

А. Б. Качинський, Д. В. Ланде, О. М. Новіков

ТЕОРЕТИКО-ГРАФОВІ МОДЕЛІ ІНФОРМАЦІЙНО-ПСИХОЛОГІЧНИХ ВІЙН

МОДЕЛЮВАННЯ МЕНТАЛЬНОЇ ВІЙНИ ІЗ ЗАСТОСУВАННЯМ
ГЕНЕРАТИВНОГО ШТУЧНОГО ІНТЕЛЕКТУ



Київ

КПІ ім. Ігоря Сікорського

2025

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ УКРАЇНИ
«КИЇВСЬКИЙ ПОЛІТЕХНІЧНИЙ ІНСТИТУТ ІМЕНІ ІГОРЯ СІКОРСЬКОГО»

А. Б. Качинський, Д. В. Ланде, О. М. Новіков

ТЕОРЕТИКО-ГРАФОВІ МОДЕЛІ ІНФОРМАЦІЙНО-ПСИХОЛОГІЧНИХ ВІЙН

**Моделювання ментальної війни із застосуванням
генеративного штучного інтелекту**

Навчальний посібник

Рекомендовано Методичною радою КПІ ім. Ігоря Сікорського
як навчальний посібник для здобувачів ступенів магістра та PhD
за спеціальністю 125 Кібербезпека та захист інформації

Київ
КПІ ім. ІГОРЯ СІКОРСЬКОГО
2025

УДК 303.732.4: 004.056

K19

Автори: *Качинський Анатолій Броніславович, д-р техн. наук, проф.
Ланде Дмитро Володимирович, д-р техн. наук, проф.
Новіков Олексій Миколайович, д-р техн. наук, проф.*

Рецензент *Данилов В. Я., д-р техн. наук, проф.*

Відповідальний редактор *Апішева А. Ш., канд. психол. наук*

*Гриф надано Методичною радою КПІ ім. Ігоря Сікорського
(протокол № 7 від 08.05.2025 р.)
за поданням вченої ради навчально-наукового фізико-технічного інституту
(протокол № 8 від 05.05.2025 р.)*

Качинський А. Б.
K19 Теоретико-графові моделі інформаційно-психологічних війн. Моделювання ментальної війни із застосуванням генеративного штучного інтелекту : навч. посіб. для здобувачів ступенів магістра та PhD за спец. 125 Кібербезпека та захист інформації / А. Б. Качинський, Д. В. Ланде, О. М. Новіков ; КПІ ім. Ігоря Сікорського. — Київ : КПІ ім. Ігоря Сікорського, 2025. — 242 с.

Навчальний посібник присвячений сучасним методам аналізу та моделювання інформаційно-психологічних війн на основі графової теорії, генеративного штучного інтелекту та мережевого аналізу. Видання орієнтоване на здобувачів ступенів магістра та PhD за спеціальністю 125 Кібербезпека та захист інформації, викладачів ЗВО, а також на аналітиків і науковців у сфері інформаційної та кібернетичної безпеки, OSINT, працівників ЗМІ.

УДК 303.732.4: 004.056

Реєстр. № НП 24/25-523. Обсяг 10,25 авт. арк.

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
проспект Берестейський, 37, м. Київ, 03056
<https://kpi.ua>

Свідectво про внесення до Державного реєстру видавців, виготовлювачів і розповсюджувачів видавничої продукції ДК № 5354 від 25.05.2017 р.

© А. Б. Качинський, Д. В. Ланде, О. М. Новіков, 2025
© КПІ ім. Ігоря Сікорського, 2025

Передмова

Навчальний посібник «Теоретико-графові моделі інформаційно-психологічних війн. Моделювання ментальної війни із застосуванням генеративного штучного інтелекту» присвячено сучасним методам аналізу та моделювання інформаційно-психологічних війн на основі графової теорії, генеративного штучного інтелекту (LLM) та мережевого аналізу.

Розглядаються основи теорії графів, включно з їхніми властивостями, типами та математичним апаратом, необхідним для моделювання складних інформаційних процесів. Особливу увагу приділено показникам центральності (degree, betweenness, closeness, eigenvector centrality тощо) як ключовим характеристикам впливовості вузлів у мережі, зокрема їхній ролі у поширенні інформації, формуванні наративів та виявленні прихованих лідерів інформаційного впливу.

Значну увагу приділено великим мовним моделям як інструменту генерації концептів та зв'язків у мережах. Описано методологію "рою віртуальних експертів", що дозволяє семантично збагачувати графові моделі, формувати зв'язки між концептами та адаптивно розширювати інформаційні структури.

Розглянуто методи побудови мережевих моделей, включаючи перехід від ієрархічних структур до складних динамічних мереж. Описано методи кластеризації, ранжування вузлів та візуалізації мережевих моделей в Gephi.

Окремий розділ присвячено моделюванню інформаційно-психологічних війн, аналізу їхньої структури, типових механізмів впливу та маніпуляцій. Зокрема, розглядаються вузли-посередники як ключові агенти інформаційних атак та операцій з дезінформації.

Наведено практичні кейси, що демонструють виявлення центрів впливу в інформаційних операціях, аналіз механізмів когнітивного тиску та прогнозування інформаційних атак.

Посібник містить лабораторний практикум із прикладами на основі промпт-інжинірингу, шаблонами документації та практичними завданнями для застосування методів на реальних кейсах.

Видання орієнтоване на студентів магістратури, аспірантів, викладачів ЗВО, а також на аналітиків і науковців у сфері інформаційної та кібернетичної безпеки, OSINT, працівників ЗМІ.

Зміст

Вступ	8
1. Теоретичні основи графових моделей	11
1.1. Основні поняття теорії графів	11
1.2. Властивості та типи графів.....	12
1.3. Математичні основи графового моделювання	14
1.4. Практична частина. Побудова простих графів.....	20
Висновки до розділу	25
Питання для самоконтролю	26
Література за розділом	27
2. Великі мовні моделі та їх застосування	28
2.1. Поняття LLM.....	28
2.2. Архітектура «Трансформер»	29
2.3. Модель шести сутностей	39
2.4. Генерація концептів та зв'язків.....	43
2.5. Методологія «рою віртуальних експертів»	47
2.4. Практична частина. Використання LLM для семантичного аналізу текстів.....	59
Висновки до розділу	63
Питання для самоконтролю	66
Література за розділом	67
3. Методи побудови мережових моделей	68
3.1. Інтеграція нових понять і зв'язків.....	68
3.2. Приклади із інформаційно-психологічних війн	69
3.3. Виклики та обмеження застосування великих мовних моделей ..	70
3.4. Розширення ієрархічних структур до мереж.....	71
3.5. Семантичне розширення мережових моделей.....	76
3.6. Практична частина. Формування семантичних мереж на основі аналізу текстів.....	82
Висновки до розділу	89

Питання для самоконтролю	90
Література за розділом	92
4. Аналіз та візуалізація даних.....	93
4.1. Кластеризація мереж.....	93
4.2. Ранжирування вузлів мережі.....	95
4.3. Візуалізація мереж	96
4.4. Практична частина. Завантаження, аналіз і візуалізація мережевих структур в Gephi	97
Висновки до розділу	101
Питання для самоконтролю	102
Література за розділом	103
5. Ознаки центральності у графах	104
5.1. Degree Centrality (Ступенева центральність).....	104
5.2. Closeness Centrality (Близькість до інших вузлів)	106
5.3. Eigenvector Centrality (Центральність власного вектора)	108
5.4. Центральність посередництва (betweenness centrality)	110
5.5. Виявлення впливів в мережах	113
5.5.1. Семантична модель.....	113
5.5.2. Модель перколяції	116
5.5.3. Модель розповсюдження інформації: SIR-модель	118
5.6. Приклади застосування у реальних ситуаціях	121
5.7. Практична частина. Розрахунок центральності	122
Висновки до розділу	126
Питання для самоконтролю	128
Література за розділом	128
6. Моделювання інформаційно-психологічних війн	130
6.1. Структура й особливості інформаційно-психологічної війни....	130
6.2. Практична реалізація результатів моделювання	141
6.3. Аналіз результатів.....	149

6.4. Практична частина. Аналіз достовірності повідомлень в інтернет через мережевий нетворкінг	153
Висновки до розділу	157
Питання для самоконтролю	158
Література за розділом	159
7. Динамічні аспекти інформаційно-психологічних війн.....	162
7.1. Основні характеристики динамічних мереж ІПВ	163
7.2. Методологія дослідження динаміки	164
7.3. Еволюція мереж у часі.....	164
7.4. Моделі еволюції мереж	166
7.5. Моделі розповсюдження інформації	167
7.6. Приклади еволюції мереж у реальних ситуаціях	172
7.7. Виявлення атак.....	173
7.8. Математична формалізація атак.....	174
7.9. Стратегії контрзаходів	174
7.10. Прогнозування інформаційних атак	177
7.11. Практична частина. Моделювання динаміки інформаційних графів.....	184
Висновки до розділу	188
Питання для самоконтролю	190
Література за розділом	191
8. Автоматизація аналізу та прогнозування	192
8.1. Використання ML та AI для аналізу графів	192
8.2. Graph Neural Networks	194
8.3. Node Embedding	198
8.4. Recurrent Neural Networks (RNN) та LSTM	205
8.5. Random Forests та Gradient Boosting	209
8.6. Приклади застосування ML та AI у графових моделях ІПВ	214
8.7. Розробка автоматизованих аналітичних інструментів.....	216

8.8. Практична частина. Побудова моделі прогнозування змін у мережі.....	219
Висновки до розділу	223
Питання для самоконтролю	225
Література за розділом	226
9. Практичні кейси та застосування	227
9.1. Аналіз реальних випадків ІПВ.....	227
9.2. Практичні сценарії для досліджень	230
9.3. Використання методів у кібербезпеці	231
9.4. Приклади застосування графових моделей у кібербезпеці	235
9.5. Практична частина. Проведення самостійного дослідження...	236
Висновки до розділу	237
Питання для самоконтролю	239
Література за розділом	240
Висновки	241

Вступ

Сучасний світ переживає безпрецедентне зростання інформаційних потоків, що створює нові виклики для національної безпеки, суспільства та окремих громадян. Особливо гостро ця проблема постає у контексті інформаційно-психологічних війн, які стали одним із ключових інструментів впливу на свідомість людей, формування суспільної думки та нав'язування певних політичних, соціальних або культурних наративів. Наш ворог, росія, зокрема, активно використовує методи ментальної війни проти України, спрямовані на дестабілізацію внутрішньої ситуації, руйнування національної ідентичності та маніпулювання масовою свідомістю. Маніпулятивні технології, дезінформація та інформаційні атаки активно використовуються в медіа, соціальних мережах і кіберпросторі.

Для ефективного протистояння таким викликам необхідні методи аналізу, що дозволяють ідентифікувати інформаційні впливи, оцінювати їхню структуру та прогнозувати розвиток подій.

У таких умовах зростає потреба у фахівцях, які володіють методиками виявлення та аналізу інформаційно-психологічних ментальних впливів, можуть ідентифікувати ключові вузли поширення дезінформації та оцінювати наслідки таких кампаній.

Навчальний посібник «Теоретико-графові моделі інформаційно-психологічних війн. Моделювання ментальної війни із застосуванням генеративного штучного інтелекту» спрямований на формування у читачів практичних навичок аналізу інформаційного середовища та моделювання впливів із використанням сучасних методів мережевого аналізу, штучного інтелекту та алгоритмічного прогнозування. Він орієнтований на підготовку фахівців у галузі інформаційної безпеки, OSINT, кібербезпеки, соціального моделювання та когнітивних наук, які зможуть використовувати сучасні методи аналізу інформаційних процесів для виявлення загроз та оцінки ризиків.

Метою навчального курсу є формування у слухачів глибоких теоретичних знань та практичних навичок у галузі аналізу інформаційно-психологічних війн. Цей курс спрямований на підготовку фахівців, які зможуть ефективно протистояти сучасним інформаційним загрозам, використовуючи передові методи графового аналізу, штучного інтелекту та кіберрозвідки.

Для досягнення цієї мети визначено наступні конкретні завдання:

1. Вивчення теоретичних основ інформаційно-психологічних війн: аналіз історичного розвитку поняття «інформаційна війна», її типів та особливостей.
2. Опанування методами графового моделювання: навчання студентів виявляти ключові вузли впливу, аналізувати структуру інформаційних мереж та прогнозувати динаміку інформаційних атак.
3. Застосування великих мовних моделей: вивчення можливостей штучного інтелекту для аналізу текстових масивів, ідентифікації дезінформації та генерації контрнарративів.
4. Розробка стратегій протидії інформаційним загрозам: формування навичок моделювання сценаріїв інформаційних атак та розробки ефективних контрзаходів.
5. Практична підготовка фахівців: проведення аналізу реальних випадків ІПВ, розробка практичних рекомендацій та застосування отриманих знань у сфері кібербезпеки та OSINT.

Курс поєднує теоретичні та практичні заняття, що охоплюють основи інформаційно-психологічного впливу, методи графового аналізу, роботу з великими мовними моделями (LLM) та застосування кіберрозвідки для оцінки інформаційних загроз. Він побудований за модульним принципом і складається з дев'яти основних розділів, кожен з яких включає теоретичний матеріал, практичні заняття та самостійну роботу студентів. До основних напрямків курсу відносяться:

1. Теоретичні основи графових моделей: вивчення базових понять теорії графів, властивостей та типів графів, математичних основ моделювання.
2. Великі мовні моделі та їх застосування: аналіз принципів роботи LLM, генерація концептів та зв'язків, методологія «рою віртуальних експертів».
3. Методи побудови мережевих моделей: вивчення способів формування семантичних мереж, аналіз їх структури та динаміки.
4. Аналіз та візуалізація даних: застосування методів кластеризації, ранжування вузлів та візуалізації інформаційних потоків.
5. Ознаки центральності у графах: вивчення показників центральності (degree, closeness, eigenvector, тощо), виявлення ключових вузлів впливу.
6. Моделювання інформаційно-психологічних війн: аналіз структури ІПО, практичні приклади моделювання та аналіз результатів.
7. Динамічні аспекти інформаційних війн: вивчення еволюції мереж у часі, виявлення атак та розробка стратегій контрзаходів.

8. Автоматизація аналізу та прогнозування: застосування машинного навчання та штучного інтелекту для передбачення інформаційних атак.
9. Практичні кейси та застосування: аналіз реальних випадків ІПВ, розробка практичних сценаріїв та проведення самостійних досліджень.

До очікуваних результатів від вивчення курсу можна віднести:

- теоретичні знання, а саме, здобувачі освіти отримають глибоке розуміння механізмів інформаційно-психологічних війн, їх структури та методів впливу;
- практичні вміння та навички застосовувати графовий аналіз, великі мовні моделі та інструменти кіберрозвідки для аналізу інформаційних загроз;
- аналітичні можливості, у тому числі, здатність виявляти ключові вузли впливу, прогнозувати динаміку інформаційних атак та розробляти ефективні контрзаходи;
- проектні компетенції здобувачів, що включають проведення самостійних досліджень, аналіз реальних ІПВ та розробку практичних рекомендацій.

По завершенню курсу здобувачі зможуть:

- аналізувати інформаційні атаки та розпізнавати механізми психологічного впливу;
- визначати ключові інформаційні вузли та агентів впливу за допомогою мережевого аналізу;
- використовувати інструменти штучного інтелекту для аналізу текстів, побудови семантичних карт та виявлення прихованих патернів дезінформації;
- моделювати розвиток інформаційних атак і розробляти контрзаходи;
- застосовувати засоби OSINT-аналізу інформаційного середовища для прогнозування загроз та оцінки ризиків.

Курс орієнтовано на студентів магістратури, аспірантів, викладачів ЗВО а також на аналітиків і науковців у сфері інформаційної та кібернетичної безпеки, OSINT, працівників ЗМІ. Він є важливим кроком у підготовці фахівців, здатних ефективно протистояти сучасним інформаційним загрозам, захищати національну ідентичність та забезпечувати інформаційну безпеку держави.

1. Теоретичні основи графових моделей

Методи графового аналізу є ключовим інструментом для дослідження інформаційно-психологічних впливів, операцій, війн, зокрема ментальної війни, оскільки дозволяють описувати складні мережі взаємозв'язків між джерелами інформації, поширювачами та реципієнтами. Використовуючи теорію графів, можна моделювати процеси поширення дезінформації, виявляти найбільш впливові вузли та оцінювати ефективність інформаційних атак.

У цьому розділі розглядаються основи теорії графів, властивості графів та математичний апарат, що дозволяє будувати та аналізувати інформаційні мережі.

1.1. Основні поняття теорії графів

Граф – це математична структура, що складається з множини вузлів (вершин) та множини зв'язків (ребер) між ними. Він може використовуватися для моделювання взаємодій між об'єктами, що мають певні відносини.

До основних понять теорії графів відносяться:

- вершина (вузол, node) – основний елемент графа, що представляє сутність (наприклад, користувача соціальної мережі, ЗМІ, бота);
- ребро (зв'язок, edge) – з'єднання між двома вершинами, яке відображає взаємодію між об'єктами (наприклад, ретвіт, поширення інформації);
- ступінь вершини (degree) – кількість зв'язків, що виходять з вершини;
- компоненти зв'язності (connected components) – підграфи, у яких будь-які дві вершини з'єднані шляхом.

Графи можуть класифікуватися за різноманітними ознаками, які визначають їхню структуру та властивості. Нижче наведено найпоширеніші види графів.

За типом зв'язків між вершинами графи поділяються на орієнтовані, неорієнтовані та змішані. Орієнтований граф (або орграф) – це граф, у якому ребра мають напрямок і називаються дугами. У неорієнтованому графі ребра не мають напрямку. Змішаний граф містить як орієнтовані, так і неорієнтовані ребра.

За наявністю додаткових елементів графи бувають мультиграфами, псевдографами або зваженими. Мультиграф допускає кратні ребра, тобто кілька ребер між однією парою вершин. Псевдограф, окрім

кратних ребер, може мати петлі – ребра, які починаються і закінчуються в одній і тій самій вершині. У зваженому графі кожному ребру або вершині присвоюється певна вага, яка може символізувати відстань, вартість чи інші параметри.

За структурою графа розрізняють повні, порожні, двочасткові, деревоподібні, лісові, циклічні та ациклічні графи. Повний граф – це такий граф, у якому кожна пара вершин з'єднана ребром. Порожній граф не містить жодного ребра. Двочастковий граф характеризується тим, що його вершини можна поділити на дві множини таким чином, що всі ребра з'єднують вершини з різних множин. Дерево – це зв'язний граф без циклів, а ліс – сукупність дерев, тобто граф без циклів, який може бути незв'язним. Циклічний граф містить хоча б один цикл, тоді як ациклічний граф таких циклів не має.

За особливостями впорядкування графи можуть бути планарними або непланарними. Планарний граф можна намалювати на площині так, щоб його ребра не перетиналися. У реальних інформаційних мережах більшість графів непланарні, тому ця властивість графів для рішення задач інформаційно-психологічних війн не досліджується.

За числом компонент зв'язності графи поділяються на зв'язні та незв'язні. У зв'язному графі будь-які дві вершини можна з'єднати шляхом, тоді як незв'язний граф складається з кількох окремих частин, які називаються компонентами зв'язності.

За регулярністю графи бувають регулярними та нерегулярними. Регулярний граф має однаковий степінь для всіх своїх вершин, тоді як в нерегулярному графі степені вершин можуть відрізнятися.

У контексті інформаційно-психологічних ментальних війн графова модель дозволяє виявляти центри впливу, поширення наративів та оцінювати ефективність комунікаційних стратегій.

1.2. Властивості та типи графів

Моделювання інформаційно-психологічних війн передбачає аналіз структури та динаміки інформаційних потоків, маніпуляцій і поширення впливу через мережі.

Використання графових моделей дає змогу відстежувати взаємодію акторів, аналізувати стратегії маніпуляції та прогнозувати розвиток інформаційних кампаній. Нижче подано класифікацію основних типів графів, які можуть бути використані для цього аналізу.

Графи можуть мати різні властивості, що визначають їхню структуру та можливості аналізу, наприклад, щільність графа (density) як відношення кількості наявних зв'язків до максимально можливої

кількості. Щільні графи можуть свідчити про високу взаємодію між агентами інформаційного впливу. Важливою ознакою для вивчення графів у аспекті ІПВ є центральність (centrality) вершин як міра їх впливовості в графі, яка може визначати, наприклад, ключові джерела інформації або агентів дезінформації.

Крім того, в рамках дослідження ІПВ розглядаються окреми різновиди графів, серед яких найпоширеніші:

1. Випадкові графи Ердеша-Реньї (Erdős–Rényi). У таких графах кожне ребро між парами вершин створюється з певною фіксованою ймовірністю. Для них властива відсутність чіткої структури. Такі графи застосовуються для моделювання неструктурованих комунікацій, наприклад, хаотичного поширення інформації у випадкових інформаційних середовищах.

2. Масштабно-інваріантні графи (scale-free networks), а саме, графи, де розподіл ступенів вершин відповідає степеневому закону, при якому більшість вузлів мають мало зв'язків, а деякі – дуже багато, спостерігається наявність хабів (вузлів із високим ступенем), що виконують роль центрів впливу. Ці графи характерні для соціальних мереж, де основні актори (лідери думок, медіа-ресурси) відіграють ключову роль у розповсюдженні інформації.

3. Графи «малого світу» (small-world networks), тобто такі графи, що поєднують високу кластеризацію та малі середні відстані між вузлами. Наприклад, локальні групи акторів (діючих осіб, фігурантів) тісно взаємодіють, але між групами існують короткі зв'язки, що забезпечують швидке поширення інформації. Такі графи моделюють реальні мережі комунікацій, такі як внутрішні спільноти та інформаційні бульбашки.

4. Динамічні графи – графові структури, які змінюються з часом (змінюються вершини, зв'язки, ваги ребер). Вони відображають еволюцію інформаційного простору і широко використовуються для аналізу змін у поширенні інформації та впливу під час інформаційно-психологічних операцій.

5. Гіперграфи – це узагальнення «звичайних» графів, охоплює структури, у яких ребро може з'єднувати більше ніж дві вершини. Вони дозволяють моделювати складні взаємозв'язки між багатьма агентами. В галузі інформаційної безпеки застосовуються для моделювання одночасних координаційних атак або взаємодії багатьох інформаційних каналів.

6. Мультиграфи, тобто такі графи, у яких між двома вершинами може існувати більше одного ребра. Вони відображають паралельні

комунікаційні канали та можуть аналізувати різні шляхи маніпуляцій, таких як одночасне поширення правдивої та фейкової інформації.

Графові моделі є сучасним інструментом аналізу інформаційно-психологічних війн. Вибір типу графа залежить від особливостей комунікаційної мережі та поставленої задачі: від випадкових взаємодій до ієрархічного контролю над інформаційними потоками.

Комбінація різних типів графів дозволяє створити складніші моделі для розуміння та прогнозування інформаційних впливів у сучасному інформаційному середовищі.

1.3. Математичні основи графового моделювання

Графове моделювання є засобом аналізу структурних властивостей мереж, зокрема інформаційних систем. Для формального опису та аналізу графових структур використовуються різні математичні підходи, які дозволяють вивчати властивості цілої мережі, окремих вузлів та ребер.

Граф G можна формально визначити як пару множин $G = (V, E)$, де V – множина вершин (вузлів), E – множина ребер (зв'язків між вершинами).

Ребра можуть бути:

- направлені (орієнтовані): ребро (u, v) вказує на зв'язок від вершини u до вершини v ;
- ненаправлені: ребро (u, v) вказує на взаємний зв'язок між вершинами u і v .

Графи також можуть мати ваги на ребрах, що визначають силу зв'язку між вершинами. У цьому випадку граф називається зваженим.

Для практичного аналізу графів використовуються різні способи їх представлення, зокрема, *матриці суміжності, інцидентності, списки суміжності*.

Матриця суміжності A є квадратною матрицею розміру $n \times n$, де n – кількість вершин у графі. Елементи матриці визначаються так:

$$A_{ij} = \begin{cases} w_{ij}, & \text{якщо існує ребро між вершинами } i \text{ та } j \text{ з вагою } w_{ij}; \\ 0, & \text{якщо ребра немає.} \end{cases}$$

Для ненаправлених графів матриця A є симетричною ($A_{ij} = A_{ji}$).

Матриця інцидентності I має розмір $n \times t$, де n – кількість вершин, t – кількість ребер. Елементи матриці визначаються так:

$$I_{ij} = \begin{cases} 1_{ij}, & \text{вершина } i \text{ пов'язана з ребром } j; \\ 0, & \text{якщо вершина } i \text{ не пов'язана з ребром } j. \end{cases}$$

Для направленої графу значення I_{ij} може бути -1 , 0 або $+1$, що вказує на напрямок ребра.

Список суміжності представляє граф як набір списків, де кожен список містить вершини, суміжні з даною вершиною. Цей метод ефективний для розріджених графів (графів з невеликою кількістю ребер порівняно з максимально можливою).

