

prof. UAM. dr hab. Krzysztof Dyczkowski
Zakład Sztucznej Inteligencji
Wydział Matematyki i Informatyki
Uniwersytet im. Adama Mickiewicza w Poznaniu

Poznań, 20 lipca 2026 r.

RECENZJA

rozprawy doktorskiej **mgr. inż. Jana Edwarda Baumgarta**
„Zastosowanie starożytnych strategii bitew w nowych algorytmach rojowych
do rozwiązywania problemów optymalizacyjnych”

Promotor rozprawy: dr hab. inż. Jacek M. Czerniak, prof. Politechniki Bydgoskiej
Promotor pomocniczy: dr inż. Dawid Ewald, prof. Politechniki Bydgoskiej

1 CEL I TEMATYKA PRACY DOKTORSKIEJ

Recenzowana rozprawa doktorska mgr. inż. Jana Edwarda Baumgarta podejmuje problem projektowania i oceny metaheurystycznych algorytmów optymalizacji ciągłej. Główną ideą pracy jest systematyczne przełożenie wybranych elementów starożytnych strategii bitewnych na mechanizmy obliczeniowe wykorzystywane w algorytmach rojowych. Autor opracował rodzinę algorytmów Ancient Battlefield Optimizer (ABO), w której zbiór tworzą heterogeniczne typy agentów, odpowiadające różnym sposobom przeszukiwania przestrzeni rozwiązań. Algorytm uzupełniono o system taktyk i formacji, moduł rozpoznania przestrzeni, mechanizm honoru oraz moduł Commander AI, wykorzystujący Q-learning do adaptacyjnego wyboru taktyki.

Tematyka pracy wpisuje się w aktualny obszar badań nad metaheurystykami wielostrategicznymi oraz adaptacyjnym doбором operatorów. Na podkreślenie zasługuje fakt, że doktorant nie ogranicza się do przedstawienia atrakcyjnej metafory wojskowej, lecz świadomie konfrontuje swoje podejście z krytyką algorytmów „inspirowanych naturą”. W rozprawie przedstawia również alternatywną, pozbawioną metafory interpretację ABO jako repozytorium sześciu operatorów ruchu sterowanych meta-kontrolerem Q-learning. Takie pozycjonowanie pracy uważam za trafne i metodologicznie dojrzałe.

Główną tezę rozprawy można streścić następująco: systematyczna integracja heterogenicznych operatorów przeszukiwania, inspirowanych strategiami bitewnymi, z hierarchicznym mechanizmem koordynacji oraz adaptacyjnym doбором taktyk pozwala uzyskać



RPW/2960/2026 P
Data:2026-07-22

algorytm konkurencyjny wobec współczesnych metaheurystyk na standardowych funkcjach benchmarkowych w szerokim zakresie wymiarowości.

Z tezy tej wynikają dwie podstawowe hipotezy badawcze:

1. Algorytm ABO oraz jego wariant CommanderABO będą osiągały konkurencyjne lub lepsze wyniki niż standardowe algorytmy optymalizacji na zróżnicowanym zbiorze wielowymiarowych funkcji testowych.
2. Adaptacyjny wybór taktów za pomocą Q-learningu, bez ręcznego harmonogramowania ich sekwencji, będzie konkurencyjny wobec najlepszych wariantów fazowych w kryterium mediany, a jednocześnie podniesie jakość najlepszych uzyskiwanych rozwiązań w kryterium best-of-run.

Cele badawcze zostały jasno sformułowane i odpowiadają zakresowi przeprowadzonych badań. Szczególnie dobrze zdefiniowana jest relacja pomiędzy dwoma poziomami wkładu: podstawową architekturą heterogenicznej populacji ABO oraz dodatkowym mechanizmem adaptacyjnego wyboru operatora. Należy jednak zaznaczyć, że hipoteza druga ma charakter wielokryterialny, a jej ocena zależy od przyjętej miary jakości. Jak wynika z rezultatów, w kryterium mediany CommanderABO nie przewyższa ręcznie strojonego harmonogramu fazowego, natomiast jego przewaga ujawnia się głównie w kryterium najlepszego wyniku. W dalszej części recenzji odnoszę się do konsekwencji tej obserwacji.

