



Filip Wójcik

Graflowe sieci neuronowe

Teoria i praktyka

Wydawnictwo
Naukowe
Helion

Wszelkie prawa zastrzeżone. Nieautoryzowane rozpowszechnianie całości lub fragmentu niniejszej publikacji w jakiegokolwiek postaci jest zabronione. Wykonywanie kopii metodą kserograficzną, fotograficzną, a także kopiowanie książki na nośniku filmowym, magnetycznym lub innym powoduje naruszenie praw autorskich niniejszej publikacji.

Wszystkie znaki występujące w tekście są zastrzeżonymi znakami firmowymi bądź towarowymi ich właścicieli.

Autor oraz wydawca dołożyli wszelkich starań, by zawarte w tej książce informacje były kompletne i rzetelne. Nie biorą jednak żadnej odpowiedzialności ani za ich wykorzystanie, ani za związane z tym ewentualne naruszenie praw patentowych lub autorskich. Autor oraz wydawca nie ponoszą również żadnej odpowiedzialności za ewentualne szkody wynikłe z wykorzystania informacji zawartych w książce.

Redaktor prowadzący: Małgorzata Kulik

Projekt okładki: Studio Gravite/Olsztyn
Obarek, Pokoński, Pazdrijowski, Zaprucki

Grafika na okładce została wykorzystana za zgodą AdobeStock.com.

Helion S.A.
ul. Kościuszki 1c, 44-100 Gliwice
tel. 32 230 98 63
e-mail: helion@helion.pl
WWW: helion.pl (księgarnia internetowa, katalog książek)

Drogi Czytelniku!

Jeżeli chcesz ocenić tę książkę, zajrzyj pod adres

helion.pl/user/opinie/siegra

Możesz tam wpisać swoje uwagi, spostrzeżenia, recenzję.

Kolorowe wersje wybranych rysunków dostępne są pod adresem:

<https://ftp.helion.pl/przyklady/siegra.zip>.

ISBN: 978-83-289-3390-3

Copyright © Filip Wójcik 2026

Printed in Poland.

- [Kup książkę](#)
- [Poleć książkę](#)
- [Oceń książkę](#)

- [Księgarnia internetowa](#)
- [Lubię to! » Nasza społeczność](#)

Spis treści

Wstęp	5
Notacja i oznaczenia	11
ROZDZIAŁ 1. Narzędzia analizy grafów w środowisku Pythona	19
1.1. Biblioteki do klasycznej analizy grafów	19
1.2. Grafowe bazy danych	20
1.3. Biblioteki do budowania grafowych sieci neuronowych	20
1.3.1. PyTorch Geometric	20
1.3.2. Deep Graph Library	21
1.3.3. Porównanie bibliotek	21
ROZDZIAŁ 2. Wybrane zagadnienia teorii grafów	24
2.1. Podstawowe definicje i oznaczenia	25
2.2. Reprezentacja grafów i sąsiedztwa	27
2.2.1. Macierze sąsiedztwa	28
2.2.2. Listy sąsiedztwa	30
2.3. Liczbowe własności wierzchołków i grafów	32
2.3.1. Miary centralności wierzchołków	33
2.3.2. Numeryczna reprezentacja grafu	36
2.3.3. Problem izomorfizmu grafów i test Weisfeilera-Lehmana	39
2.4. Grafy heterogeniczne	43
2.4.1. Podstawowe pojęcia	43
2.4.2. Reprezentacja grafów heterogenicznych	45
ROZDZIAŁ 3. Grafowe sieci neuronowe GNN — wprowadzenie	47
3.1. Zadania realizowane przez sieci GNN	47
3.2. Podstawowe zasady działania sieci GNN	49
3.3. Model przekazywania wiadomości — MPNN	51
3.4. Implementacja modelu MPNN w PyG	56
3.5. Modele MPNN jako część większej sieci	60
ROZDZIAŁ 4. Przegląd wybranych warstw spłotu grafowego	67
4.1. Splot GNN	70
4.1.1. Formalizacja i działanie	70
4.1.2. Implementacja	72
4.1.3. Podsumowanie	74
4.2. Splot GCN	74
4.2.1. Formalizacja i działanie	75
4.2.2. Implementacja	76
4.2.3. Podsumowanie	80
4.3. Splot SAGE	80
4.3.1. Formalizacja i działanie	81
4.3.2. Implementacja	85
4.3.3. Podsumowanie	90

4.4. Splot GAT	90
4.4.1. Formalizacja i działanie	90
4.4.2. Implementacja	99
4.4.3. Podsumowanie	105
4.5. Splot GIN	106
4.5.1. Formalizacja i działanie	106
4.5.2. Implementacja	110
4.5.3. Podsumowanie	114
4.6. Warstwy splotu dla grafów heterogenicznych	114
4.7. Podsumowanie omówionych warstw splotu	125

ROZDZIAŁ 5. Wybrane zagadnienia procesu szkolenia sieci grafowych 131

5.1. Podział danych grafowych na treningowe, walidacyjne i testowe	131
5.1.1. Indukcja i transdukcja	134
5.1.2. Podziały wierzchołków	134
5.1.3. Podziały krawędzi	140
5.1.4. Podziały grafów	150
5.2. Proces szkolenia na dużych zbiorach danych	151
5.2.1. Próbkowanie w oparciu o sąsiedztwo	152
5.2.2. Próbkowanie w oparciu o społeczności	160
5.2.3. Podsumowanie procesu szkolenia	162
5.3. Trudności i wyzwania w uczeniu warstw splotu grafowego	164
5.3.1. Problem nadmiernego wygładzania	164
5.3.2. Ograniczenie testem Weisfeilera-Lehmana	171
5.4. Dodatkowe modyfikacje warstw splotu	
usprawniające proces predykcji	175
5.4.1. Network in graph neural network	175
5.4.2. Agregacje wielokrotne	178
5.4.3. Mechanizm połączeń skokowych	180
5.4.4. Porównanie wyników przy zastosowaniu rozszerzeń	187

ROZDZIAŁ 6. Przykłady zastosowań grafowych sieci neuronowych 189

6.1. Klasyfikacja wierzchołków	189
6.1.1. Postać formalna	189
6.1.2. Znaczenie problemu i zastosowania	190
6.1.3. Klasyfikacja tematyczna stron na Facebooku	191
6.2. Klasyfikacja krawędzi	194
6.2.1. Postać formalna	194
6.2.2. Znaczenie problemu i zastosowania	195
6.2.3. Badanie oddziaływania pomiędzy lekami — klasyfikacja krawędzi	196
6.3. Klasyfikacja grafów	199
6.3.1. Postać formalna	199
6.3.2. Znaczenie problemu i zastosowania	203
6.3.3. Badania toksyczności cząsteczek — klasyfikacja grafów	203
6.4. Systemy rekomendacyjne	207
6.4.1. Postać formalna	208
6.4.2. Znaczenie problemu i zastosowania	209
6.4.3. Rekomendacje filmów MovieLens	210

Zakończenie	215
--------------------------	------------

Bibliografia	216
---------------------------	------------

Skorowidz	222
------------------------	------------

ROZDZIAŁ 6.

Przykłady zastosowań grafowych sieci neuronowych

Niniejszy rozdział zostanie poświęcony omówieniu praktycznych zastosowań grafowych sieci neuronowych w kontekście rozwiązywania wybranych problemów biznesowych. W każdej kategorii zadań, takich jak klasyfikacja wierzchołków, predykcja połączeń czy klasyfikacja całych grafów, przedstawiony zostanie konkretny problem wraz z odpowiadającym mu zbiorem danych. Dodatkowo zaprezentowane zostanie przykładowe rozwiązanie wraz z jego implementacją. Taki układ treści ma na celu ułatwienie utrwalenia wiedzy teoretycznej oraz zrozumienia architektury GNN, omówionych w poprzednich rozdziałach.

