

Poznań, 25-05-2026

dr hab. inż. Paweł Śniatała, prof. PP
Wydział Informatyki i Telekomunikacji
Politechniki Poznańskiej
ul. Piotrowo 2, 60-965 Poznań
pawel.sniatala@put.poznan.pl



RECENZJA ROZPRAWY DOKTORSKIEJ

mgr inż. Piotra Kopy Ostrowskiego

zatytułowanej:

Efficient Deep Neural Networks for Video Denoising

wykonana na podstawie Umowy zawartej z Wydziałem Elektroniki, Telekomunikacji i Informatyki Politechniki Gdańskiej.

1. Problem badawczy i jego znaczenie

Głównym problemem badawczym zdefiniowanym w rozprawie jest opracowanie technik, które umożliwią poprawę wydajności modeli AI używanych do usuwania szumów z obrazów wideo poprzez lepsze wykorzystanie kontekstu czasowego.

W recenzowanej pracy nacisk położono na filmy nagrane w spektrum widzialnym, tj. wstępnie przetworzone dane RGB lub RAW. Są one najpopularniejsze i mają najszerszy zakres zastosowań. Proponowane techniki można potencjalnie rozszerzyć na inne modalności, takie jak filmy hiperspektralne, ultrasonograficzne, filmy przedstawiające głębię czy też filmy oparte na zdarzeniach.

Wideo głębokie (Depth Video) to sekwencja obrazów (podobnie jak w przypadku zwykłego wideo), w której każdy piksel odzwierciedla odległość, a nie kolor. Zamiast rejestrować światło czerwone, zielone i niebieskie, czujnik głębi oblicza dokładnie, ile metrów lub milimetrów dzieli obiekt od kamery. Ten rodzaj wideo ma zastosowanie w rzeczywistości rozszerzonej (AR), robotyce, skanowaniu 3D oraz efektach trybu portretowego w smartfonach. Dzięki najnowszym rozwiązaniom w dziedzinie sztucznej inteligencji oprogramowanie może przekształcać zwykłe filmy 2D w filmy z głębią.

Wideo oparte na zdarzeniach to materiał filmowy zarejestrowany lub wygenerowany przy użyciu neuromorficznych kamer zdarzeniowych, które naśladują działanie ludzkiego oka. W odróżnieniu od tradycyjnych kamer, które rejestrują serię pełnych obrazów ze stałą częstotliwością klatek, kamery zdarzeniowe rejestrują jedynie zmiany natężenia światła na poziomie pikseli, przysyłając dane asynchronicznie z dokładnością do mikrosekund. Ponieważ kamera zdarzeniowa ignoruje statyczne tło i rejestruje wyłącznie ruch, generuje nawet 1000 razy mniej danych niż konwencjonalne czujniki.

Jak wspomniano powyżej, Autor w swojej pracy skupił się głównie na technikach usuwania szumów z obrazów i filmów w zakresie pasma widzialnego. Proces ma na celu wyeliminowanie szumów przy jednoczesnym zachowaniu szczegółów obrazu. Jednak w trakcie redukcji szumów może dojść do utraty części tekstury obrazu, a resztkowe artefakty szumów mogą pozostać. Aktualne rozwiązania wykorzystują do tego zadania modele AI, które w kolejnych pracach badawczych są ulepszone. Jedno z takich podejść, na przykład, wykorzystuje połączenie architektury konwolucyjnej sieci neuronowej (CNN) i generatywnej sieci przeciwstawnej (GAN), które pozwoliło zachować strukturę obrazu wraz ze wzrostem głębokości modelu. Dodatkowo, poprzez włączenie modelu GAN po sieci CNN poprawia się jakość wizualna obrazu.

Praca Autora wpisuje się w ten trend ulepszania istniejących modeli AI i uzyskania efektywniejszego odszumiania obrazu poprzez wykorzystanie informacji z poprzednich, w stosunku do aktualnej, klatek filmu wideo. Przedstawione rezultaty z pewnością wzbogacają aktualny stan wiedzy w obszarze przetwarzania sygnałów wideo.

2. Wkład autora

Doktorant zaproponował dwie uzupełniające się strategie usuwania szumów z obrazów wideo z wykorzystaniem kontekstu czasowego.