2 OCENA FORMY I STRUKTURY

Przedstawiona do recenzji rozprawa została napisana w języku polskim i liczy 293 strony wraz z bibliografią, listą publikacji autora i obszernymi załącznikami. Tekst zasadniczy składa się z siedmiu rozdziałów. Bibliografia obejmuje 128 pozycji z zakresu optymalizacji, inteligencji rojowej, uczenia ze wzmocnieniem, analizy statystycznej oraz historii wojskowości. W załącznikach zamieszczono kompletne pseudokody głównych komponentów, zestaw parametrów oraz szczegółowe wyniki indeksu trudności funkcji benchmarkowych.

Rozdział pierwszy zawiera wprowadzenie do problematyki optymalizacji, omówienie podstawowych pojęć i ograniczeń metaheurystyk oraz motywację zastosowania strategii wojskowych. Autor formułuje główną tezę, dwa cele badawcze i odpowiadające im hipotezy. Istotnym elementem jest wprowadzenie pojęcia Strategic Swarm Intelligence jako ogólnej ramy koncepcyjnej oraz rozróżnienie jej od konkretnych implementacji ABO i CommanderABO.

Rozdział drugi przedstawia przegląd literatury: od metod klasycznych i ewolucyjnych, przez algorytmy rojowe i podejścia interdyscyplinarne, po wykorzystanie uczenia ze wzmocnieniem w adaptacyjnym doborze operatorów. Autor omawia również wcześniejsze algorytmy inspirowane konfliktem i działaniami wojennymi. Bardzo wartościowa jest sekcja poświęcona krytyce metaheurystyk. Doktorant przywołuje zarzuty dotyczące nadmiarowości matematycznej, braku ablacji i asymetrycznych porównań, a następnie próbuje precyzyjnie określić, gdzie znajduje się rzeczywisty wkład ABO.

Rozdział trzeci ma charakter interdyscyplinarny i przedstawia wybrane starożytne typy jednostek, formacje, taktyki, mechanizmy dowodzenia, morale oraz rozpoznania. Materiał historyczny służy zbudowaniu mapowania pomiędzy elementami strategii militarnej a mechanizmami algorytmicznymi. Rozdział ten jest ciekawy i dobrze napisany, choć miejscami zbyt rozbudowany w stosunku do jego znaczenia dla formalnej części rozprawy.

Rozdział czwarty stanowi zasadniczą część projektową pracy. Przedstawiono w nim modułową architekturę ABO, sześć typów jednostek i odpowiadające im operatory: przeszukiwanie współrzędnościowe, wariant ewolucji różnicowej, wielopunktowe próbkowanie zdalne, loty Lévy'ego z pamięcią kierunku, ruch z pędem Cauchy'ego oraz przeskoki między dolinami z kryterium akceptacji. Opisano także sześć taktyk, formacje, system honoru, Reconnaissance Intelligence i Commander AI. Pięciowymiarowy, zdyskretyzowany stan meta-kontrolera obejmuje fazę optymalizacji, stagnację, momentum, różnorodność populacji i odsetek bohaterów, co daje 243 stany i sześć możliwych akcji.

Rozdział piąty opisuje implementację, środowisko programistyczne, funkcje testowe i metodykę eksperymentalną. Autor rozdziela funkcje używane do treningu tablicy Q od funkcji benchmarkowych i deklaruje zamrożenie tablicy podczas ewaluacji. Plan badań obejmuje 33 funkcje, pięć wymiarowości, 100 niezależnych uruchomień oraz 46 algorytmów, w tym cztery warianty ABO. Zastosowano testy Friedmana, Nemenyiego i Wilcoxona, korekty wielokrotnych porównań oraz analizę wielkości efektu.