Ważne: różnice w wynikach

Biblioteki wykorzystane w przykładzie i dołączonych materiałach korzystają z obliczeń opartych na kartach graficznych (CUDA).

Pomimo zastosowania mechanizmów zapewniania deterministyczności procesu szkolenia (m.in. ustawianie ziarna losowego za pomocą biblioteki PyG, deterministyczne uczenie w PyTorch Lightning itp.) nie ma gwarancji, że uzyskane wyniki będą za każdym razem identyczne. Różnice mogą wynikać m.in. **z wersji używanego sterownika CUDA oraz innych systemowych sterowników karty graficznej**, kompilacji bibliotek dla konkretnego systemu operacyjnego (Linux, Windows, macOS).

6.1. Klasyfikacja wierzchołków

6.1.1. Postać formalna

Klasyfikacja wierzchołków w grafach to jedno z głównych zagadnień w analizie danych strukturalnych, które znajduje zastosowanie w wielu dziedzinach, takich jak sieci społecznościowe, biologia obliczeniowa czy analiza finansowa.

Formalnie zadanie to polega na przypisaniu każdemu wierzchołkowi $v_i \in V$ z grafu $G = (V, E)$ etykiety $y_i \in \mathcal{Y}$ spośród zbioru możliwych klas $\mathcal{C}$.

Aby wykonać zadanie klasyfikacji wierzchołków, niezbędne są, oprócz grafu, następujące elementy:

- zbiór wierzchołków $V = \{v_1, v_2, \dots, v_N\}$, gdzie $N = |V|$;
- zbiór krawędzi $E \subseteq V \times V$, które reprezentują relacje między wierzchołkami;

- macierz cech $\mathbf{X} \in \mathbb{R}^{N \times d}$, gdzie każdy wiersz $x_i \in \mathbb{R}^d$ zawiera wektor cech opisujących wierzchołek v_i ;
- zbiór etykiet $Y = \{y_1, y_2, \dots, y_N\}$, gdzie $y_i \in Y = \{1, 2, \dots, C\}$.

Celem klasyfikacji jest nauczenie modelu:

$$\text{GNN}: (G, \mathbf{X}) \rightarrow Y, \quad (45)$$

który przewiduje etykiety dla wszystkich lub wybranych wierzchołków. W przypadku częściowo nadzorowanego uczenia (ang. *semi-supervised learning*) etykiety są znane jedynie dla podzbioru wierzchołków $V_L \subset V$, a zadaniem modelu jest przewidzenie etykiet dla pozostałych $V_U = V \setminus V_L$. Problem ten można formalnie zapisać jako minimalizację funkcji kosztu:

$$\mathcal{L} = \frac{1}{|V_L|} \sum_{v_i \in V_L} \text{loss}(\text{GNN}(v_i; G, \mathbf{X}), y_i), \quad (46)$$

gdzie:

- $\mathcal{L}(\cdot, \cdot)$ to globalna funkcja kosztu (np. entropia krzyżowa);
- $\text{loss}(\cdot)$ to funkcja kosztu dla pojedynczego przykładu uczącego;
- $\text{GNN}(v_i; G, \mathbf{X})$ to model splotu grafowego dokonujący predykcji klasy dla wierzchołka v_i ,
- y_i to rzeczywista etykieta przypisana do tego wierzchołka.

W przeciwieństwie do klasycznych modeli uczenia maszynowego, grafowe sieci neuronowe stanowią efektywne podejście do tego problemu, gdyż wykorzystują zarówno lokalne informacje o sąsiedztwie każdego wierzchołka (struktura grafu), jak i jego cechy.

6.1.2. Znaczenie problemu i zastosowania

Klasyfikacja wierzchołków ma szerokie zastosowania praktyczne.

Jednym z najważniejszych jest wykrywanie botów w sieciach społecznościowych. W tym przypadku modele analizują grafy znajomości oraz wzorce interakcji, aby identyfikować konta automatyczne lub prowadzące złośliwą działalność. Techniki te są skuteczne nawet przy ograniczonych danych wejściowych i pozwalają na rozpoznawanie botów o różnych strategiach kamuflażu (Kalameyets, 2021).

W sektorze finansowym klasyfikacja wierzchołków wspiera wykrywanie oszustw poprzez analizę relacji między transakcjami, punktami sprzedaży i użytkownikami kart kredytowych. Sieci GNN umożliwiają identyfikację złożonych wzorców oszustw znacznie skuteczniej niż tradycyjne metody (Johannessen, Jullum, 2023; Pereira, Murai, 2021).

W biologii obliczeniowej klasyfikacja genów na podstawie sieci molekularnych pozwala przewidywać funkcje nieznanych genów oraz ich związki z chorobami. Metody te wykorzystują podejście „wina przez skojarzenie” (ang. *guilt-by-association*), zakładające, że geny silnie połączone w sieci mają podobne funkcje (R. Liu, Mancuso, Yannakopoulos et al., 2020; Mancuso, Johnson, Liu et al., 2024).

W bieżącym podrozdziale zajmiemy się klasyfikacją tematyczną stron internetowych na Facebooku na podstawie ich sąsiedztwa (wzajemnych linków) oraz zawartości. Analiza ta wspiera personalizację treści i optymalizację kampanii reklamowych. Przykładowo klasyfikacja interakcji na stronach marek pozwala zrozumieć wzorce zaangażowania użytkowników oraz segmentować strony według ich popularności i aktywności fanów (Chiu, 2021).

6.1.3. Klasyfikacja tematyczna stron na Facebooku



Eksperyment 11. Klasyfikacja tematyczna stron na Facebooku

Implementacja eksperymentu znajduje się w notatniku *rozdzial_6_zastosowania_sieci_gnn/01_facebook_klasyfikacja_wierzchołkow.ipynb*.

Ze względu na stopień skomplikowania i objętość omawianego przykładu w części książkowej zostaną omówione założenia, parametry i wyniki. Detale implementacyjne znajdują się w dołączonych materiałach (notatniku).

W tym podrozdziale przedstawiamy szczegółową analizę studium przypadku opartego na zbiorze danych *Facebook Large Page-Page Network*, który jest częścią repozytorium **Stanford SNAP**¹. Zbiór ten ilustruje zastosowanie grafowych sieci neuronowych do wieloklasowej klasyfikacji wierzchołków i koncentruje się na rzeczywistych danych z sieci społecznościowej (Rozemberczki, Allen, Sarkar, 2021).

Opisywany zbiór danych to graf reprezentujący zweryfikowane strony na Facebooku, gdzie:

- **Wierzchołki** odpowiadają oficjalnym stronom.
- **Krawędzie** reprezentują wzajemne „polubienia” między nimi.

Każdy wierzchołek należy do jednej z czterech kategorii zdefiniowanych przez Facebook:

1. Firmy (~30% wartości).
2. Organizacje rządowe (~28% wartości).
3. Politycy (~25% wartości).
4. Programy telewizyjne (~14% wartości).

Problemem badawczym jest **wieloklasowa klasyfikacja wierzchołków** — przewidzenie kategorii każdej strony na podstawie jej cech oraz struktury grafu.

