie z zasadami określonymi w załączniku nr 1 do uchwały nr 157 Senatu Uniwersytetu Warszawskiego. Jednocześnie rekomenduję nadanie mu stopnia doktora w dziedzinie nauk ścisłych i przyrodniczych, w dyscyplinie informatyka.

Mając na uwadze oryginalność uzyskanych rezultatów, ich rangę potwierdzoną publikacjami w wysoko punktowanych, recenzowanych międzynarodowych konferencjach naukowych, a także spójność i dojrzałość całego programu badawczego, uważam, że rozprawa wyraźnie przekracza standardowy poziom rozpraw doktorskich. Wobec powyższego **wnioskuję o wyróżnienie rozprawy doktorskiej mgr Piotra Tempczyka**.

