

INTEGRACJE DEZINTEGRACJE

Efektywność pracy zespołowej w nauce

Granice stosowalności AI – regulacje prawne i etyczne

Nowe okoliczności – kierunki badań, oczekiwania i zagrożenia

INTEGRACJE – DEINTEGRACJE

Efektywność pracy zespołowej w nauce

Granice stosowalności AI – regulacje prawne i etyczne

Nowe okoliczności – kierunki badań, oczekiwania i zagrożenia

Symposium Stowarzyszenia Singularis
Centrum Studiów Zaawansowanych PW

Ośrodek Wypoczynkowy PW w Wildze
ul. Wiewiórek 6, 08-470 Wilga

10-12 października 2025



Centrum Studiów
Zaawansowanych
Politechnika Warszawska



SPOTKANIA
OTWARTYCH UMYŚŁÓW

Zespół prowadzący:

Przemysław Bieчек (PW), Leon Gradoń (PW), Janusz Hołyst (PW),
Stanisław Janeczko (PW), Agnieszka Jastrzębska (PW), Artur Kasprzak (PW),
Julia Kuczak (PW), Marek Kuś (PAN), Mariusz Malinowski (PW), Jacek Mańdziuk (PW),
Dariusz Plewczyński (PW), Jacek Rogala (UW), Bartłomiej Skowron (PW),
Adam Woźniak (PW)

Organizacja:

Wanda Borkowska (CSZ)

Harmonogram

piątek » 10 października

17:00 – 18:30	zakwaterowanie
19:00 – 21:00	kolacja powitanie gości Stanisław Janeczko – dyrektor Centrum Studiów Zaawansowanych, Politechnika Warszawska

sobota » 11 października

09:00 – 10:00	śniadanie
10:30 – 10:50	<i>Warunki intensywnej współpracy zespołu</i> – Stanisław Janeczko – Centrum Studiów Zaawansowanych, PW
10:50 – 11:00	dyskusja
11:00 – 11:20	<i>Centrum AI jako przykład i wyzwanie do organizacji centrów badawczych w strukturze Uczelni</i> – Leon Gradoń – Wydział Inżynierii Chemicznej i Procesowej, PW
11:20 – 11:30	dyskusja
11:30 – 11:50	<i>Kilka bioetycznych problemów biomedycznych</i> – Ewa Bartnik – Wydział Biologii, UW
11:50 – 12:00	dyskusja
12:00 – 12:20	<i>Kontrola jakości wspierana sztuczną inteligencją w sektorze MŚP: wnioski z badań empirycznych</i> – Krzysztof Ejsmont – Wydział Mechaniczny Technologiczny, PW
12:20 – 12:30	dyskusja
12:30 – 13:30	przerwa kawowa
13:30 – 13:50	<i>2441</i> – Bartłomiej Skowron – Wydział Administracji i Nauk Społecznych, PW
13:50 – 14:00	dyskusja

14:00 – 14:20	<i>Abstrakcyjne rozumowanie wizualne: wciąż poza zasięgiem modeli sztucznej inteligencji</i> – Jacek Mańdziuk – Wydział Matematyki i Nauk Informatycznych, PW
14:20 – 14:30	dyskusja
14:30 – 16:00	obiad
16:00 – 16:20	<i>Wielopoziomowe przeciążenie informacyjne: dlaczego jest niebezpieczne i czy możemy je pokonać?</i> – Janusz Hołyst – Wydział Fizyki, PW
16:20 – 16:30	dyskusja
16:30 – 16:50	<i>AI w badaniach w zastosowaniach między dziedzinowych szansa i zagrożenie</i> – Jacek Rogala – Centrum Badania Ryzyka Systemowego, Artes Liberales, UW
16:50 – 17:00	dyskusja
17:00 – 17:30	przerwa kawowa
17:30 – 17:50	<i>Arystoteles, duże zbiory danych i społeczeństwo cyfrowe: jak integrować badania wokół koncepcji filozoficznych?</i> – Marcin Koszowy – Wydział Administracji i Nauk Społecznych, PW
17:50 – 18:00	dyskusja
19:00	uroczysta kolacja

niedziela » 12 października

09:00 – 10:00	śniadanie
10:30 – 10:50	<i>Modelologia – nowe i ciekawe wyzwania związane z analizą modeli</i> – Przemysław Biecek – Wydział Matematyki i Nauk Informatycznych, PW
10:50 – 11:00	dyskusja
11:00 – 11:20	<i>Zespoły interdyscyplinarne w trójwymiarowej genomice: od sekwencji DNA do struktury chromatyny w skali populacji</i> – Dariusz Plewczyński – Wydział Matematyki i Nauk Informatycznych, PW
11:20 – 11:30	dyskusja