Analizowany graf składa się z 22 470 wierzchołków, 171 002 krawędzi, jest strukturą spójną (tzn. z każdego wierzchołka można dotrzeć do każdego innego), a średni stopień wierzchołka wynosi 15,22, co wskazuje na intensywną siatkę połączeń. Dla porównania maksymalny stopień wierzchołka to 709. Na zbiór treningowy/walidacyjny i testowy wybrano losowo odpowiednio 15 760, 3370 i 3370 wierzchołków.

¹ <https://snap.stanford.edu/data/facebook-large-page-page-network.html> [dostęp: 20 listopada 2025].

Na zbiór cech wierzchołków składa się macierz 31 zanonimizowanych wartości numerycznych, reprezentujących rozmaite cechy poszczególnych witryn.

Ze względu na wielkość grafu i wektora cech oraz społecznościowy charakter problemu do treningu wybrano metodę próbkowania wycinków uczących w oparciu o społeczność (omawianą w podrozdziale 5.2.2).

Pierwszą fazą eksperymentu był wybór hiperparametrów i optymalizacja modeli. Optymalizator dokonywał wyboru spośród:

- architektur splotu SAGE, GAT, GCN oraz ClusterGCN (będącego odmianą GCN przeznaczoną dla problemów społecznościowych);
- agregacji pojedynczych (suma/średnia/maksimum) oraz ich kombinacji dwu- i trzelementowych;
- mechanizmu połączeń skokowych z różnymi sposobami łączenia poprzednich warstw (konkatenacja/maksimum/LSTM);
- ewentualnie dodatkowych hiperparametrów przeznaczonych dla poszczególnych modeli splotu (takich jak normalizacja SAGE lub liczba „głów” uwagi).

Optymalizacja hiperparametrów była prowadzona za pomocą biblioteki Optuna z wykorzystaniem zbiorów treningowego i walidacyjnego. W jej wyniku wybrane zostały najlepsze hiperparametry (tabela 6.1).

TABELA 6.1. Hiperparametry dla modelu klasyfikacji stron na Facebooku

Hiperparametr	Wartość
Architektura splotu	SAGE
Liczba neuronów w pojedynczej warstwie	32
Liczba warstw	3
Agregacja	Maksimum
Połączenia skokowe	Konkatenacja
Normalizacja	Tak
Projekcja	Tak
Normalizacja próbki treningowej	Tak

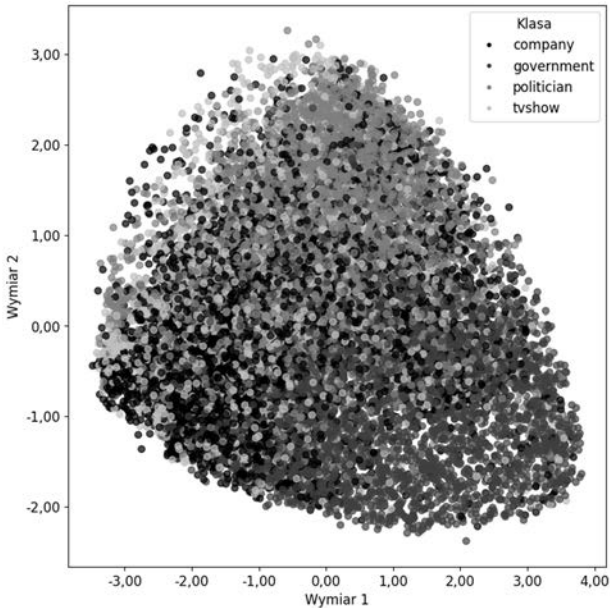
Po zakończeniu tego procesu przystąpiono do szkolenia modelu z pomocą biblioteki PyTorch Lightning. Trening prowadzony był przez 50 epok z uruchomionym mechanizmem wczesnego przerywania po 3 epokach w przypadku pogorszenia wyników oraz bieżącym zapisywaniem najlepszego modelu.

Tabela 6.2 podsumowuje uzyskane wyniki. Widać ewidentną poprawę zarówno dla metryki F1, jak i AUC dla zbioru walidacyjnego — przed treningiem i po nim. Ostateczny wynik dla zbioru testowego nie odbiega znacząco od wyników walidacyjnych.

TABELA 6.2. Wyniki klasyfikacji zbioru danych stron na Facebooku

Metryka	Zbiór	Faza	
		Przed treningiem	Po treningu
F1	Walidacyjny	0,164	0,686
	Testowy	–	0,671
AUC	Walidacyjny	0,493	0,839
	Testowy	–	0,841

Oprócz klasyfikacji za pomocą modelu można go także wykorzystać do wizualizacji zbioru danych. Budując wektory reprezentacji wierzchołków, a następnie poddając je redukcji wymiarowości, można stworzyć reprezentację wektorów osadzenia na płaszczyźnie. Jest to niezwykle przydatne w sytuacji, gdy chcemy sprawdzić jakość wykonania zadania i separację poszczególnych klas. Rysunek 6.1 przedstawia uzyskany rezultat. Widać wyraźnie, że klasy takie jak government, company oraz politician zostały prawidłowo odseparowane i tworzą odróżniające się klastry na przeciwległych krańcach chmury punktów. Klasa tvshow jest wymieszana z pozostałymi. Model ewidentnie miał trudności z odróżnieniem jej od pozostałych.



RYSUNEK 6.1. Wizualizacja wyniku klasyfikacji stron na Facebooku

Przeprowadzony eksperyment pokazuje, że model sieci grafowej relatywnie dobrze poradził sobie z klasyfikacją stron na Facebooku. Polem do poprawy jest zdecydowanie jakość predykcji dla klasy tvshow, która zaniża wynik ogólny i (co jest widoczne w wizualizacji na płaszczyźnie) nie odróżnia się wyraźnie od pozostałych. Możliwe jest przeprowadzenie dalszego strojenia parametrów i wykorzystanie bardziej złożonej architektury.

6.2. Klasyfikacja krawędzi

6.2.1. Postać formalna

Klasyfikacja krawędzi w grafach to zadanie, w którym celem jest przypisanie każdemu połączeniu (krawędzi) e_{uv} między wierzchołkami v oraz u etykiety y_{uv} należącej do zbioru Y . W zależności od specyfiki problemu etykieta może być binarna (0 — krawędź nie występuje, 1 — krawędź występuje) lub mieć wiele wartości odzwierciedlających określone zjawiska. Dla uproszczenia oraz w związku z prezentowanym dalej przykładem omówiona zostanie formalizacja tego działania dla klasyfikacji binarnej.

W rozwiązywaniu tego zadania stosuje się podejście dwuetapowe:

- 1. Krok 1.** Budowanie reprezentacji wierzchołków z wykorzystaniem grafowych sieci neuronowych:

$$\mathbf{h}_v = \text{GNN}(G, \mathbf{X}), v \in V, \quad (47)$$

$$\mathcal{H} = \{\mathbf{h}_v \mid \forall v \in V\}, \quad (48)$$

gdzie G to graf z wierzchołkami i krawędziami, $\mathbf{X}$ to macierz cech wierzchołków, a GNN to model spłotu grafowego.