Wykorzystanie informacji czasowych stanowi wyzwanie w realistycznych scenariuszach ze względu na zróżnicowany ruch między klatkami, co wymaga przeznaczenia znacznej części mocy obliczeniowej modelu na kompensację ruchu. W rezultacie podczas uczenia modele mogą wybierać łatwiejszą, ale nieoptymalną ścieżkę, skupiając się przede wszystkim na klatce centralnej. Aby zachęcić do większego wykorzystania sąsiednich klatek, proponowana jest strategia wstępnego uczenia, w której model usuwania szumów z wideo jest najpierw trenowany na powiązonym zadaniu, w którym klatka centralna nie jest dostępna, a mianowicie na interpolacji klatek. Ponadto, aby zmniejszyć różnicę między zadaniami interpolacji klatek a usuwaniem szumów z wideo, wprowadzono pośredni etap uczenia, w którym informacje z klatki centralnej są stopniowo ponownie wprowadzane.

Zaproponowaną metodę można więc określić jako trenowanie sieci neuronowych z wymuszeniem niejawnej estymacji ruchu, wykorzystując uproszczone, jawne dopasowanie. Wstępne trenowanie modeli sieci neuronowych odbywa się w warunkach, w których sygnał z klatki po usunięciu szumu jest ograniczony. Sprzyja to niejawnej estymacji ruchu w dużej mierze opierając się na informacji pochodzących z sąsiednich ramek obrazu. Przedstawiono dwa warianty implementacji tego podejścia, tj. klatka po usunięciu szumu jest maskowana podczas wstępnego trenowania sieci i drugie podejście, gdy szum w klatce po usunięciu szumu jest wzmacniany. Aby zwiększyć efektywność uczenia modelu z wykorzystaniem kontekstu czasowego, podczas wstępnego uczenia redukuje się szum w sąsiednich klatkach. Procedury te są stosowane wyłącznie podczas trenowania sieci, i nie mają one wpływu na proces wnioskowania, co skutkuje lepszą równowagą między jakością a wydajnością.

W celu walidacji opracowanego podejścia Autor wykonał eksperymenty na czterech zbiorach danych. Do celów syntetycznego odszumiania gaussowskiego wykorzystano dwa standardowe zbiory danych RGB: DAVIS oraz Set8. Zbiór Set8 zawiera sekwencje o wyższej

rozdzielczości przestrzennej i większym zakresie ruchu w porównaniu do sekwencji zawartych w zbiorze danych DAVIS. Dodatkowo, modele były trenowane i oceniane na zbiorach danych RAW zawierających rzeczywisty szum, tj. CRVD i ReCRVD. Jakość usuwania szumów z obrazu wideo za pomocą opracowanych modeli była oceniana przy użyciu trzech standardowych miar. Pierwszym wskaźnikiem jakości był PSNR (peak signal-to-noise ratio) podający stosunek maksymalnej mocy sygnału do mocy szumu zakłócającego ten sygnał. Drugą metryką był indeks SSIM (Structural Similarity Index Measure). SSIM służy do matematycznego określenia stopnia podobieństwa między dwoma obrazami. Trzecim wskaźnikiem jakości modeli był LPIPS (Learned Perceptual Image Patch Similarity). Miara ta ocenia wizualne podobieństwo między dwoma obrazami poprzez pomiar odległości między ich cechami wyodrębnionymi za pomocą głębokiego uczenia. W odróżnieniu od tradycyjnych miar opartych na pikselach (takich jak MSE czy SSIM), wskaźnik LPIPS jest bardzo zbliżony do ludzkiej oceny, co sprawia, że stał się podstawowym narzędziem do oceny treści generowanych przez sztuczną inteligencję. Do oceny spójności czasowej użyto wskaźnika E(W). Jest to miara stosowana w wizji komputerowej i w głębokim uczeniu do pomiaru niespójności wizualnej, zniekształceń geometrycznych lub błędów topologicznych między klatkami lub obrazami.