11:30 – 11:50	<i>Granice przewidywania</i> – Agnieszka Jastrzębska – Wydział Matematyki i Nauk Informacyjnych, PW
11:50 – 12:00	dyskusja
12:00 – 13:00	przerwa kawowa
13:00 – 13:20	<i>Od fragmentacji do synergii: sztuczna inteligencja wobec barier interdyscyplinarności</i> – Błażej Żyliński – Wydział Fizyki, PW
13:20 – 13:30	dyskusja
13:30 – 13:50	<i>Wykorzystanie dużych modeli językowych w programowaniu graficznym sterowników PLC: możliwości, ograniczenia i rekomendacje dla edukacji przemysłowej</i> – Łukasz Żukowski – Wydział Mechatroniki PW
13:50 – 14:00	dyskusja
14:00 – 15:00	obiad

Spis abstraktów

6

Stanisław JANECKO, <i>Warunki intensywnej współpracy zespołu</i>	7
Leon GRADOŃ, <i>Centrum AI jako przykład i wyzwanie do organizacji centrów badawczych w strukturze Uczelni</i>	8
Ewa BARTNIK, <i>Kilka bioetycznych problemów biomedycznych</i>	10
Krzysztof EJSFONT, <i>Kontrola jakości wspierana sztuczną inteligencją w sektorze MŚP: wnioski z badań empirycznych</i>	11
Bartłomiej SKOWRON, <i>2441</i>	12
Jacek MAŃDZIUK, <i>Abstrakcyjne rozumowanie wizualne: wciąż poza zasięgiem modeli sztucznej inteligencji</i>	13
Janusz HOŁYST, <i>Wielopoziomowe przeciążenie informacyjne: dlaczego jest niebezpieczne i czy możemy je pokonać?</i>	14
Jacek ROGALA, <i>AI w badaniach w zastosowaniach między dziedzinowych szansa i zagrożenie</i>	16
Marcin KOSZOWY, <i>Arystoteles, duże zbiory danych i społeczeństwo cyfrowe: jak integrować badania wokół koncepcji filozoficznych?</i>	17
Przemysław BIECEK, <i>Modelologia – nowe i ciekawe wyzwania związane z analizą modeli</i>	18
Dariusz PLEWCZYŃSKI, <i>Zespoły interdyscyplinarne w trójwymiarowej genomice: od sekwencji DNA do struktury chromatyny w skali populacji</i>	19
Agnieszka JASTRZĘBSKA, <i>Granice przewidywania</i>	20
Błażej ŻYLIŃSKI, <i>Od fragmentacji do synergii: sztuczna inteligencja wobec barier interdyscyplinarności</i>	21
Łukasz ŻUKOWSKI, Jakub MOŻARYN, <i>Wykorzystanie dużych modeli językowych w programowaniu graficznym sterowników PLC: możliwości, ograniczenia i rekomendacje dla edukacji przemysłowej</i>	22

Stanisław JANECKO

Wydział Matematyki i Nauk Informatycznych
Centrum Studiów Zaawansowanych
Politechnika Warszawska

Warunki intensywnej współpracy zespołu

Zespół badawczy w nauce współpracuje najbardziej intensywnie kiedy przestaje być zbiorem rozseparowanych elementów wykonujących nakazowe zadania do-
raźne i tworzy jeden podmiot poznający. W systemie zbyt daleko posuniętej indy-
widualizacji i szczegółowego rozliczania taki stan jest trudny do osiągnięcia. Jeśli
poziom kolektywnej podmiotowości członków zespołu jest dostatecznie wysoki to
poziom zaangażowania twórczego jest bardzo wysoki i słabo zależy od zewnętrz-
nych deformacji. Przeanalizujemy pewne rodzaje zespołów badawczych i wpływ
sportowej rywalizacji na wynikowy rezultat zespołu.

Leon GRADON

Wydział Inżynierii Chemicznej i Procesowej
Politechnika Warszawska

Centrum AI jako przykład i wyzwanie do organizacji centrów badawczych w strukturze Uczelni

Tematyka Sympozjum planowanego na październik w Wildzie jest wielowątkowa. Jednym z ważnych, moim zdaniem, zagadnień, które warto poruszyć, jest wykorzystanie do pewnych działań przyznania Zespołowi prof. Przemysława Biecka grantu w programie MAB-FENG. Przy umiejętnym i dobrze zaplanowanym działaniu może to być źródłem utworzenia, znaczącego co najmniej w skali krajowej, naukowego centrum doskonałości (CD).