- 2. Krok 2.** Predykcja etykiety dla par wierzchołków $\{(u, v) \mid u, v \in V\}$ za pomocą wyspecjalizowanego modelu LinkPred. Jako wejście wykorzystuje on wektory reprezentacji skonstruowane przez sieć GNN:

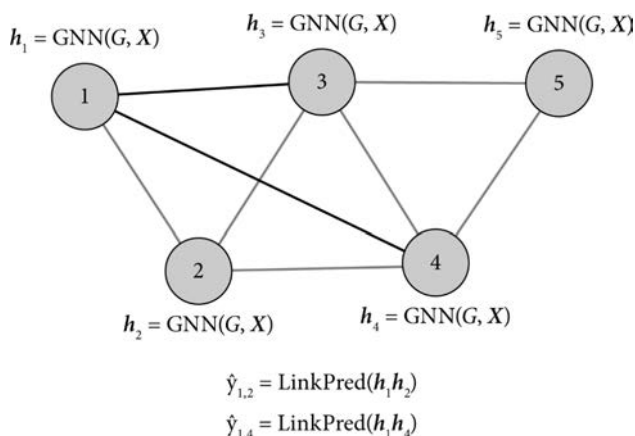
$$\hat{y}_{uv} = \text{LinkPred}(\mathbf{h}_v, \mathbf{h}_u), \hat{y}_{uv} \in [0, 1]. \quad (49)$$

Wynik predykcji, będący prawdopodobieństwem istnienia/nieistnienia krawędzi, jest następnie podstawą do obliczenia funkcji kosztu, najczęściej entropii krzyżowej lub binarnej:

$$\mathcal{L}_{BCE} = \frac{1}{|E|} \sum_{(u,v) \in E} [y_{uv} \log \hat{y}_{uv} + (1 - y_{uv}) \log(1 - \hat{y}_{uv})]. \quad (50)$$

Opisane zadanie można przedstawić graficznie jak na ilustracji poniżej (rysunek 6.2). Początkowo model GNN tworzy reprezentację wektorową każdego wierzchołka. Następnie pary takich reprezentacji (dla wierzchołków połączonych krawędziami) są wykorzystywane do klasyfikacji połączeń (prawdziwe/istniejące — jasne; nieprawdziwe/nieistniejące — ciemne).

W tak przedstawionej formalizacji problemu kryje się jednak pewna pułapka. Standardowo bowiem graf G zawiera tylko krawędzie realnie istniejące — „prawdziwe”. Aby przeprowadzić proces nauki, konieczne jest sztuczne dodanie krawędzi „negatywnych” — oznaczonych etykietą „0” (Hamilton, 2020; Hamilton, Ying, Leskovec, 2017). Taki proces nosi nazwę „próbkiowania negatywnego” (ang. *negative sampling*) i jest szeroko stosowany w zadaniach takich jak systemy rekomendacyjne (Liu, Zhong, Che et al., 2022), a nawet analiza języka naturalnego (Yang, Ding, Huang et al., 2024). Generowane negatywne próbki mogą być całkowicie losowe lub uwzględniać bardziej zaawansowane kryteria (np. odległość topologiczną między wierzchołkami).



RYSUNEK 6.2. Wizualizacja dwustopniowego procesu klasyfikacji wierzchołków

Biblioteka PyG oferuje kilka narzędzi wspomagających proces próbkowania negatywnego:

- `RandomLinkSplit` — umożliwia jednorazowe i statyczne dodanie krawędzi negatywnych do zbiorów treningowego, walidacyjnego i testowego. Narzędzie to pozwala na kontrolowanie proporcji negatywnych próbek względem pozytywnych poprzez parametr `neg_sampling_ratio`.
- `LinkNeighborLoader` — stosowany podczas konstruowania próbek uczących (ang. *batch*). Dynamicznie generuje negatywne próbki dla każdej iteracji treningowej, co pozwala na bardziej efektywne wykorzystanie danych i zwiększenie różnorodności próbek.

Mając na uwadze powyższe, po zastosowaniu próbkowania negatywnego zbiór krawędzi grafu E zawiera zarówno krawędzie „prawdziwe” (realnie istniejące, których etykieta klasy to 1), jak i krawędzie „fałszywe” (nieistniejące, których etykieta klasy to 0): $E = E^+ \cup E^-$, co umożliwia przeprowadzenie procesu uczenia i klasyfikacji binarnej.

6.2.2. Znaczenie problemu i zastosowania

Klasyfikacja krawędzi odgrywa kluczową rolę w wielu dziedzinach nauki i biznesu, umożliwiając analizę relacji między obiektami w sieciach grafowych. Jednym z istotnych zastosowań jest wykrywanie interakcji między lekami (ang. *Drug-Drug Interaction*, DDI), co ma szczególne znaczenie w farmakologii i opiece zdrowotnej. Identyfikacja potencjalnie niebezpiecznych interakcji pozwala na zwiększenie bezpieczeństwa terapii oraz opracowanie skuteczniejszych rekomendacji medycznych (Han, Xie, Li et al., 2022). Problem ten zostanie szerzej omówiony w kolejnym podrozdziale.

W analizie sieci społecznościowych klasyfikacja krawędzi służy do wykrywania powiązań między użytkownikami, co znajduje zastosowanie w systemach rekomendacji znajomości, identyfikacji osób o kluczowym wpływie w danej społeczności oraz detekcji fałszywych kont i botów (Silva, Correia, Maziero, 2023). W sektorze finansowym natomiast umożliwia wykrywanie anomalii poprzez analizę nietypowych

transakcji i powiązań między podmiotami, co wspomaga przeciwdziałanie oszustwom oraz nieprawidłowościom w systemach bankowych i płatniczych (Johannessen, Jullum, 2023; Pereira, Murai, 2021).

Biologia obliczeniowa wykorzystuje klasyfikację krawędzi do badania zależności między białkami i genami. Grafowe sieci neuronowe pozwalają na identyfikację nowych połączeń w sieciach biologicznych, co może prowadzić do głębszego zrozumienia mechanizmów chorób oraz wskazania potencjalnych celów terapeutycznych (Muzio, O'Bray, Borgwardt, 2021; Wu, Gao, Zeng et al., 2022).

6.2.3. Badanie oddziaływania pomiędzy lekami — klasyfikacja krawędzi



Eksperyment 12. Klasyfikacja interakcji pomiędzy lekami



Implementacja eksperymentu znajduje się w notatniku *rozdzial_6_zastosowania_sieci_gnn/02_interakcje_lekow_klasyfikacja_wierzchoлков.ipynb*.

Ze względu na stopień skomplikowania i objętość omawianego przykładu w części książkowej zostaną omówione założenia, parametry i wyniki. Detale implementacyjne znajdują się w dołączonych materiałach (notatniku).

W tym podrozdziale analizujemy praktyczne zastosowanie klasyfikacji krawędzi w grafach na przykładzie sieci interakcji leków (DDI). Studium przypadku opiera się na zbiorze danych *ogbl-ddi*, który pochodzi z platformy Open Graph Benchmark (OGB)² i został zaprojektowany do oceny metod uczenia maszynowego na grafach dla celów branży farmaceutycznej (Guney, 2017). Zbiór ten przedstawia jednorodny, nieskierowany i nieważony graf, w którym:

- 1. Wierzchołki** reprezentują leki zatwierdzone przez FDA lub będące w fazie eksperymentalnej.
- 2. Krawędzie** odzwierciedlają interakcje między lekami, czyli sytuacje, w których wspólne stosowanie dwóch leków prowadzi do efektu istotnie różniącego się od oczekiwanego działania każdego z nich osobno.