Властивості графів

До основних властивостей всього графа відносяться:

1. Порядок графа – кількість вершин $|V|$.
2. Розмір графа – кількість ребер $|E|$.
3. Густина графа – для ненаправлених графів відношення кількості існуючих ребер до максимально можливої кількості ребер:

$$D = \frac{2 \cdot |E|}{|V| \cdot (|V| - 1)}.$$

4. Діаметр графа – максимальна довжина найкоротшого шляху між будь-якими двома вершинами.
5. Середня довжина шляху – середнє значення найкоротших шляхів між усіма парами вершин:

$$L = \frac{1}{n \cdot (n - 1)} \sum_{i \neq j} d(i, j),$$

де $d(i, j)$ – довжина найкоротшого шляху між вершинами i та j .

До основних властивостей окремих вершин графа відносяться:

1. Ступінь вершини: кількість ребер, що виходять з вершини v :

$$k(v) = \sum_{i=1}^n A_{vi}.$$

Для направленої графу розрізняються вхідний ступінь – кількість ребер, що входять у вершину, і вихідний ступінь – кількість ребер, що виходять з вершини.

2. Коефіцієнт кластеризації – показує, наскільки сусіди вершини з'єднані між собою. Для вершини v :

$$C(v) = \frac{2}{k(v) \cdot (k(v) - 1)} T(v),$$

де $T(v)$ – кількість трикутників, що містять вершину v .

3. Центральність вершини. Існую декілька десятків різних критеріїв центральності вершин, серед яких при дослідженні ІПВ найчастіше розглядаються:

- ступінь центральності – пропорційна ступеню вершини;
- посередництво (betweenness centrality) – вимірює частоту, з якою вершина знаходиться на найкоротшому шляху між іншими вершинами;
- близькість (closeness centrality) – обернено пропорційна сумарній відстані до всіх інших вершин.

Розподіл ступенів вершин графу $P(k)$ визначає ймовірність того, що випадково обрана вершина має ступінь k . Це одна з ключових характеристик графа, яка впливає на його структуру та поведінку.

При дослідженні ІПВ розглядаються такі основні типи розподілу ступенів вузлів як рівномірний, експоненціальний, степеневий, Гаусів.

Рівномірний розподіл: всі вершини мають приблизно однаковий ступінь.

$$P(k) = \text{const.}$$

Експоненціальний розподіл: ступені вершин зменшуються експоненціально:

$$P(k) \propto e^{-\lambda k},$$

де $P(k)$ – ймовірність того, що вершина має ступінь k , а $\lambda > 0$ – параметр, що визначає швидкість спадання. Експоненціальні розподіли характерні для випадкових графів.

Степеневий розподіл: характерний для безмасштабних мереж:

$$P(k) \propto k^{-\gamma},$$

де γ – параметр, що визначає крутизну розподілу. Зазвичай $2 < \gamma < 3$.

Мережі зі степеневим розподілом також називають безмасштабними. Такі графові структури (наприклад, соціальні мережі) демонструють степеневий розподіл, що свідчить про наявність "гігантських хабів".

Гаусів розподіл – ступені вершин зосереджені навколо середнього значення:

$$P(k) \propto e^{-\frac{(k-\mu)^2}{2\sigma^2}},$$

де μ – середнє значення ступенів, σ^2 – дисперсія, яка визначає ширину розподілу.

При подальшому вивченні зв'язків в мережевих структурах до слід прийняти такі властивості ребер графів, як їх вага, тобто, числове значення, що характеризує силу зв'язку між вершинами, а також пропускна здатність – максимальна кількість "ресурсів", які можуть передаватися через ребро.

Серед багатьох алгоритмів аналізу графів розглянемо алгоритм ранжирування вузлів мережі за впливами PageRank і алгоритми виявлення спільнот Louvain і Girvan-Newman.

PageRank визначає важливість вершин у графі на основі їх зв'язків. Основна формула:

$$PR(v_i) = \frac{1-d}{N} + d \sum_{v_j \in M(v_i)} \frac{PR(v_j)}{L(v_j)},$$

де:

- $PR(v_i)$ – ранг вершини v_i ;
- d – фактор демпфування ($0 < d < 1$);
- N – загальна кількість вершин у графі;
- $M(v_i)$ – множина вершин, що посиляються на v_i (тобто множина вершин, які мають зв'язок до v_i);
- $L(v_j)$ – кількість вихідних ребер вершини v_j (тобто кількість вершин, на які посиляється v_j).

Ця формула враховує дві компоненти:

- базова компонента $\frac{1-d}{N}$ забезпечує мінімальне значення рангу для кожної вершини, навіть у випадку відсутності посилянь;

- ітеративна компонента $d \sum_{v_j \in M(v_i)} \frac{PR(v_j)}{L(v_j)}$ враховує внесок інших вершин, які посилаються на v_i . Вага кожного посилення залежить від рангу вершини v_j та кількості її вихідних зв'язків.

Таким чином, PageRank розподіляє "ранг" або "важливість" між вершинами графу на основі їхньої структури зв'язків.

Алгоритми виявлення спільнот призначені для виділення груп вершин у графі, які мають більш щільну внутрішню зв'язність порівняно із зв'язками між групами. Такі структури часто називають *спільнотами* або *кластерами*. Розглянемо два популярних алгоритми: Louvain та Girvan-Newman, враховуючи додатковий параметр для Louvain – розподільну здатність (гамма).

Алгоритм Louvain – це ієрархічний метод виявлення спільнот, який максимізує функціонал модулярності графа.

Модулярність (Q) визначається як міра того, наскільки сильно структура графа відрізняється від випадкового розподілу зв'язків. Алгоритм працює ітеративно, покращуючи модулярність у спосіб перегрупування вершин між спільнотами. Важливим параметром є розподільна здатність (гамма), яка контролює чутливість до розміру спільнот.

Формальне визначення модулярності задається формулою:

$$Q = \frac{1}{2m} \sum_{i,j} \left[A_{ij} - \gamma \frac{k_i k_j}{2m} \right] \delta(c_i, c_j),$$

де

A_{ij} – елемент матриці суміжності графа (1, якщо є ребро між вершинами i та j , інакше 0);

k_i, k_j – ступені вершин i та j (кількість ребер, що виходять із кожної вершини);

m – загальна кількість ребер у графі;

c_i, c_j – спільноти, до яких належать вершини i та j ;

$\delta(c_i, c_j)$ – функція, яка дорівнює 1, якщо $c_i = c_j$, і 0 інакше;

γ – параметр розподільної здатності (гамма), який регулює важливість випадкового розподілу зв'язків. Зазвичай $\gamma > 0$.

Роль параметра γ визначається його значенням:

- якщо $\gamma=1$, формула модулярності відповідає класичному визначенню;
- якщо $\gamma < 1$, алгоритм надає перевагу більшим спільнотам (менша чутливість до маленьких кластерів);
- якщо $\gamma > 1$, алгоритм формує менші спільноти (більша чутливість до деталей).

Для визначення спільнот у відповідності із наведеним формальним визначенням застосовується алгоритм, що складається із послідовності кроків:

1. Ініціалізація. Кожна вершина графа розглядається як окрема спільнота.
2. Локальна оптимізація. Для кожної вершини перевіряється, чи переміщення її до сусідньої спільноти покращить значення модулярності з урахуванням γ . Якщо так, вершина переміщується.
3. Агрегація. Після завершення локальної оптимізації граф агрегується: кожна спільнота стає новою «супервершиною», а ребра між спільнотами перераховуються.
4. Ітерація. Кроки 2–3 повторюються до тих пір, поки модулярність не перестане зростати.

Переваги цього підходу полягають у швидкості – алгоритм має невелику складність $O(n \log(n))$, можливості працювати з великими графами, а також можливості налаштування до розміру спільнот завдяки параметру γ .

До недоліків слід віднести можливість алгоритму зупинитись на локальних оптимумах, при цьому не знайти глобального максимуму модулярності, а також чутливість до параметрів – неправильний вибір γ може призвести до невірної розбиття.

Алгоритм Girvan-Newman базується на ідеї поступового видалення ребер з найвищим значенням центральності за посередництвом (edge betweenness centrality). Міжцентральність ребра визначається як кількість найкоротших шляхів, які проходять через це ребро. Видалення таких ребер поступово розділяє граф на компоненти, які вважаються спільнотами.

Формальне визначення міжцентральності ребра:

$$C_B(e) = \sum_{s \neq t} \frac{\sigma_{st}(e)}{\sigma_{st}},$$

де σ_{st} – загальна кількість найкоротших шляхів між вершинами s та t , $\sigma_{st}(e)$ – кількість найкоротших шляхів між s та t , які проходять через ребро e .

Етапи алгоритму Girvan-Newman визначаються наступним чином:

1. Обчислення міжцентральності. Для кожного ребра обчислюється його міжцентральність.
2. Видалення ребра. Ребро з найвищим значенням міжцентральності видаляється.
3. Перерахунок компонент. Після видалення ребра граф розбивається на компоненти зв'язності.
4. Ітерація. Кроки 1–3 повторюються до тих пір, поки граф не буде повністю розділений на ізольовані вершини.

Переваги алгоритму Girvan-Newman полягають у його простоті (алгоритм легко зрозуміти та реалізувати) і точності, він добре працює для невеликих графів з чітко вираженими спільнотами.

До недоліків слід віднести складність обчислення міжцентральності для великих графів є обчислювально дорогим ($O(n^3)$), а також неможливість паралельної обробки: алгоритм вимагає послідовного видалення ребер.

1.4. Практична частина. Побудова простих графів

Графовий аналіз – це не лише теоретичний інструмент, а й практичний метод, який дозволяє наочно відобразити та дослідити структуру інформаційних мереж. Практична робота з графами у Gephi – це потужний інструмент для дослідження мережевих структур. Вона дозволяє не лише візуалізувати складні мережі взаємодій, а й проводити глибокий аналіз їх структури. Завдяки цьому ми можемо виявляти ключові вузли, кластери та прогнозувати можливі сценарії розвитку інформаційних атак. Розглянемо, як можна застосувати теоретичні знання до реальних даних, щоб побудувати графову модель інформаційного простору.

Підготовка до роботи з графами

Перш ніж почати роботу з Gephi, важливо правильно підготувати дані та зрозуміти, як вони будуть імпортовані та візуалізовані. Робота з графами передбачає створення та аналіз мережевих структур, де основними елементами є вузли (nodes) та ребра (edges). У цьому контексті підготовка даних є ключовим кроком для успішного аналізу.

Gephi підтримує різні формати даних, але найчастіше використовуються наступні:

- CSV (Comma-Separated Values);
- GEXF (Graph Exchange XML Format);
- GraphML – XML-формат для представлення графів;
- JSON – популярний формат для веб-додатків.

Для подальшої роботи ми будемо використовувати дані представлені у найпростішому вигляді – у форматі CSV, придатному для імпорту в Gephi. Варіант такого формату – це список рядків, що містить назви пар зв'язаних вузлів, розділених крапкою з комою, що можна вважати позначенням зв'язку (ребра).

Імпорт даних у Gephi

Коли дані підготовлені, першим кроком є їх імпорт у Gephi. Програма пропонує інтуїтивно зрозумілий інтерфейс для завантаження даних. Відкрийте меню **File** → **Open**, оберіть підготовлений CSV-файл та налаштуйте параметри імпорту. Наприклад, вкажіть, що ми імпортуємо Adjacency list, вказуємо, який символ розподіляє вузли, кодування (рис. 1.1).

Після імпорту Gephi автоматично створить базовий граф. На перший погляд він може здаватися хаотичним, оскільки вершини та ребра розташовані без певного порядку. Але саме тут починається найцікавіша частина роботи – структурування та аналіз графа.

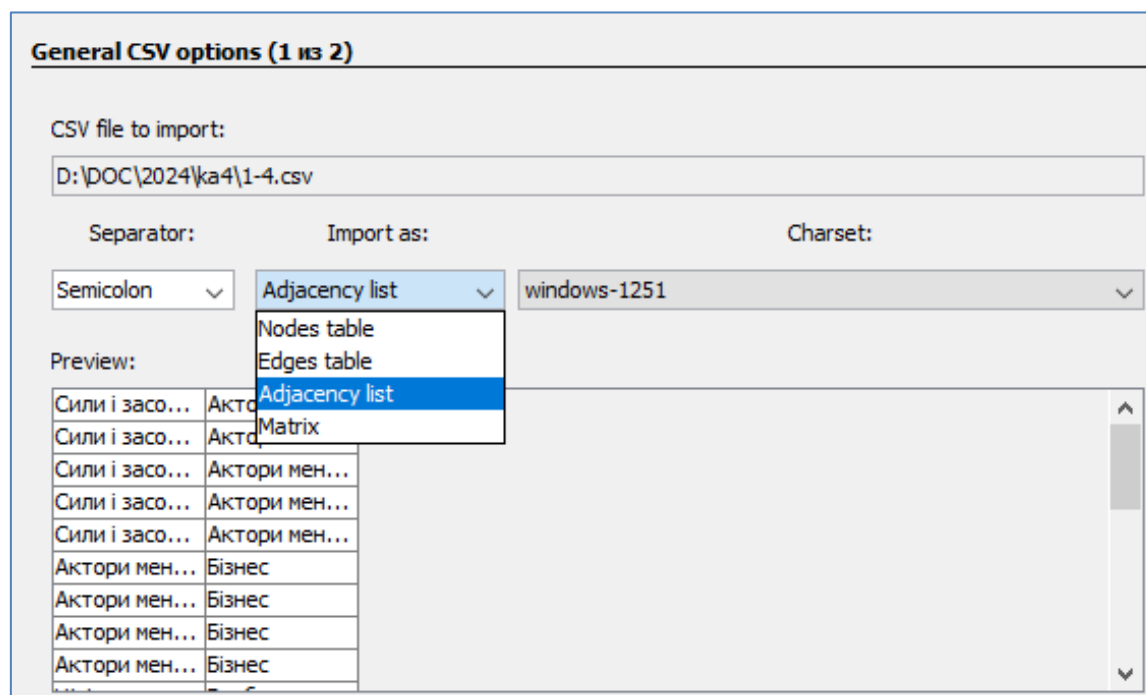


Рисунок 1.1: Налаштування формату даних, що вводяться

Візуалізація графа

Один із ключових аспектів роботи з графами – це їх візуалізація. Gephi пропонує декілька алгоритмів розташування вершин, які допомагають зробити граф більш зрозумілим. Наприклад, алгоритм **ForceAtlas2** імітує фізичні взаємодії між вершинами, розташовуючи їх таким чином, щоб схожі або сильно пов'язані вершини згрупувалися разом.

Щоб застосувати алгоритм, перейдіть до вкладки **Layout** та оберіть **ForceAtlas2**. Налаштуйте параметри, такі як сила відштовхування між вершинами або гравітація, і запустіть алгоритм. Через кілька секунд граф набуде більш структурованого вигляду, що дозволяє проаналізувати. Тепер можна помітити окремі кластери – групи вершин, які щільно пов'язані між собою. Ці кластери часто відображають спільноти користувачів, які діляться подібними інтересами або поглядами.

Приклад

Як перший приклад візуалізації, розглянемо візуалізацію ієрархічного графа, що відповідає ментальній війні.

Перший рівень графового представлення – «Цілі ментальної війни», другий рівень – «Сили і засоби ментальної війни», третій рівень – «Актори ментальної війни», четвертий рівень – «Цілі окремих акторів ментальної війни», п'ятий рівень – «Політики реалізації цілей ментальної війни», які у свою чергу приводять до «Результатів ментальної війни».

Відповідно, кожний рівень характеризується визначеними попередньо експертами поняттями, а саме:

1. Цілі ментальної війни: зміна ідентичності народу.
2. Сили і засоби ментальної війни: література, мистецтво, ЗМІ, соціальні медіа.
3. Актори ментальної війни: митці, уряд, художники тощо.
4. Цілі окремих акторів ментальної війни: зміна уявлень, зміна цінностей, зміна установок, зміна національної свідомості.
5. Політики реалізації цілей ментальної війни: мова, віра, культура, історична пам'ять.

Результати ментальної війни мають привести до перекодування уявлень людей щодо держави, суспільства, особи.

Для завантаження цих даних у програму Gephi формується файл у форматі CSV такого змісту:

Цілі ментальної війни; Зміна ідентичності народу
 Цілі ментальної війни; Сили і засоби ментальної війни
 Сили і засоби ментальної війни; Література
 Сили і засоби ментальної війни; Мистецтво
 Сили і засоби ментальної війни; ЗМІ
 Сили і засоби ментальної війни; Соціальні медіа
 Сили і засоби ментальної війни; Актори ментальної війни
 Актори ментальної війни; Митці
 Актори ментальної війни; Уряд
 Актори ментальної війни; Художники
 Актори ментальної війни; Цілі окремих акторів ментальної війни
 Цілі окремих акторів ментальної війни; Зміна уявлень
 Цілі окремих акторів ментальної війни; Зміна цінностей
 Цілі окремих акторів ментальної війни; Зміна установок
 Цілі окремих акторів ментальної війни; Зміна національної свідомості
 Цілі окремих акторів ментальної війни; Політики реалізації цілей ментальної війни
 Політики реалізації цілей ментальної війни: Мова
 Політики реалізації цілей ментальної війни: Віра
 Політики реалізації цілей ментальної війни: Культура
 Політики реалізації цілей ментальної війни: Історична пам'ять
 Політики реалізації цілей ментальної війни; Результати ментальної війни
 Результати ментальної війни; Перекодування уявлень українців щодо держави
 Результати ментальної війни; Перекодування уявлень українців щодо суспільства
 Результати ментальної війни; Перекодування уявлень українців щодо особи

Після завантаження цих даних у програму Gephi і укладання графу у відповідності з до алгоритму **ForceAtlas2** отримується представлення, наведене на Рис. 1.2. Цей рисунок буде розглядатись як схема для розширення за алгоритмами, наведеними у наступних розділах.

Аналіз графа

Наступний крок – це аналіз структури графа. Gephi надає широкий набір інструментів для обчислення різних метрик, таких як центральність, щільність та коефіцієнт кластеризації. Розглянемо основні показники:

- Ступінь вершини (Degree) – кількість зв'язків, що виходять з вершини. Вершини з високим ступенем часто є ключовими джерелами інформації.
- Центральність за посередництвом (Betweenness Centrality) – міра того, наскільки часто вершина знаходиться на найкоротшому шляху між іншими вершинами. Такі вершини є важливими для передачі інформації між різними кластерами.
- PageRank. Алгоритм, який визначає важливість вершини на основі її зв'язків з іншими вершинами. PageRank особливо корисний для виявлення впливових акторів у мережі.



Рисунок 1.2: Ієрархія базового графу «Ментальна війна»

Щоб обчислити ці метрики, скористайтесь вкладкою **Statistics**. Наприклад, оберіть **Average Degree** для оцінки загальної щільності зв'язків у графі або **Modularity** для виявлення спільнот. Після обчислення метрик вони автоматично додаються як атрибути вершин, що дозволяє візуалізувати їх за допомогою різних кольорів або розмірів.

Виявлення ключових вузлів та кластерів

Однією з головних цілей графового аналізу є виявлення ключових вузлів та кластерів у мережі. За допомогою Gephi це можна зробити наочно. Наприклад, застосуйте градієнтну заливку для вершин, щоб виділити їхню центральність за розміром або кольором. Вершини з високим значенням PageRank можна зробити більшими або виділити яскравими кольорами, що допоможе ідентифікувати впливових акторів. Крім того, модуль **Partition** дозволяє розфарбувати вершини за приналежністю до певних кластерів. Це особливо корисно для виявлення інформаційних «бульбашок» – замкнених спільнот, де

інформація циркулює між однаково налаштованими учасниками. Наприклад, у соціальних мережах такі кластери часто формуються навколо політичних, культурних або ідеологічних поглядів.

Прогнозування та практичне застосування

Побудований граф не лише допомагає зрозуміти поточну структуру інформаційної мережі, а й може слугувати основою для прогнозування майбутніх сценаріїв. Наприклад, аналізуючи ключові вузли, ми можемо передбачити, які актори будуть найбільш впливовими у поширенні інформації. Аналіз кластерів дозволяє ідентифікувати групи, які найбільше піддаються інформаційним атакам, що є важливим для розробки стратегій протидії.

Також графовий аналіз можна застосовувати для виявлення ботів у соціальних мережах. Боти часто мають характерні ознаки, такі як висока кількість зв'язків з іншими вершинами або певні шаблони поведінки. Виявивши такі вершини, ми можемо ізолювати їх від основної мережі, що допоможе зменшити ефективність інформаційних атак.

Висновки до розділу

Графові моделі є необхідним інструментом для сучасного аналізу інформаційних мереж. Вони дозволяють не лише зрозуміти поточну структуру комунікацій, а й розробляти стратегії протидії дезінформації, прогнозувати поведінку акторів та оптимізувати поширення достовірної інформації. Сучасні методи та інструменти графового аналізу стають все більш актуальними у контексті інформаційно-психологічних війн, де успіх залежить від глибини розуміння та ефективності використання мережевих структур.

У цьому розділі було розглянуто теоретичні основи графових моделей, які є ключовим інструментом для аналізу інформаційно-психологічних впливів. Графовий аналіз дозволяє формалізувати складні мережеві взаємодії між джерелами, поширювачами та одержувачами інформації, надаючи можливість виявляти центри впливу, прогнозувати поширення наративів та оцінювати ефективність комунікаційних стратегій.

Основні поняття теорії графів, такі як вершини, ребра, ступінь вершин, напрямленість зв'язків та вага ребер, утворюють базис для побудови моделей, які адекватно відображають реальні інформаційні мережі. Зокрема, важливими характеристиками графів є щільність, коефіцієнт кластеризації, центральність вершин та розподіл ступенів, які допомагають аналізувати структуру та поведінку мережі. Ці метрики особливо корисні для виявлення ключових акторів у поширенні інформації, таких як лідери думок або медіа-ресурси.

Особливу увагу приділено класифікації типів графів, кожен з яких має свої переваги та застосування. Наприклад, випадкові графи Ердеша-Реньї використовуються для моделювання хаотичного поширення інформації, масштабно-інваріантні графи підходять для аналізу соціальних мереж із «гігантськими хабами», а графи малого світу забезпечують можливість вивчати швидке поширення інформації через локальні групи та короткі зв'язки між ними. Динамічні графи, гіперграфи та мультиграфи дозволяють враховувати часову еволюцію мережі, складні взаємодії між багатьма агентами та паралельні канали комунікації.

Математичний апарат графового моделювання, включаючи матриці суміжності та інцидентності, списки суміжності, а також алгоритми PageRank та виявлення спільнот (Louvain, Girvan-Newman), надає потужні інструменти для аналізу структури та функціональності графів. Особливо важливим є використання параметра γ у методі Louvain, який дозволяє налаштовувати чутливість до розміру спільнот, що є критичним для точного розбиття мережі.

Практичне застосування графового аналізу, наприклад, у програмі Gephi, демонструє, як теоретичні знання можуть бути реалізовані для візуалізації та аналізу реальних даних. Це дозволяє не лише виявляти ключові вузли та кластери, а й прогнозувати сценарії розвитку інформаційних атак, ідентифікувати ботів та оцінювати вразливість певних груп до маніпулятивних впливів.

Питання для самоконтролю

1. Що таке граф у контексті теорії графів, і які основні елементи він містить?
2. Які основні поняття теорії графів використовуються для аналізу інформаційних мереж?
3. Що таке ступінь вершини, і як вона враховується при аналізі графа?
4. Як визначається щільність графа, і що вона показує у контексті інформаційних мереж?
5. Які типи графів можуть бути використані для моделювання інформаційно-психологічних війн?
6. Що таке масштабно-інваріантні графи, і яку роль вони відіграють для аналізу у соціальних мережах?
7. У чому полягають особливості динамічних графів стосовно інформаційних впливів?
8. Що таке гіперграфи, і як вони використовуються для моделювання складних взаємодій?

9. Які способи представлення графів використовуються для практичного аналізу?
10. Що таке центральність вершини, і які існують типи центральності?
11. Як визначається коефіцієнт кластеризації, і що він показує у графовому аналізі?
12. Як працює алгоритм PageRank, і яке значення він має для аналізу впливових вершин у графі?
13. У чому полягає методика алгоритму Louvain для виявлення спільнот, і яку роль відіграє параметр γ ?
14. Як працює алгоритм Girvan-Newman, і чому він може бути менш ефективним для великих графів?

Література за розділом

1. Основи теорії складних мереж: навч. пос. / Снарський А.О., Ланде Д.В., Субач І.Ю. ІСЗІ КПІ ім. Ігоря Сікорського. – Київ: ТОВ "Інжиніринг", 2023. – 225 с. ISBN 978-9662344950.
2. Візуалізація та аналіз мережових структур : навч. пос. / Ланде Д.В., Субач І.Ю. – Київ : КПІ ім. Ігоря Сікорського, Вид-во "Політехніка", 2021. – 80 с. ISBN 978-9662577143.
3. Amy N. Langville, Carl D. Meyer. Google's PageRank and Beyond: The Science of Search Engine Rankings. Princeton University Press, 2006. – 240 p. ISBN 978-0691122021
4. Mark Newman. Networks (2nd edn). Oxford University Press, 2018. – 800 p. ISBN 978-0198805090
5. Ken, Cherven (2015). Mastering Gephi Network Visualization. Packt Publishing, 378. ISBN 978-1783987344
6. Douglas A. Luke A User's Guide to Network Analysis in R. Springer; 1st ed. 2015 edition (December 21, 2015). 250 pp. ISBN: 978-3319238821

2. Великі мовні моделі та їх застосування

Штучний інтелект і великі мовні моделі відкрили нові можливості для аналізу та моделювання інформаційно-психологічних операцій і ментальних війн. Вони дозволяють виявляти ключові наративи, прогнозувати розвиток інформаційних атак і навіть створювати симуляції для оцінки ефективності впливових впливу кампаній.