8

Od szeregu lat, w różnych konfiguracjach osobowych, dyskutowaliśmy o potrzebie zmian w strukturze organizacyjnej Uczelni dla podniesienia poziomu naukowego oraz jakości absolwenta, a także zapewnienia stabilności działania na różnego rodzaju kryzysy finansowe, polityczne i demograficzne. Jednym z pomysłów było rozdzielenie pionów dydaktycznego i naukowego. Ten drugi miał dać swobodę w powstaniu interdyscyplinarnych zespołów, wyróżnionych w formie laboratoriów badawczych (LB). LB mogłyby łączyć się w CD dostosowane do realizacji programów wynikających ze strategii gospodarczej państwa. Integracja pionu dydaktycznego i badawczego miałyby odbywać się na etapie wykonywania prac magisterskich i doktoratów w laboratoriach wybieranych przez studentów. Tak wykształceni absolwenci byłiby cennymi fachowcami, atrakcyjnymi dla gospodarki. Niestety, z wielu względów, nie doszło nawet z symulacyjnych rozmów na te tematy. Mam nadzieję, że nastąpią jakieś ruchy wymuszające ukierunkowany rozwój gospodarczy kraju skutkujące zdefiniowanymi aktywnościami Uczelni z programami dostosowanymi do tych celów.

Zacznijmy więc od innej strony, od konstrukcji przykładu który mógłby stać się wzorcem dla podobnych przedsięwzięć. Budowa CD z obszaru AI będzie możliwa przez przyznanie finansowania na ten cel funduszy ze wspomnianego programu MAB. Stabilność i trwałość takiego centrum zależeć będzie od jego struktury i podejmowanej tematyki badań. Kluczowym elementem tej struktury jest Lider, na

którym spoczywa odpowiedzialność za osiągnięcie wyników w zaplanowanym programie badań. To Lider odpowiada za działanie CD przed podmiotami finansującymi zgodnie z programem i regulaminem MAB.

Biorąc pod uwagę wielowątkowość tematyki badawczej w programie grantu, struktura zarządzania będzie wymagać silnej integracji podmiotów merytorycznych podporządkowanych realizacji głównego celu działalności centrum. Będzie to związane ze zdefiniowaniem i określeniem zadań dla części administracyjnej i podkreśleniem służebnej roli urzędników. Różnorodność tematyczna działalności centrum w sposób naturalny wymuszać będzie dezintegrację całej struktury w zakresie szczegółowych badań zapewniając swobodę aktywności kierowników części zadaniowej dla osiągnięcia zaplanowanych celów badawczych. Dynamika całego systemu w jego działaniu powinna doprowadzić do osiągnięcia stabilnej równowagi prowadzącej do spełnienia wymagań dotyczących całego projektu.

Zdaję sobie sprawę, że zespół badawczy sam już określił sposób działania i ma sprecyzowane plany, o których wie najlepiej jak je realizować. Powyższe refleksje dotyczą przyszłości i wskazują na te elementy, które zapewnią trwałość istnienia centrum i spowodują, że będzie to nowa jakość w strukturze Uczelni.

AI staje się ważnym obszarem aktywności naukowej, a wyniki badań odbierane przez zewnętrzne otoczenie gospodarcze i społeczne mają pozytywne, ale i negatywne skutki. Często nie mamy wpływu na to jak wykorzystane będą uzyskane wyniki. Ukierunkowanie badań i wzmocnienie wektora pozytywnych skutków jest trudne ze względu na moralną i merytoryczną jakość odbiorców. Dominująca postnowoczesność rozszerza nieuchwytność sensu w kontrze do postaw konserwatywnych, gdzie kluczowe epistemologicznie jest wykorzystanie związków przyczynowo-skutkowych zjawisk, logiki i empirii. Te narzędzia dominująco zastąpione zostały labiryntem sugestii i przypuszczeń. System zaufków, z którego często nie ma wyjścia, ani potrzeby, żeby w nim cokolwiek poszukiwać, a wszystko można zakwestionować, jest groźny przede wszystkim dla nauki.

W epoce natychmiastowej globalnej łączności i stale rosnącej jednorodności idei, coraz trudniej nadać nowym pomysłom wartościowy sens i kierunek. Ogłupiające i powszechnie obecne w życiu obywateli mass-media zdominowały drogi dotarcia do odbiorcy. Rozchwiane zostały system wartości i postawy etyczne. Nad AI wisi groźne niebezpieczeństwo szkodliwego użycia nowych osiągnięć i narzędzi.

Tym bardziej osadzenie projektu MAB-FENG w ramach powstającego centrum, we właściwych ryzach pozytywnych ograniczeń i zdefiniowania użyteczności wyników badań jest dużym wyzwaniem i wymaga szczególnej troski wszystkich uczestników tego przedsięwzięcia. Dobry przykład właściwego funkcjonowania Centrum może być załącznikiem do naśladownictwa prowadzącego do ewolucyjnych zmian w Uczelni.