Aby porównać modele pod kątem ich wydajności, obliczano parametr FPS (frames-per-second, liczbę klatek na sekundę) przy przetwarzaniu na procesorze graficznym NVIDIA Quadro RTX 5000, liczbę operacji mnożenia i akumulacji (GMAC) oraz opóźnienie klatek. W celu dalszej analizy zachowania modeli pod kątem wykorzystania kontekstu czasowego zaproponowano dwa dodatkowe wskaźniki: różnicę pojedynczej klatki (SFD - Single Frame Difference) oraz procent poprawności przepływu wewnętrznego (PCIF - Percentage of Correct Intrinsic Flow). Wskaźnik SFD mierzy spadek wydajności spowodowany zastąpieniem elementów sekwencji ramką centralną. Autor założył, że modele w większym stopniu opierające się na informacjach z sąsiednich klatek są bardziej wrażliwe na utratę kolejnych klatek. Takie modele będą osiągać wyższe wyniki w skali SFD. Głównymi zaletami miary SFD są jej prostota oraz możliwość obliczenia jej na dużej liczbie próbek w rozsądnym czasie. Drugi zaproponowany wskaźnik – PCIF – związany jest z optymalizacją doboru odpowiadające sobie punktów w sąsiednich klatkach. Stanowi to szczególne wyzwanie dla modeli pozbawionych jawnego szacowania ruchu. Wykazano, że modele oparte na uczeniu głębokim potrafią wykorzystywać informacje czasowe w takich scenariuszach. Aby ocenić, które punkty mają największy wpływ na wynik odsumowania, zaproponowany przez Autora wskaźnik PCIF wyznaczany jest na podstawie technik opartych na gradientach.

W eksperymentach z zaproponowaną architekturą wykorzystywano pięć różnych sieci: Unet, NAFnet, ADFNet, Restormer oraz CGNet. Na początku każdej sieci umieszczono warstwę konwolucji grupowej, do której wprowadzane są trzy klatki. W celu zwiększenia wydajności rozmiary modeli zostały zmniejszone (NAFnet z domyślnych 17 mln do 3,9 mln, Restormer z 26 mln do 1,4 mln, CGNet z 167 mln do 22 mln). Wykorzystano oryginalne implementacje sieci, a wszystkie modele są trenowane przez $E = 80$ epok. Powstałe modele są określane przy użyciu nazwy sieci szkieletowej oraz głębokości architektury, np. Unet-D2. W przypadku modeli trenowanych z wykorzystaniem FIP zachowana jest ta sama łączna liczba epok, przy czym początkowe EFIP = 10 epok jest przeznaczone na wstępne trenowanie, po którym następuje pojedyncza epoka Post-FIP.

Autor wykazał, że przy użyciu nowatorskich miar (tj. PCIF, SFD) oraz techniki wizualizacji, że w różnych architekturach stosujących technikę FIP (Frame Interpolation Pre-training) uzyskano lepsze wykorzystanie informacji czasowych. Przekłada się to na poprawę jakości usuwania szumu w poszczególnych klatkach oraz zwiększoną spójność czasową. Proponowana technika jest odporna nie tylko na zmianę architektur modeli, ale także w na różne typy szumów, ponieważ zapewnia poprawę zarówno w przypadku sztucznego szumu Gaussa, jak i szumu rzeczywistego. Ta część wyników została opublikowana w prestiżowym czasopiśmie IEEE Transactions On Circuits and Systems For Video Technology.

Drugi obszar badań, który oczywiście obejmuje również główne zagadnienie badawcze, dotyczy wykorzystania przepływu optycznego w rekonstrukcji wideo. Informacja o przepływie optycznym nie była powszechnie integrowana z wydajnymi filtrami usuwającymi szumy wideo ze względu na wysokie koszty obliczeniowe zarówno tradycyjnych, jak i opartych na głębokim uczeniu się estymatorów przepływu optycznego. Podejście zaproponowane przez Doktoranta wykorzystuje przepływ optyczny oszacowany na podstawie klatek o znacznie zmniejszonej rozdzielczości. Zmniejszenie rozdzielczości danych wejściowych do estymatora przepływu optycznego znacznie obniża nakłady obliczeniowe, podczas gdy uzyskane mapy przepływu nadal dostarczają użytecznych informacji o ruchu. Wynikające z tego zgrubne dopasowanie pozwala na znaczną poprawę jakości usuwania szumów bez znaczącego wpływu na wydajność. Zaproponowane podejście do problemu skupia się na jawnej estymacji ruchu, co może okazać się konieczne w przypadku niewielkich modeli, gdy występują znaczne przemieszczenia. Autor postawił tezę, że uzyskana mapa przepływu optycznego o niskiej rozdzielczości LROF (Low-Resolution Optical Flow), nadal może uchwycić wystarczającą ilość informacji o ruchu, aby znacząco poprawić jakość usuwania szumów z materiału wideo.