LLM здатні не лише обробляти величезні обсяги текстової інформації, а й генерувати нові концепти та зв'язки між ними. Це робить їх ефективним інструментом для моделювання динаміки інформаційних операцій, визначення тактик дезінформації та розробки контрзаходів.

2.1. Поняття LLM

У цьому розділі розглядаються принципи роботи LLM, методи генерації концептів, підхід «рою віртуальних експертів» для побудови семантичних мереж, а також практичне застосування цих технологій для аналізу інформаційно-психологічних впливів.

Великі мовні моделі – це штучні нейронні мережі, навчені на великих обсягах текстових даних для розпізнавання, аналізу та генерації мовного контенту. Їхня робота базується на декількох ключових принципах:

- *Токенізація* – розбиття тексту на окремі складові (слова, фрази, символи).
- *Трансформерна архітектура* – використання механізмів самоуваги (self-attention) для розуміння взаємозв'язків між словами у великому контексті.
- *Контекстне прогнозування* – визначення найбільш імовірного наступного слова або фрази на основі попереднього контексту.
- *Файнтюнінг* (донавчання) – адаптація LLM до специфічних задач, наприклад, аналізу інформаційних атак чи виявлення фейкових новин.

LLM можуть використовуватися для аналізу текстових масивів у сфері інформаційно-психологічних війн, зокрема для автоматичного визначення ключових тем і наративів, ідентифікації дезінформації та інформаційних атак, аналізу динаміки поширення певних ідей у медіапросторі, генерації відповідей та контрнативів тощо.

Великі мовні моделі, такі як GPT-4, Qwen2.5, Llama-4, та інші, є генеративними мовними моделями, які здатні створювати та аналізувати тексти на основі введених запитів та контексту. Ці системи

розвиваються на основі сучасних технологій штучного інтелекту (ШІ, Artificial Intelligence, AI) і стали ключовими інструментами у розв’язанні завдань обробки природної мови (Natural Language Processing, NLP) і машинного навчання. Вони дозволяють генерувати текст, робити переклад з різних мов, інформативно відповідати на запитання.

2.2. Архітектура «Трансформер»

Основи та принципи роботи LLM базуються на глибокому навчанні та архітектурі «Трансформер» (Transformer), яка дозволяє ефективно моделювати довгострокові залежності в тексті. Ця архітектура вперше була представлена у 2017 році в роботі "Attention is All You Need" (<https://arxiv.org/pdf/1706.03762>) і стала технологічною базою багатьох сучасних NLP моделей.

Архітектура «Трансформер» стала справжнім проривом у світі штучного інтелекту, особливо в обробці природної мови. Вона дозволила в багатьох випадках відмовитись від традиційних рекурентних (RNN) та згорткових (CNN) мереж, які обробляли дані послідовно, і запропонувала новий підхід, заснований на механізмах уваги (Attention Mechanisms). В рамках цієї архітектури елементи послідовності аналізуються одночасно, що значно прискорює навчання та покращує якість результатів.

«Трансформер» не використовує рекурентні або згорткові шари, тому при його роботі невідомо, в якому порядку йдуть слова. Щоб вирішити цю проблему, до кожного елемента мови – токена, додається позиційне кодування – спеціальний вектор, який вказує на позицію слова в реченні. Це дозволяє моделі враховувати порядок слів, що є важливим для розуміння контексту.

Для пояснення роботи архітектури «Трансформер» звернемося до схеми, наведеної на рис. 2.1, де детально показано, як працюють основні складові трансформера: кодувальник, декодувальник і механізми уваги:

- Ліва частина схеми – кодувальник (Encoder), який складається з кількох шарів ($N \times$).
- Права частина – декодувальник (Decoder), який також складається з кількох шарів ($N \times$).
- Стрілки позначають зв’язки між різними компонентами, включаючи залишкові з’єднання та потоки даних між кодувальником і декодувальником.

Обидві частини складаються з кількох однакових шарів (зазвичай $N = 6$), які мають подібну структуру, але виконують різні функції.

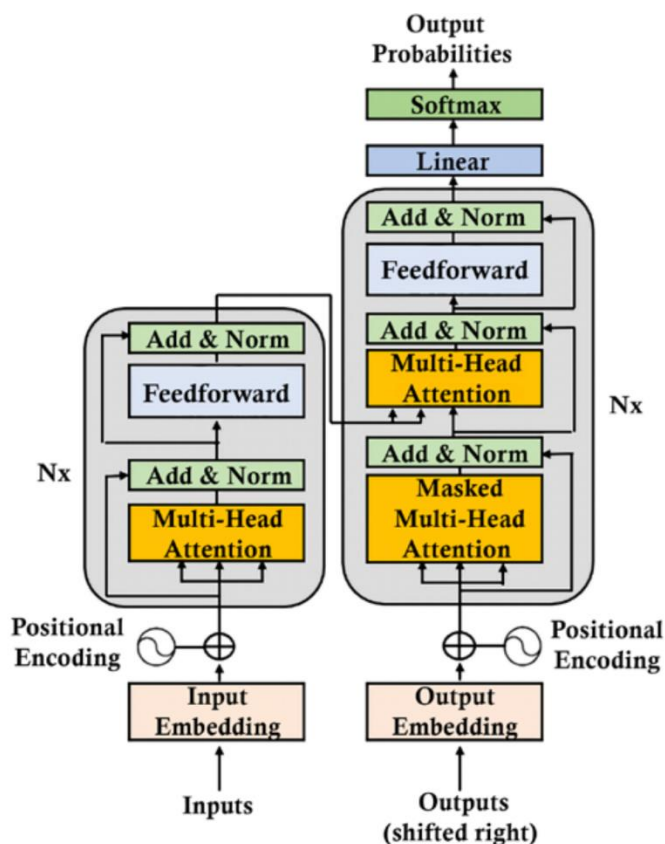


Рисунок 2.1: Архітектура моделі «Трансформер»
(<https://arxiv.org/pdf/1706.03762>)

Нижче детально розглядаються всі елементи схеми.

«Мозоком» моделі Трансформер є *Кодувальник (Encoder)*, який аналізує вхідний текст. Він складається з кількох шарів, кожен з яких виконує дві основні операції (як показано на схемі):

При роботі кодувальника кожен токен у реченні «зосереджується» на інших токенах, використовуючи так званий механізм багатоголової уваги. Це схоже на те, як ви читаете речення і зосереджуєте увагу на ключових словах, щоб зрозуміти сенс. Наприклад, у реченні «Кіт сидить на столі» модель може зосередитися на слові «кіт», щоб зрозуміти, хто виконує дію.

Кожна «голова» уваги навчається зосереджуватися на різних аспектах тексту, наприклад, на синтаксичних зв'язках (хто і що робить) або семантиці (значення слів). Результати всіх голів об'єднуються, щоб створити єдине представлення тексту.

Механізм багатоголовної уваги (Multi-Head Attention) дозволяє моделі «зосереджуватися» на різних частинах вхідної послідовності. Кожний токен (слово або частина слова) може взаємодіяти з будь-яким іншим токеном у послідовності, незалежно від відстані між ними. Механізм працює наступним чином:

- Кожен токен генерує три вектори: запит (Query), ключ (Key) та значення (Value).
- Запити порівнюються з ключами, щоб визначити, наскільки важливий кожен токен для поточного.
- На основі цих порівнянь обчислюються вагові значення, які визначають, яку частину інформації зі значень (Value) слід враховувати.
- Цей процес повторюється декілька разів (зазвичай 8) з різними наборами параметрів (це і є «багатоголова увага»), що дозволяє моделі враховувати різні аспекти даних (наприклад, синтаксичні та семантичні зв'язки).

Після шару уваги дані проходять через повнозв'язний шар (Feed-Forward Network, FFN), який додає нелінійність і дозволяє моделі краще засвоювати інформацію. Цей шар обробляє дані, отримані від механізму уваги.

Таким чином, кодувальник приймає вхідну послідовність (наприклад, речення) і перетворює її на набір векторів, які кодують інформацію про кожне слово та його контекст. Кожен шар кодувальника складається з двох лінійних перетворень і функції активації ReLU (Rectified Linear Unit) між ними: Він додає нелінійність і дозволяє моделі краще засвоювати інформацію.

$$FFN(x) = \max(0, xW_1 + b_1)W_2 + b_2,$$

де W_1 , W_2 – вагові матриці, а b_1 , b_2 – зміщення.

Завдяки цьому додається нелінійність, що дозволяє моделі краще засвоювати інформацію.

Навколо кожного підшару (уваги та повнозв'язного шару) додаються залишкові з'єднання (Residual Connections). Це означає, що вихід підшару додається до його входу перед нормалізацією.

До кожного підшару додається залишкове з'єднання, після чого застосовується нормалізація шару (Layer Normalization), яка стабілізує навчання. Таким чином, вихід кожного підшару обчислюється за формулою функції, яку реалізує підшар:

$$\text{LayerNorm}(x + \text{Sublayer}(x)).$$

Щоб спростити залишкові з'єднання, всі підшари моделі, як і шари вбудовування, генерують виходи розмірності $d_{model} = 512$.

Декодувальник (Decoder) – це частина моделі, яка генерує вихідний текст (наприклад, переклад речення). Декодувальник генерує вихідну послідовність (наприклад, переклад речення) на основі виходу кодувальника.

Він також складається з кількох шарів, але має додаткові функції.

Під час навчання декодувальник використовує *масковану увагу* (Masked Multi-Head Attention), яка запобігає «підгляданню» в майбутні токени. Це забезпечує те, що система не може зчитувати наступні токени в тексті, поки не прочитає поточний. Маскована увага – це механізм уваги, який запобігає «підгляданню» в майбутні токени під час навчання. Це важливо для забезпечення коректної генерації послідовностей.

Під час навчання модель має доступ лише до поточного та попередніх токенів, а не до майбутніх. Це досягається за допомогою маски, яка приховує майбутні токени.

Механізм уваги можна описати як перетворення запиту (query) та набору пар ключ-значення (key-value) у вихід, де всі ці елементи представлені векторами. Вихід обчислюється як зважена сума значень, де ваги визначаються функцією сумісності запиту з відповідними ключами.

Особливий вид уваги, використаний у трансформері, називається «масштабована увага на основі скалярного добутку».

Вхід включає:

- Запити Q і ключі K розмірності d_k .
- Значення V розмірності d_v . Це значення, які відповідають ключам і містять власне інформацію.

Обчислення включають такі кроки:

1. Знаходиться скалярний добуток запиту з кожним ключем.
2. Результати діляться на $\sqrt{d_k}$.
3. Застосовується softmax для нормалізації вагових значень.
4. Вихід – це зважена сума значень.

Формула для обчислення уваги має вигляд:

$$Attention(Q, K, V) = softmax\left(\frac{QK^T}{\sqrt{d_k}}\right)V.$$

Декодувальник також "зосереджується" на виході кодувальника, щоб враховувати контекст вхідного тексту. Наприклад, під час перекладу з однієї мови на іншу, декодувальник використовує інформацію з кодувальника, щоб правильно перекласти кожне слово.

Багатоголова увага над виходом кодувальника – це механізм уваги, який дозволяє декодувальнику "зосереджуватися" на виході кодувальника. Це допомагає моделі враховувати контекст вхідного тексту під час генерації вихідного.

Декодувальник використовує запити (Query) зі свого попереднього шару, а ключі (Key) та значення (Value) беруться з виходу кодувальника. Це дозволяє моделі враховувати інформацію з вхідного тексту під час генерації вихідного.

Замість використання одного механізму уваги, багатоголовкова увага проєктує запити, ключі та значення h разів у різні підпростори, після чого виконує паралельну увагу. Це дозволяє моделі одночасно обробляти інформацію з різних підпросторів представлення.

Формально багатоголову увагу над виходом кодувальника можна записати наступним чином :

$$MultiHead(Q, K, V) = Concat(head_1, \dots, head_h)W^O,$$

де $head_i = Attention(QW_i^Q, KW_i^K, VW_i^V)$.

У цьому записі проєкції є параметричними матрицями $W_i^Q \in \mathbb{R}^{d_{model} \times d_k}$, $W_i^K \in \mathbb{R}^{d_{model} \times d_k}$, $W_i^V \in \mathbb{R}^{d_{model} \times d_v}$ і $W^O \in \mathbb{R}^{d_{model}}$.

У первинній моделі трансформера використовувалось $h=8$ голів, причому $d_k = d_v = d_{model} / h = 64$.

Аналогічно до кодувальника, декодувальник також має повнозв'язний шар (Feed-Forward Network), який додає нелінійність і обробляє дані.

Як і в кодувальнику, навколо кожного підшару додаються залишкові з'єднання та виконується нормалізація шару.

Декодер також складається з $N=6$ ідентичних шарів. Окрім двох підшарів, аналогічних енкодеру, декодер додає третій підшар, який виконує багатоголовкову увагу до виходу стеку енкодера.

Як і в енкодері, до кожного підшару додається залишкове з'єднання, за яким слідує нормалізація шару. Однак у підшарі самоуваги декодера

змінюється спосіб роботи: він запобігає доступу до майбутніх позицій. Це досягається за допомогою маскування, яке гарантує, що прогноз для позиції i залежить лише від відомих виходів для позицій, менших за i .

Оскільки трансформер не використовує рекурентні або згорткові шари, він не має вбудованого розуміння порядку елементів у послідовності. Щоб вирішити цю проблему, до вхідних векторів додаються *позиційні кодування*. Це спеціальні вектори, які кодують інформацію про позицію кожного токена в послідовності. Формула для позиційного кодування виглядає так:

$$PE_{(pos,2i)} = \sin\left(\frac{pos}{10000^{2i/d_{model}}}\right),$$
$$PE_{(pos,2i+1)} = \cos\left(\frac{pos}{10000^{2i/d_{model}}}\right),$$

де pos – позиція токена, а i – індекс розмірності.

Після останнього шару декодувальника дані проходять через лінійний шар та функцію `softmax`, яка перетворює вектори на ймовірності для кожного токена в словнику. Це дозволяє моделі вибирати найбільш ймовірний наступний токен під час генерації тексту.

Токенізація

Перший крок навчання трансформера – це токенизація. Токенізатор – це як ножиці для слів. Він розрізає текст на маленькі шматочки (токени), які трансформер потім аналізує окремо. Текст розбивається на менші частини, які називають токенами [Ovalle, 2024]. Це можуть бути слова, частини слів або навіть символи. Наприклад, речення "Transformer is amazing!" може бути розбите на токени: ["Transform", "er", " is", " amazing", "!"]. Інший приклад – під час аналізу речення: "Шкідлива програма WannaCry використовує вразливість CVE-2023-1234." Можуть бути виділені такі токени: ["Шкідлива", "програма", "WannaCry", "використовує", "вразливість", "CVE-2023-1234"], або токени як мітки сутностей: ["Malware", "WannaCry", "Action", "Vulnerability", "CVE-2023-1234"].

Токенізація допомагає моделі працювати з текстом як з набором «токенів», а не довгими реченнями. Після цього кожен токен перетворюється на вектор за допомогою ембеддингів – числових

представлень, які допомагають моделі працювати з текстом як з послідовністю чисел.

Токенізація є основним етапом попередньої обробки тексту в процесі семантичного індексування. Вона забезпечує структуроване представлення текстових даних і підготовлює їх для подальших етапів, таких як екстракція сутностей і зв'язків. Точна і правильно налаштована токенізація є ключовим елементом для досягнення ефективності на наступних етапах індексування, аналізу та пошуку в великих масивах текстових даних.

Токенізація служить базовою операцією, яка підготовлює текст для екстракції сутностей, зв'язків і подальшого семантичного аналізу. Зупинимось детальніше на цьому процесі, розглянувши його різновиди, важливість і приклади застосування.

Існує декілька видів токенізації, серед яких токенізація на рівні слів, токенізація на рівні слів символів, токенізація на рівні слів фраз або N-грам тощо.

Токенізація на рівні слів – це найпоширеніший підхід, де текст розбивається на окремі слова. Він підходить для більшості природних мов, де слова чітко розділені пробілами або знаками пунктуації. Розглянемо приклад токенізації речення: «Атака на сервер була проведена вночі». У результаті можна отримати токени: [«Атака», «на», «сервер», «була», «проведена», «вночі»]. Слід зазначити, що при такому процесі виникають виклики при роботі з такими елементами як аббревіатури, багатокomпонентні назви (наприклад, CVE-2023-1234), та складні слова.

При токенізації на рівні символів текст розбивається на окремі символи. Така токенізація зазвичай використовується у випадках, коли важливо зберегти всі символи або при роботі з мовами, де немає чітких розділювачів між словами (наприклад, китайська), наприклад, фрагмент тексту "CVE-2023-1234" розбивається на токени: ["C", "V", "E", "-", "2", "0", "2", "3", "-", "1", "2", "3", "4"]. Цей підхід дозволяє працювати з окремими символами та спеціальними знаками, що може бути корисним при аналізі технічних текстів (наприклад, у кібербезпеці при роботі з назвами вразливостей чи версій програм).

При токенізації на рівні фраз або N-грам текст розбивається на групи слів або фрази, що дозволяє зберегти більше контексту між словами. N-грами – це послідовність із N слів, які розглядаються як єдиний токен. Наприклад, текст: «Шкідлива програма WannaCry використовує вразливість.» Розбивається на токени (біграми, N=2): [«Шкідлива програма», «програма WannaCry», «WannaCry використовує»,

«використовує вразливість»]. Такий підхід часто використовується для поліпшення точності аналізу при витягуванні сутностей і зв'язків між ними, особливо для складних термінів, імен або фраз.

Навчання LLM

Великі мовні моделі навчаються на величезних обсягах тексту, найчастіше з Інтернету. Вони запам'ятовують, як люди зазвичай будують речення, і можуть писати текст так, ніби його написала людина.

Особливістю таких моделей є здатність до «навчання без вчителя» (unsupervised learning). Це означає, що модель не потребує розмічених даних (наприклад, пар речень для перекладу). Натомість вона навчається на великих обсягах неструктурованого тексту, використовуючи завдання маскування. Вони можуть самостійно вивчати структуру мови та її закономірності, опрацьовуючи текстові дані без попередньо визначених пар «запит–відповідь». Це дозволяє LLM бути надзвичайно гнучкими й адаптованими до виконання різних завдань, зокрема автоматичного перекладу, створення резюме текстів, генерації відповідей, синтезу контенту, класифікації даних тощо. Наприклад, модель може отримати речення з прихованим словом і навчитися передбачати це слово на основі контексту. Цей підхід дозволяє Transformer засвоювати загальні закономірності мови, які потім можна використовувати для різних задач.

Послідовність навчання великої мовної моделі можна представити як послідовність етапів (рис. 2.2):

1. **Збір даних.** Спочатку збирають велику кількість текстових даних з різних джерел, таких як книги, статті, вебсайти. Це потрібно, щоб модель могла "читати" багато інформації і навчитися розуміти, як будується мова.
2. **Токенізація.** Наступний крок навчання трансформера – це токенизація. Токенізатор – це як ножиці для слів. Він розрізає текст на маленькі шматочки (токени), які трансформер потім аналізує окремо. Це допомагає моделі працювати з текстом як з набором «токенів», а не довгими реченнями. Після цього кожен токен перетворюється на вектор за допомогою ембеддингів – числових представлень, які допомагають моделі працювати з текстом як з послідовністю чисел.
3. **Тренування моделі.** На основі зібраних даних та токенів модель навчається розуміти структуру мови. Вона вчиться передбачати, яке слово або фраза має з'явитися наступним у тексті. Це схоже на тренування мозку вгадувати продовження речень.

4. **Вдосконалення (Fine-tuning).** Модель навчають під конкретні задачі. Наприклад, її можна навчити писати тексти, відповідати на питання або створювати оповідання.
5. **Тестування.** Модель перевіряють на нових текстах, яких вона раніше не бачила. Це потрібно, щоб переконатися, що вона правильно розуміє завдання і дає точні відповіді.
6. **Реальне використання (Deployment).** Коли модель добре працює, її можна використовувати в реальних застосуваннях – наприклад, у чатах, для написання текстів або допомоги в пошуку інформації.



Рисунок 2.2: Послідовність навчання великої мовної моделі

Методи претренування дозволяють моделям вивчати мовні структури на великих об'ємах тексту перед тим, як її використовувати для конкретних завдань.

Модель навчається передбачати наступне слово в реченні чи текстовому фрагменті, використовуючи контекст з попередніх слів. LLM використовують авторегресійний підхід до генерації тексту, де кожне наступне слово генерується одне за одним з використанням контексту попередніх слів. Це дозволяє моделі генерувати довгі та змістовні висловлювання.

Після претренування модель може бути налаштована на конкретне завдання, таке як сентимент-аналіз, машинний переклад чи відповіді на питання. Це дозволяє моделі використовувати здобуті знання для конкретних сценаріїв.

Чому трансформер став революцією

Трансформер – це архітектура, яка змінила правила гри в світі ШІ. Вона дозволила створити моделі, здатні ефективно обробляти великі обсяги даних, навчатися без вчителя та застосовуватися до різноманітних задач. Її механізми уваги та паралелізація стали основою для таких відомих моделей, як GPT, BERT та багато інших, які сьогодні використовуються в мільйонах додатків по всьому світу.

Трансформер можна уявити як розумного помічника, який вміє читати та розуміти тексти. Коли трансформер читає текст, він спочатку розбиває його на частини, потім дивиться на кожну частинку і намагається зрозуміти, які слова найважливіші. Далі він переводить ці частинки у зрозумілі для нього числа, обробляє їх і створює відповідь або новий текст, дотримуючись правильного порядку слів.

Фактично, трансформер має багато очей і може одночасно дивитися на всі слова у реченні, звертаючи особливу увагу на ті, що важливі для розуміння. Це як «суперсила», яка дозволяє розуміти, про що йде мова, навіть якщо речення складне.

Енкодери виступають як розумні перекладачі, які перетворюють кожен шматочок тексту в набір чисел, щоб трансформер міг їх зрозуміти. Декодери – це такі ж перекладачі, але вони перетворюють числа назад у текст, щоб створити відповідь або новий текст.

Позиційне кодування використовується як номерки для кожного слова, щоб не загубити їхній порядок. Багатоголовий механізм уваги – це як команда з кількох спеціалістів, кожен з яких звертає увагу на різні частини тексту одночасно.

Цей механізм дозволяє фокусуватись на різних словах чи фразях. Наприклад, одне "око" може звертати увагу на імена, інше – на дії, а ще одне – на місця. Завдяки цьому трансформер отримує більше інформації і краще розуміє сенс тексту.

Коли трансформер читає текст, він використовує багатоголовий механізм уваги, щоб одночасно звертати увагу на різні частини тексту. Це допомагає йому краще зрозуміти контекст і значення кожного слова, що в свою чергу дозволяє йому створювати більш точні і корисні відповіді.

На питання, завдяки чому ця архітектура стала революцією можна відповісти, що вона відрізняється такими властивостями, як паралелізація, глобальна увага, масштабованість і універсальність:

1. На відміну від рекурентних мереж, які обробляють дані послідовно, трансформер може обробляти всі елементи послідовності одночасно. Це значно прискорює навчання.
2. Механізм уваги дозволяє моделі враховувати залежності між будь-якими елементами послідовності, незалежно від їхньої відстані.
3. Трансформер може ефективно навчатися на великих обсягах даних, що дозволило створити такі потужні моделі, як GPT, BERT та інші.
4. Архітектура трансформер виявилася придатною не лише для NLP, але й для обробки зображень, генерації музики та інших задач.

Таким чином, трансформери можуть аналізувати та розуміти мову краще, ніж попередні методи. Вони паралельно аналізують всю послідовність даних, що значно пришвидшує їх та дозволяє обробляти довгі тексти. Ці нейронні мережі на цей час вже використовуються в багатьох сучасних програмах, зокрема для перекладу текстів, написання статей та навіть створення музики.

2.3. Модель шести сутностей

Дослідження суті і можливостей великих мовних моделей, як основи побудови систем LLM, дозволило узагальнено розглядати ці моделі як сукупність шести сутностей, а саме, токенів, слів, сутностей (понять), наративів, подій і документів (рис. 2.3).

При цьому різні сутності відповідають різним фрагментам документального масиву, на якому відповідна модель навчалася. Так всьому документальному масиву відповідає множина токенів, яких у сучасних системах може бути до сотень мільярдів.

Слова можуть відповідати одному або декільком токенам, а відповідно сутності складаються із слів. Наративи формулюються словами, і їм відповідає менша множина документів. Деякі наративи охоплюють події, що їм відповідають. Не всім наративам відповідають події (особливо в епоху постправди), відповідно, події співвідносяться ще з меншим обсягом документів.