Ewa BARTNIK

Wydział Biologii
Uniwersytet Warszawski

Kilka bioetycznych problemów biomedycznych

Szybki rozwój nauk biomedycznych doprowadził do ogromnych osiągnięć, ale także do wielu pytań. Sekwencjonowanie DNA stało się dostępne, i coraz częściej jest stosowane, ale jak i o czym należy informować osobę, od której pochodzi DNA? Czy każdy powinien mieć zsekwencjonowany genom tuż po narodzinach? Czy jeśli dla danej choroby nie ma terapii a występuje u dorosłych należy określać obecność powodującej ją mutacji u nieletnich? Terapia genowa zaczyna mieć ogromne osiągnięcia, ale też kosztuje miliony dolarów - jak i czy możliwe jest sprawiedliwe podejście to jej stosowania? Czy możemy i czy powinniśmy modyfikować DNA zarodków by uniknąć chorób? Jak traktować embrioidy – modelowe zarodki ludzkie powstające bez udziału plemnika i komórki jajowej? To tylko kilka przykładów.

Krzysztof EJSMONT

Wydział Mechaniczny Technologiczny
Politechnika Warszawska

Kontrola jakości wspierana sztuczną inteligencją w sektorze MŚP: wnioski z badań empirycznych

Współczesna nauka i przemysł rozwijają się w warunkach dynamicznych przemian technologicznych i społecznych, w których sztuczna inteligencja (AI) staje się jednym z kluczowych narzędzi transformacji. Jej wdrażanie w sektorze małych i średnich przedsiębiorstw (MŚP) odślania zarówno potencjał, jak i ograniczenia. Prezentacja przedstawia wyniki badań empirycznych nad zastosowaniem AI w obszarze kontroli jakości w europejskich MŚP, koncentrując się na procesach integracji nowych technologii z pracą ludzką oraz na zjawiskach dezintegracyjnych wynikających z barier kompetencyjnych, oporu organizacyjnego i trudności w zarządzaniu danymi. Szczególną uwagę poświęcono granicom stosowalności AI, uwzględniając aspekty regulacyjne i etyczne, które warunkują odpowiedzialne wdrażanie tej technologii. Pokazane przykłady wdrożeń wskazują nie tylko na wymierne korzyści – poprawę jakości, bezpieczeństwa i efektywności – lecz także na kluczową rolę interdyscyplinarnej współpracy badaczy i praktyków. Dyskusja obejmuje również nowe okoliczności kształtujące kierunki badań, oczekiwania użytkowników i potencjalne zagrożenia, wyznaczając ścieżki odpowiedzialnego rozwoju AI w przemyśle i nauce.

Bartłomiej SKOWRON

Wydział Administracji i Nauk Społecznych
Politechnika Warszawska

2441

2441 lat minęło, odkąd w Atenach w domu tragika Agatona odbyła się uczta, aby uczcić zwycięstwo Agatona w zawodach poetyckich. Zebrani filozofowie, artyści i arystokraci postanowili, że będą chwalić Erosa, boga miłości. Pośród wygłaszających mowy na część Erosa byli m.in. młody uczeń Sokratesa Fajdros, sofista Pausaniasz, lekarz Eryksimachos, słynny komediopisarz Arystofanes, Agaton i sam Sokrates. O treści owych pochwał dowiadujemy się z klasycznego dialogu Platona o niezaskakującym tytule „Uczta”. W swoim wystąpieniu przypomnę zasadnicze charakterystyki Erosa z poszczególnych mów, a w szczególności z mowy pochwalnej Sokratesa, a następnie przedstawię eksperyment myślowy: co byłoby, gdybyśmy Erosa na powrót zaprosili do PW.

Jacek MAŃDZIUK

Wydział Matematyki i Nauk Informatycznych
Politechnika Warszawska

Abstrakcyjne rozumowanie wizualne: wciąż poza zasięgiem modeli sztucznej inteligencji

Abstrakcyjne rozumowanie wizualne (ang. *Abstract Visual Reasoning, AVR*) obejmuje problemy wymagające zdolności odkrywania wspólnych pojęć leżących u podstaw zbioru obrazów poprzez proces analogii, podobny do sposobu, w jaki ludzie rozwiązują testy IQ. W wystąpieniu krótko scharakteryzuję główne typy problemów AVR, ze szczególnym uwzględnieniem tzw. Problemów Bongarda (ang. *Bongard Problems*), które stanowią fundamentalne wyzwanie w obszarze AVR, głównie ze względu na konieczność połączenia rozumowania wizualnego z opisem werbalnym oraz ograniczoną liczbę przykładów uczących.