Rozdział szósty jest najobszerniejszą częścią rozprawy i przedstawia wyniki eksperymentalne. Zawiera ranking algorytmów, analizę statystyczną, porównanie wariantów ABO, analizę wpływu objętości treningu, eksperyment CommanderABO-OnTheFly, analizę zbieżności, kosztu czasowego i wyników dla poszczególnych funkcji. Rozdział siódmy podsumowuje wyniki, omawia wkład pracy, ograniczenia oraz kierunki dalszych badań. Autor otwarcie wskazuje m.in. brak pełnej ablacji, wykorzystanie funkcji w formie kanonicznej, brak walidacji na problemach rzeczywistych oraz nierówną liczbę wywołań funkcji celu.

Praca jest przygotowana bardzo starannie, a jej układ jest logiczny. Liczne tabele, wykresy, schematy i pseudokody ułatwiają analizę zaproponowanego rozwiązania. Szczególnie pozytywnie oceniam obszerną dokumentację implementacyjną i jawne przedstawienie ograniczeń. Tekst jest jednak nadmiernie rozbudowany, miejscami powtarzalny, a część narracji historycznej, przykładów i refleksji końcowych mogłaby zostać istotnie skrócona bez szkody dla wartości naukowej. W kilku miejscach wymienne używanie terminów „uruchomienie”, „epizod” i „ewaluacja funkcji” może prowadzić do niejednoznaczności, zwłaszcza przy omawianiu budżetu obliczeniowego.

3 OCENA ROZPRAWY

Za główny wkład rozprawy uznaję opracowanie i implementację spójnej, modułowej architektury wielostrategicznego algorytmu populacyjnego. Poszczególne operatory ruchu nie są całkowicie nowymi metodami numerycznymi, co autor uczciwie przyznaje. Oryginalność rozwiązania polega przede wszystkim na ich integracji w heterogenicznej populacji, połączeniu z mechanizmami formacji, rozpoznania i honoru oraz zastosowaniu meta-kontrolera Q-learning do wyboru taktyki. Tak rozumiany wkład jest wystarczająco wyraźny, aby traktować ABO jako oryginalne rozwiązanie problemu naukowego, a nie jedynie zmianę nazewnictwa istniejącego algorytmu.

Istotną zaletą jest oddzielenie warstwy metaforycznej od formalnej. Autor przedstawia ABO również jako instancję hyper-heuristic z Adaptive Operator Selection. Pozwala to oceniać rozwiązanie w ustalonej nomenklaturze naukowej i wskazuje, że wojskowe nazwy pełnią przede wszystkim funkcję porządkującą i mnemoniczną. W pracy podano równania, pseudokody, struktury stanów, funkcję nagrody, sposób dyskretyzacji oraz konfigurację treningu. Dzięki temu podstawowe elementy metody są możliwe do odtworzenia.

Na bardzo wysoką ocenę zasługuje skala eksperymentów. Końcowy benchmark obejmuje 759 tys. uruchomień, a badania uzupełniające dodatkowe eksperymenty treningowe i analizę wariantu uczącego się od zera. Cztery warianty ABO zajęły cztery pierwsze miejsca w globalnym rankingu Friedmana dla median wyników. Ręcznie strojony ABO-ManualPhases osiągnął średnią rangę 6,21, CommanderABO 6,77, natomiast najlepszy algorytm spoza rodziny GWO 10,02. Wyniki wskazują, że zaproponowana architektura dobrze radzi sobie szczególnie w wymiarach od 10 do 100.

Weryfikację hipotezy pierwszej oceniam pozytywnie, ale wyłącznie w granicach przyjętego protokołu. Rodzina ABO uzyskała bardzo dobre wyniki na wybranym zestawie 33 funkcji kanonicznych. Rezultat ten potwierdza konkurencyjność metody na analizowanej klasie problemów, nie stanowi jednak dowodu jej ogólnej przewagi nad nowoczesnymi metodami optymalizacji ciągłej. Sam autor słusznie odwołuje się do twierdzenia No Free Lunch i ogranicza zakres wniosków do badanego zbioru.