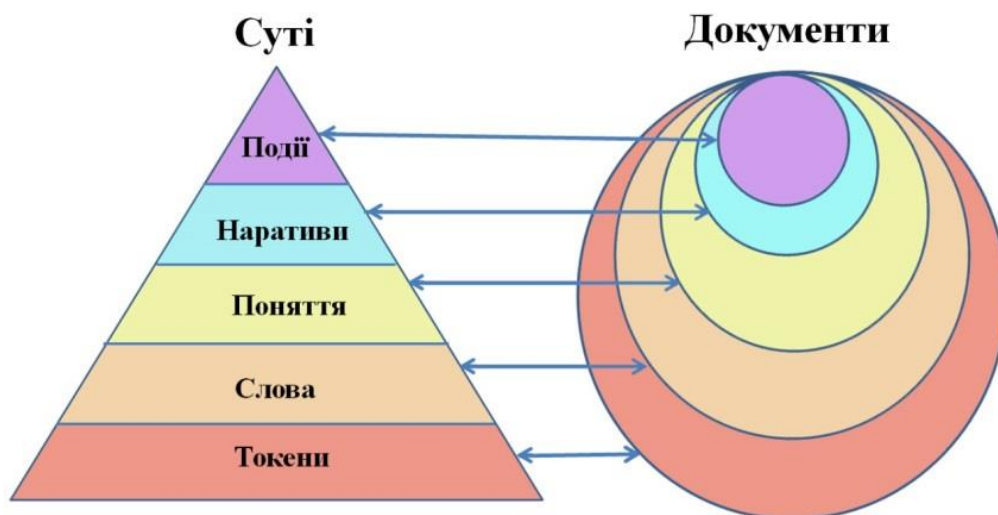


Рисунок 2.3: Елементи великої мовної моделі

Як показано на рис. 2.3 елементами моделі, що розглядається, є:

Токени, тобто найменші одиниці тексту, з яких складаються дані, на яких навчаються великі мовні моделі. Це можуть бути окремі символи, частини слів або цілі слова.

Слова – групи токенів, які мають певне значення. Для ефективного аналізу тексту слова часто нормалізуються до їхньої нормальної форми (леми). Наприклад:

«відбулася» → «відбутися» (дієслово в інфінітиві);

«конференції» → «конференція» (іменник у називному відмінку однини).

Використання лем дозволяє зменшити кількість унікальних слів у моделі, що спрощує обробку тексту та покращує розуміння семантики.

Сутності (поняття) – іменовані об'єкти, такі як люди, місця, організації, події тощо. Сутності можуть бути ідентифіковані через аналіз слів у їхній нормальній формі. Наприклад:

«Києві» → «Київ» (іменник у називному відмінку);

«2023 року» → «2023» (число).

Наративи – це історії або оповіді, які складаються зі слів та сутностей. Нормалізація слів допомагає виділити основні теми та ідеї, що лежать в основі наративу.

Події, тобто дії, які відбуваються в часі. Вони часто пов'язані з дієсловами, які також нормалізуються до інфінітиву для спрощення аналізу. Наприклад:

«відбулася» → «відбутися»;

«обговорювалися» → «обговорюватися».

Документи, або зібрання тексту, такі як статті новин, звіти або законопроекти. У рамках документів слова та фрази також можуть бути приведені до нормальної форми для покращення пошуку та аналізу.

Приклад. Речення «У 2023 році в Києві відбулася конференція з штучного інтелекту, на якій обговорювалися нові підходи до навчання великих мовних моделей.»

Це речення можна розкласти на основні сутності, враховуючи нормальні форми слів:

Токени:

[«У», «2023», «року», «в», «Києві», «відбулася», «конференція», «з», «штучного», «інтелекту», «,»,...]

Слова (нормальні форми):

«У» → «у»;

«2023 року» → «2023»;

«Києві» → «Київ»;

«відбулася» → «відбутися»;

«конференція» → «конференція»;

«обговорювалися» → «обговорюватися»;

«навчання» → «навчання».

Сутності (поняття):

Локація: «Київ».

Подія: «відбутися конференція».

Технології: «великі мовні моделі».

Час: «2023».

Наративи:

Оповідь про конференцію з штучного інтелекту, де обговорювалися нові підходи до навчання ВММ.

Події:

«Відбутися конференція з штучного інтелекту».

Документи:

Речення є частиною більшого документа, наприклад, статті чи звіту.

Запропонована модель, яка включає шість основних сутностей (токени, слова, сутності, наративи, події та документи), є більш деталізованою

та універсальною порівняно з попередніми спрощеними моделями. Розгляд такої моделі має кілька важливих переваг для аналізу текстових даних, особливо у контексті великих мовних моделей (ВММ) та інформаційного моделювання.

Раніше, наприклад, у системах типу Palantir, використовувалися більш прості моделі, які базувалися лише на трьох основних сутностях: слово, подія та документ. Це дозволяло ефективно працювати з конкретними завданнями, такими як аналіз новин або розслідування подій, але мало певні обмеження:

- обмежена глибина аналізу, тому що така модель не враховувала семантичні зв'язки між словами, не виділяла окремі сутності (наприклад, іменовані об'єкти) і не враховувала наративи;
- низька адаптивність, оскільки спрощені моделі погано справлялися зі складними текстами, де потрібно було аналізувати взаємодію між різними рівнями текстової інформації;
- відсутність нормалізації слів. У таких моделях не враховувалися нормальні форми слів (леми), що ускладнювало обробку тексту та призводило до надлишковості даних.

Запропонована модель вирішує ці проблеми, додаючи нові рівні абстракції та деталізації.

До переваг запропонованої моделі можна віднести універсальність, глибокий семантичний аналіз, ефективну обробку тексту, покращену інтерпретацію подій, підтримку складних завдань.

Модель охоплює всі ключові аспекти текстового аналізу: від найменших одиниць (токени) до найвищого рівня організації (документи). Вона може бути застосована до будь-якого типу текстових даних: від коротких повідомлень до довгих наукових статей чи законодавчих документів.

Виділення сутностей (понять) дозволяє точно ідентифікувати ключові об'єкти, про які йдеться у тексті.

Аналіз наративів допомагає зрозуміти загальну тему або історію, що розповідається.

Використання нормальних форм слів (лем) спрощує обробку тексту, зменшує кількість унікальних слів і покращує семантичне розуміння.

Модель враховує взаємозв'язки між різними рівнями тексту, що дозволяє краще аналізувати контекст.

Події пов'язуються з часом, місцем та учасниками, що полегшує їх ідентифікацію та аналіз.

Модель враховує, що не всі наративи відповідають реальним подіям (особливо в епоху постправди), що робить її більш гнучкою.

Модель може бути використана для широкого кола завдань, таких як генерація тексту, аналіз наративів, виявлення дезінформації, прогнозування трендів тощо.

Таким чином, запропонована модель, яка включає шість сутностей (токени, слова, сутності, наративи, події та документи), є значним кроком вперед у порівнянні з попередніми спрощеними моделями. Вона забезпечує більш глибокий та комплексний аналіз текстових даних, що особливо важливо у сучасному світі, де обсяги інформації постійно зростають. Основні переваги моделі полягають у її універсальності, семантичній глибині та здатності ефективно обробляти великі масиви тексту.

Такий підхід дозволяє ліпше розуміти структуру та зміст тексту, виявляти ключові сутності, наративи та події, генерувати високоякісний текст, який відповідає контексту, аналізувати інформаційні потоки та виявляти потенційні загрози.

Отже, запропонована модель створює понятійну основу для подальшого розвитку технологій обробки природної мови та інформаційного моделювання.

2.4. Генерація концептів та зв'язків

Генерація концептів та виявлення зв'язків між ними у контексті аналізу текстових даних при дослідженні ІПВ є одним із ключових завдань сучасних великих мовних моделей, особливо, побудови семантичних моделей та прогнозування стратегій інформаційного впливу. Це завдання включає екстракцію нових сутностей (понять) та виявлення різних типів зв'язків між ними з акцентом на промпт-інжиніринг як один із найефективніших інструментів для таких завдань.

Для кращого розуміння нижче будуть наведені приклади, включаючи великий приклад з області інформаційно-психологічних війн.

Екстракція понять з текстів – це процес автоматичного виявлення значущих об'єктів, подій, характеристик та взаємозв'язків, які є ключовими для розуміння тексту. Це можуть бути іменовані сутності (наприклад, люди, організації, локації), абстрактні поняття (наприклад, "свобода", "безпека") або конкретні терміни (наприклад, "штучний інтелект").

Для виявлення нових сутностей в рамках цього курсу використовується промпт-інжиніринг (Prompt Engineering) – це техніка, при якій

формулюється специфічне запитання або інструкція для LLM, щоб отримати бажаний результат.

Наприклад, для виявлення нових понять можна використовувати такий запит (промпт):

Виділи з тексту всі ключові поняття, пов'язані з темою інформаційної безпеки.

Модель аналізує текст і повертає список понять, таких як «кібератаки», «фейкові новини», «пропаганда».

Приклад. Текст: «У 2023 році спостерігалось зростання кібератак на критичну інфраструктуру. Зловмисники використовували фейкові новини для маніпуляції суспільною думкою.»

Результат: Ключові поняття: «кібератаки», «критична інфраструктура», «зловмисники», «фейкові новини», «маніпуляція», «суспільна думка».

В рамках методів іменованої екстракції сутностей (Named Entity Recognition, NER) здійснюється виявлення іменованих сутностей у тексті (наприклад, імен людей, назв організацій, дат).

У комбінації з LLM цей метод стає більш гнучким, оскільки модель може враховувати контекст і виділяти абстрактні поняття.

Приклад. Текст: «Організація 'Кіберщит' провела тренінг для захисту від кібератак у 2023 році.»

Результат охоплює такі аспекти:

1. Іменовані сутності: «Кіберщит» (організація), «2023 рік» (час).
2. Абстрактні поняття: «тренінг», «захист», «кібератаки».
3. Тематичне моделювання (Topic Modeling): Тематичне моделювання (наприклад, LDA) дозволяє виявляти загальні теми у великому масиві текстів.

Наприклад, у текстах про інформаційну безпеку можуть бути виділені теми "кіберзахист", "соціальна інженерія" тощо.

Кластеризація семантичних векторів:

Слова та фрази перетворюються у вектори за допомогою методів, таких як Word2Vec, FastText або BERT embeddings. Після цього застосовується кластеризація (наприклад, K-means), щоб виділити групи пов'язаних понять.

Виявлення зв'язків та їх типів між поняттями

Після того як поняття виявлені, важливо встановити зв'язки між ними. Це дозволяє створити семантичну модель тексту, яка показує, як різні поняття пов'язані між собою.

Зв'язки між поняттями можуть бути різними за своєю природою:

Семантичні зв'язки: Зв'язки, що відображають схожість або протилежність між поняттями. Наприклад, «безпека» → «захист», «небезпека» → «атака».

Контекстуальні зв'язки: Зв'язки, що виникають у певному контексті. Наприклад, у текстах про інформаційну війну «пропаганда» може бути пов'язана з «маніпуляціями».

Каузальні зв'язки: Зв'язки, що вказують на причину та наслідок.

Наприклад, «фейкові новини» → «дезінформація» → «суспільна паніка».

Хронологічні зв'язки: Зв'язки, що відображають послідовність подій.

Наприклад, «атака» → «реакція» → «результат».

Для виявлення зв'язків між поняттями використовуються такі методи:

- 1) Промпт-інжиніринг, який дозволяє задавати питання про зв'язки між поняттями. Наприклад, при вирішенні завдання виявлення зв'язків понять «фейкові новини» та «політична нестабільність» модель повертає пояснення та може навести приклади. Інший промпт: «Як пов'язані поняття 'фейкові новини' та 'політична нестабільність'?» привдить до результату: «Фейкові новини можуть спричиняти політичну нестабільність через маніпуляцію суспільною думкою та зниження довіри до влади.»
- 2) Аналіз семантичних векторів при якому семантична близькість між словами визначається за допомогою косинусної відстані між їхніми векторами. Наприклад, слова «безпека» та «захист» матимуть високий показник близькості.
- 3) Графове моделювання, при якому поняття представляються як вершини графа, а зв'язки між ними – як ребра. Це дозволяє візуалізувати структуру тексту та аналізувати її властивості (наприклад, кластери, центральні вершини).
- 4) Реляційний аналіз дозволяє виявляти зв'язки між поняттями на основі граматичних конструкцій у тексті. Наприклад, аналіз речень може показати, що «фейкові новини» часто викликають «недовіру».

Екстракція концептів та зв'язків має широке практичне застосування для таких задач, як:

- 1) Аналіз інформаційних кампаній, у тому числі:
 - виявлення центральних тем і понять у пропагандистських матеріалах;

- аналіз змін у наративах з часом.
- 2) Прогнозування сценаріїв, а саме:
- моделювання можливих сценаріїв поширення інформації та контрзаходів;
 - виявлення потенційних загроз на основі зв'язків між поняттями.
- 3) Виявлення нових методів впливу шляхом:
- ідентифікації нових стратегій маніпуляції через аналіз зв'язків між поняттями.
 - прогнозування зміни акцентів у пропагандистських кампаніях.
- 4) Оцінка ефективності інформаційного впливу, тобто
- аналіз того, як змінюються зв'язки між поняттями у текстах до та після інформаційної кампанії.

Наведемо приклад з області інформаційно-психологічних війн, розглянемо текст для аналізу:

«У 2023 році спостерігалось активне використання фейкових новин для маніпуляції суспільною думкою. Основними цілями були політична нестабільність та економічна паніка. Зловмисники використовували соціальні мережі для поширення дезінформації, що призвело до масових протестів у деяких регіонах.»

Наведемо найменування етапів дослідження та їх результатів:

Етап 1. Екстракція понять:

- Іменовані сутності: "2023 рік" (час), "соціальні мережі" (платформа).
- Абстрактні поняття: "фейкові новини", "маніпуляція", "суспільна думка", "політична нестабільність", "економічна паніка", "дезінформація", "масові протести".

Етап 2. Виявлення зв'язків:

- Семантичні зв'язки:
 - «фейкові новини» → «дезінформація».
 - «маніпуляція» → «суспільна думка».
- Каузальні зв'язки:
 - «фейкові новини» → «маніпуляція суспільною думкою».
 - «маніпуляція суспільною думкою» → «політична нестабільність».
 - «політична нестабільність» → «масові протести».

- Хронологічні зв'язки:
 - «2023 рік» → «активне використання фейкових новин».

Етап 3. Візуалізація (графове моделювання):

- Створюється граф, де вершинами є поняття, а ребрами – зв'язки між ними. Наприклад:
 - Вершини: «фейкові новини», «маніпуляція», «суспільна думка», «політична нестабільність», «масові протести».
 - Ребра: «фейкові новини» → «маніпуляція», «маніпуляція» → «суспільна думка», «політична нестабільність» → «масові протести».

Етап 4. Висновки:

- Головна тема: використання фейкових новин для маніпуляції суспільною думкою з метою створення політичної нестабільності.
- Ключові зв'язки: фейкові новини → маніпуляція → політична нестабільність → масові протести.
- Потенційні контрзаходи: блокування джерел дезінформації, просвітницькі кампанії.

Таким чином, генерація концептів і зв'язків є ключовим інструментом для аналізу текстових даних та побудови семантичних моделей. Запропонований підхід, який включає екстракцію понять та виявлення зв'язків, забезпечує глибоке розуміння тексту та дозволяє ефективно вирішувати складні завдання, такі як аналіз інформаційних кампаній, прогнозування сценаріїв та виявлення нових методів впливу.

Промпт-інжиніринг відіграє ключову роль у цьому процесі, оскільки дозволяє точно керувати LLM для отримання бажаних результатів. Комбінація промпт-інжинірингу з іншими методами, такими як семантичний аналіз, векторизація та графове моделювання, робить цей підхід максимально ефективним.

2.5. Методологія «рою віртуальних експертів»

Метод «рою віртуальних експертів» (Swarm of Virtual Experts) – це підхід до автоматизованого семантичного аналізу, в якому LLM використовуються для багаторазового вилучення концептів і зв'язків із тексту, що дозволяє сформулювати зважену семантичну мережу.

LLM стали інструментом для витягу сутностей (концептів) з текстів і зв'язків між ними. Цей процес лежить в основі створення семантичних мереж, які можуть застосовуватися в різних сферах – від аналізу великих даних до побудови систем знань. Однак при використанні цих

моделей виникає кілька суттєвих проблем, таких як галюцинації, зайві або несуттєві зв'язки, а також пропуск важливих понять. Це особливо критично в задачах, пов'язаних із безпекою, де пропущена загроза може призвести до серйозних наслідків.

Для часткового подолання цих проблем було запропоновано підхід, який ґрунтується на множинних викликах LLM (великих мовних моделей) з різними промптами, ролями та моделями, а також на агрегуванні результатів. Тобто, запропонований підхід базується на багаторазовій взаємодії з LLM через різні запити, з різними ролями та на різних моделях. Це дозволяє збільшити ймовірність знаходження релевантних концептів та зв'язків, але після цього виникає питання щодо оптимальної кількості промптів і ролей, а також необхідності людської верифікації результатів.

Важливу роль у процесі семантичного нетворкінгу відіграє віртуальний експерт – образ, який представляє LLM у відповіді на кожен промпт користувача. Кожен запит можна розглядати як думку віртуального експерта, який надає вузли мережі (поняття), а також зв'язки між цими вузлами та їхніми атрибутами. Віртуальні експерти працюють на основі знань, вбудованих у моделі, та текстів або документів, наданих користувачем, що перетворює LLM на своєрідний лінгвістичний процесор, здатний аналізувати, класифікувати та розширювати семантичні мережі.

Справжній прорив у семантичному мережевому аналізі досягається завдяки можливості отримання багатьох відповідей на запити, зроблені в різний час, на різних моделях LLM, з використанням різних підходів до формування промптів. Рій віртуальних експертів об'єднує різноманітність і варіативність думок, забезпечуючи семантичну мережу додатковими шарами точності та повноти. Цей рій формує мережу з концептів, взаємодій або взаємозв'язків, які людина потім може фільтрувати, агрегувати та впорядковувати, підвищуючи достовірність і повноту знань. Кожен запит, кожен «експерт» у рої – це частка загального знання, яка, як краплі в океані, об'єднується для досягнення спільної мети – побудови більш точної та повної картини досліджуваної області.

Додаткове нововведення, запропоноване в цій роботі, полягає у введенні етапу валідації за участю експертів-людей. Це допомагає забезпечити, щоб отримані семантичні мережі були не лише статистично обґрунтованими, але й практично значущими. Такий гібридний підхід – поєднання штучних і людських експертів – може значно підвищити точність і надійність результатів.

Роль людини в цьому процесі, попри потужність і ерудицію LLM, залишається унікальною та незамінною. Саме людина, як ментор, визначає мету, задає напрямки та контролює якість створюваної мережі. Її досвід і знання контексту дозволяють відрізнити значуще від другорядного, визначати ключові вузли та їх зв'язки. У цьому унікальному симбіозі людина не лише використовує LLM для формування мережі, але й направляє рій віртуальних експертів до глибшого розуміння теми.

Рой віртуальних експертів у семантичному нетворкінгу можуть бути використані для аналізу текстів, звітів і даних при вирішенні таких завдань, як виявлення прихованих смислів і взаємозв'язків, моніторингу та оцінка тенденцій, створення рекомендацій і прийняття рішень, ідентифікація та боротьба з дезінформацією, створення навчальних матеріалів і контенту, аналіз великих масивів документів.

Рій віртуальних експертів допомагає виявити смислові зв'язки, які не є очевидними на перший погляд. «Віртуальні експерти» аналізують тексти на різних рівнях абстракції та можуть пропонувати нові вузли та зв'язки, які відображають приховані залежності в даних.

Принципи методології «рою віртуальних експертів» полягають у наступному:

- Мультизапитування LLM – для кожного тексту робиться декілька незалежних запитів щодо ключових концептів і зв'язків між ними.
- Агрегація результатів – кожне поняття та зв'язок отримують вагу залежно від частоти їхньої появи у відповідях LLM.
- Формування семантичної мережі – отримані дані використовуються для побудови графової структури інформаційного простору.
- Ітераційне оновлення – можливість додавання нових зв'язків і коригування структури на основі нових даних.

Метод дозволяє:

- Виявляти приховані закономірності в інформаційних війнах.
- Розпізнавати ключових гравців і агентів впливу.
- Моделювати поширення ідей у суспільстві.
- Автоматизувати аналіз великих масивів даних в рамках OSINT.

Такий підхід значно перевершує традиційні методи семантичного аналізу, оскільки дозволяє отримати більш комплексну картину інформаційного середовища.

Використовуючи масив текстів, наприклад, звіти, новинні статті чи наукові публікації, рій віртуальних експертів допомагає виділяти ключові тенденції та динаміку змін, представляючи їх у вигляді семантичної мережі. Це дозволяє відстежувати розвиток теми та прогнозувати зміни в майбутньому.

Рій експертів може використовуватися при аналізі численних джерел даних, виявляючи патерни та зв'язки, корисні для прийняття рішень. Наприклад, у бізнесі чи політиці можна обґрунтовувати стратегії та рекомендації, спираючись на аналіз звітів і прогнозів, підготовлений штучним інтелектом на основі історичних даних.

Семантичний нетворкінг на основі рою віртуальних експертів дозволяє аналізувати фрагменти тексту, щоб виявляти суперечливу або потенційно неправдиву інформацію. Це особливо важливо для роботи з соціальними медіа, де рій віртуальних експертів може оцінювати достовірність публікацій і пропонувати найбільш ймовірні істинні зв'язки.

Рої віртуальних експертів можуть обробляти та спрощувати складні матеріали, виділяючи ключові концепти та зв'язки, які легше сприймаються аудиторією.

Під час обробки та аналізу великих обсягів інформації, таких як звіти або дослідницькі публікації, рій віртуальних експертів дозволяє систематизувати знання, виділяючи основні теми та концепти, створюючи зв'язки та структури для подальшого аналізу.

В рамках семантичного нетворкінгу рій віртуальних експертів може включати нову інформацію в мережі та оперативно перебудовувати структуру понять, покращуючи мережу при надходженні нових даних. Це робить її гнучкою та актуальною, здатною адаптуватися до змінних умов.

Використання рою віртуальних експертів у семантичному нетворкінгу дозволяє вирішувати ці задачі, витягуючи максимальну цінність із складних інформаційних структур і допомагаючи робити обґрунтовані висновки в різних галузях.

Впровадження великих мовних моделей, таких як GPT, здійснило революцію в задачах обробки природної мови, включаючи витягнення понять та зв'язків із неструктурованих даних. Однак залишається значна проблема – схильність LLM до "галюцинацій", коли модель генерує нерелевантні концепти або не може виявити ключові зв'язки між важливими сутностями у сфері кібербезпеки.

Метод рою віртуальних експертів покликаний вирішити ці проблеми за рахунок використання кількох ітерацій запитів до однієї або кількох

LLM для витягнення концептів та взаємозв'язків. Нещодавні дослідження показують, що подібні методи можуть використовуватися для підвищення точності в інших галузях, таких як медична діагностика та фінансовий аналіз.

Важливою темою є також методика використання ролевих моделей для розширення різноманітності концепцій. Дослідження пропонує методики, схожі на «рій експертів», проте без застосування багаторазового виконання промптів із різними ролями. Також важливим напрямком є включення людського фактора на етапі верифікації.

Підхід, заснований на використанні «рою віртуальних експертів», який надає широкі можливості для гнучкої адаптації до складних сценаріїв.

Таким чином, під «роєм віртуальних експертів» мається на увазі група відповідей, отриманих від великої мовної моделі на серію запитів (промптів), спрямованих на дослідження певної галузі знань. Кожна окрема відповідь від LLM виступає як думка «віртуального експерта», відображаючи інтерпретацію даних моделі на конкретний запит. Однак одна відповідь може містити лише частину можливої інформації, особливо якщо запит вимагає суб'єктивної інтерпретації або комплексного аналізу.

Щоб підвищити точність і повноту семантичної мережі, фахівці застосовують підхід «рою», багаторазово звертаючись до LLM з різними промптами і навіть на різних моделях, щоб накопичити широкую сукупність точок зору та інтерпретацій. Ці відповіді разом можна агрегаційно обробляти, щоб виявляти ключові вузли (поняття) та зв'язки між ними, а також їхні властивості, які найчастіше збігаються серед усіх відповідей. Статистична обробка множини відповідей і залучення «людини-ментора» дозволяє покращити семантичну мережу, виявляючи найважливіші концепти та зв'язки. Таким чином, «рій віртуальних експертів» є потужним інструментом у створенні повноцінних і точних семантичних мереж, пропонуючи симбіоз людського аналізу та можливостей ШІ.

У контексті терміна «рій віртуальних експертів» узгодженість дій віртуальних експертів – важливий аспект, який дозволяє створювати цілісні та достовірні семантичні мережі. Незважаючи на те, що кожен «експерт» є окремим запуском промпта, існує кілька механізмів, які сприяють узгодженості їхньої роботи та допомагають створювати злагоджені результати, а саме:

- Частота повторюваності вузлів і зв'язків у відповідях на схожі запити дозволяє виділяти більш значущі та стійкі концепти. Чим

частіше зустрічається певний вузол або зв'язок, тим вища ймовірність його значущості для побудови мережі.

- Структурування результатів на основі ієрархії концептів допомагає вибудовувати мережу з пріоритетом ключових понять. Це дозволяє виділити вузли вищого рівня, які об'єднують під собою пов'язані, але менш значущі елементи.
- У сучасних LLM часто зберігається інформація про попередні запити в межах однієї сесії, що допомагає моделям підтримувати певну узгодженість у відповідях. Крім того, використання кумулятивних запитів, які уточнюють або розвивають попередні, дозволяє уточнювати зв'язки та корегувати мережеві вузли.
- Людина виступає як ментор рою віртуальних експертів, задаючи послідовність запитів і відстежуючи відповіді. Це дозволяє корегувати питання та додавати нові, за потреби перевіряючи гіпотези. Людський аналіз дозволяє уникнути випадкових відповідей, відкидаючи нерелевантні зв'язки та підкріплюючи найбільш значущі.
- Для різних вузлів і зв'язків можуть призначатися вагові значення, які відображають їхню значущість на основі повторюваності або унікальності. При цьому зв'язки з високими ваговими значеннями розглядаються як надійніші, що сприяє створенню більш точної та узгодженої семантичної мережі.