13

Janusz HOŁYST

Wydział Fizyki
Politechnika Warszawska

Wielopoziomowe przeciążenie informacyjne: dlaczego jest niebezpieczne i czy możemy je pokonać?

Obecnie każdego dnia otrzymujemy więcej informacji, niż jesteśmy w stanie przetworzyć, co wiąże się z poważnymi kosztami w gospodarce i prowadzi również do zagrożenia naszego zdrowia. Będę chciał Państwa przekonać, że przestrzeń informacyjna powinna być obecnie postrzegana jako część naszego naturalnego środowiska i dlatego niezbędne jest prowadzenie badań nad skutkami zarządzaniem przeciążenia informacyjnego i jego ograniczeniem.

We współczesnym świecie dostęp do informacji jest często traktowany jako dobro lub wręcz niezbywalne prawo człowieka. Istnieje jednak również „ciemna strona” informacji: nadmiar danych przekraczający możliwości ich przetwarzania prowadzi do tzw. przeciążenia informacyjnego (ang. information overload, IOL). Pojęcie to niepokoiło ludzkość na długo przed wynalezieniem druku i było badane z różnych perspektyw – od neurobiologii po dziennikarstwo.

Zazwyczaj przeciążenie informacyjne rozpatrywane jest jednak na poziomie jednostkowym poprzez analizę pojedynczych czynników lub jednego poziomu, które ostatecznie prowadzą do ograniczenia aktywności danej jednostki. Nie są natomiast znane skutki IOL występującego jednocześnie na wielu poziomach, czyli tzw. wielopoziomowego przeciążenia informacyjnego.

Powyższe obserwacje stanowią punkt wyjścia do głównego celu międzynarodowego projektu finansowanego z grantu UE: OMINO – Overcoming Multilevel Information Overload. Projekt OMINO ma na celu pomiar wielopoziomowego przeciążenia informacyjnego w różnych systemach, opracowanie modeli IOL oraz metod przeciwdziałania temu zjawisku. W czasie referatu omówię m.in. problemy przecią-

żenia informacyjnego w nauce, zarządzaniu, procesach poznawczych oraz relacje między przeciążeniem informacyjnym i rozwojem sztucznej inteligencji.

Referencje

1. <https://ominoproject.eu>
2. J.A. Holyst et al, *Protect our environment from information overload*, *Nature Human Behaviour* (2024), vol. 8, s.402-403. DOI:10.1038/s41562-024-01833-8

Jacek ROGALA

Centrum Badania Ryzyka Systemowego
Wydział *Artes Liberales*
Uniwersytet Warszawski

AI w badaniach w zastosowaniach międzydziedzinowych szansa i zagrożenie

Sztuczna inteligencja jest dziś siłą napędową nauki i wielu gałęzi przemysłu. Jednak wiele jej zastosowań w nauce i biznesie wydaje się nie trafiać w cel – szczególnie w obszarze, który znam z własnych doświadczeń badawczych, czyli biologii i medycyny. Część z tych prac prezentuje znakomite umiejętności programistyczne czy metody, ale stosowane przez inżynierów i programistów bez dogłębnej znajomości problemu prowadzą do trywialnych, a nawet mylących wyników. Z drugiej strony wielu specjalistów dziedzinowych sięga po gotowe rozwiązania, również bez głębszego zrozumienia narzędzi, co także skutkuje błahymi lub myłymi rezultatami. Wysiłek i energia tysięcy badaczy są w ten sposób zwyczajnie marnowane. Jedynym – choć banalnym – rozwiązaniem jest ścisła współpraca ekspertów AI i specjalistów dziedzinowych. Okazuje się to jednak bardzo trudne – język i metodologia obu stron są często odmienne, a zbudowanie wspólnego zespołu wymaga czasu.

Marcin KOSZOWY

Wydział Administracji i Nauk Społecznych
Politechnika Warszawska

Arystoteles, duże zbiory danych i społeczeństwo cyfrowe: jak integrować badania wokół koncepcji filozoficznych?