Problemem badawczym jest przewidywanie interakcji między lekami na podstawie już znanych powiązań. W tym celu wykorzystuje się zarówno miary F1, AUC i czułość jako metryki ewaluacji, jak i K najlepszych trafień (ang. *Hits@K*), która ocenia zdolność modelu do budowania poprawnego rankingu: prawdziwe interakcje (krawędzie) powinny uzyskiwać wyższe prawdopodobieństwa niż losowo wybrane próbki negatywne. Dla tego zbioru danych przyjęto specyficzny podział na zestawy treningowy, walidacyjny i testowy, bazujący na oddziaływaniach leków na białka. Dzięki temu testowane są możliwości modeli w generowaniu praktycznie użytecznych prognoz dla leków o odmiennych mechanizmach biologicznego oddziaływania w porównaniu do substancji obecnych w zbiorze treningowym.

² <https://ogb.stanford.edu/docs/linkprop/#ogbl-ddi> [dostęp: 20 listopada 2025].

Zbiór ogbl-ddi cechuje się dużą objętością: zawiera **4267 wierzchołków** oraz **1 334 889 krawędzi**, co czyni go wymagającym pod względem obliczeniowym. Na potrzeby tej książki oraz towarzyszącego przewodnika zdecydowano się na zmniejszenie jego wielkości. W tym celu podjęto następujące kroki:

1. Przeanalizowano spójne składowe grafu, czyli połączone elementy sieci.
2. Spośród wszystkich spójnych składowych wybrano tę, która zawiera około 20% wszystkich wierzchołków grafu.
3. Ograniczono graf do tej wybranej składowej, zachowując jedynie około 20% wierzchołków oraz krawędzie łączące je ze sobą.
4. Podzielono dane na zestawy: zbiory treningowy (80%), walidacyjny (10%) oraz testowy (10%). Dodatkowo dodano krawędzie negatywne (brak interakcji) w proporcji 1:1 względem krawędzi pozytywnych (rzeczywistych interakcji).

Po przeprowadzeniu powyższych kroków uzyskano następujące próbki:

1. **Liczba wierzchołków:** 845 we wszystkich zbiorach (treningowym, walidacyjnym i testowym).
2. **Zbiór treningowy:** 136 384 krawędzi (pozytywnych i negatywnych).
3. **Zbiór walidacyjny:** 17 046 krawędzi (pozytywnych i negatywnych).
4. **Zbiór testowy:** 17 046 krawędzi (pozytywnych i negatywnych).

Pomimo redukcji rozmiaru danych proces treningu pozostaje czasochłonny ze względu na dużą liczbę krawędzi i złożoność modelu.

Szkolenie modelu odbywało się z wykorzystaniem mechanizmów próbkowania w oparciu o sąsiedztwo (opisanych w podrozdziale 5.2.1), dobranych w taki sposób, aby zapewnić kompromis między efektywnością uczenia a jego czasem. Szczególnie warto zwrócić uwagę na fragment kodu zaprezentowany poniżej, ukazujący konfigurację klasy `LinkNeighborLoader`:

```
1. train_loader = pyg.loader.LinkNeighborLoader(  
2.     train_data,  
3.     batch_size=1024,  
4.     num_neighbors=[-1, 15, 10, 5],  
5.     shuffle=True,  
6.     edge_label_index=train_data.edge_label_index,  
7.     edge_label=train_data.edge_label,  
8.     num_workers=3)
```

Przedstawiona metoda próbkowania zbioru treningowego rekurencyjnie pobiera sąsiadów w czterech krokach (wiersz 4.):

- wszystkich bezpośrednich sąsiadów danego wierzchołka (wartość -1);
- 15 sąsiadów dla każdego sąsiada bezpośredniego;
- kolejnych 10 w odległości dwóch skoków;
- kolejnych 5 w odległości trzech skoków.

Takie próbkowanie może szybko doprowadzić do gwałtownego wzrostu liczby wierzchołków. Z tego względu ograniczono całkowitą wielkość próbki do 1024 elementów (wiersz 3.).

W przypadku zbiorów walidacyjnego i testowego, ze względu na mniejsze rozmiary, możliwe jest ich całościowe załadowanie do pamięci od razu.

Podobnie jak wcześniej, pierwszą fazę eksperymentu stanowił wybór hiperparametrów i optymalizacja modeli. Tym razem (głównie ze względów wydajnościowych) wykorzystana została tylko architektura SAGE, po której następował model dokonujący klasyfikacji wierzchołka. Hiperparametrybrane podczas optymalizacji to:

- liczba warstw i neuronów dla modelu SAGE;
- sposoby agregacji sąsiednich wierzchołków (suma/średnia/maksimum oraz ich kombinacje 2- i 3-elementowe);
- mechanizm połączeń skokowych z różnymi sposobami łączenia poprzednich warstw (konkatenacja/maksimum/LSTM);
- architektura sieci klasyfikującej połączenia (liczba warstw i neuronów).

Optymalizacja hiperparametrów była prowadzona za pomocą biblioteki Optuna z wykorzystaniem zbiorów treningowego i walidacyjnego.

Tabela 6.3 przedstawia końcowe hiperparametry.

TABELA 6.3. Przykładowe zoptymalizowane hiperparametry dla predykcji połączeń

Hiperparametr	Wartość
Liczba neuronów w warstwie spłotu	128
Liczba warstw	3
Agregacja	Suma, maksimum, średnia
Połączenia skokowe	Konkatenacja
Normalizacja	Tak
Projekcja	Tak
Normalizacja próbki treningowej	Tak
Liczba warstw klasyfikatora połączeń	3
Liczba neuronów w warstwie klasyfikatora	256

Po zakończeniu tego procesu przystąpiono do szkolenia modelu. Trening prowadzony był przez 25 epok z uruchomionym mechanizmem wczesnego przerywania (po 3 epokach) w przypadku pogorszenia wyników oraz bieżącym zapisywaniem najlepszego modelu.

Tabela 6.4 podsumowuje uzyskane wyniki. Widać ewidentną poprawę zarówno dla metryki F1, jak i AUC dla zbioru walidacyjnego — przed treningiem i po nim. Ostateczny wynik dla zbioru testowego nie odbiega znacząco od wyników walidacyjnych.

Model wykazał znaczącą poprawę wyników w procesie uczenia, co jest widoczne we wszystkich analizowanych metrykach. W zadaniu predykcji istnienia lub braku krawędzi, opierającym się na rozróżnianiu próbek negatywnych od rzeczywistych, kluczowe znaczenie mają współczynniki trafień. Określają one, jak często spośród K predykcji o najwyższym prawdopodobieństwie rzeczywiste krawędzie uzyskują wyższe wskazania niż obiekty nieistniejące. W przeprowadzonym eksperymencie zauważalna jest istotna różnica między wskaźnikami trafień dla top 10 i top 20 predykcji dla zbioru testowego. Średnio 55% spośród 10 najwyżżej ocenionych predykcji poprawnie wskazuje istniejące połączenia. W przypadku rozszerzenia analizy do 20 najlepszych predykcji odsetek ten wzrasta do około 78%.

TABELA 6.4. Wyniki klasyfikacji połączeń w modelu interakcji leków

Metryka	Zbiór	Faza	
		Przed treningiem	Po treningu
F1	Walidacyjny	0,0	0,861
	Testowy	–	0,665
AUC	Walidacyjny	0,395	0,93
	Testowy	–	0,92
Śr. trafień w top 10	Walidacyjny	0,0	0,57
	Testowy	–	0,55
Śr. trafień w top 20	Walidacyjny	0,0	0,79
	Testowy	–	0,78

Dłuższy trening oraz bardziej agresywne strojenie hiperparametrów mogłyby istotnie poprawić wyniki w zakresie top 10.