Ці механізми дозволяють «рою віртуальних експертів» бути не просто випадковим набором думок, а цілеспрямованим процесом формування надійних знань. Завдяки узгодженості дій рій віртуальних експертів у підсумку будує повну та стійку семантичну мережу, де людина та ШІ діють як єдина система.

Щоб збільшити різноманітність і точність витягнутих зв'язків, запропоновано декілька методик, серед яких вилучення сутностей і зв'язків за допомогою LLM, використання ролей, застосування різних моделей LLM, агрегація результатів, верифікація людиною.

На рис. 2.4 наведено схему, яка ілюструє взаємодію між агентами, ролями, клієнтом і LLM у межах рою віртуальних експертів.

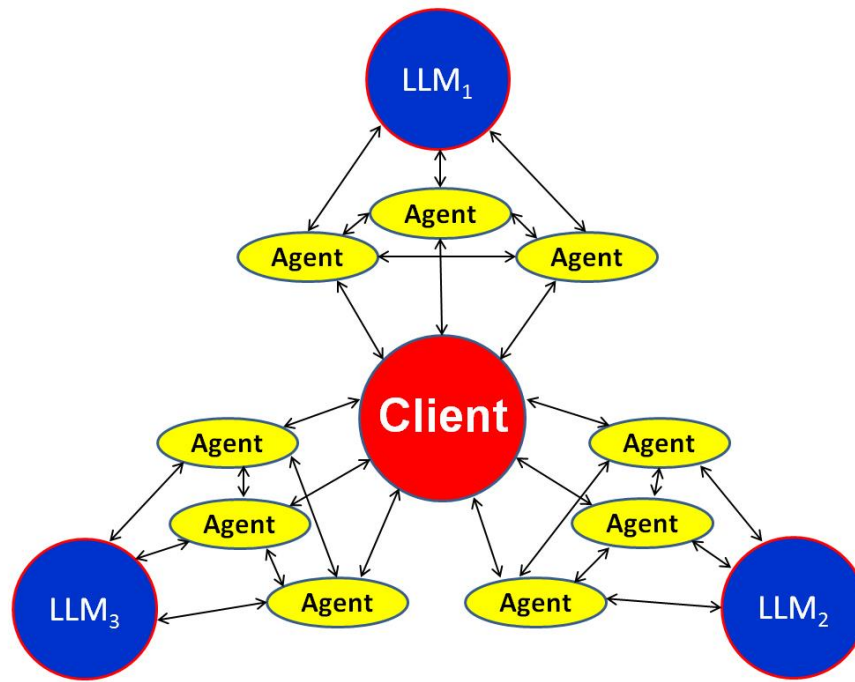


Рисунок 2.4: Схема рою віртуальних агентів

При цьому розглядаються такі процеси:

1. Вхідний запит q_k передається клієнтом до рою агентів.
2. Кожен агент отримує запит і генерує відповідь із урахуванням ролі.
3. Відповіді агентів обмінюються між собою (взаємодія).
4. Клієнт синхронізує та агрегує відповіді від різних LLM.
5. Результат узгоджується із цілями.

Як бачимо, рій утворюється також завдяки зв'язку між агентами та клієнтом (людиною), який координує процес. Основними аспектами зв'язку є обмін інформацією між агентами, загальний клієнтський контекст, агрегація результатів людиною. Один агент може враховувати відповідь іншого для уточнення свого результату. Агенти працюють із загальною метою, заданою клієнтом, але виконують різні ролі. Результати кожного агента аналізуються та узгоджуються на рівні клієнта для формування єдиного рішення.

Рій віртуальних агентів можна формалізувати як множину взаємозв'язаних агентів:

$$A = \{a_1, a_2, \dots, a_n\},$$

де кожен елемент a_i позначає окремого віртуального агента.

Кожен агент має здатність реагувати на зовнішні запити середовища. Позначимо через q_k k -й запит, який надходить у систему. Відповідь агента a_i на цей запит позначимо як r_{ik} .

Особливістю моделі є те, що відповідь агента не є автономною, а формується з урахуванням результатів інших агентів. Інакше кажучи, відповідь r_{ik} є функцією не лише самого запиту q_k , але й відповідей решти агентів на цей запит.

Позначимо через $R_{-i,k}$ сукупність відповідей усіх агентів, окрім агента a_i , на k -й запит. Тоді відповідь конкретного агента може бути записана як:

$$r_{ik} = f(a_i, q_k, R_{-i,k}),$$

де a_i це характеристики або внутрішній стан агента, q_k – поточний запит, $R_{-i,k}$ – результати інших агентів на цей самий запит, $f(\cdot)$ – функція, що визначає логіку прийняття рішення агентом.

Така модель моделює групове прийняття рішень, де кожен агент «слухає» інших перед тим, як висловити свою думку. Це нагадує колективне мислення, де кінцева відповідь системи є результатом консенсусу або узгодження між учасниками.

Цей принцип можна порівняти з роботою експертної комісії: коли кожен спеціаліст аналізує питання окремо, але перед остаточним висновком знайомиться з думками своїх колег, щоб уточнити або підтвердити власну позицію. Такий підхід забезпечує вищу якість рішень завдяки взаємному контролю та корекції помилок.

Основою методу є послідовне вилучення концептів і відносин між ними з тексту. Для цього використовується один і той самий промпт, який багаторазово переформулюється і запускається, щоб мінімізувати вплив повторюваних відповідей. Повторне виконання промпта допомагає доповнити систему, що збільшує ймовірність знаходження більш релевантних зв'язків. Perezавантаження системи дозволяє генерувати більш різноманітні результати, уникнувши повторень і посиливши ймовірність вилучення унікальних і релевантних зв'язків.

Для різноманітності відгуків і мінімізації помилок була додана методика зміни ролей. Кожен запит до LLM формулюється з урахуванням ролі, наприклад, від імені експерта в певній галузі, що сприяє екстрагування нових понять і зв'язків. Різноманітність моделей дозволяє зменшити вплив особливостей конкретної моделі на результат і підвищити достовірність витягуваних даних. Кожна модель генерує свої унікальні сутності та зв'язки, що підвищує повноту отримуваних семантичних мереж.

Після збору даних від різних моделей і в різних ролях проводиться їх агрегація. Зв'язкам приписуються ваги, що відповідають частоті їх

появи в різних мережах. Зв'язки з низькими частотами можуть бути виключені, щоб уникнути хибних або малозначущих даних.

На завершальному етапі пропонується включення людини-експерта для верифікації мереж. Цей крок необхідний для мінімізації помилок, допущених LLM, особливо у випадках, коли важливі зв'язки пропущені, а неіснуючі – додані.

Визначення ролей – важливий аспект методології "рою віртуальних експертів", який дозволяє витягувати сутності та зв'язки в різних контекстах і з різних точок зору. Виділення ролей забезпечує багатоаспектне сприйняття даних, а також допомагає створювати семантичну мережу, здатну інтегрувати різноманітні інтерпретації та підходи до однієї й тієї самої інформації.

Застосування різних ролей дозволяє віртуальним експертам формувати мережу сутностей та зв'язків, які включають різноманітні підходи та точки зору. Такий багатоаспектний аналіз робить семантичну мережу більш цілісною та гнучкою. Людина, як диригент, визначає загальну спрямованість та корегує ролі за необхідності, але LLM, виступаючи як віртуальні експерти, доповнюють процес, ініціюючи та пропонуючи нові, потенційно значущі ролі.

Таким чином, розподіл ролей між людиною та LLM не лише забезпечує глибину аналізу, але й дозволяє створювати семантичні мережі, здатні до більш точного та різнобічного відображення складних взаємозв'язків у досліджуваному контексті. Одним із ключових питань є визначення оптимальної кількості ролей віртуальних експертів. Надмірна кількість ролей може призводити до зниження якості мережі через інформаційний шум, тоді як недостатня кількість ролей не дозволить виявити всі можливі зв'язки. Для цього може використовуватися підхід статистичного аналізу приросту нових зв'язків. У запропонованій моделі "рою віртуальних експертів" оцінка кількості ролей та їх підтвердження людиною є важливими елементами для підвищення точності та надійності семантичних мереж, створюваних за допомогою LLM.

Людина, як ментор, може визначати ключові ролі, ґрунтуючись на цілях аналізу та специфіці предметної області. Ці ролі можуть представляти різні перспективи, підходи та контексти, такі як експертна, аналітична та оціночна ролі.

- Експертна роль полягає у вилученні сутностей та зв'язків із акцентом на наукові чи технічні терміни, підходи та концепції.
- Аналітична роль полягає в аналізі взаємозв'язків з точки зору причинно-наслідкових залежностей, послідовності подій тощо.

- Оціночна роль полягає у виявленні сутностей та зв'язків із фокусом на пріоритетах, значущості та вагах.

Задаючи конкретні ролі, людина формує спрямованість аналізу, вказуючи, які концепти та відносини є найважливішими для виділення. Це допомагає рою концентруватися на вилученні інформації в межах заданих перспектив та інтерпретацій.

Великі мовні моделі також можуть пропонувати власні ролі та контексти, які людина могла б не враховувати. LLM можуть розпізнавати та пропонувати ролі на основі структури даних або тексту, аналізу частотності понять або навіть на основі вбудованих семантичних категорій. Наприклад, модель може запропонувати історичні, соціальні та міждисциплінарні ролі.

- Історична роль інтерпретує інформацію з акцентом на історичний контекст та розвиток понять.
- Соціальна роль передбачає акцент на зв'язках між соціальними аспектами даних, аналіз інформації в рамках взаємодій та їх значущості в суспільстві.
- Міждисциплінарна роль забезпечує вилучення даних з точки зору взаємозв'язку різних дисциплін або підходів.

LLM можуть виявляти такі ролі автоматично, використовуючи свої вбудовані знання та аналітичні можливості, що відкриває для рою віртуальних експертів можливість бачити дані з несподіваних та інноваційних перспектив.

Ролі експертів у моделі слугують для збільшення різноманітності відповідей. Кожна роль задає інший контекст або точку зору, з якої LLM (велика мова моделей) може розглядати концепти та їх зв'язки. Це допомагає мінімізувати вплив можливих помилок або пропусків, які можуть виникнути при запитах від однієї ролі. Важливо знайти оптимальну кількість ролей. Надмірне збільшення кількості ролей може призвести до генерації зайвих, нерелевантних зв'язків, тоді як замала кількість не забезпечить достатнього різноманіття відповідей. Оптимум можна шукати через статистичний аналіз: після кількох ітерацій промптів із різними ролями можна спостерігати, за якої кількості ролей з'являється найбільш стабільна та релевантна мережа зв'язків.

Метод підрахунку ролей може базуватися на концепції збіжності. Якщо при додаванні нової ролі частота появи нових концептів і зв'язків починає зменшуватися, це може сигналізувати про те, що більшість релевантних концептів уже враховано, і подальше додавання ролей стає малоефективним. Це можна формалізувати через спостереження за приростом нових концептів з кожною додатковою роллю та

встановленням порогу, при якому додавання нових ролей стає статистично незначущим.

Роль людини в моделі є фінальною перевіркою агрегованої семантичної мережі. Після того як LLM обробить численні промпти від різних ролей, людина-експерт виконує верифікацію, аналізуючи зв'язки між концептами, їх вагу та релевантність. Людина-експерт допомагає усунути галюцинації LLM, неважливі або малозначущі зв'язки, а також заповнити пропуски важливих зв'язків. Для зручності валідації результатів можна впровадити інтерактивні інтерфейси, які дозволять людині спостерігати за вагами зв'язків (отриманими на основі частоти їх виявлення віртуальними експертами) та корегувати їх за потреби.

Таким чином, система отримує додатковий рівень надійності завдяки включенню людини-експерта, який підтверджує або коригує результати, надані віртуальними експертами.

Надамо коротку математичну формалізацію концепції «рою віртуальних експертів», для чого введемо позначення:

- $S = \{s_1, s_2, \dots, s_n\}$ – множина сутностей, що екстраговані з тексту;
- $R = \{r_{ij}\}$ – множина зв'язків між сутностями, де r_{ij} позначає зв'язок між сутностями s_i та s_j ;
- P – набір промптів для витягу зв'язків;
- $M = \{M_1, M_2, \dots, M_k\}$ – набір різних моделей LLM;
- $W(r_{ij})$ – вага зв'язку між сутностями s_i та s_j , яка залежить від частоти повторення цього зв'язку.

Для кожного промпту $p \in P$ і кожної моделі $M_i \in M$ виконується витягнення множини сутностей S_p^l та множини зв'язків R_p^l . Крім того, поновлюється вага зв'язку: $W(r_{ij}) = W(r_{ij}) + 1$, якщо $r_{ij} \in R_p^l$.

Зв'язок r_{ij} включається до підсумкової мережі, якщо його вага перевищує певний поріг $W(r_{ij}) \geq \theta$, де θ поріг визначається як статистичний критерій, заснований на кількості запусків віртуальних експертів і частоті зв'язків.

Для визначення оптимального порогу θ пропонуються наступні методи:

- Нехай $f(r_{ij})$ – частота появи зв'язку r_{ij} в результатах різних промптів і моделей.
- Нехай T – загальна кількість викликів (промптів і моделей).

- Визначимо частотний розподіл зв'язків $P(f)$, який показує ймовірність появи зв'язків з різною частотою.

Щоб вибрати оптимальний поріг Θ , можна використовувати критерій значимості, заснований на кількості запусків і дисперсії частот зв'язків:

$$\theta = \mu f + \alpha \sigma_f,$$

де μf – середнє значення частоти появи зв'язків, σ_f – стандартне відхилення частоти, α – коефіцієнт значимості (наприклад, $\alpha = 1.5$, щоб відсіяти менш значимі зв'язки).

Цей поріг дозволяє відсіювати зв'язки, які є випадковими і не мають значного ваги.

Кількість викликів віртуальних експертів N прямо впливає на достовірність побудови мережі. Для оцінки оптимальної кількості запусків можна використовувати метод стабілізації частот зв'язків.

Нехай N – кількість запусків, $f(r_{ij})$ – частота зв'язку r_{ij} після N запусків.

Якщо при збільшенні N приріст кількості нових зв'язків $\Delta f(r_{ij})$ стає незначним (наприклад, $\Delta f(r_{ij}) < \varepsilon$, де ε – мале число), це вказує на досягнення стабільності. Таким чином, можна визначити оптимальне значення N :

$$N = \min \left\{ N' : \frac{1}{|R|} \sum_{r_{ij} \in R} \Delta f(r_{ij}) < \varepsilon \right\}.$$

Цей критерій дозволяє вибрати N , при якому модель стабілізується, і нові експерти не додають значущих зв'язків.

Наведемо формалізацію для оцінки кількості ролей та підтвердження:

1. Нехай R – кількість ролей, від яких запитуються дані у LLM. Для кожної ролі i , отримуємо множину концептів C_i та множину зв'язків між ними S_i .
2. Агрегована множина зв'язків: $S = \bigcup_{i=1}^R S_i$, кожному зв'язку призначається вага $w(s)$, яка дорівнює частоті появи зв'язку серед ролей, а саме:

$$w(s) = (\text{кількість ролей, у яких зв'язок } s \text{ зустрічається } w) / R$$
3. Вводиться порогове значення τ , які зв'язки вважаються значущими. Наприклад, якщо зв'язок зустрічається у менш ніж τ долі ролей, він відкидається:

$$S^* = \{s \in S \mid w(s) \geq \tau\}.$$

Для оцінки необхідної кількості ролей R^* можна спостерігати за приростом нових зв'язків із додаванням нових ролей. Якщо приріст кількості нових концептів і зв'язків стає незначним (наприклад, $\frac{\Delta S_i}{\Delta R} \rightarrow 0$), можна зупинити збільшення кількості ролей.

4. Після побудови мережі S^* , людина-експерт проводить фінальну верифікацію, підтверджуючи або відхиляючи запропоновані зв'язки, що забезпечує фінальну корекцію мережі.

Для визначення оптимального порогу частотності зв'язків можна здійснити наступним чином:

- Проводиться серія експериментів із різними значеннями порогу, починаючи з мінімальних значень (врахування всіх зв'язків) до збільшених порогів.
- Зв'язок із вагою 1 може виявитися найчастішим при малій кількості експертів. Щоб уникнути цього, у вибірці експериментів із різною кількістю ролей і різними порогамі τ проводиться аналіз, який дозволяє оцінити, при якому значенні порогу приріст нових зв'язків стає незначним.
- Це значення і вважається оптимальним.

Людина-експерт підтверджує або коригує результати, надані віртуальними експертами.

2.4. Практична частина. Використання LLM для семантичного аналізу текстів

Ця практична частина демонструє, як великі мовні моделі можуть бути застосовані для аналізу текстів у контексті інформаційно-психологічних війн. Буде показано, як виконати семантичний аналіз текстів за допомогою відкритих API, таких як Hugging Face Inference API або Groq API. Цей підхід дозволяє працювати з великими мовними моделями без необхідності завантаження та налаштування локальних моделей.

Підготовка середовища

Для початку потрібно налаштувати середовище для роботи з онлайн-API. Для цього:

1. Зареєструйтесь на платформі, наприклад, Hugging Face або Groq, і отримайте API-ключ.
2. Встановіть необхідні бібліотеки:

pip install requests

Формування промптів

Для аналізу текстів ми будемо використовувати різні мовні моделі. Наприклад, GPT-4.

Приклад 1: Визначення ключових тем

Ми маємо повідомлення із мережі Інтернет:

Alibaba випустила нову версію штучного інтелекту – стверджує, що краща за DeepSeek
29.01.2025 19:20

Китайська технологічна компанія Alibaba у середу випустила нову версію своєї моделі штучного інтелекту (ШІ) Qwen 2.5-Max.

Про це повідомляє Reuters, передає Укрінформ.

У компанії стверджують, що нова версія Qwen 2.5-Max перевершує "майже за всіма параметрами" інший китайський ШІ-стартап DeepSeek, а також передові моделі OpenAI та Meta – GPT-4o і Llama-3.1-405B.

Випуску Qwen 2.5-Max відбувся у перший день Нового року за місячним календарем, коли більшість китайців не працюють. Зазначається, що це свідчить про тиск стрімкої популярності стартапу DeepSeek не лише на іноземних, а й на внутрішніх конкурентів.

Як повідомляв Укрінформ, китайський стартап DeepSeek зі своєю моделлю генеративного штучного інтелекту DeepSeek-V3 випередив у Сполучених Штатах свого конкурента ChatGPT і став найпопулярнішим безкоштовним застосунком, доступним в американському Apple App Store.

DeepSeek став першим стартапом, який отримав високі оцінки від технологічної індустрії США за рівень продуктивності, що відповідає і навіть перевершує найсучасніші американські моделі штучного інтелекту.

Укрінформ 2025.01.29 19:31
<https://www.ukrinform.ua/rubric-world/3954239-alibaba-vipustila-novu-versiu-stucnogo-intelektu-stverdvue-so-krasa-za-deepseek.html>

Нехай перед нами стоїть задача визначення основних тем із завантаженого тексту. Для цього до великої мовної моделі формується запит (промпт):

Визнач ключові теми у введеному тексті

Після завантаження тексту вихідного файлу в систему ChatGPT (позначка «+») і наведеного промпту, велика мовна модель надає відповідь:

У тексті з документа виділяються такі ключові теми:

1. **Випуск нової моделі ШІ від Alibaba** – компанія Alibaba презентувала нову версію своєї моделі штучного інтелекту Qwen 2.5-Max.
2. **Порівняння з конкурентами** – стверджується, що Qwen 2.5-Max перевершує за більшістю параметрів інші моделі, зокрема DeepSeek, GPT-4o (OpenAI) та Llama-3.1-405B (Meta).
3. **Конкуренція на ринку ШІ** – випуск у святковий день символізує тиск конкуренції, особливо з боку стартапу DeepSeek.

4. **Успіх DeepSeek у США** – модель DeepSeek-V3 стала найпопулярнішим безкоштовним застосунком в Apple App Store США, обійшовши ChatGPT.
5. **Визнання китайських ШІ-технологій** – DeepSeek отримав високу оцінку у США за свою продуктивність, яка відповідає або перевершує західні моделі.

Якщо потрібно, можу сформулювати ці теми як теги або розгорнути кожен з них для подальшого аналізу.

Приклад 2: Виявлення персон

Із наданого тексту повідомлення виявити персони, що згадуються. Ось це повідомлення:

У "Коаліції охочих" визначили чотири чіткі цілі для України

Ці країни разом виступають як єдине ціле і готові забезпечити майбутнє для України після досягнення мирної угоди про припинення війни.

Плани є добре розробленими і коаліція має чіткі цілі для України / фото facebook.com/NATO

У "Коаліції охочих" визначили чіткі чотири цілі для якнайсильнішого посилення України і формування "гарантійних сил" безпеки, які складатимуть контингент іноземних сил. Як передає кореспондент УНІАН, про це міністр оборони Великої Британії Джон Гілі сказав на початку засідання "Коаліції охочих" на рівні міністрів оборони у Брюсселі в штаб-квартирі НАТО.

Гілі розпочав з подяки двом сотням військових планувальників з 30 країн, які в останні тижні наполегливо працювали. За його словами, ці країни разом виступають як єдине ціле і готові забезпечити майбутнє для України після досягнення мирної угоди про припинення війни.

...

УНІАН 2025.04.10 18:13

<https://www.unian.ua/world/u-koaliciji-ohochih-viznachili-chotiri-chitki-cili-dlya-ukrajini-12973086.html>

Для вирішення поставленої задачі формується промпт:

Визнач всіх персон, що згадуються у повідомленні

У результаті опрацювання цього промпту система LLM, зокрема, Qwen-2.5 надає відповідь:

У цьому повідомленні згадуються наступні персони:

- Джон Гілі (John Healey) - міністр оборони Великої Британії
- Президент Трамп (Donald Trump) - колишній президент США (згаданий в контексті)
- Президент Путін (Володимир Путін) - президент Росії
- Олександр Сирський - Головнокомандувач Збройних сил України
- Володимир Зеленський - президент України

Також стаття згадує "дві сотні військових планувальників з 30 країн" і "багато міністрів", які беруть участь у зустрічах, але конкретно не називає їх.

Приклад 3: Виявлення зв'язків між поняттями

Для вирішення цієї задачі аналізуємо повідомлення з мережі Інтернет:

Шукати шляхи нової пропаганди

20.11.2024

Пропаганда була, є і буде. Це "мрія", яку продають суспільству. Але вона завжди спирається на емоції, лише мімікруючи під серйозну розмову. Пропаганда сильна лише тоді, коли за нею стоять реальні дії.

Сьогоднішня війна багато в чому "знищила" статус старих інструментів управління, але коли вона закінчиться, вони знову прийдуть. З цієї причини треба бути готовим до них. І одним із таких інструментів була і буде пропаганда. Масова свідомість відрізняється від індивідуальної, вона потребує свого підходу. І пропаганда якраз є розмовою з масовою свідомістю. Пропаганда завжди супроводжувала людину. Саме слово, як і формальний інститут пропаганди, з'явилося давним-давно в надрах католицької церкви для навернення невірників до віри. Ігнатій Лойола шукав і знаходив, як робити це найефективніше.

...

Академія Української Преси 2024.11.20 12:12

<https://www.aup.com.ua/news/shukati-shlyakhi-novoi-propagandi/>

Промпт до LLM:

Вияви зв'язки між основними поняттями поняттями у введеному тексті, наведи їх у вигляді пар зв'язаних понять у форматі «поняття;поняття», наприклад, «хмара;дощ». Ось текст: [текст повідомлення]

Після завантаження тексту через промпт у систему ChatGPT, ця велика мовна модель надає відповідь:

пропаганда;мрія
пропаганда;емоції
пропаганда;реальні дії
пропаганда;масова свідомість
пропаганда;віртуальний простір
пропаганда;революції
масова свідомість;емоції
масова свідомість;візуальність
масова свідомість;управління
масова свідомість;реклама
масова свідомість;публік рілейшнз
війна;інформаційна війна
війна;когнітивна війна
війна;ментальна війна
війна;свідомість
війна;віртуальність
війна;емоції
війна;зміна мрії
війна;реальність
емоції;свідомість
емоції;управління
емоції;революція

...

Візуалізація результатів

Для кращого розуміння результатів аналізу важливо візуалізувати їх. Наприклад, можна побудувати граф взаємозв'язків понять, для чого

скористуємось сервісом Semantic Networking (<https://bigsearch.space/uli.html>), стартову сторінку якого наведено на рис. 2.5.

Після опрацювання сервісом Semantic Networking (<https://bigsearch.space>), на екрані користувача (рис. 2.6) виводяться поняття як вузли мережі, кількість зв'язків кожного, поняття, деякі статистичні показники і сама мережа в інтерактивному поданні (натискання на вузли або ребра) призводять до переходу за гіперпосиланнями на результати пошуку зовнішніх пошукових систем.