Arystotelesowska koncepcja etosu mająca wymiary wiedzy i umiejętności (*phronesis*), moralności (*areté*) i życzliwości (*eunoia*) stanowi źródło inspiracji do opracowania programu integrującego badania nad przemianami komunikacji w społeczeństwie cyfrowym. Zostanie ona omówiona jako przykład teoretycznej podstawy systemu wykorzystującego duże zasoby danych do identyfikowania tendencji kształtujących percepcję wiarygodności i zaufania względem uczestników procesu komunikacji. W tym kontekście strategie atakowania, podważania lub wzmacniania autentyczności mówców lub użytkowników mediów społecznościowych zyskują zarówno teoretyczną ramę modelową, jak i solidne ugruntowanie empiryczne. Ten przykład nie zawsze oczywistego połączenia klasycznych idei filozoficznych z technikami tworzenia korpusów typowymi dla lingwistyki obliczeniowej oraz z metodami eksploracji danych czyni z arystotelesowskiego ujęcia etosu narzędzie systematycznego badania wielkoskalowych strategii i taktyk komunikacyjnych, inspirując m.in. do refleksji nad wiedzotwórczym i moralnym wymiarem komunikacji w dynamicznym środowisku cyfrowym.

17

Przemysław BIECEK

Wydział Matematyki i Nauk Informatycznych
Politechnika Warszawska

Modelologia – nowe i ciekawe wyzwania związane z analizą modeli

Rosnąca popularność modeli fundamentalnych toruje ścieżkę dla nowej dziedziny badań – badań nad modelami, tzw. Model Science. W przeciwieństwie do podejść skoncentrowanych na danych, Model Science umieszcza wyszkolony model w centrum analizy, mając na celu interakcję, weryfikację, wyjaśnienie i kontrolę jego zachowania w różnych kontekstach operacyjnych. Niniejsza prezentacja wprowadza ramy koncepcyjne dla nowej dziedziny, wraz z propozycją jej czterech kluczowych filarów: Weryfikacja, która wymaga ścisłych, świadomych kontekstu protokołów oceny; Wyjaśnienie, które jest rozumiane jako różne podejścia do badania wewnętrznych operacji modelu; Kontrola, która integruje techniki dostosowywania w celu kierowania zachowaniem modelu; oraz Interfejs, który rozwija interaktywne i wizualne narzędzia wyjaśniające w celu poprawy ludzkiej kalibracji i podejmowania decyzji. Proponowane ramy mają na celu kierowanie rozwojem wiarygodnych, bezpiecznych i dostosowanych do człowieka systemów sztucznej inteligencji.

Dariusz PLEWCZYŃSKI

Wydział Matematyki i Nauk Informatycznych
Politechnika Warszawska

Zespoły interdyscyplinarne w trójwymiarowej genomice: od sekwencji DNA do struktury chromatyny w skali populacji

W referacie pokażę wizję, jak zespoły interdyscyplinarne łączące biologów molekularnych, informatyków, fizyków, chemików i matematyków – wspierane agentami AI — mogą przejść od sekwencji DNA do predykcji i wyjaśniania trójwymiarowej struktury chromatyny w skali populacyjnej. Punktem wyjścia jest koncepcja Wirtualnego Laboratorium (Virtual Lab), w której zespół badaczy współpracuje z zespołem agentów-naukowców (PI, krytyk, specjaliści) prowadzących iteracyjne „narady” i realizujących pełne cykle projektowe – podejście, które przyniosło już zweryfikowane eksperymentalnie wyniki i które adaptujemy do 3D genomiki. Przedstawię, jak modele sekwencyjne – transformatory i hybrydy CNN transformer (np. Enformer, Borzoi, C.Origami/Orca) – mapują DNA na sygnały regulacyjne i/lub macierze kontaktów Hi C, podkreślając ich mocne strony (globalny kontekst setek kb) i ograniczenia (pamięć/megabajowy zasięg), a także kiedy łączyć je z fizyką polimerów. Pokażę również, jak agentowe modele podstawowe dla pojedynczych komórek (np. Geneformer, scGPT) integrują wielomodalne dane, przewidują skutki perturbacji i pomagają priorytetyzować hipotezy oraz edycje sekwencji do testów. W warstwie eksperymentalnej omówię zestaw protokołów dla walidacji w pętli: Hi C/Micro C, obrazowanie super rozdzielcze, CRISPRi/a i skale pojedynczej komórki. Zakończę pokazaniem, jak takie zespoły – z jasno zdefiniowanymi rolami ludzi i agentów AI – umożliwiają biobankową (setki tysięcy – miliony genomów) imputację stanów chromatyny i ocenę wariantów funkcjonalnych w tkankach, wraz z omówieniem wyzwań interpretowalności, solidności i etyki współpracy człowiek-AI.

19

Agnieszka JASTRZĘBSKA

Wydział Matematyki i Nauk Informatycznych
Politechnika Warszawska

Granice przewidywania

Prognostowanie przyszłości na podstawie danych historycznych jest jedną z popularniejszych dziedzin analizy danych. W wystąpieniu przyjrzymy się popularnonaukowemu aspektowi prognozowania szeregów czasowych: różnym modelom analizy danych oraz pojęciom ważnym dla tej dziedziny. Zastanowimy się również nad ograniczeniami tych metod – zwłaszcza w obliczu tzw. czarnych łabędzi, czyli rzadkich, nieprzewidywalnych zdarzeń o ogromnych konsekwencjach. Dlaczego nawet najlepsze modele zawodzą w obliczu pandemii, kryzysów finansowych czy przełomów technologicznych? Czy prognozowanie to sztuka, nauka, czy może iluzja kontroli nad niepewnością?