Opisane podejście znajduje szerokie zastosowanie w systemach rekomendacyjnych, gdzie kluczowe jest określenie, jaka część proponowanych użytkownikowi sugestii rzeczywiście spełnia jego oczekiwania.

6.3. Klasyfikacja grafów

6.3.1. Postać formalna

Klasyfikacja całych grafów jest kluczowym zagadnieniem w analizie danych strukturalnych, znajdującym szerokie zastosowanie w wielu gałęziach nauki i biznesu. W przeciwieństwie do klasyfikacji wierzchołków czy krawędzi, gdzie predykcje dotyczą lokalnych komponentów grafu, klasyfikacja grafów koncentruje się na jego globalnych właściwościach.

Celem tego zadania jest przypisanie pojedynczej etykiety $y \in \mathcal{Y}$ do całego grafu $G = (V, E)$, gdzie $\mathcal{Y}$ oznacza zbiór możliwych klas. Proces ten można podzielić na dwa główne etapy:

1. **Konstrukcja reprezentacji grafu**, w której globalne właściwości strukturalne i atrybuty wierzchołków są transformowane do zwartych reprezentacji numerycznych.
2. **Klasyfikacja oparta na tych reprezentacjach**, gdzie model uczący dokonuje predykcji etykiety na podstawie osadzenia całego grafu.

Pierwszym krokiem jest wyznaczenie reprezentacji wierzchołków przy użyciu dowolnej grafowej sieci neuronowej (ang. *Graph Neural Network*, GNN). Proces ten jest analogiczny do tego stosowanego w klasyfikacji krawędzi i wierzchołków:

$$\begin{aligned} \mathbf{h}_v &= \text{GNN}(G, \mathbf{X}), v \in V, \\ \mathcal{H} &= \{ \mathbf{h}_v \mid \forall v \in V \}. \end{aligned}$$

Następnie uzyskane osadzenia wierzchołków są agregowane do jednej reprezentacji opisującej cały graf:

$$\mathbf{h}_G = \text{GlobalPool}(\{\mathbf{h}_v \mid \forall v \in V\}). \quad (51)$$

Funkcja $\text{GlobalPool}(\cdot)$ pełni rolę agregatora i może przyjmować różne formy, takie jak średnia, maksimum, suma czy mechanizmy bardziej zaawansowane, np. uwzględniające stopnie wierzchołków.

Zdefiniowana reprezentacja $\mathbf{h}_G$ stanowi wejście do modelu klasyfikacyjnego, który przypisuje grafowi etykietę:

$$\hat{y}_G = \text{GraphClf}(\mathbf{h}_G), \quad (52)$$

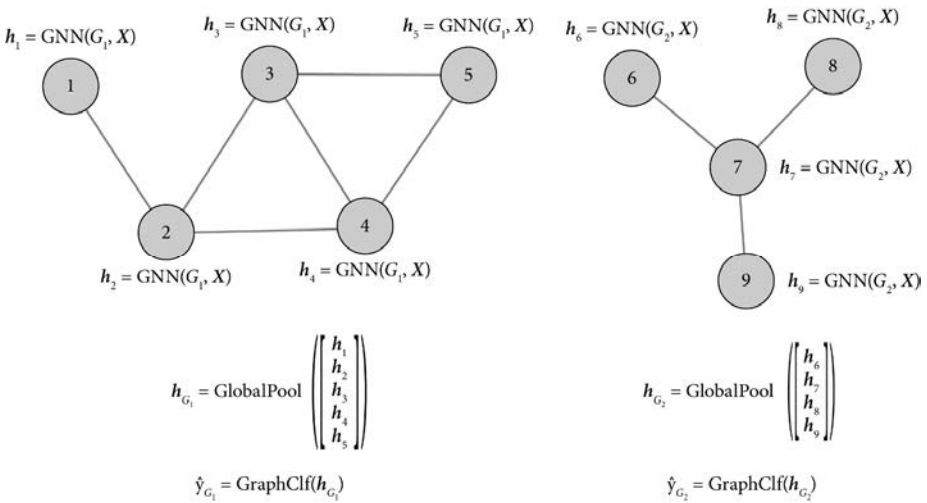
gdzie $\text{GraphClf}(\cdot)$ to dowolny klasyfikator, np. wielowarstwowa sieć neuronowa MLP. Proces optymalizacji przebiega poprzez minimalizację funkcji kosztu, np. entropii krzyżowej:

$$\mathcal{L} = \frac{1}{N} \sum_{i=1}^N \text{loss}(\text{GraphClf}(\mathbf{h}_{G_i}), y_i), \quad (53)$$

gdzie:

- $\text{loss}(\cdot, \cdot)$ to funkcja kosztu dla pojedynczego przykładu uczącego,
- $\mathbf{h}_{G_i}$ to wektor reprezentacji grafu G_i ,
- y_i to rzeczywista etykieta przypisana do grafu G_i ,
- N to liczba grafów w zbiorze uczącym.

Rysunek 6.3 ilustruje opisany proces.



RYСУNEK 6.3. Ilustracja procesu budowania reprezentacji całego grafu

W zadaniach klasyfikacji grafów często pracujemy z wieloma grafami, tj. zbiorem $\{G_1, G_2, \dots, G_N\}$, gdzie każdy graf $G_i = (V_i, E_i)$ ma własny zestaw wierzchołków V_i i krawędzi E_i . Aby umożliwić efektywne uczenie modeli, grafy te muszą zostać

scalone w jedną wspólną strukturę. PyG oferuje wyspecjalizowane mechanizmy umożliwiające ich efektywne połączenie, obejmujące:

1. Konkatenację macierzy cech wierzchołków:

$$X = \text{concat}(X_1, X_2, \dots, X_N), \quad (54)$$

gdzie $X_i \in \mathbb{R}^{|V_i| \times d}$ to macierz cech wierzchołków dla grafu G_i .

2. Łączenie macierzy sąsiedztwa w globalną strukturę:

$$A = \begin{bmatrix} A_1 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & 0 & A_N \end{bmatrix}, \quad (55)$$

gdzie każda macierz A_i jest macierzą sąsiedztwa w grafie G_i .

3. Dodanie identyfikatora przynależności każdego wierzchołka do konkretnego grafu:

$$\mathbf{batch} = \left[\underbrace{1, 1, \dots, 1}_{|V_1|}, \underbrace{2, 2, \dots, 2}_{|V_2|}, \dots, \underbrace{N, N, \dots, N}_{|V_N|} \right], \quad (56)$$

gdzie wartość w wektorze **batch** wskazuje numer grafu źródłowego dla danego wierzchołka.

Posługując się grafem z przywołanego wcześniej przykładu (patrz rysunek 6.3), możemy przyjąć, że macierze cech wierzchołków dla grafu 1 (wierzchołki 1 – 5) oraz grafu 2 (wierzchołki 6 – 9) wyglądają następująco:

$$X_1 = \begin{bmatrix} 1 & 1 \\ 2 & 2 \\ 3 & 3 \\ 4 & 4 \\ 5 & 5 \end{bmatrix},$$

$$X_2 = \begin{bmatrix} 6 & 6 \\ 7 & 7 \\ 8 & 8 \\ 9 & 9 \end{bmatrix}.$$

Dla uproszczenia przyjęto, że cechy wierzchołków to dwuelementowe wektory o wartościach odpowiadających indeksowi wierzchołka.