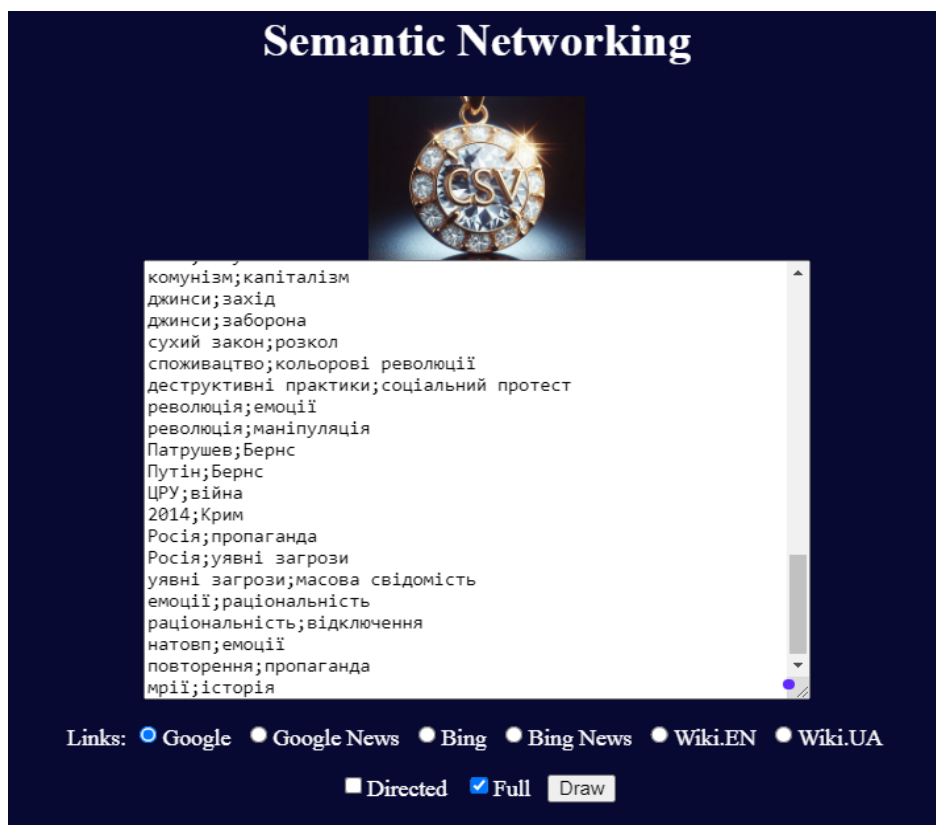


Рисунок 2.5: Форма введення результату роботи LLM для подальшого аналізу і візуалізації

Висновки до розділу

Цей розділ присвячено великим мовним моделям та їхній ролі у аналізі, генерації та моделюванні інформаційно-психологічних війн (ІПВ). Тут були детально розглянуті архітектура LLM, методи їх навчання, механізми обробки текстової інформації, а також практичне застосування для виявлення концептів, зв'язків та наративів. Особливу увагу приділено можливостям та обмеженням цих моделей у контексті аналізу інформаційних потоків.

Великі мовні моделі, такі як GPT-4, Qwen, DeepSeek-V3, Llama-4 та інші, базуються на архітектурі Transformer, яка дозволяє ефективно обробляти великі обсяги текстових даних.

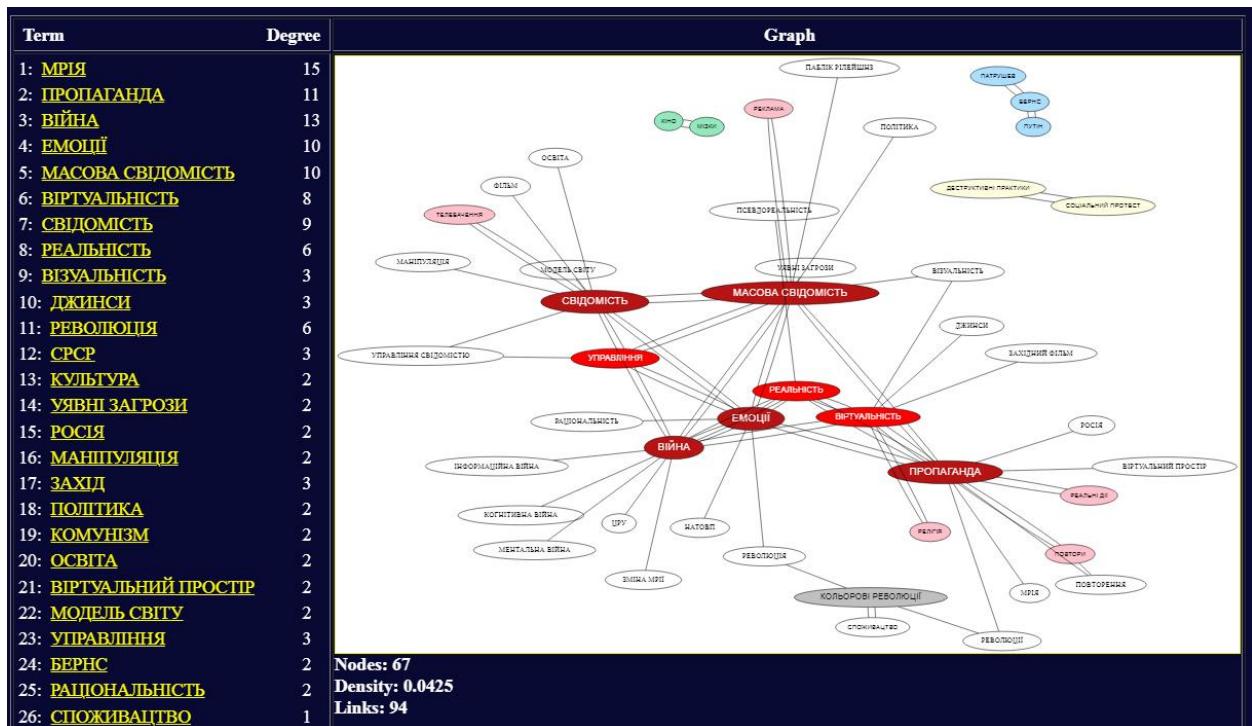


Рисунок 2.6: Результат візуалізації

Основні характеристики цих моделей:

- Навчання без вчителя – моделі навчаються на великих обсягах неструктурованого тексту, використовуючи завдання маскування чи прогнозування слів.
- Механізм уваги – багатоголовий механізм уваги дозволяє моделям одночасно враховувати різні частини тексту, що забезпечує краще розуміння контексту.
- Авторегресійний підхід – кожне наступне слово генерується на основі попередніх слів, що дозволяє створювати довгі та змістовні тексти.

Ці особливості роблять LLM гнучкими інструментами для широкого спектру завдань, включаючи автоматичний переклад, створення резюме, класифікацію текстів та генерацію контенту.

Одним із ключових завдань LLM є генерація нових концептів та виявлення зв'язків між ними. Це завдання включає два основні аспекти:

- екстракція сутностей. Виділення іменованих сутностей (наприклад, організацій, часу, локацій) та абстрактних понять (наприклад, "захист", "тренінг");

- виявлення зв'язків. Визначення семантичних зв'язків між поняттями, наприклад, через аналіз граматичних конструкцій або графове моделювання.

Для покращення точності генерації концептів запропоновано використовувати множинні промпти (запити) з різними параметрами та агрегувати результати. Це дозволяє збільшити ймовірність знаходження релевантних концептів, хоча людська верифікація залишається важливою.

LLM мають широке застосування у вивченні інформаційно-психологічних війн:

- виявлення ключових наративів. Моделі можуть аналізувати великі масиви тексту для виявлення домінуючих тем та наративів;
- прогнозування інформаційних атак. LLM допомагають прогнозувати розвиток інформаційних кампаній, виявляючи закономірності у поширенні дезінформації;
- створення симуляцій. Моделі можуть моделювати сценарії інформаційних операцій для оцінки їхньої ефективності.

Однак існує проблема "галюцинацій" – ситуацій, коли модель генерує нерелевантні концепти або не може виявити ключові зв'язки. Для подолання цієї проблеми запропоновано комбінувати автоматизовані методи з людською експертизою.

У розділі також представлено модель, яка включає шість основних сутностей для аналізу текстових даних:

1. Токени. Найменші одиниці тексту (символи, частини слів, слова).
2. Слова. Леми або нормальні форми слів.
3. Сутності. Конкретні поняття, такі як локації, події, технології.
4. Наративи. Системи оповіді, які формулюються словами.
5. Події. Реальні або уявні події, пов'язані з часом, місцем та учасниками.
6. Документи. Збірки тексту, такі як статті, звіти або книги.

Ця модель забезпечує комплексний аналіз текстових даних, враховуючи взаємозв'язки між різними рівнями тексту. Вона дозволяє:

- ліпше розуміти структуру та зміст тексту;
- виявляти ключові сутності, наративи та події;
- генерувати високоякісний текст, який відповідає контексту;
- аналізувати інформаційні потоки та виявляти потенційні загрози.

До основних переваг LLM відносяться:

- моделі можуть застосовуватися до широкого кола завдань, включаючи аналіз наративів, виявлення дезінформації та прогнозування трендів;
- здатність аналізувати контекст та виявляти складні зв'язки між поняттями;
- моделі можуть обробляти великі обсяги даних за короткий час.

Основні обмеження застосування LLM наступні:

- галюцинації. Моделі іноді генерують нерелевантні концепти або не можуть виявити ключові зв'язки;
- потреба в людській верифікації. Результати моделей часто потребують перевірки експертами;
- залежність від якості даних. Точність моделей залежить від якості та різноманітності навчальних даних.

Великі мовні моделі є потужним інструментом для аналізу та моделювання інформаційно-психологічних війн. Вони дозволяють:

- ефективно обробляти великі обсяги текстових даних;
- генерувати нові концепти та виявляти зв'язки між ними;
- прогнозувати розвиток інформаційних атак та створювати симуляції для оцінки їхньої ефективності.

Однак для досягнення найкращих результатів важливо комбінувати автоматизовані методи з людською експертизою. Це дозволить мінімізувати ризики "галюцинацій" та забезпечити надійний аналіз інформаційних потоків.

Питання для самоконтролю

1. Що таке трансформерна архітектура, і як вона використовується для аналізу текстових даних?
2. Як працює механізм самоуваги (self-attention), і чому він важливий для обробки природної мови?
3. Що таке позиційне кодування, і чому воно необхідне для трансформерних моделей?
4. Як працює повнозв'язний шар (Feed-Forward Network) у трансформерній архітектурі?
5. У чому полягає концепція "рою віртуальних експертів"?
6. Як визначається оптимальний поріг для відсіву незначущих зв'язків у методі рою віртуальних експертів?
7. Як людська верифікація допомагає покращити точність результатів, отриманих від LLM?

8. Які основні етапи процесу побудови семантичної мережі за допомогою LLM?
9. Як використовується промпт-інжиніринг для екстракції нових сутностей із тексту?
10. Які існують методи для зменшення впливу повторюваних відповідей при багаторазовому виконанні промптів?
11. Як визначається вага зв'язку між сутностями у методі рою віртуальних експертів?
12. Як можна виявити ключові наративи у великому масиві текстових даних за допомогою LLM?
13. Як можна оцінити стабільність результатів при збільшенні кількості запусків LLM?
14. Які існують методи для виявлення "галюцинацій" у відповідях LLM?

Література за розділом

1. Dmytro Lande, Leonard Strashnoy. GPT Semantic Networking: A Dream of the Semantic Web – The Time is Now – Kyiv: Engineering, 2023. – 168 p. ISBN 978-966-2344-94-3, DOI: 10.5281/zenodo.14278893
2. Ланде Д.В., Страшной Л.Л. Семантичний нетворкінг на основі великих мовних моделей : монографія. Київ: Інжиніринг, 2025. – 274 с. ISBN 978-617-8180-01-0
3. OSINT у кібербезпеці : навч. пос. / Ланде Д.В. – Київ: ТОВ "Інжиніринг", 2024. – 522 с. ISBN 978-966-2344-97-4
4. Stephen Wolfram. What Is ChatGPT Doing ... and Why Does It Work? Wolfram Media, Inc., March 9, 2023. 112 p. ISBN 978-1-57955-081-3
5. Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, Illia Polosukhin. Attention Is All You Need. 2017. arXiv: 1706.03762. ArXiv Preprint. DOI: 10.48550/arXiv.1706.03762
6. Srinivasan Ramanujam. The LLM Revolution: Transforming Industries with Large Language Models. Kindle Edition, 2024. – 160 p. ASIN B0D7VS5BNK

3. Методи побудови мережевих моделей

Для дослідження інформаційно-психологічних операцій і ментальних війн передусім необхідно моделювання інформаційного середовища, для чого застосовуються графові структури – мережі, які дозволяють описувати взаємодію між суб'єктами інформаційного впливу, простежувати поширення наративів і виявляти ключові вузли, що визначають динаміку інформаційних процесів.

Методи побудови мережевих моделей розвивалися від жорстких ієрархічних графових структур до динамічних адаптивних мереж, що здатні змінювати свою конфігурацію в залежності від потоку інформації, застосовуються методи семантичного розширення графів та інтеграції нових зв'язків. Ці підходи дозволяють створювати гнучкі аналітичні системи для виявлення та прогнозування впливів у цифровому просторі.

3.1. Інтеграція нових понять і зв'язків

В реальності інформаційні мережі є динамічними – їхня структура постійно змінюється, тому важливою задачею є автоматизована інтеграція нових понять і зв'язків у мережу для підтримки актуальності її змісту.

У контексті інформаційно-психологічних війн великі мовні моделі, такі як GPT-4, LLaMA, або Qwen можуть допомогти виявити приховані патерни, стратегії ворожих дій та методи протидії. Далі розглянемо, як LLM допомагають інтегрувати нові поняття і зв'язки, а також наведемо конкретні приклади із реальних інформаційно-психологічних війн.

Великі мовні моделі можуть аналізувати величезні масиви текстової інформації та запропонувати нові поняття і зв'язки на основі закономірностей у даних. Наприклад, при аналізі соціальних мереж LLM може виявити, що певні типи наративів (наприклад, «дискредитація влади») часто поширюються через конкретні платформи (наприклад, Telegram). Також LLM може запропонувати нові механізми впливу, наприклад, в задачі використання штучного інтелекту для створення персоналізованих фейкових новин.

Великі мовні моделі мають здатність «розуміти» контекст та надавати детальні пояснення нових зв'язків. Наприклад, якщо модель виявляє зв'язок між «генерацією штучних мемів» та «соціальними мережами», вона може пояснити, як алгоритми рекомендацій сприяють їх поширенню серед цільової аудиторії.

Великі мовні моделі можуть одночасно перевіряти дані з різних джерел, наприклад, порівняти інформацію з наукових статей, новин та

експертних звітів, або запропонувати застосувати додаткові види джерел для перевірки достовірності нового поняття або зв'язку.

LLM можуть допомогати знаходити нові поняття та формулювати нові зв'язки, інтегрувати їх у існуючу мережу, враховуючи тип зв'язків, напрямки та вагові значення зв'язків і понять. Це дозволяє прискорити процес розширення мережі.

3.2. Приклади із інформаційно-психологічних війн

Розглянемо деякі приклади із сфери ІПВ і відповідні процеси виявлення понять та інтеграції нових зв'язків.

Приклад 1: Російська дезінформаційна кампанія проти України

Задача: Виявити нові поняття і зв'язки у семантичній мережі, пов'язані з російською дезінформаційною кампанією.

Процес:

1. Аналіз даних. Мовна модель проаналізувала тексти з російських медіа, соцмереж та офіційних документів.
2. Генерація гіпотез. Модель запропонувала новий зв'язок: «використання ботнетів → поширення наративів про "нацистську владу в Україні"».
3. Перевірка достовірності. Експерти підтвердили, що ботнети активно використовувалися для поширення таких наративів у Twitter та Facebook.
4. Інтеграція у мережу. Зв'язок був доданий до мережі як «ботнети» → «поширюють наративи про "нацистську владу"» з вагою 0,9.

Результат: Ці поняття і зв'язок допомагають зрозуміти механізм масштабного поширення пропаганди та розробити стратегії протидії (наприклад, блокування ботнетів).

Приклад 2: Дезінформаційна кампанія під час президентських виборів у США

Задача: Виявити нові поняття і зв'язки, пов'язані з втручанням іноземних акторів у виборчий процес.

Процес виявлення понять та інтеграції нових зв'язків:

1. Аналіз даних. Мовна модель проаналізувала тексти з соціальних мереж, новин та офіційних звітів.
2. Генерація гіпотез: Модель запропонувала новий зв'язок: «використання мікротаргетингу → посилення поляризації суспільства».
3. Перевірка достовірності. Аналіз даних показав, що мікротаргетинг дійсно використовувався для поширення

політичної реклами серед конкретних груп населення, що призвело до зростання поляризації.

4. Інтеграція у мережу. Зв'язок був доданий як «мікротаргетинг → посилює поляризацію суспільства» з вагою 0,85.

Результат: Ці поняття і зв'язок допомагають розробити рекомендації щодо регулювання рекламних платформ для запобігання подібним втручанням у майбутньому.

Приклад 3: Дезінформація про вакцини проти COVID-19 (2020–2022)

Задача: Виявити нові поняття і зв'язки, пов'язані з поширенням дезінформації про вакцини.

Процес виявлення понять та інтеграції зв'язків:

1. Аналіз даних. Мовна модель проаналізувала тексти з соцмереж, YouTube-каналів та форумів.
2. Генерація гіпотез. Модель запропонувала новий зв'язок: «використання алгоритмів рекомендацій → поширення теорій змов про вакцини».
3. Перевірка достовірності. Дослідження показали, що алгоритми рекомендацій дійсно сприяли поширенню теорій змов серед користувачів.
4. Інтеграція у мережу. Зв'язок був доданий як «алгоритми рекомендацій → поширюють теорії змов про вакцини» з вагою 0,8.

Результат: Цей зв'язок допоміг розробити стратегії для покращення алгоритмів рекомендацій та зменшення поширення дезінформації.

3.3. Виклики та обмеження застосування великих мовних моделей

Хоча LLM є потужним інструментом, їх застосування має свої обмеження, а саме:

- ці моделі можуть виявляти поняття та генерувати зв'язки, які не завжди відповідають реальним процесам;
- значення понять та якості понять зв'язків залежить від якості та різноманітності даних, які модель аналізує;
- використання мовних моделей для аналізу чутливих тем може викликати етичні питання.

Великі мовні моделі є новим інструментом для екстракції понять та інтеграції нових зв'язків у семантичні мережі, особливо у контексті інформаційно-психологічних війн. Вони дозволяють генерувати гіпотези, перевіряти адекватність понять та достовірність зв'язків та

автоматизувати процес оновлення мережі. Приклади із реальних ІПВ демонструють, як ці моделі можуть допомогти виявити нові механізми впливу та розробити стратегії протидії. Однак важливо враховувати можливі виклики та обмеження, щоб забезпечити точність та ефективність інтеграції нових зв'язків.

3.4. Розширення ієрархічних структур до мереж

У той час як загально прийняти ієрархічна модель дає змогу структурувати проблему ІПВ на основі суворого підпорядкування рівнів, вона має обмежену гнучкість у відображенні складних взаємозв'язків між елементами. Тому для глибшого аналізу необхідно цю ієрархічну структуру розширити, увівши нові сутності та зв'язки, які формують мережеву модель. Це дозволить врахувати не тільки вертикальні, а й горизонтальні залежності, а також динамічну взаємодію між акторами, їхніми цілями та засобами, що є критично важливим для розуміння механізмів ментальної війни в реальних умовах.

Категорію «ментальна війна» спочатку будемо розглядати як складну систему, яка має ієрархічну структуру. Дана модель є, звичайно, її математичною абстракцією – рівні й елементи визначають суть проблеми.

Надамо її формальний опис: нехай H – скінчена частково упорядкована множина, що має найбільший елемент b . H є ієрархією, якщо виконуються наступні умови:

- Існує розбиття H на підмножини L_k , $k = 1, 2, \dots, h$, де $L_1 = \{b\}$.
- Із $x \in L_k$, випливає, що $x^- \in L_{k+1}$, $k = 1, 2, \dots, h-1$.
- Із $x \in L_k$, випливає, що $x^+ \in L_{k-1}$, $k = 2, \dots, h$.

Тут x^- – елемент, який слідує за елементом x , а x^+ – елемент, що передує елементу x за рівнями ієрархії.

В такому разі дана проблема піддається декомпозиції на простіші складові з подальшим оцінюванням відносного ступеня взаємодії елементів даної ієрархічної структури. У складній системі «ментальна війна» можна виділити п'ять рівнів L_i , де $i = 1, 2, \dots, 5$.

У нашому випадку, при традиційному розгляді ментальної війни як ієрархічної структури, зміст рівнів має таке призначення: на першому рівні (L_1 «Цілі ментальної війни») розглядають єдиний елемент – фокус, і розміщують його у вершині ієрархії (війна за зміну ідентичності); на другому рівні (L_2 «Сили і засоби ментальної війни») відображають економічні, політичні та соціальні сили, що впливають на результат (фінанси, література, мистецтво загалом, ЗМІ, Інтернет тощо); третій

рівень (L_3 «Актори ментальної війни») утворюють актори, що маніпулюють цими силами (уряд, митці, меценати тощо); четвертий рівень (L_4 «Цілі окремих акторів ментальної війни») – це цілі кожного актору (зміна уявлень, цінностей, настанов, стереотипів, традицій, історичної пам'яті); п'ятий рівень (L_5 «Політики реалізації цілей акторів ментальної війни») описує можливі сценарії, або результати, яких прагнуть досягнути кожен з акторів, застосовуючи свої політики, в першу чергу, йдеться про перекодування, як цивілізаційної ідентичності держави, так і перекодування культурних цінностей суспільства й індивіду.

Ієрархічна декомпозиція дозволяє структурувати систему на рівні підсистем, де кожна підсистема відповідає за конкретні цілі. Для кожного рівня визначаються концепти, які утворюють кластери, які взаємодіють між собою для досягнення загальної мети. Математична модель допомагає формалізувати взаємодії між підсистемами через графи, матриці зв'язків та цільові функції, що описують ефективність системи в цілому.

Рівні L_1 і L_4 ієрархії включають певну множину цілей та підцілей. Нехай на рівні i знаходиться множина цілей $T_i = \{ T_{i1}, T_{i2}, \dots, T_{jmi} \}$, де $j = \{1, 4\}$, а mi – кількість цілей на рівні i . У свою чергу, кожній цілі T_{ij} відповідає множина підцілей $F_{ij} = \{ F_{ij1}, F_{ij2}, \dots, F_{ijn} \}$, де ijn – кількість підцілей цілі T_{ij} . Так в умовах російської агресії глобальною ціллю першого рівня T_{1i} є остаточне й безповоротне розчинення української ідентичності у “загальноросійській” та відмова українців як політичної нації від претензій на власну державність за відсутності усвідомлення власної самотності. У той час цілі акторів T_{4i} , залишаючись підпорядкованими глобальній цілі “ментальної війни” росії проти України можуть видозмінюватися залежно від епохи передовсім за формою та певними ознаками, залишаючись незмінними за своєю суттю.

Елементи і підцілі на кожному рівні виконують певні функції (інформаційні, економічні, організаційні тощо) та можуть бути пов'язані з іншими елементами і підцілями F_{ijk} на тому ж або інших рівнях, формуючи мережу взаємозв'язків. Вона описується множиною функціональних зв'язків між елементами та підцілями:

$$\text{Зв'язки} = \{ (F_{ijk}, F_{i'j'k'}) | F_{ijk} \text{ пов'язана з } F_{i'j'k'} \}.$$

Множина функціональних зв'язків між елементами, цілями та підцілями утворює орієнтований граф, де вершинами є елементи, підцілі і цілі F_{ijk} , а ребрами – зв'язки між ними. Позначимо цей $G=(V, E_G)$, де V – множина вершин, а $E_G \subseteq V \times V$ – множина

орієнтованих ребер, які описують функціональні залежності між елементами, підцілями та цілями.

Елементи і підцілі на кожному рівні виконують певні функції (інформаційні, економічні, організаційні тощо) та можуть бути пов'язані з іншими елементами і підцільми F_{ijk} на тому ж або інших рівнях, формуючи мережу взаємозв'язків. Вона описується множиною функціональних зв'язків між елементами та підцільми:

Для кожного елемента системи F_{ijk} визначається функція $f(F_{ijk})$, яка описує її роль у системі. Функції можна представити як сукупність відношень або формул, які описують, як ці елементи взаємодіють між собою в межах мережі.

Взаємозв'язки між елементами, цілями та підцільми різних рівнів можуть бути представлені через матрицю зв'язків між рівнями. Нехай $A_{ii'}$ – це матриця розміру $m_i \times m_{i'}$, де кожен елемент a_{ij} , $a_{i'j'}$ описує наявність або відсутність зв'язку між цілями і підцільми T_{ij} і $T_{i'j'}$.

Для моделювання процесу досягнення основних цілей на кожному рівні, можна визначити цільову функцію системи «ментальна війна» Φ , яка залежить від виконання цілей і підцілей на різних рівнях:

$$\Phi = \alpha \cdot f(T_{ij}) + (1-\alpha) \cdot f(F_{ijk}),$$

де α – ваговий коефіцієнт у інтервалі із значеннями у інтервалі $[0,1]$.