Celem wystąpienia jest pokazanie, jak naukowcy i praktycy balansują między matematycznym porządkiem modeli a chaosem nieprzewidywalnych zdarzeń, oraz jakie wnioski możemy wyciągnąć z tej wiedzy na co dzień.

Błażej ŻYLIŃSKI

Wydział Fizyki
Politechnika Warszawska

Od fragmentacji do synergii: sztuczna inteligencja wobec barier interdyscyplinarności

Współczesne badania naukowe coraz częściej wymagają podejścia interdyscyplinarnego, jednak realizacja takiej współpracy napotyka liczne bariery. Należą do nich zarówno wyzwania organizacyjne, komunikacyjne, instytucjonalne ^[1], jak i ograniczenia systemowe, które obniżają szanse zdobycia finansowania na interdyscyplinarne badania ^[2]. Równocześnie publikacje zespołowe nie tylko dominują w tworzeniu wiedzy, ale też są częściej cytowane i mają większą szansę uzyskać status wpływowych w swoich dziedzinach ^[3]. Jednakże brak wspólnych standardów zarządzania danymi utrudnia budowanie spójnych projektów i długofalowych partnerstw ^[5].

21

Sztuczna inteligencja (SI) otwiera nowe możliwości przezwyciężania tych przeszkód. Z jednej strony wspiera analizę złożonych zbiorów danych i ułatwia przegląd obszernej literatury, z drugiej – znajduje zastosowanie w przełomowych przedsięwzięciach, takich jak AlphaFold ^[4]. Pokazują one jak SI integruje różne kompetencje i skraca drogę od postawienia hipotezy badawczej do jej weryfikacji. Warunkiem sukcesu jest jednak odpowiedzialne wdrażanie zasad otwartości i interoperacyjności (FAIR) ^[5] oraz refleksja nad ryzykami i regulacjami.

Celem wystąpienia jest omówienie głównych barier w naukowej współpracy interdyscyplinarnej oraz wskazanie, w jaki sposób odpowiedzialnie rozwijana sztuczna inteligencja może sprzyjać ich przezwyciężaniu, a w konsekwencji stymulować powstawanie nowych obszarów wiedzy.

Literatura

1. Rafols, I. et al. Quantitative Science Studies 4(3), 711–735 (2023). https://doi.org/10.1162/qss_a_00265
2. Bromham, L. et al. Nature 534, 684–687 (2016). <https://doi.org/10.1038/nature18315>
3. Wuchty, S. et al. Science 316(5827), 1036–1039 (2007). <https://doi.org/10.1126/science.1136099>
4. Jumper, J. et al. Nature 596, 583–589 (2021). <https://doi.org/10.1038/s41586-021-03819-2>
5. Wilkinson, M.D. et al. Scientific Data 3, 160018 (2016). <https://doi.org/10.1038/sdata.2016.18>

Łukasz ŻUKOWSKI, Jakub MOŻARYN

Wydział Mechatroniki
Politechnika Warszawska

Wykorzystanie dużych modeli językowych w programowaniu graficznym sterowników PLC: możliwości, ograniczenia i rekomendacje dla edukacji przemysłowej

22

W środowiskach przemysłowych sterowniki programowalne (PLC) są szeroko stosowane do zarządzania małymi i średnimi liniami produkcyjnymi, umożliwiając użytkownikom tworzenie systemów sterowania i monitorowania stanu linii. Wielu początkujących programistów PLC napotyka poważne wyzwania, ponieważ przejście od teorii do praktyki często ujawnia braki w umiejętnościach, m.in. młodszy programiści PLC, mają trudności ze zrozumieniem niuansów graficznych języków programowania, takich jak diagramy drabinkowe (LAD) i diagramy bloków funkcyjnych (FBD), powszechnie stosowanych w programowaniu PLC. Języki graficzne, takie jak LAD lub FBD, umożliwiają analizę wizualną, zrozumienie i wdrożenie logiki sterowania, stanowiąc dobre rozwiązanie w przypadku mniej złożonych systemów sterowania, natomiast tekst strukturalny (ST) nadaje się do złożonej logiki i obliczeń, przypominając programowanie wysokiego poziomu. Na Politechnice Warszawskiej (PW) prowadzone są na wielu wydziałach zajęcia wprowadzające do programowania PLC. Nasze badania są motywowane rosnącą popularnością i dostępnością dużych modeli językowych (LLM), których kluczowym zastosowaniem jest generowanie kodu w różnych językach programowania, które młodszy programiści PLC coraz częściej wykorzystują do usprawnienia rozwoju systemów sterowania, co podkreśla potrzebę wsparcia narzędziami ułatwiającymi przygotowanie programów. W związku z tym warto zbadać, czy modele LLM mogą być wykorzystywane do pisania i analizy programów dla sterowników PLC, potencjalnie automatyzując żmudne, podatne na błędy i powtarzalne zadania programistyczne, działając jako „drugi pilot inżyniera” w zakresie generowania kodu i dokumentacji.