Ocena hipotezy drugiej wymaga większej ostrożności. W kryterium mediany CommanderABO nie poprawia wyników względem ręcznie strojonego harmonogramu: ABO-ManualPhases wygrywa na 75 ze 165 konfiguracji, a różnica jest istotna statystycznie. CommanderABO osiąga natomiast najlepszą rangę w kryterium best-of-run i przewyższa warianty QTable4K oraz QTable10K. Wynik potwierdza, że adaptacyjny dobór taktyk może zwiększać prawdopodobieństwo uzyskania wyjątkowo dobrego rozwiązania, ale nie poprawia typowego wyniku pojedynczego uruchomienia. Z tego względu uznałbym hipotezę H2 za potwierdzoną warunkowo, zgodnie z jej dwuczęściowym sformułowaniem, a nie w sensie ogólnej przewagi adaptacyjnego sterowania.

Dobór testów statystycznych jest zasadniczo prawidłowy i odpowiada charakterowi wyników algorytmów stochastycznych. Autor stosuje metody nieparametryczne i przedstawia przedziały ufności oraz wielkości efektu. Jednocześnie należy pamiętać, że test Nemenyiego przy 46 algorytmach jest bardzo konserwatywny, a traktowanie 165 par funkcja-wymiar jako bloków może częściowo ignorować zależności pomiędzy różnymi wymiarami tej samej funkcji. Zaletą jest uzupełnienie analizy globalnej testami parowymi, jednak interpretacja licznych porównań powinna pozostać ostrożna.

Ważnym elementem jest analiza ograniczeń. Autor sam wskazuje, że wszystkie funkcje wykorzystano w formie kanonicznej, co może faworyzować operator przeszukiwania współrzędnościowego. Zwraca także uwagę na brak pełnej ablacji systemu honoru, rekonesansu i formacji oraz na brak testów na problemach rzeczywistych. Są to ograniczenia istotne, ponieważ główny rezultat pracy dotyczy kompozycji wielu mechanizmów, a bez ich niezależnego wyłączenia trudno określić, które z nich odpowiadają za obserwowaną przewagę.

Najpoważniejsze zastrzeżenie metodologiczne dotyczy budżetu obliczeniowego. Wszystkie algorytmy otrzymały tę samą liczbę epok i ten sam rozmiar populacji, ale warianty ABO z lokalnym przeszukiwaniem wykonują dodatkowe wywołania funkcji celu. Oznacza to, że porównanie nie jest

prowadzone przy równym budżecie Function Evaluations. W optymalizacji black-box liczba wywołań funkcji celu jest zazwyczaj bardziej miarodajną jednostką kosztu niż liczba epok. Porównanie czasów wykonania częściowo uzupełnia obraz, lecz nie zastępuje eksperymentu przy identycznym limicie FE, szczególnie gdy funkcja celu jest kosztowna.

Drugim ważnym problemem jest dobór algorytmów odniesienia. Zestaw 42 metod jest szeroki, ale brakuje w nim m.in. klasycznego PSO, silnego algorytmu CMA-ES oraz bezpośrednich konkurentów z obszaru hyper-heuristics i Adaptive Operator Selection. Duża liczba porównywanych metaheurystyk nie musi oznaczać wysokiej jakości punktów odniesienia. Dla pełnej oceny naukowej bardziej przekonujące byłoby porównanie z mniejszą liczbą starannie dobranych, silnych i aktualnych metod, dostrojonych według wspólnego protokołu.