Перехід від ієрархічних до мережових структур дозволяє більш точно моделювати інформаційні війни. Семантичне розширення мереж дає змогу не тільки аналізувати їхню топологію, але й розуміти зміст інформаційних впливів. Інтеграція нових зв'язків допомагає оновлювати модель та адаптувати її до реального інформаційного середовища.

Інформаційні системи традиційно будувалися за ієрархічним принципом, де потоки даних рухалися від центрального джерела до нижчих рівнів. Така структура характерна для традиційних засобів масової інформації, державних органів та централізованих пропагандистських механізмів.

Основні характеристики ієрархічних структур, це централізоване управління інформацією, фіксовані ролі агентів комунікації, лінійний рух інформації (зверху вниз), висока передбачуваність взаємодій.

У той же час сучасне інформаційне середовище має переважно мережовий характер, якому притаманне:

- децентралізоване поширення інформації (наприклад, соціальні мережі, форуми, месенджери);

- наявність множинних каналів зв'язку, що сприяє швидкому розповсюдженню контенту;
- виникнення саморегульованих інформаційних кластерів, що можуть формувати автономні наративи.

Таким чином, основні характеристики мережевих структур забезпечують децентралізацію впливу, адаптацію до змін інформаційного середовища, високу варіативність каналів комунікації, можливість швидкого формування нових спільнот навколо ідей.

Перехід до мережевих моделей дозволяє більш точно описувати динаміку інформаційних війн, оскільки взаємодія суб'єктів у них відбувається не за жорстко заданими правилами, а залежно від їхньої ролі у мережі та ступеня зв'язності.

У контексті ієрархічної структури ментальної війни, як підцілі, так і поняття відповідного рівня можуть бути визначені експертами, зокрема віртуальними експертами (штучним інтелектом), які допомагають створювати і вдосконалювати систему цілей і відповідних їм понять (концептів). Якщо окремі концепти можуть одночасно належати кільком різним рівням ієрархії, система зв'язків, що раніше утворювала ієрархію, перестає бути просто ієрархічною, а перетворюється у мережеву. Це дозволяє ефективніше переходити від початкового стану (цілі) до наслідків, використовуючи мережеву структуру зв'язків.

Формалізуємо цю думку для подальшого розширення моделі за допомогою LLM, як і раніше, нехай $T = \{T_1, T_2, \dots, T_n\}$ – множина рівней первинної ієрархії, а $F = \{f_1, f_2, \dots, f_k\}$ – множина концептів. Експерти, включаючи віртуальних, визначають ці концепти на основі аналізу рівней первинної ієрархії моделі ментальної війни. Концепти можуть належати одразу кільком рівням (у тому числі, цілям, визначеним вище), що робить зв'язки між рівнями моделі складнішими. Якщо концепт f_i належить кільком рівням одночасно, це вводить нові зв'язки між рівнями. Тепер структура системи переходить від чистої ієрархії до мережевої моделі.

Для формалізації цього перетворення використовуємо граф $G=(T,E)$, де T – це вершини, які представляють рівні, а E – це множина ребер, де кожне ребро між двома рівнями означає наявність спільного концептів між ними.

Якщо концепт f_i належить рівням T_a і T_b , тоді між ними є зв'язок:

$$E = \{(T_a, T_b) \mid f_i \in F, f_i \in T_a \cap T_b\}.$$

Вага кожного ребра (T_a, T_b) визначається кількістю спільних концептів між цими рівнями:

$$w(T_a, T_b) = |T_a \cap T_b|.$$

Таким чином, чим більше спільних концептів відповідають різним рівням, тим сильніший зв'язок між ними.

Використовуючи мережеву структуру, можна ефективніше досягати кінцевих результатів. Найкоротші шляхи в мережі між рівнями (зокрема, між первинною цілю і наслідками) дозволяють скоротити час і ресурси для досягнення цих наслідків. Це можна виразити через функцію мінімізації шляху в графі:

$$\tau(T_i) = \min_{P \subseteq G} \sum_{(T_a, T_b) \in P} w(T_a, T_b),$$

де P – це шлях від початкової цілі до необхідного результату через інші рівні/концепти.

Віртуальні експерти можуть допомагати не тільки у визначенні концептів, а й у динамічному оновленні мережі. Це означає, що експерти можуть додавати нові зв'язки (new links) або змінювати існуючі залежно від контексту ментальної війни. Формально це описується через динамічну зміну графу моделі:

$$G' = \{G \cup (T_i, T_j) \mid f_i \in F, f_i \in T_i \text{ I } T_j, \text{new links}\}.$$

При цьому перетворення ієрархії на мережу дозволяє більш гнучко та швидко адаптуватися до змін, а також краще використовувати наявні ресурси для досягнення наслідків. Зв'язки між рівнями первинної ієрархії через спільні концепти стають основою для вибору найефективнішого шляху від цілі до наслідків. Цей підхід надає можливість використовувати спільні концепти кількох рівнів, оптимізувати процеси за допомогою віртуальних експертів, і динамічно адаптуватися до нових умов.

Вищенаведеній формалізації відповідає методика, яка передбачає можливість повторення процесу у відповідності із концепцією «рою віртуальних експертів» ([3], [8]) під наглядом експерта до досягнення остаточного розуміння стану предметної області. Реалізація методики передбачає послідовність дій:

1. Представлення початкової схеми, що включає формування початкової схеми в форматі CSV, де відображені базові семантичні зв'язки між поняттями.

2. Формування промптів для великих мовних моделей, тобто створення промптів для LLM з метою генерації нових концептів та зв'язків.
3. Інтеграція нових зв'язків з початковою схемою, використання формули для посилення зв'язків.
4. Лінгвістична обробка даних для оцінки нових і існуючих зв'язків та ранжирування вузлів.
5. Аналіз і візуалізація даних, а саме, завантаження даних в систему аналізу графів (наприклад, Gephi), проведення кластеризації за класами модулярності.
6. Формування і уточнення кластерів, визначення їхніх назв за допомогою LLM, перевірка узгодженості і видалення непотрібних елементів.
7. Остаточна перевірка і верифікація розширеної моделі, забезпечення її узгодженості та коректності.

3.5. Семантичне розширення мережевих моделей

У контексті інформаційно-психологічних війн (ІПВ), де ключовим є розуміння складних соціальних, психологічних та інформаційних процесів, семантичні мережі можуть допомогти виявити приховані патерни, стратегії ворожих дій та методи протидії. Однак для ефективного застосування таких мереж необхідно постійно оновлювати та розширювати їх, враховуючи нові дані, концепції та зв'язки. Цей процес називається семантичним розширенням мережевих моделей.

Семантичне розширення полягає у пошуку нових вузлів (понять) та ребер (зв'язків) для існуючої мережі. У випадку дослідження ІПВ це дозволяє:

- виявляти нові поняття, наприклад, виявлення нових методів маніпуляції інформацією або нових каналів поширення пропаганди;
- уточнювати існуючі зв'язки, пропозити додавання деталей до вже наявних зв'язків, наприклад, для уточнення механізмів впливу на цільову аудиторію;
- проводити багаторівневий аналіз, враховувати як мікро-, так і макrorівні (від окремих повідомлень до глобальних інформаційних кампаній).

3.5.1. Приклад семантичного розширення графа ментальної війни

Для розширення первинної ієрархічної моделі на основі застосування великих мовних моделей пропонуються такі кроки:

1. Початкова схема «ментальних війн» подається у форматі CSV, де кожен рядок відображає семантичний зв'язок між первинними рівнями у форматі зв'язку "Concept 1; Concept 2", наприклад:

Цілі ментальної війни; Сили і засоби ментальної війни
Сили і засоби ментальної війни; Актори ментальної війни
Актори ментальної війни; Цілі окремих акторів ментальної війни
Цілі окремих акторів ментальної війни; Політики реалізації цілей ментальної війни
Політики реалізації цілей ментальної війни; Результати ментальної війни

2. Створюються промпти для розширення всіх рівнів початкової моделі з урахуванням ментальних війн. Відповіді отримуються від систем ChatGPT, Gemini та Llama-3. Наведемо перелік промптів для 6 рівнів безової моделі:

Промпт 1:

Основна ціль ментальної війни проти України – зміна ідентичності народу. Які ще є цілі саме ментальної війни, що стосуються ментальності людей?

Промпт 2:

Основні сили і засоби ментальної війни проти України – це література, мистецтво, ЗМІ, соціальні медіа. Які ще є сили і засоби саме ментальної війни, впливу на ментальність українців?

Промпт 3:

Основні дійові особи (актори) ментальної війни проти України – це митці, уряд, художники тощо. Які ще є актори саме ментальної війни, хто впливає на ментальність українців?

Промпт 4:

Цілями окремих акторів ментальної війни є зміна у українців уявлень, цінностей, установок, національної свідомості. Які ще є цілі у акторів ментальної війни проти ментальності українців?

Промпт 5:

Політиками реалізації цілей окремих акторів ментальної війни є мова, віра, культура, історична пам'ять. Які ще є політиками реалізації цілей ментальної війни у ворогів?

Промпт 6:

Результатами ментальної війни вороги бачуть перекодування уявлень українців щодо держави, суспільства, особи. Які ще можуть бути результати ментальної війни за думкою ворогів?

У відповідь на наведені промпти, системи генеративного штучного інтелекту надають відповіді. Приклад відповіді системи ChatGPT на Промпт 1 наведено нижче:

- Окрім зміни ідентичності народу, ментальна війна має й інші стратегічні цілі, що безпосередньо стосуються ментальності людей:
1. **Делегітимізація державних інститутів**
 - Підриг довіри до влади, судової системи, армії, правоохоронних органів.
 - Створення враження, що держава не здатна ефективно діяти.
 2. **Послаблення національної єдності**
 - Розпалювання регіональних, мовних, релігійних та історичних суперечностей.
 - Просування наративів про «штучність» української державності.
 3. **Формування комплексу меншовартості**
 - Насадження думки, що українська культура, мова та історія менш значущі.
 - Створення відчуття безсилля перед «великими імперіями» чи «цивілізаційними центрами».
 4. **Підміна історичної пам'яті**
 - Спотворення або затирання важливих історичних подій.
 - Формування альтернативних версій історії, вигідних агресору.
 5. **Знищення довіри між людьми**
 - Посилення розколу в суспільстві через конспірологію, теорії змови, нав'язування ворожнечі між групами.
 - Дискредитація волонтерського та громадського рухів.
 6. **Нав'язування фаталізму та апатії**
 - Створення відчуття безвиході («від нас нічого не залежить»).
 - Просування наративу, що боротьба марна, а майбутнє визначене зовнішніми силами.
 7. **Мілітаризація свідомості або навпаки – пацифізм і капітуляція**
 - Формування думки, що війна неминуха і безкінечна.
 - Або ж просування ідей «миру за будь-яку ціну», закликів до здачі без опору.
 8. **Підриг довіри до науки, освіти та критичного мислення**
 - Популяризація псевдонауки, теорій змови, антиінтелектуальних рухів.
 - Дискредитація експертних спільнот та незалежних аналітичних центрів.
 9. **Фрагментація суспільства через цифрові технології**
 - Створення «інформаційних бульбашок» та радикалізація окремих груп через соціальні мережі.
 - Використання алгоритмів для підсилення ворожих наративів.
 10. **Дегуманізація ворога або власного народу**
 - Створення образу «зайвих людей», які не заслуговують на повноцінне життя (військових, переселенців, волонтерів тощо).
 - Формування образу «ворога серед своїх».

У відповідності до відповідей систем LLM формуються додаткові файли у форматі CSV, для наведеного прикладу він буде мати вигляд:

Цілі ментальної війни; Делегітимізація державних інститутів
 Цілі ментальної війни; Послаблення національної єдності
 Цілі ментальної війни; Формування комплексу меншовартості
 Цілі ментальної війни; Підміна історичної пам'яті
 тощо

3. Отримані нові зв'язки інтегруються з початковою схемою. Мережа складається зі споріднених понять. Якщо нове поняття або зв'язок вже

представлені у мережі, вага цього вузла або зв'язку збільшуються. Нові зв'язки, отримані від LLM, інтегруються з початковою схемою (рис. 3.1). Для урахування значимості зв'язків використовується формула ітеративної корекції ваги відповідних зв'язків:

$$S_{new} = \alpha \cdot S_{old} + \beta \cdot S_{new},$$

де S_{new} – новий ваговий коефіцієнт зв'язку, S_{old} – вага існуючого зв'язку, α – ваговий коефіцієнт для старих зв'язків, β – ваговий коефіцієнт для наявних зв'язків.

Старий зв'язок S_{old} представляє вагу існуючого зв'язку між поняттями в семантичній мережі.

Новий зв'язок S_{new} – це ваговий коефіцієнт, що обчислюється для нових зв'язків, що додаються на основі нових даних від LLM.

Коефіцієнти α і β дозволяють контролювати вплив кожного типу зв'язку на загальну мережу.

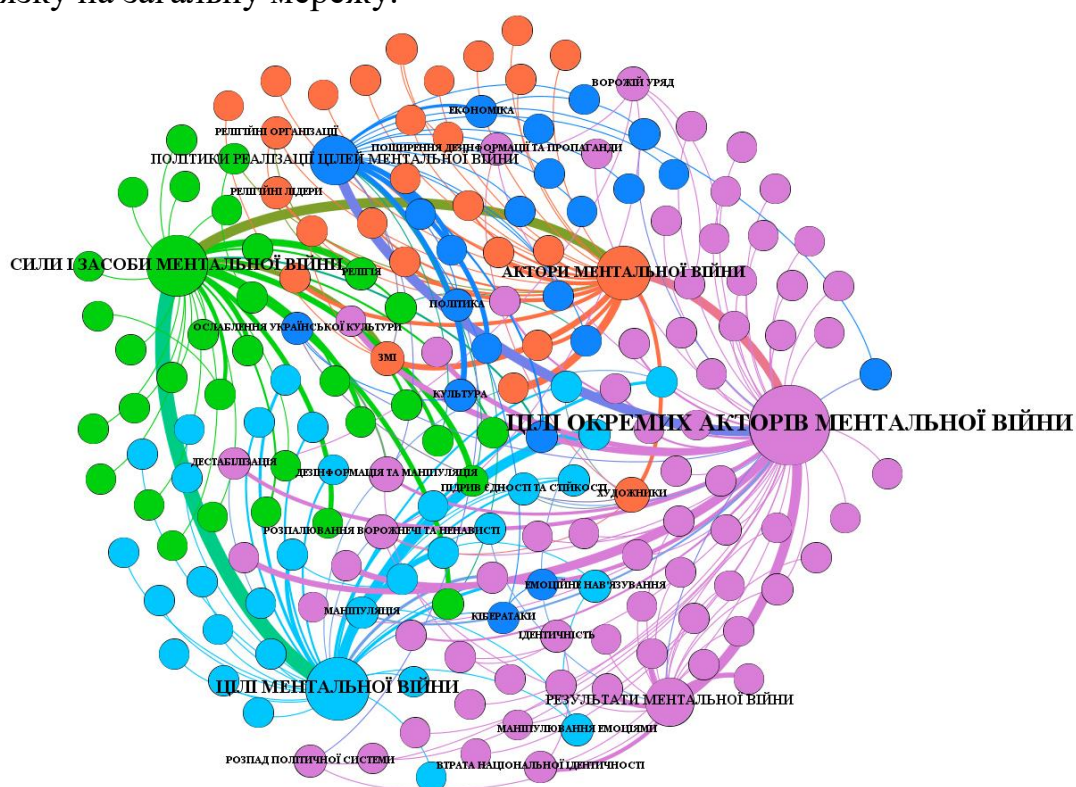


Рисунок 3.1: Укрупнена мережа ментальної війни

Обчислення нового вагового коефіцієнта: Формула дозволяє обчислювати новий ваговий коефіцієнт як комбінацію вагових значень старих і нових зв'язків, де коефіцієнти α і β допомагають налаштувати баланс між існуючими та новими даними.

Перевірка на конфлікти, в рамках цього підходу реалізується залученням людини-експерта, який у спосіб перевірки схожості або відмінностей у зв'язках між концептами встановлює, чи нові зв'язки не вводять суперечностей у контексті вже існуючих.

3.5.2. Джерела для семантичного розширення

Для розширення семантичної мережі можна використовувати такі джерела:

- Мовні моделі такі як великі мовні моделі (GPT, Qwen, DeepSeek або LLaMA), можуть допомогти генерувати нові ідеї, зв'язки та гіпотези. Наприклад, модель може запропонувати нові методи інформаційного впливу, які ще не були описані в літературі.
- Релевантні документи з інтернету щодо аналізу новин, блогів та соціальних мереж допомагають виявити актуальні тренди та нові терміни. Наприклад, аналіз постів у соцмережах може виявити нові меми або наративи, які використовуються для маніпуляції масовою свідомістю.
- Корпоративні документи, внутрішні бази знань, звіти та аналітичні матеріали можуть надати специфічну інформацію про конкретні ІПВ.
- Експертні знання, тобто результати консультацій з фахівцями з ІПВ, психологами, соціологами та аналітиками, які допомагають перевірити достовірність нових зв'язків та понять.

3.5.3. Приклади семантичного розширення

Розглянемо кілька прикладів, як семантичне розширення може бути застосовано до тематики ІПВ.

Приклад 1. Виявлення нових методів інформаційного впливу

Початкова семантична мережа містить поняття "пропаганда" та "соціальні мережі". Під час аналізу даних з інтернету та корпоративних документів було виявлено новий метод – "генерація штучних мемів". Це поняття додається як новий вузол, а зв'язок між ним і "соціальними мережами" встановлюється через механізм "поширення через алгоритми рекомендацій".

Приклад 2. Уточнення зв'язків

У мережі є поняття «маніпуляція інформацією» та «психологічний вплив». Звернення до великої мовної моделі дозволило виявити новий тип зв'язку: «використання мікротаргетингу для посилення

психологічного впливу». Цей зв'язок додається до мережі, що дозволяє точніше моделювати процеси інформацій-психологічних впливів.

Приклад 3. Багаторівневий аналіз

На мікрорівні мережа містить поняття «фейкові новини» та «довіра до джерел». На макрорівні аналіз показує, що "фейкові новини" впливають на «суспільну стабільність». Звернення до експертів дозволило виявити новий вузол – «стратегічна дезінформація», який поєднує обидва рівні через зв'язок «систематична кампанія».

3.5.4. Переваги семантичного розширення

До переваг семантичного розширення можна віднести:

- покращення точності аналізу, зокрема, розширення мережі дозволяє точніше моделювати складні процеси ІПВ;
- прогнозування нових загроз, виявлення нових методів інформаційного впливу допомагає розробляти стратегії протидії;
- інтеграція різних джерел знань, що забезпечує більш повне розуміння проблеми.

Окремо можна виділити лінгвістичну обробку даних та ранжування вузлів, що дозволяє поєднувати концепти, які, є синонімами, похідними одне від одного, тощо. Для цього, зокрема, використовується спеціальне програмне забезпечення, вбудоване в сервіс аналізу і візуалізації графів CSV2Graph (<http://bigsearch.space/uli.html>). Ось приклад нових зв'язків, отриманих після лінгвістичної обробки:

Освіта;Освіта та навчання
Культура;Культурне нав'язування
Культура;Ослаблення української культури
Культура;Знищення української культури
Релігія;Релігійне нав'язування
Релігія;Релігійні організації
Релігія;Релігійні лідери
тощо

Семантичне розширення мережевих моделей є важливим інструментом для аналізу інформаційно-психологічних війн. Воно дозволяє постійно оновлювати знання, враховуючи нові дані, концепції та зв'язки. Застосування великих мовних моделей, аналіз релевантних документів та консультації з експертами забезпечують комплексний підхід до розширення мережі. У результаті такий підхід дозволяє покращити прогнозування загроз, розробку стратегій протидії та загальне розуміння складних процесів ІПВ.

3.6. Практична частина. Формування семантичних мереж на основі аналізу текстів

Розглянемо, як на практиці створити семантичні мережі з текстових даних. Цей процес включає декілька етапів: підготовку даних, аналіз текстів для виявлення ключових концептів та зв'язків між ними, побудову графової моделі та її візуалізацію. Також буде показано, як використовувати доступні інструменти (Hugging Face API, Groq API тощо) для автоматизації цього процесу.

Ця практична частина демонструє, як формувати семантичні мережі з текстів для подальшого аналізу інформаційно-психологічних впливів. Для цього екстрагуємо з тексту пари найбільш впливових пов'язаних між собою понять, які сформують утворюють первинну мережу. Після цього на основі попередніх знань LLM для кожного поняття із первинної мережі знайдемо поняття, яке найбільш зв'язано з ним в контексті тематики первинного тексту. Таким чином розширимо первинну мережу. Розширену мережу візуалізуємо за допомогою сервісу Semantic Networking (CSV2Graph). Задачу будемо вирішувати на основі фреймворку «безкодового програмування», суть якого викладемо у наступному пункті.

3.6.1. Фреймворк безкодового програмування

Суть фреймворку безкодового програмування полягає у представленні промптів як аналогів програмних конструкцій (умовні оператори, цикли, функції) шляхом математичної формалізації їх логіки та взаємодій. Для формування промптів застосовуються такі основні примітиви («цеглинки»), такі як «Умова», «Цикл» і «Функція», а також методи композиції цих примітивів для побудови складних систем, зокрема семантичних мереж.

Опишемо фреймворк, на основі якого у подальшому будуть формуватись промпти великими мовними моделями і застосовуватись для формування семантичних мереж, виконуючись, у свою чергу, у середовищі великих мовних моделей.

Кожен примітив у фреймворку безкодового промпт-інжинірингу описується як математичний об'єкт із чіткими правилами трансформації вхідних даних у вихідні. Ці примітиви є основними "цеглинками" для створення складних логічних конструкцій, аналогічних до операторів мов програмування, але заснованих на природній мові.

Примітив «Умова» (If-Else)

Формалізація. Нехай:

- $Input$ – вхідні дані (наприклад, текст, числовий параметр тощо);
- C – умова (предикат), що повертає значення True або False;
- A_1, A_2 – дві можливі дії, які виконуються у випадках, коли умова дорівнює True або False відповідно.

Промпт-функція P визначається як:

$$P(Input) = \begin{cases} A_1(Input), & \text{якщо } C(Input) = True, \\ A_2(Input), & \text{якщо } C(Input) = False. \end{cases}$$

Це формальне визначення аналогічне класичному оператору if-else у мовах програмування, але застосовується через текстові інструкції.

Механізм роботи примітива «Умова» полягає в тому, щоб передати LLM задачу з двома можливими сценаріями, один з яких виконується в залежності від результату перевірки умови. Умова може бути будь-якою логічною конструкцією, наприклад:

- перевірка наявності певного слова у тексті,
- порівняння чисел,
- аналіз контексту або категорії даних.

Приклад. «Якщо текст містить термін «кібербезпека», поверни його визначення; інакше виведи список споріднених термінів»

Тут:

- C – наявність слова «кібербезпека» у тексті.
- A_1 – генерація визначення терміну «кібербезпека».
- A_2 – пошук асоціацій до терміну «кібербезпека».

Цей примітив дозволяє моделі приймати рішення на основі заданих критеріїв, що ефективно моделює логічну умову.

Примітив «Цикл» (For-Loop)

Формалізація. Нехай:

- $S = \{s_1, s_2, \dots, s_n\}$ – множина елементів, які потрібно обробити.
- F – операція, яка застосовується до кожного елемента множини.

Промпт-функція $P(S)$ визначається як об'єднання результатів застосування операції F до кожного елемента множини:

$$P(S) = \bigcup_{i=1}^n F(s_i).$$

Це формальне визначення аналогічне класичному циклу `for` у мовах програмування, де одна і та ж операція повторюється для всіх елементів множини.

Механізм роботи примітива "Цикл" полягає в тому, щоб передати LLM список елементів разом із операцією, яку потрібно застосувати до кожного елемента. Результатом є об'єднання всіх отриманих відповідей. Такий підхід корисний для масового аналізу даних, виділення паттернів або генерації відповідей для декількох входів одночасно.

Приклад. "Для кожного терміну зі списку ['фішинг', 'брандмауер'] знайди 3 приклади його використання."

Тут:

- $S = \{\text{'фішинг'}, \text{'брандмауер'}\}$ – множина термінів;
- F – операція знаходження прикладів використання терміну.

Результатом буде об'єднання прикладів для обох термінів.

Примітив «Функція» (Abstraction)

Формалізація. Нехай:

- $F: X \rightarrow Y$ – функція, яка перетворює елементи з множини X у елементи множини Y ;
- x – вхідний елемент;
- *параметр* – набір параметрів, які регулюють поведінку функції.
- Промпт-функція $F_{\text{extract}}(x, \text{параметр})$ реалізується через інструкцію, яка передається моделі:

$$F_{\text{extract}}(x, \text{параметр}) = \text{Promt}(x, \text{інструкція з параметр}).$$

Це формальне визначення аналогічне абстракції у мовах програмування, де функція може приймати параметри для регулювання своєї логіки.

Механізм роботи примітива "Функція" полягає в тому, щоб передати LLM конкретну задачу з параметрами, які визначають деталі її виконання. Це дозволяє створювати гнучкі інструкції, які можуть бути повторно використані для різних входів.

Загальні принципи вибору примітивів

Чіткість формулювання. Кожен примітив має бути точно описаною конструкцією, щоб забезпечити однозначне розуміння моделлю.