Lista sąsiedztwa dla obu grafów przedstawia się następująco:

$$A_1 = \begin{bmatrix} 1 & 2 & 2 & 3 & 3 \\ 2 & 3 & 4 & 4 & 5 \end{bmatrix},$$

$$A_2 = \begin{bmatrix} 6 & 8 & 9 \\ 7 & 7 & 7 \end{bmatrix}.$$

Mimo że oba grafy są nieskierowane, dla uproszczenia przedstawiono jednokierunkowe listy sąsiedztwa.

Podczas budowania próbki treningowej PyG skonstruowałby struktury zbliżone do przedstawionych poniżej, łącząc ze sobą grafy:

$$X = \begin{bmatrix} 1 & 1 \\ 2 & 2 \\ 3 & 3 \\ 4 & 4 \\ 5 & 5 \\ 6 & 6 \\ 7 & 7 \\ 8 & 8 \\ 9 & 9 \end{bmatrix}$$

$$A = \begin{bmatrix} 1 & 2 & 2 & 3 & 3 & 6 & 8 & 9 \\ 2 & 3 & 4 & 4 & 5 & 7 & 7 & 7 \end{bmatrix}$$

$$\text{batch} = [1 \ 1 \ 1 \ 1 \ 1 \ 2 \ 2 \ 2]$$

Indeksy w macierzy **batch** odpowiadają przypisaniu przykładów do oryginalnego grafu. W ten sposób nie trzeba w żaden sposób modyfikować standardowego kodu służącego do uczenia modeli sieci grafowych.

Fragment kodu poniżej ilustruje przykład skonstruowania opisywanej próbki (pamiętajmy o indeksowaniu numerów wierzchołków od zera).

```

1. g1 = pyg.data.Data()
2. g1.x = th.tensor([[i, i] for i in range(1, 6)], dtype=th.float)
3. g1.edge_index = th.tensor([
4.     [0, 1, 1, 2, 2, 3],
5.     [1, 2, 3, 3, 4, 4]
6. ])
7.
8. g2 = pyg.data.Data()
9. g2.x = th.tensor([[i, i] for i in range(6, 10)], dtype=th.float)
10. g2.edge_index = th.tensor([
11.     [0, 2, 3],
12.     [1, 1, 1]
13. ])
14.
15. small_data_loader = pyg.loader.DataLoader([g1, g2], batch_size=10)
16. batch = next(iter(small_data_loader))
17. batch
18. >> DataBatch(x=[9, 2], edge_index=[2, 9], batch=[9], ptr=[3])
19. batch.x
20. >> tensor([[1., 1.],
21. >>         [2., 2.],
22. >>         [3., 3.],
23. >>         [4., 4.],
24. >>         [5., 5.],
25. >>         [6., 6.],
26. >>         [7., 7.],
27. >>         [8., 8.],
28. >>         [9., 9.]])
29. batch.edge_index
30. >> tensor([[0, 1, 1, 2, 2, 3, 5, 7, 8],
31. >>         [1, 2, 3, 3, 4, 4, 6, 6, 6]])
32. batch.batch
33. >> tensor([0, 0, 0, 0, 0, 1, 1, 1, 1])

```


Skorowidz

A

agregacje wielokrotne, 178
algorytm
 mechanizmu Jumping Knowledge, 183
 MPNN, 53
 splotu GAT, 92
 splotu GIN, 109
 splotu GraphSAGE, 81, 84
 testu WL, 41
 uczenia sieci GNN, 163
analiza
 danych strukturalnych, 189, 199
 sieci społecznościowych, 5
 spektralna, 30
analizy molekularne, 5
architektury przestrzenne i czasowe, STGNNs, 51

B

badania toksyczności
 cząsteczek, 203
bazy danych grafowe, 20
biblioteka
 Deep Graph Library, 21
 graph-tool, 19
 iGraph, 19
 NetworkX, 19, 24
 Open Graph Benchmark, 151
 PyTorch Geometric, 20, 56
 SNAP, 19
 Stanford DeepSNAP, 141
bliskość, 35

C

cechy, features, 32, 36
centralność
 sąsiedztwa, closeness centrality, 35
 stopnia, degree centrality, 33
ChebNet, Chebyshev Spectral CNN, 75

D

dane
 testowe, 131, 163
 treningowe, 131, 163
 walidacyjne, 131, 163
Deep Graph Library, DGL, 21, 22
droga w grafie, 26
drzewo rekurencyjnego sąsiedztwa, 174
duże zbiory danych, 151
działania
 na poziomie całego grafu, 47
 na poziomie krawędzi, 48, 140
 na poziomie wierzchołków, 48

E

ekwiwariancja, 50
epoka, epoch, 152, 163

F

funkcja
 AGGREGATE, 53
 betweenness_centrality, 34

closeness_centrality, 35
degree_centrality, 33
degree_histogram, 37
MESSAGE, 52
UPDATE, 53

G

GAT, Graph Attention Network, 90
GCN, Graph Convolution Network, 5, 74–76
generowanie grafów, 48
GIN, Graph Isomorphism Network, 106, 172
głowa predykcyjna, prediction head, 60
GNN, Graph Neural Networks, 5, 47, 70
graf
 „Karate Club”, 132, 133, 138, 148–150
 Cora, 63–65
 MovieLens, 211
 MUTAG, 151
graflety, graphlets, 36
grafowe
 autoenkodery, GAEs, 51
 bazy danych, 20
 sieci neuronowe, GNN, 5, 47, 70, *Patrz także*
 warstwy splotu grafowego
 model MPNN, 51
 zadania, 47
 zasady działania, 49
 zastosowania, 189
grafy
 długość drogi, 26
 generowanie, 48

heterogeniczne, 43, 44
 reprezentacja, 45
 warstwy spłotu, 114
 klasyfikacja, 199
 krawędzi, 194
 wierzchołków, 189
 listy sąsiedztwa, 30
 macierze sąsiedztwa, 38
 nadmierne wygładzanie, 165
 nieskierowane, 25, 27
 niespójne, 26
 podgrafy, 26
 podział indukcyjny, 151
 problem izomorfizmu, 39
 regresja krawędzi, 208
 reprezentacja
 numeryczna, 36, 38
 skierowane, digrafy, 25, 27
 spójne, 26
 ścieżka, 26
 średnia długość
 najkrótszej ścieżki, 38
 średnica, 37
 wiedzy, KG, 6
 graph-tool, 19

H

heterogeniczne
 atrybuty, 45
 sieci informacyjne, HIN, 43
 hiperparametry, 192, 198, 205, 213

I

iGraph, 19
 iloczyn
 Hadamarda, 12
 skalarny, 12
 implementacja
 konwersji warstwy
 spłotu, 121
 modelu MPNN, 56
 naiwna JK, 183
 NGNN, 177, 179
 spłotu GAT, 99
 spłotu GCN, 76
 spłotu GIN, 110
 spłotu GNN, 72

splotu SAGE, 85
 testu WL, 41
 wbudowana JK, 185
 indeksy, 12
 indukcja, 134
 indukcyjny podział
 grafu, 151
 krawędzi, 141
 wierzchołków, 135
 inwariancja, 50
 izomorfizm grafów, 39

K

klasyfikacja, 47, 48, 61, 63
 całych grafów, 199
 badanie toksyczności
 cząsteczek, 203
 zastosowania, 203
 interakcji pomiędzy
 lekami, 196
 krawędzi, 194
 binarna, 208
 wieloklasowa, 208
 zastosowania, 195, 196, 207
 tematyczna stron, 191
 wierzchołków, 168, 189
 wieloklasowa, 191
 zastosowania, 190
 Knowledge Graphs, KG, 6
 konkatenacja, 13
 konwolucyjne/splotowe
 sieci grafowe, ConvGNNs, 51
 krawędzie, 140, 191, 196
 negatywne, negative samples, 140
 podział indukcyjny, 141
 podział transdukcyjny, 145–148
 przekazujące
 wiadomości, message passing links, 140, 145
 uczące, supervision links, 140, 145