Przedstawiona analiza zbieżności ma w znacznej części charakter intuicyjny i empiryczny. Autor uczciwie zaznacza, że klasyczne warunki zbieżności Q-learningu nie są spełnione i że pełny algorytm nie posiada formalnej gwarancji. Sformułowana propozycja zbieżności globalnej wymagałaby jednak ścisłego dowodu i precyzyjniejszych warunków dotyczących rozkładu próbkowania, nieskończonej liczby wizyt w każdym otoczeniu oraz wpływu projekcji na granice. Obecny fragment należy traktować jako hipotezę teoretyczną, a nie udowodnione twierdzenie.

Mimo wskazanych ograniczeń rozprawa pokazuje bardzo dobre przygotowanie autora do prowadzenia badań naukowych. Doktorant potrafi połączyć projektowanie algorytmów, programowanie, eksperymenty wielkoskalowe, analizę statystyczną i krytyczną interpretację wyników. Za szczególnie wartościowe uważam transparentne pokazanie, że głównym źródłem jakości jest architektura ABO, natomiast wkład Q-learningu zależy od kryterium oceny. Jest to bardziej wiarygodne niż próba przedstawienia każdego elementu proponowanej metody jako jednoznacznie korzystnego.

4 UWAGI I WNIOSKI KOŃCOWE

Rozprawa stanowi wartościowy i oryginalny wkład w obszar metaheurystycznej optymalizacji ciągłej. Autor zaproponował rozbudowaną architekturę wielostrategiczną, zaimplementował ją w postaci działającego systemu oraz przeprowadził bardzo szeroką weryfikację eksperymentalną. Jednocześnie charakter pracy i sposób ewaluacji rodzą pytania, które powinny zostać przedyskutowane podczas publicznej obrony:

1. Na stronach 132 oraz 179 autor wyraźnie stwierdza, że warianty ABO wykonują dodatkowe wywołania funkcji celu, a porównanie oparto na równej liczbie epok i rozmiarze populacji, nie na równym budżecie FE. Jaka jest rzeczywista liczba wywołań funkcji celu dla poszczególnych wariantów ABO i najważniejszych konkurentów? Czy autor przeprowadził lub może oszacować ranking przy identycznym limicie FE? W jakim stopniu przewaga rodziny ABO może wynikać z większej liczby ocen funkcji celu?
2. Na stronie 146 podano, że harmonogram ABO-ManualPhases został wybrany po przetestowaniu wielu sekwencji taktyk i progów przejść, a następnie wybraniu najlepszego wariantu „na zestawie funkcji benchmarkowych”. Czy harmonogram ten był strojony wyłącznie na zbiorze treningowym rozłącznym z 33 funkcjami ewaluacyjnymi, czy także na

końcowym zbiorze benchmarkowym? Jeżeli wykorzystano dane z benchmarku końcowego, jak wpływa to na interpretację przewagi ManualPhases i porównanie z CommanderABO?