Kolejnym nowatorskim rozwiązaniem zaproponowanym w pracy jest zaproponowanie techniki augmentacji NI (Noise Imbalancing). Ma ona na celu rozwiązanie problemów pojawiających w opisanej uprzednio metodzie FIP (Frame Interpolation Pre-training), takich jak zamiana klatek wejściowych oraz dodatkowa epoka pośrednia po zastosowaniu FIP. NI wzmacnia szum w klatce referencyjnej, jednocześnie redukując go w pozostałych klatkach. Szybkie zmiany scen (ruchliwość obiektów), które często występują w filmach o wysokiej rozdzielczości, są trudne do przetworzenia w sposób niejawny przez sieci neuronowe. Aby uprościć to zadanie, sąsiednie klatki są często dopasowywane względem klatki odniesienia za pomocą mapy przepływu optycznego LROF. W trakcie początkowych epok szkolenia poziomy szumu we wszystkich klatkach są płynnie równoważone. Zakłada się, że przepływ optyczny obliczony z klatek o zrównoważonym szumie będzie lepszej jakości niż ten obliczony na klatkach z rozszerzeniem i zachęci to model do wykorzystania informacji z sąsiednich klatek, dlatego też wyrównanie podczas wstępnego szkolenia powinno być jak najdokładniejsze. W tabeli 4.1 przedstawiono wyniki ablacji proponowanego modelu wzbogaconego o przepływ optyczny, w których oceniono wpływ współczynnika skalowania

Doktorant wykazał, że przepływ optyczny obliczony na klatkach o radykalnie zmniejszonej rozdzielczości (nawet dwunastokrotnie) może poprawić jakość usuwania szumów. Efekt ten jest szczególnie wyraźny w przypadku danych zawierających znaczny ruch, ale poprawia on również odporność w przypadku bardziej statycznych treści. Metoda NI okazała się szczególnie efektywna w przypadku danych głównie statycznych.

Wykazano, że obie metody są kompatybilne, a ich połączenie daje najlepsze wyniki w różnych zbiorach danych. Model uwzględniający obie proponowane techniki, określany przez Autora jako LVD (Lightweight Video Denoiser) został przetestowany na zbiorach danych CRVD oraz ReCRVD. CRVD, stanowi podstawowy zbiór danych obejmujący różne ustawienia ISO, uzyskany metodą „stop-and-motion capture”. ReCRVD (Recaptured Raw Video Denoising) jest zbiorem danych testowych, który powstał poprzez ponowne zarejestrowanie obrazów wideo w rozdzielczości 4K przy różnych ustawieniach czułości ISO w celu utworzenia par obrazów z szumami i obrazów bez szumów, dopasowanych pikselowo. Ułatwia to szkolenie modeli na scenach dynamicznych. Dla obu zbiorów testowych opracowany model LVD wykazał się bardzo dobrą efektywnością mierzoną parametrami PSNR oraz FPS.

Należy stwierdzić, że niewątpliwie wyniki prac badawczych przeprowadzonych przez Doktoranta mają istotny wkład w przedmiotowy obszar badawczy w dyscyplinie Automatyka, Elektronika, Elektrotechnika i Technologie Kosmiczne.

3. Poprawność

Przedstawiona do oceny rozprawa została napisana w języku angielskim, składa się z pięciu rozdziałów i bibliografii, zajmując 113 stron. We wstępie, w jasny sposób, Autor przedstawia zagadnienie objęte pracą badawczą, określa cel pracy i formułuje dwie hipotezy, które omówiłem w poprzednich sekcjach recenzji.

Kolejny, rozdział drugi, prezentuje aktualny stan wiedzy w zakresie tematyki poruszanej w pracy. W szczególności omówiono aspekty uczenia głębokiego, uwzględniając problemy związane z wydajnością sieci neuronowych, zaczynając od przeglądu podejść architektonicznych do tworzenia wydajnych modeli, następnie omawiając techniki optymalizacji modeli, a na koniec przedstawiając strategie uczenia, które poprawiają wydajność modelu bez utraty jego efektywności. Kolejna sekcja przedstawia aktualny stan wiedzy w zakresie odtwarzania obrazów i filmów, zaczynając od ogólnych metod odtwarzania obrazów, następnie omawiając postępy w odtwarzaniu filmów, a kończąc na szczegółowym przeglądzie metod usuwania szumów zarówno z obrazów, jak i filmów. Podsumowaniem rozdziału jest sformułowanie problemów badawczych, które nie są dotychczas rozwiązane w istniejącej literaturze, a które podejmuje niniejsza rozprawa doktorska.