Nasze prace stanowią rozszerzenie dotychczasowych badań nad modelami LLM w programowaniu graficznym, wykraczając poza zadania logiki dyskretnej i skupiając się na systemach sterowania ciągłego oraz ich walidacją. Praca podejmuje

problem braku badań nad oceną logiki sterowania w językach LAD i FBD generowanej przez LLM w przypadku złożonych systemów, takich jak np. kaskadowa regulacja temperatury powietrza. Celem pracy było sprawdzenie użyteczności dużych modeli językowych (LLM), na przykładzie ChatGPT-4o, w programowaniu sterowników PLC za pomocą języków graficznych, takich jak Ladder Diagram (LAD) lub Function Block Diagram (FBD). Obecnie powstające rozwiązania (LLM-4PLC, Agents4PLC) skupiają się na wykorzystaniu języków tekstowych, takich jak tekst strukturalny (ST), które ze względu na większe podobieństwo do języków programowania ogólnego przeznaczenia, takich jak C++ lub Java, jednak ich struktura jest trudniejsza do analizy przez człowieka. W przypadku języków graficznych zachodzi odwrotna zależność: są bardziej problematyczne dla LLM w analizie, natomiast łatwiejsze dla człowieka.

Podczas interakcji z modelem przyjęto dwie strategie budowy logiki: rozwój inkrementalny, gdzie program stopniowo rozbudowywano o kolejne funkcjonalności oraz pełna prezentacja zadania, gdzie po przedstawieniu wymagań model był proszony o dokładne opisywanie poszczególnych elementów programu. Do weryfikacji uzyskanych programów wykorzystano stanowisko wyposażone w sterownik PLC do regulacji temperatury powietrza w rurociągu. Wyniki wskazują na ogólną poprawność modelu w programowaniu regulacji jednoobwodowej, natomiast w przypadku regulacji kaskadowej pojawiają się liczne błędy, w tym niespójności logiczne. By uzyskać poprawny kod, niezbędne były uwagi ze strony programisty. Zauważono również większą skuteczność modelu w diagnozowaniu błędów w kodzie w sytuacji, gdy ma wybrać najbardziej prawdopodobną przyczynę spośród podanych, niż gdy ma sam takową zaproponować.

Wnioski z badań wskazują, że ChatGPT-4o może wspierać programowanie PLC w językach graficznych (LAD, FBD), dostarczając przydatnych instrukcji, tłumacząc działanie niektórych funkcji i tworząc podstawowe szkielety, jednak ograniczenia obejmują fakt, że wygenerowany kod często zawiera błędy niskiego poziomu, generowanie odpowiedzi jest czasochłonne i podatne na błędną interpretację pytania operatora, a ślepe poleganie na wynikach LLM jest niebezpieczne.

Zalecenia obejmują wykorzystywanie LLM w roli doradczej oraz edukacyjnej, a także do tworzenia wstępnych ram programu, zwracanie uwagi na treść sformułowanych zapytań, aby uniknąć błędnej interpretacji przez LLM, oraz dokładną weryfikację zaproponowanego kodu pod kątem zarówno jego wewnętrznej spójności logicznej, jak i zgodności z rozwiązaniami dostępnymi w używanym środowisku. Rekomendujemy również dalsze badania nad automatyzacją "promptów" i oceną programów ST (LLM4PLC, Agents4PLC, AutoPLC), a także opracowanie dedykowanych offline modeli LLM, aby zminimalizować ryzyko cyberataków i szpiegostwa.

Praktyczne implikacje wskazują, że młodzi inżynierowie mogą szybciej tworzyć akceptowalny kod i uzyskać cenne wsparcie w nauce, przy obserwowanym wzroście popularności „przemysłowych kopilotów”, czyli zintegrowanych narzędzi opartych na LLM, wspomagających programistów PLC. Manualna weryfikacja kodu pozostanie niezbędna w przypadku systemów sterowania o krytycznym zagrożeniu dla bezpieczeństwa.