3. Wśród 42 algorytmów porównawczych nie znalazły się PSO i CMA-ES, a także bezpośrednie metody z obszaru hyper-heuristics lub Adaptive Operator Selection, do których autor sam odnosi alternatywne sformułowanie tezy. Jakie były kryteria pominięcia tych metod? Czy wstępny benchmark 58 algorytmów, wspomniany w rozdziale siódmym, pozwala ocenić stabilność głównych wniosków po włączeniu silniejszych punktów odniesienia?
4. Najważniejszym argumentem na rzecz modułu Commander AI jest najlepszy wynik w kryterium best-of-run, podczas gdy mediana jest gorsza od ABO-ManualPhases. Minimum ze 100 uruchomień jest silnie zależne od liczby powtórzeń i może premiować algorytm o większej wariancji. Jak autor uzasadnia praktyczne znaczenie tego kryterium? Czy podobne wnioski uzyskano by dla kwantyla 10%, oczekiwanej jakości przy ustalonym budżecie kilku uruchomień lub miary anytime performance?
5. Na stronach 175-177 sformułowano propozycję asymptotycznej zbieżności CommanderABO. Samo istnienie niezerowego prawdopodobieństwa wygenerowania punktu w otoczeniu rozwiązania oraz elityzm nie są wystarczające do wykazania zbieżności z prawdopodobieństwem jeden, jeżeli nie zapewni się odpowiednio częstego i pełnego próbkowania przestrzeni. Ponadto wzór dla taktyki Scouting generuje punkty wokół aktualnego najlepszego rozwiązania i nie w każdej konfiguracji musi pokrywać całą dziedzinę. Jakie dodatkowe założenia są potrzebne do sformułowania poprawnego dowodu?
6. Autor trafnie zauważa, że funkcje kanoniczne, nierotowane i nieprzesunięte mogą faworyzować przeszukiwanie współrzędnościowe używane przez Heavy Infantry. Jakiego spadku pozycji ABO autor spodziewa się na rotowanych i przesuniętych zestawach CEC? Czy przeprowadzono choćby pilotażową analizę, która pozwala oddzielić rzeczywistą odporność architektury od wykorzystania separowalności osiowej funkcji?
7. Indeks trudności funkcji jest przedstawiany jako jeden z wkładów pracy, lecz wagi jego komponentów mają charakter ekspercki, a w analizie wrażliwości uwzględniono jedynie pięć z siedmiu składowych. Co więcej, przy wagach równych zgodność listy dziesięciu najtrudniejszych konfiguracji spada do 1/10 pomimo wysokiej korelacji globalnej. Jak należy interpretować stabilność tego indeksu w skrajnym ogonie rankingu i czy powinien on być traktowany jako samodzielny rezultat naukowy, czy raczej narzędzie pomocnicze do doboru benchmarku?

W pracy występują również drobne uchybienia redakcyjne i terminologiczne. Tekst można byłoby skrócić, ograniczając powtórzenia pomiędzy rozdziałem wynikowym i podsumowaniem. Warto także konsekwentnie rozróżnić liczbę uruchomień algorytmu, liczbę epok treningowych oraz liczbę wywołań funkcji celu. Uwagi te nie wpływają jednak na zasadniczo pozytywną ocenę rozprawy.

W mojej opinii rozprawa doktorska mgr. inż. Jana Edwarda Baumgarta zawiera bardzo szeroki i dobrze udokumentowany materiał badawczy. Stanowi oryginalne rozwiązanie problemu naukowego oraz wykazuje ogólną wiedzę teoretyczną autora w dyscyplinie informatyka techniczna i telekomunikacja, a także jego umiejętność samodzielnego prowadzenia pracy naukowej. Wskazane przeze mnie ograniczenia dotyczą przede wszystkim zakresu i rygoru porównań eksperymentalnych, nie podważają natomiast oryginalności zaproponowanej architektury ani wartości wykonanej pracy implementacyjnej.

Na uwagę zasługuje dorobek publikacyjny doktoranta obejmujący 13 prac, w tym artykuł w czasopiśmie Electronics, rozdziały w monografiach Springer oraz publikacje konferencyjne. Co najmniej pięć prac jest bezpośrednio związanych z rozwojem koncepcji algorytmów inspirowanych działaniami wojennymi i dokumentuje kolejne etapy prowadzące do powstania ABO.

Stwierdzam, że rozprawa spełnia wymagania stawiane rozprawom doktorskim, w tym wymagania określone w art. 187 ustawy z dnia 20 lipca 2018 r. - Prawo o szkolnictwie wyższym i nauce (t.j. Dz. U. z 2024 r. poz. 1571 z późn. zm.).

W związku z powyższym wnioskuję do Rady Naukowej Dyscypliny Informatyka Techniczna i Telekomunikacja Politechniki Bydgoskiej im. Jana i Jędrzeja Śniadeckich o dopuszczenie rozprawy doktorskiej mgr. inż. Jana Edwarda Baumgarta pt. „Zastosowanie starożytnych strategii bitew w nowych algorytmach rojowych do rozwiązywania problemów optymalizacyjnych” do publicznej obrony i dalszego procedowania.


/prof. UAM. dr hab. Krzysztof Dyczkowski/