Pierwsza luka badawcza dotyczy faktu, że aktualne techniki proponowane w najnowocześniejszych metodach odsumiania obrazu nie są dostosowane do zastosowań, które wymagają wydajnego przetwarzania. W przypadku modeli bez wyraźnego wyrównania czasowego, zazwyczaj nie ma mechanizmu zachęcającego do wykorzystania sąsiednich klatek, co prowadzi do tendencji skupiania się na klatce referencyjnej. Aby rozwiązać ten problem w pracy zaproponowano nowy paradygmat uczenia. Zmusza on model do skupienia się na sąsiednich klatkach poprzez wstrzymywanie informacji z centralnej klatki podczas początkowego etapu wstępnego uczenia. Zachowanie wyuczone na tym etapie jest następnie rozwijane, gdy udostępniane są pełne informacje.

Dru ga luka badawcza w istniejących rozwiązaniach dotyczy jawnego szacowania ruchu. W przypadku wystąpienia znacznych ruchów ich skuteczne przetworzenie zazwyczaj wymaga specjalnych mechanizmów wyrównywania, szczególnie w przypadku modeli o niewielkiej wydajności obliczeniowej. Metody te zazwyczaj nie są dostosowane do takich modeli, z wyjątkiem BasicVSR++, w kontekście usuwania szumów i rozmycia. Jednak w przypadku BasicVSR++ wpływ zmniejszenia rozmiaru klatek przed oszacowaniem przepływu optycznego nie jest analizowany osobno, a zakres rozważanych współczynników zmniejszenia jest ograniczony do czterokrotności. W przedstawionym w pracy rozwiązaniu zbadano wpływ zmiany rozdzielczości wejściowej sieci przepływu optycznego na wydajność usuwania szumów z wideo. Wykazano, że nawet zgrubne dopasowanie uzyskane z przepływu optycznego obliczonego na mocno zmniejszonych klatkach (nawet 12-krotnie) może nadal zapewniać znaczną poprawę jakości i ma zastosowanie w modelach lekkich. Doktorant umiejętnie sformułował tezy, które stanowią uzupełnienie tych brakujących rozwiązań. Rozwiązania wypełniające te luki badawcze zostały bardziej szczegółowo opisane w części recenzji dotyczącej wkładu własnego autora. Rozdziały trzeci i czwarty przedstawiają autorskie rozwiązania Autora, udowadniając postawione hipotezy. Rozdziały te zajmują w sumie 38 stron, 19 rysunków i 19 tabel, które w przejrzysty sposób prezentują wyniki pracy doktoranta. Rozdział piąty stanowi podsumowanie pracy oraz określa przyszłe plany badawcze w rozważnej tematyce. W treści pracy Autor odwołuje się do obszernej, liczącej 247 pozycji, bibliografii. Całość, przedstawionej do oceny rozprawy, stanowi spójnie, jasno i poprawnie zaprezentowany materiał opisujący wyniki badań Doktoranta.

4. Wiedza kandydata

W przedstawionej rozprawie Autor wykazał się szeroką wiedzą i sprawnym warsztatem naukowym. Obszerna, obejmująca 247 pozycji bibliografia, potwierdza przeprowadzoną przez Doktoranta szeroką analizę aktualnego stanu wiedzy w zakresie podjętej problematyki badawczej. Według informacji zawartej w rozprawie, Doktorant jest autorem/współautorem 5 publikacji w renomowanych czasopismach z wysokim IF oraz liczbą tzw. pkt ministerialnych (2x100, 2x140, 200). Ponadto Doktorant jest współtwórcą baz danych:

“Vident-lab: a dataset for multi-task video processing of phantom dental scenes” oraz
“Vident-real: an intra-oral video dataset for multi-task learning”.

Należy również wspomnieć o udziale Doktoranta w trzech projektach badawczych w tym jednym finansowanym przez NCBiR oraz trzy miesięcznym stażu w the Centre for Visual Computing, Paris-Saclay, CentraleSupélec.

Z pewnością można stwierdzić, że Kandydat wykazał się odpowiednią wiedzą z zakresu dyscypliny Automatyka, Elektronika, Elektrotechnika i Technologie Kosmiczne.