L

listy sąsiedztwa, 30

M

macierz, 11
 macierze sąsiedztwa, 28
 mechanizm połączeń
 skokowych, JK, 180
 metaścieżki, meta-path, 44
 miary centralności
 wierzchołków, 33, 36
 MLP, multilayer perceptron, 49
 model przekazywania
 wiadomości, MPNN, 51, 178, 203
 implementacja, 56
 jako część sieci, 60
 zasady działania, 54
 zastosowania, 61
 MPNN, Message Passing
 Neural Network, 51, 178, 203

N

NetworkX, 19, 24
 NGNN, network in graph
 neural network, 175, 177, 179, 187

O

odkrywanie leków, 5, 203
 optymalizacja struktur, 6

P

perceptron
 wielowarstwowy, MLP, 49
 podgrafy, 26
 pośrednictwo, betweenness centrality, 33
 predykcja, 168, 175, 209
 problem nadmiernego
 wygładzania, oversmoothing, 62, 164
 projektowanie nowych
 materiałów, 203
 próbki uczące, mini-batch, 152, 163
 próbkowanie
 negatywne, 195
 w oparciu o sąsiedztwo, 152
 w oparciu o społeczność, 160

przewidywanie
funkcji białek, 203
właściwości molekuł,
203
PyTorch Geometric, PyG,
20, 22, 56

R

regresja, 47, 48, 61
krawędzi, 208
rekomendacje filmów
MovieLens, 210
rekurencyjne
grafy sąsiedztwa, 173
sieci grafowe, RecGNNs,
51
repozytorium Stanford
SNAP, 191
rozkład częstości stopni
wierzchołków, 37

S

SAGE, 80
sąsiedztwo, 26, 152
rekurencyjne, 174
sieci
MPNN, 60
spektralne, Spectral
Graph Convolution, 51
w sieci grafowej, NGNN,
175
skalar, 11
SNAP, Stanford Network
Analysis Platform, 19
splot

GAT, 90
działanie, 90
implementacja, 99
GCN, 74
działanie, 75
implementacja, 76
GIN, 106
działanie, 106
implementacja, 110
GNN, 70
działanie, 70
implementacja, 72

SAGE, 80
działanie, 81
implementacja, 85
z mechanizmem
NGNN, 187
z mechanizmem
połączeń
skokowych, 187
z wieloma
agregacjami, 187
splotowa sieć grafowa, GCN,
5, 74–76
sploty grafu, graph
convolutions, 53
społeczności, 160
spójność grafu, 26
Stanford
DeepSNAP, 141
SNAP, 191
stopień
średni wierzchołka, 36
wierzchołka, 29, 32
systemy rekomendacyjne, 6,
207, 209

T

test Weisfeilera-Lehmana,
39, 171
implementacja, 41
transdukcja, 134
transdukcyjny podział
krawędzi, 145–148
wierzchołków, 137

U

uczenie na dużych zbiorach,
151
próbkiowanie w oparciu
o sąsiedztwo, 152
o społeczności, 160

W

walidacja krzyżowa,
cross-validation, 131
warstwy
splotu dla grafów
heterogenicznych, 114

splotu grafowego, graph
convolution networks,
53, 60, 67, 125
agregacje
wielokrotne, 178
GAT, 90
GCN, 74
GIN, 106
GNN, 70
mechanizm połączeń
skokowych, 180
ograniczenie testem
WL, 171
podejście NGNN, 175
problem nadmiernego
wygładzania, 164
SAGE, 80

wektor, 11
wierzchołki, 25, 191, 196
miary centralności, 33, 36
podział indukcyjny, 135
podział transdukcyjny,
137
rozkład częstości stopni,
36
stopnie, 29, 32
stopnie średnie, 36
wykrywanie
botów, 5
interakcji między lekami,
195

Z

zbiór
danych Iris, 132
danych MovieLens, 208,
210
danych MUTAG, 203
elementów, 12
testowy, 163
treningowy, 163
walidacyjny, 163

PROGRAM PARTNERSKI

— GRUPY HELION —

- 
1. ZAREJESTRUJ SIĘ
 2. PREZENTUJ KSIĄŻKI
 3. ZBIERAJ PROWIZJĘ

Zmień swoją stronę WWW w działający bankomat!

Dowiedz się więcej i dołącz już dzisiaj!

<http://program-partnerski.helion.pl>

GRUPA
Helion

Cicha rewolucja, która nadeszła

Grafowe sieci neuronowe (GNN, ang. *graph neural networks*) to klasa modeli uczenia głębokiego przeznaczona do analizy danych o strukturze grafowej. W początkowym okresie ich rozwój był ograniczony przez brak efektywnych metod projektowania i optymalizacji; w ostatnich latach bariery te w dużej mierze zostały pokonane, co przełożyło się na dynamiczny postęp teorii i praktyki. Modele GNN znajdują zastosowanie między innymi w analizie sieci społecznościowych, optymalizacji procesów logistycznych, marketingu i pracy z bazami wiedzy.

Ta książka zawiera kompleksowe opracowanie tematyki sieci grafowych w kontekście uczenia maszynowego. Tym samym wypełnia istotną lukę na polskim rynku wydawniczym, oferując połączenie solidnych podstaw teoretycznych z praktycznym zastosowaniem GNN. To przewodnik, który systematycznie przeprowadza przez kolejne zagadnienia związane z sieciami grafowymi:

- od narzędzi klasycznej analizy grafów w środowisku Pythona i wybranych zagadnień teorii grafów
- przez wprowadzenie do grafowych sieci neuronowych, a także przegląd wybranych warstw spłotu grafowego i dobrych praktyk ich projektowania
- po zagadnienia związane ze szkoleniem sieci GNN i praktyczne przykłady ich zastosowań

Dr inż. Filip Wójcik

— analityk danych, zajmuje się systemami uczenia maszynowego i sztucznej inteligencji od 2012 roku. Projektował i wdrażał rozwiązania ML/AI dla międzynarodowych i polskich przedsiębiorstw w sektorach finansowym, telekomunikacyjnym i logistycznym, pracując jako analityk danych (data scientist), kierownik zespołu badawczego i lider techniczny. Łączy działalność komercyjną z pracą akademicką; od 2017 roku związany z Uniwersytetem Ekonomicznym we Wrocławiu, gdzie specjalizuje się w optymalizacji procesów decyzyjnych z wykorzystaniem modeli uczenia maszynowego, w szczególności głębokich sieci neuronowych. Chętnie dzieli się wiedzą i doświadczeniem ze słuchaczami w ramach studiów podyplomowych, kursów, a także Biznesowego Indywidualnego Programu Studiów (BIPS). Autor licznych recenzowanych publikacji naukowych, prelegent międzynarodowych konferencji poświęconych najnowszym trendom w dziedzinie uczenia maszynowego.

Wydawnictwo
Naukowe
Helion

helion.pl

HELION S.A.
ul. Kościuszki 1c
44-100 Gliwice
tel.: 32 230 98 63
helion@helion.pl

KOD KORZYŚCI
Sięgnij po więcej!



ISBN 978-83-289-3390-3



9 788328 933903

Cena: 79,00 zł