5. Inne uwagi

Praca jest napisana poprawnym i zrozumiałym językiem. Poniżej przedstawiam jedynie kilka drobnych uwag edytorskich:

wykres Fig. 3.4: brak opisu wartości przedstawionych na wykresie,
tabela 4.9, 4.10, 4.11: brak informacji, że tabela dotyczy wartości PSNR.

W ramach komentarzy i pytań do Doktoranta kilka uwag/pytań:

1. Do wyznaczania przepływu optycznego wykorzystywane są sieci neuronowe, takie jak RAFT (Recurrent All-Pairs Field Transforms) i BM3C (Bilateral Motion Estimation with a Dynamic Blend Filter). W pracy wykorzystano RAFT. Jakie kryteria były brane pod uwagę przy doborze tego modelu?
2. Czy istnieje parametr (Figure-of-Merit), który pozwoliłby na porównywanie różnych algorytmów odsumiania, uwzględniając więcej parametrów (poza PSNR. FPS), np. pobór mocy, wymagania dotyczące mocy obliczeniowej itp.?
3. Jak przedstawione w pracy rozwiązania są ułożone w stosunku do istniejących rozwiązań komercyjnych. Czy mogą zostać dołączone jako dodatkowy „akcelerator odsumiania video”? Wspominają istniejące narzędzia odwołuję się przykładowo do następujących narzędzi: „Neat Video”, które pozwala na dokładne sprofilowanie szumu z konkretnego aparatu i usunięcie go z zachowaniem detali; „Topaz Video AI”, które uważane jest za standard w kategorii odsumiania i upscalingu opartego na sztucznej inteligencji, „DaVinci Resolve Studio”, które posiada wbudowane, referencyjne narzędzia redukcji szumów (Temporal i Spatial NR), wykorzystujące zaawansowane wykrywanie ruchu w kadrze.
4. Doktorant przeprowadził test na dostępnych zbiorach danych testowych. Czy wykonano również testy na danych pochodzących z akwizycji własnej video? Czy i jakie w takim przypadku należy przeprowadzić procedury przygotowania danych wejściowych?

6. Podsumowanie

Biorąc pod uwagę opinie zaprezentowane w poprzednich punktach i wymagania zdefiniowane przez art. 187 Ustawy z dnia 20 lipca 2018 r. Prawo o szkolnictwie wyższym i nauce (z późniejszymi zmianami), moja ocena rozprawy pod względem trzech podstawowych kryteriów jest następująca:

A. Czy rozprawa zawiera oryginalne rozwiązanie problemu naukowego? (wybierz jedną opcję stawiając znak X)

☒	☐	☐	☐	☐
<i>Zdecydowanie TAK</i>	<i>Raczej TAK</i>	<i>Trudno powiedzieć</i>	<i>Raczej NIE</i>	<i>Zdecydowanie NIE</i>

B. Czy po przeczytaniu rozprawy zgadzasz się, że kandydat posiada ogólną wiedzę teoretyczną w dyscyplinie Automatyka, Elektronika, Elektrotechnika i Technologie Kosmiczne?

☒	☐	☐	☐	☐
<i>Zdecydowanie TAK</i>	<i>Raczej TAK</i>	<i>Trudno powiedzieć</i>	<i>Raczej NIE</i>	<i>Zdecydowanie NIE</i>

C. Czy kandydat posiada umiejętność samodzielnego prowadzenia pracy naukowej?

☒	☐	☐	☐	☐
<i>Zdecydowanie TAK</i>	<i>Raczej TAK</i>	<i>Trudno powiedzieć</i>	<i>Raczej NIE</i>	<i>Zdecydowanie NIE</i>

Podsumowując niniejszą recenzję stwierdzam, że rozprawa doktorska pana mgr inż. Piotra Kopy Ostrowskiego zatytułowana: „Efficient Deep Neural Networks for Video Denoising” spełnia wymagania stawiane rozprawom doktorskim i wnoszę o dopuszczenie doktoranta do dalszych etapów przewodu doktorskiego.

Ponadto biorąc pod uwagę jakość wykonanych badań i rozprawy doktorskiej oraz uwzględniając zasady wyróżniania rozpraw doktorskich w dyscyplinie AEEiTK wnoszę o wyróżnienie rozprawy.



podpis