Гнучкість. Примітиви можуть бути параметризовані, що дозволяє їх адаптувати під різні завдання.

Композиція. Примітиви можуть поєднуватися для створення складних систем, аналогічно до того, як програмісти пишуть код з базових конструкцій.

Ці три примітиви – «Умова», «Цикл» та «Функція», утворюють основу для безкодового створення систем через промпт-інжиніринг, дозволяючи використовувати природну мову як інструмент для управління логікою великих мовних моделей.

Композиція примітивів у систему

Система – це оркестрація примітивів, аналогічна програмам. Синтаксис промтів утворює мову з формальною граматикою, яку можна представити у вигляді AST (Abstract Syntax Tree):

$$\text{Promt} ::= \text{Примітив} | (\text{Promt} \oplus \text{Promt}) | \text{Умова}(\text{Promt}, \text{Promt})$$

де $\oplus$ – композиція.

Створення промту

Для отримання чіткого формального промту для формування розширеної мережі на основі заданого тексту наведений вище фреймворк перетворюється на документ у форматі PDF і завантажується в систему LLM (наприклад, ChaGPT, Qwen або DeepSeek). Після чого завантажуюмо промт-завдання для формування розширеної мережі наступного вигляду:

У відповідності з наведеною методикою прошу сформував загальний промт для вирішення задачі: Нехай у нас є деякий вхідний текст [вхідний текст] за тематикою [тематика]. З тексту виділяється 20 пар найбільш зв'язаних між собою понять за цією тематикою, які для утворення первинної мережі виводяться у вигляді списку у форматі "поняття;поняття", що є по суті представленням первинної мережі. Після цього для кожного поняття з цієї мережі з на основі знань LLM вибирається якесь нове поняття із заданої тематики, змістовно зв'язано з ним. Пара понять виводиться у вигляді "нове поняття; поняття з первинної мережі". Так робиться для всіх понять з первинної мережі, після чого цикл завершається. Остоточно сформована розширена мережа виводиться у вигляді списку зв'язаних понять у форматі "поняття;поняття". Надай сформований промт

Нижче наведено вигляд формалізованого промту, завдання для формування мережі переважного семантичного приєднання, сформованого у результаті виконання наведеного первинного промту:

Формальне представлення через примітиви:
Примітив "Функція" (Abstraction):
Визнач функцію Fextract(текст, тематика), яка виділяє 20 пар найбільш пов'язаних понять з тексту за заданою тематикою.
Результат: Первинна мережа.
Примітив "Цикл" (For-Loop):
Для кожного поняття з Первинної мережі:
Застосуй функцію Fextend(поняття, тематика), яка знаходить нове поняття, пов'язане з поточним поняттям за заданою тематикою.

Результат: Розширена мережа.

Примітив "Функція" (Abstraction) :

Визнач функцію Fmerge(Первинна мережа, Розширена мережа), яка об'єднує всі пари понять з обох мереж.

Результат: Остаточна розширена мережа.

Композиція примітивів :

Сформулюй загальний промпт як композицію примітивів:

Prompt = Fmerge(Fextract(текст, тематика), ForLoop(S=Первинна мережа, F=Fextend))

Також наводиться текстовий варіант промпту для LLM:

Вхідний текст: [вхідний текст]

Тематика: [тематика]

1. Виділи з тексту 20 пар найбільш пов'язаних між собою понять за тематикою [тематика]. Подай результат у форматі:

поняття;поняття

Назвемо цей список "Первинна мережа".

2. Для кожного поняття з "Первинної мережі":

– Знайди нове поняття із тематики [тематика], яке є змістовно пов'язаним з поточним поняттям.

– Подай результат у форматі:

нове поняття; поняття з первинної мережі

Назвемо цей список "Розширена мережа".

3. Об'єднай всі пари понять з "Первинної мережі" та "Розширеної мережі".

Виведи остаточну розширену мережу у вигляді списку зв'язаних понять у форматі:

поняття;поняття

3.6.2. Виконання промпту

Для виконання промпту вибирається вхідний текст – повідомлення інформаційного агентства за темою інформаційної війни:

Блокування контенту в інтернеті: де межа цензури?

Свобода слова є однією з основоположних цінностей демократичного суспільства. Водночас у цифрову епоху вона нерозривно пов'язана з можливістю поширення інформації через інтернет. Саме тому питання блокування інтернет-контенту стало однією з найгостріших правових дилем сучасності. Особливої актуальності воно набуває в умовах воєнного стану, інформаційних війн, масових кампаній дезінформації та пропаганди.

...

Юридична газета 2025.04.16 11:34

<https://yur-gazeta.com/dumka-eksperta/blokuvannya-kontentu-v-interneti-de-mezha-cenzuri.html>

Нижче наводиться протокол створення розширеної мережі за допомогою LLM Qwen-2.5 (<https://qwen.ai>):

На основі наданого тексту та вказаних вимог до методики формування мережі понять, я розпочну з виділення Первинної мережі. Ось 20 пар найбільш пов'язаних між собою понять за тематикою "Інформаційна безпека":

Первинна мережа

свобода слова; цензура
інформація; поширення інформації
блокування контенту; національна безпека
Конституція України; свободи думки
Європейський суд з прав людини; пропорційність обмежень
...

Розширена мережа

Для кожного поняття з Первинної мережі знайду нове поняття, змістовно пов'язане з ним за тематикою "Інформаційна безпека". Результат подам у форматі: нове поняття; поняття з первинної мережі

інформаційна гігієна; свобода слова
цифрова безпека; поширення інформації
кіберзахист; блокування контенту
правова система; Конституція України
права людини; Європейський суд з прав людини
медіалітературність; Закон України "Про інформацію"
психологічний захист; самоцензура
...

Остаточна розширена мережа

Об'єдную всі пари понять з Первинної мережі та Розширеної мережі. Остаточна розширена мережа має такий вигляд:

свобода слова; цензура
інформація; поширення інформації
блокування контенту; національна безпека
Конституція України; свободи думки
Європейський суд з прав людини; пропорційність обмежень
Закон України "Про інформацію"; заборона цензури
самоцензура; державна цензура
втручання у професійну діяльність; медіа
судовий контроль; обмеження доступу
санкції; інтернет-провайдери
російські соціальні мережі; блокування
недостовірна інформація; спростування
медіа; платформи спільного доступу
суспільна небезпека; заклики до насильства
...

Ця остаточна розширена мережа включає всі пари понять, що демонструють взаємозв'язки між ключовими концепціями тематики "Інформаційна безпека".

3.6.3. Візуалізація

Візуалізацію первинної мережі і остаточної мережі здійснено за допомогою сервісу Semantic Networking, описаного у Розділі 1 (рис. 3.2, рис 3.3).

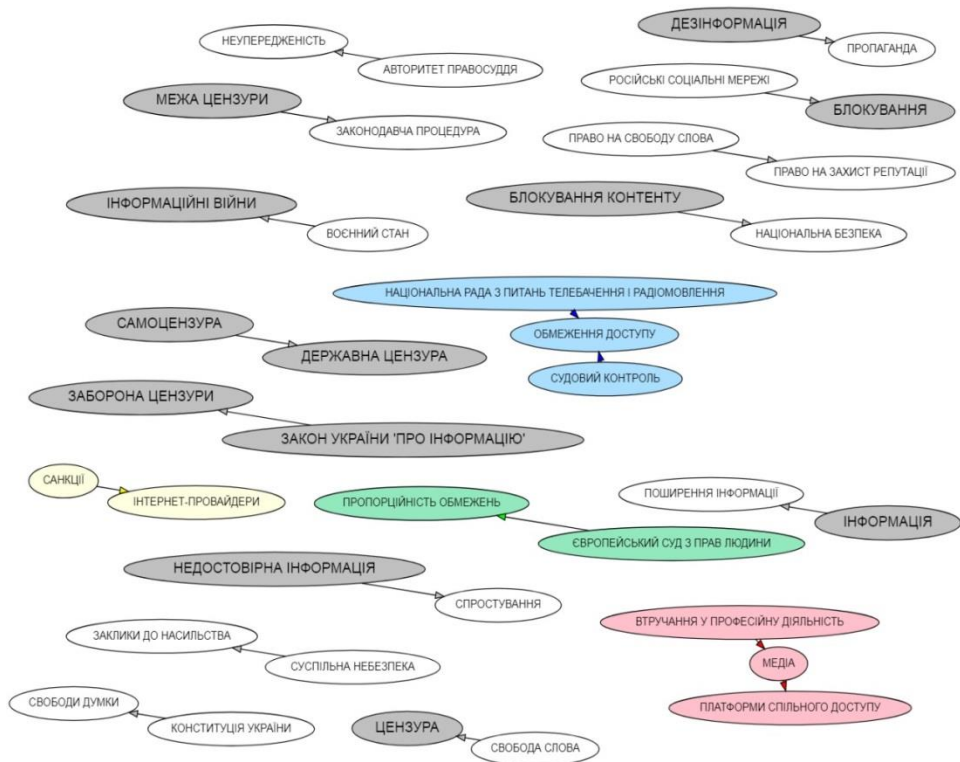


Рисунок 3.2: Візуалізація первинної мережі

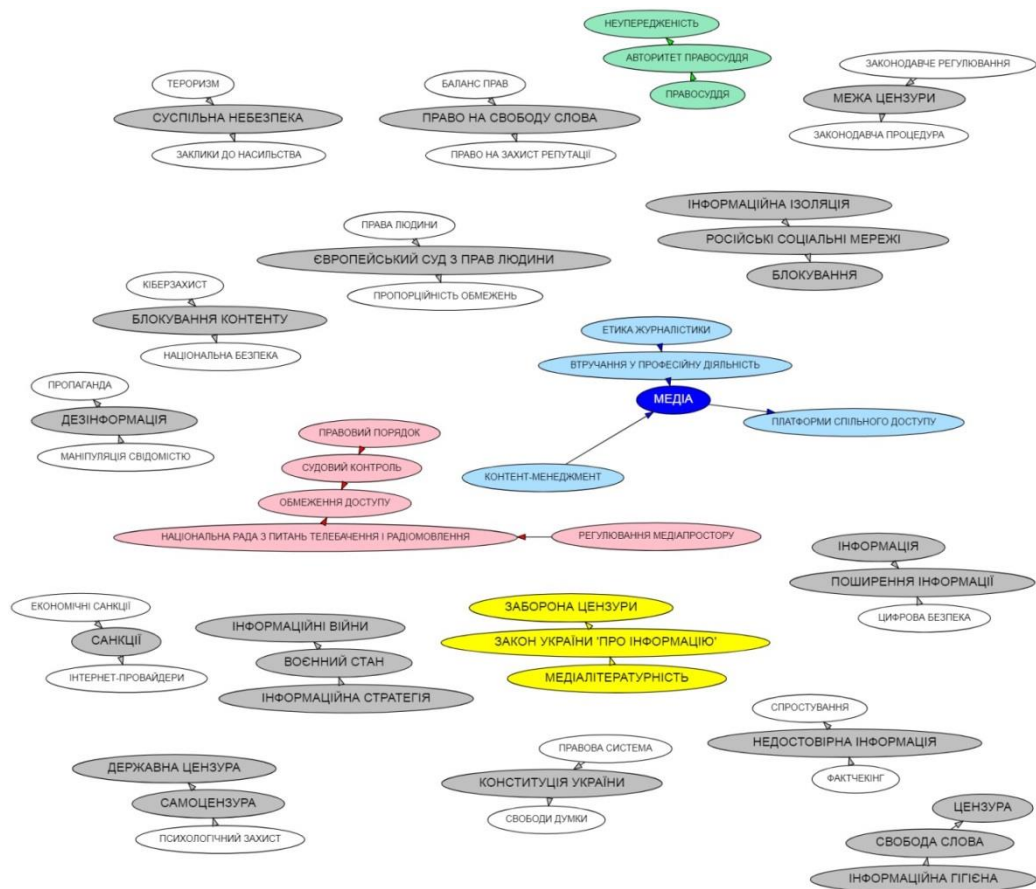


Рисунок 3.3: Візуалізація розширеної мережі

Висновки до розділу

Розділ 3 присвячений методам побудови мережевих моделей, їх еволюції від ієрархічних структур до гнучких адаптивних мереж та застосуванню цих підходів для аналізу інформаційно-психологічних війн. У цьому розділі були детально розглянуті ключові аспекти побудови мереж: семантичне розширення графів, інтеграція нових зв'язків, а також переваги та обмеження сучасних методів. Особливу увагу приділено ролі великих мовних моделей у генерації концептів та зв'язків, що дозволяє створювати динамічні та актуальні мережеві моделі.

Методи побудови мережевих моделей пройшли значний шлях розвитку:

- Ієрархічні структури. На початкових етапах моделі базувалися на жорстких ієрархіях, де кожен вузол мав чітко визначене місце. Такий підхід був корисним для опису лінійних процесів, але не враховував складності реальних інформаційних потоків;
- Мережеві структури. Сучасні моделі переходять до гнучких мереж, які враховують декілька ключових характеристик, таких як децентралізація впливу, адаптивність до змін інформаційного середовища, висока варіативність каналів комунікації, можливість швидкого формування нових спільнот навколо ідей.

Цей перехід дозволяє точніше моделювати динаміку інформаційних війн, де взаємодія суб'єктів відбувається не за жорстко заданими правилами, а залежно від їхньої ролі у мережі та ступеня зв'язності.

Цей підхід дозволяє постійно оновлювати мережу, враховуючи нові дані, концепції та зв'язки, що забезпечує її актуальність та точність.

Інтеграція нових зв'язків є важливим аспектом побудови мережевих моделей. Для цього використовуються:

- великі мовні моделі. Генерація гіпотез про нові зв'язки та їх перевірка на основі аналізу текстових даних;
- аналіз документів та джерел. Пошук підтверджуючих даних для нових зв'язків;
- консультації з експертами. Перевірка достовірності зв'язків перед їх інтеграцією.

Наприклад, виявлення зв'язку "ботнети → поширення наративів про "нацистську владу" допомогло зрозуміти механізми масового поширення пропаганди та розробити стратегії протидії.

До переваг наведених підходів відносяться:

- гнучкість. Мережеві моделі дозволяють адаптуватися до змін інформаційного середовища;
- точність. Семантичне розширення покращує аналіз складних процесів ІПВ;
- ефективність. Використання LLM та автоматизованих інструментів прискорює процес оновлення мереж.

Разом з тим, існують і деякі обмеження, а саме:

- суб'єктивність – моделі можуть генерувати зв'язки, які не завжди відповідають реальним процесам;
- залежність від якості та різноманітності даних;
- необхідність людської верифікації, яка полягає у тому, що автоматизовані методи потребують перевірки експертами для мінімізації помилок.

У контексті інформаційно-психологічних війн мережеві моделі є важливим інструментом для:

- прогнозування загроз – виявлення прихованих патернів та стратегій ворожих дій;
- розробки стратегій протидії – оптимізація механізмів захисту від інформаційних атак;
- аналізу наслідків – моделювання шляхів впливу та їх результатів.

Наприклад, інтеграція зв'язку «маніпуляція інформацією → психологічний вплив» дозволила точніше моделювати процеси ІПВ та розробляти ефективні механізми протидії.

Взагалі, методи побудови мережевих моделей є ключовим інструментом для аналізу інформаційно-психологічних війн. Вони дозволяють:

- ефективно моделювати складні процеси ІПВ;
- постійно оновлювати знання через семантичне розширення;
- інтегрувати нові зв'язки за допомогою LLM та інших інструментів.

Разом з цим, для досягнення найкращих результатів важливо комбінувати автоматизовані методи з людською експертизою, щоб забезпечити точність та надійність аналізу.

Питання для самоконтролю

1. Які основні цілі ментальної війни описані у розділі, і як вони впливають на суспільство?
2. У чому полягає методика «рою віртуальних експертів», і як вона застосовується для аналізу інформаційно-психологічних війн?

3. Які етапи включає процес побудови семантичної мережі, і яка роль LLM у цьому процесі?
4. Як формується початкова схема для аналізу зв'язків між поняттями, і який формат даних використовується?
5. Як створюються промпти для LLM, і як вони допомагають генерувати нові концепти та зв'язки?
6. Які існують способи інтеграції нових зв'язків у початкову схему, і як вони перевіряються?
7. Яка роль людини-експерта у процесі верифікації нових зв'язків, і чому це важливо?
8. Як виконується лінгвістична обробка даних для оцінки нових та існуючих зв'язків?
9. Як проводиться кластеризація даних у системах аналізу графів (наприклад, Gephi), і яке значення має цей процес?
10. Як формується та уточнюється назва кластерів, і як LLM допомагає у цьому процесі?
11. Які характеристики мережевих структур є ключовими для аналізу інформаційно-психологічних війн?
12. Як децентралізація впливу та адаптивність до змін інформаційного середовища впливають на формування спільнот?
13. Які основні методи прогнозування загроз використовуються у контексті інформаційно-психологічних війн?
14. Як формуються стратегії протидії інформаційним атакам на основі аналізу мережевих моделей?
15. Як виявляються ключові концепти та зв'язки у текстах за допомогою LLM, і які алгоритми використовуються для цього?
16. Як обчислюється новий ваговий коефіцієнт для зв'язків, і як це впливає на точність моделі?
17. Як перевіряються конфлікти між новими та існуючими зв'язками, і яка роль експертів у цьому процесі?
18. Як семантичне розширення мережевих моделей покращує точність аналізу інформаційних впливів?
19. Які етапи включає методологія семантичного розширення, і як вони допомагають виявляти нові зв'язки?
20. Як використовуються онлайн-API для аналізу великих обсягів текстових даних, і які переваги цього підходу?

Література за розділом

1. Ланде Д.В., Страшной Л.Л. Семантичний нетворкінг на основі великих мовних моделей : монографія. ТОВ «Інжиніринг», 2025. – 274 с. ISBN 978-617-8180-01-0
2. Парламентський контроль із застосуванням генеративного штучного інтелекту : монографія / Ланде Д.В., Фурашев В.М. – Київ: ТОВ «Інжиніринг», 2023. – 202 с. ISBN 978-966-2344-82-0
3. Dmytro Lande, Leonard Strashnoy. GPT Semantic Networking: A Dream of the Semantic Web – The Time is Now. Kyiv: Engineering, 2023. – 168 p. ISBN 978-966-2344-94-3.
4. Ланде Д.В., Пучков О.О., Субач І.Ю. Застосування генеративного штучного інтелекту для визначення центроїдів і сутності кластерів в мережі понять у сфері кібербезпеки // Телекомунікаційні та інформаційні технології. 2025. N 1 (86) – С. 89-97. DOI: 10.31673/2412-4338.2025.012993
5. Ланде Д., Пучков О., Субач І. Методика формування причинно-наслідкових мереж у сфері кібербезпеки засобами генеративного штучного інтелекту // Телекомунікаційні та інформаційні технології. 2024. N 4 (85). – С. 30-42. DOI: 10.31673/2412-4338.2024.040560
6. Ланде Д., Страшной Л. Формування мереж понять в галузі права за допомогою системи штучного інтелекту // Інформація і право, 2023. N 2(45)/2023. – С. 88-93. DOI: 10.37750/2616-6798.2023.2(45).282326

4. Аналіз та візуалізація даних

Для ефективного дослідження інформаційно-психологічних ментальних війн необхідно не тільки будувати мережеві моделі, але й застосовувати методи їх аналізу та інтерпретації. Аналіз мереж дозволяє виявляти основні інформаційні кластери, ідентифікувати ключові вузли впливу та оцінювати структуру інформаційного середовища.

Візуалізація мереж є важливим інструментом для розуміння взаємозв'язків між елементами інформаційної війни, оскільки вона допомагає швидко ідентифікувати осередки дезінформації, механізми маніпулятивного впливу та канали поширення інформації.

Цей розділ присвячений методам кластеризації мереж, ранжування вузлів за їхньою впливовістю та способам візуалізації інформаційних потоків.

4.1. Кластеризація мереж

Кластеризація мереж дозволяє розбивати їх на групи вузлів, які мають схожі характеристики або тісніші зв'язки між собою. У контексті аналізу ментальних війн це дає змогу виявляти інформаційні «бульбашки», де поширюється певний наратив, визначати основні інформаційні фронти (пропаганда, контрпропаганда, нейтральні групи), ідентифікувати групи акаунтів або ресурсів, які діють узгоджено;

До основних методів кластеризації мереж відносяться:

1. Метод модулярності (Louvain algorithm). Це один із найбільш ефективних алгоритмів для виявлення спільнот у великих мережах. В основі методу – максимізація модулярності – показника, що визначає, наскільки тісно пов'язані вузли в межах кластера порівняно з випадковою мережею. Це метод широко застосовується для аналізу соціальних мереж, інформаційних кампаній, бот-мереж.
2. Метод розриву ребер (Girvan-Newman algorithm) дозволяє визначати кластери, поступово видаляючи найбільш «критичні» ребра (ті, що з'єднують різні групи). Цей метод застосовується для виявлення взаємопов'язаних підмереж, наприклад, груп ботів або пропагандистських каналів.
3. Кластеризація на основі подібності контенту (Semantic Clustering) використовується, коли вузли пов'язані не тільки структурними зв'язками, а й змістовною схожістю. При цьому застосовуються алгоритми машинного навчання (BERT, Doc2Vec) для аналізу текстів і виявлення спільних тем у різних групах.

Як приклад можна розглянути кластерний аналіз соціальної мережі. Розглянемо мережу акаунтів у вибраній соціальній мережі, що поширюють певний політичний наратив. Використання алгоритму кластерного аналізу може виявити декілька кластерів групи акаунтів, що поширюють дезінформацію, які ведуть боротьбу з дезінформацією, знаходити нейтральні акаунти, що взаємодіють із різними групами.

Розглянемо кластеризацію за класами модулярності, варіанти якої можна застосовувати в Gephi, при цьому має сенс застосовувати модель Поттса, яка враховує так звану розподільну здатність, на основі якої автоматично визначається необхідна кількість класів вузлів (кластерів).

Відповідно до цієї моделі функція якості $H(G, P)$, або коротко $H(P)$, розбиття P на кластери графа G записується у вигляді:

$$H(P) = -\sum_C e_C - \gamma n_C^2,$$

де кожний кластер $C \in P$ складається з e_C ребер і n_C вузлів, а γ – це параметр розподільної здатності, який суттєво впливає на розбиття графа на кластери.

Процедура визначення кластерів – класів модулярності здійснюється за алгоритмом, що включає такі кроки:

1. Ініціалізація: початкове розподіл вузлів у початкових кластерах.
2. Оцінка: визначення модулярності для поточного розподілу вузлів.
3. Об'єднання кластерів: з'єднання вузлів або груп вузлів, щоб побачити, чи покращиться модулярність.
4. Ітерація: повторення процесу об'єднання кластерів і оцінки модулярності, поки не буде досягнута максимальна модулярність або не буде виявлено, що подальше об'єднання не покращує результат.
5. Фіналізація: отримані класи модулярності представляють собою групи вузлів, що мають високу внутрішню зв'язність та низьку зовнішню зв'язність.

Кластерізована розширена мережа може отримувати назви для кластерів, для цього може застосовуватись промпт до LLM типу:

Придумайте назву для кластеру на основі включених в нього концептів: ...

4.2. Ранжирування вузлів мережі

Ранжирування вузлів є ключовим для визначення найважливіших учасників інформаційних війн. У контексті аналізу ментальних війн важливо визначити:

- Хто є основними агентами впливу?
- Які акаунти чи медіаресурси мають найбільший охоплення?
- Які вузли виступають посередниками у розповсюдженні інформації?

До основних методів ранжування відносяться:

1. PageRank, який є класичним інтернет-алгоритмом, що використовується для визначення важливості вузлів у графі. У цьому алгоритмі враховується кількість і вага зв'язків. У контексті інформаційних війн PageRank дозволяє виявляти ключові акаунти та ресурси, що відіграють центральну роль у поширенні інформації.
2. Betweenness Centrality (Центральність посередництва) дозволяє визначати вузли, що виступають критичними сполучними точками між різними частинами мережі. Виявлення таких вузлів допомагає ідентифікувати посередників у поширенні дезінформації або інформаційних атак.
3. Closeness Centrality (Близькість вузла до інших) показує, наскільки швидко інформація може розповсюджуватися від конкретного вузла по всій мережі. Цей підхід корисний для виявлення акаунтів, що оперативно реагують на інформаційні атаки.
4. Eigenvector Centrality дозволяє визначати впливовість вузла на основі впливовості його сусідів, виявляти не лише популярні акаунти, а й ті, що взаємодіють із важливими гравцями в мережі.

Зокрема, на основі алгоритмів PageRank-TextRank [9,10] здійснюється ранжування вузлів, що відповідають вузлам мережевої моделі здійснюється за допомогою підходу до ранжування понять із текстів:

$$PR(A) = \frac{1-d}{N} + d \sum_{i \in M(A)} \frac{PR(i)}{L(i)},$$

де $PR(A)$ – ранг (PageRank) вузла A , d – коефіцієнт затухання, N – загальна кількість вузлів у графі, $M(A)$ – множина вузлів, що посилаються на A , $L(i)$ – кількість виходів з вузла i .

Перший член формули для розрахунку PageRank множник $\frac{1-d}{N}$ забезпечує базовий рівень PageRank для всіх вузлів, де чисельник

відповідає ймовірності того стану, коли користувач випадково вибирає вузол графа, а N – загальна кількість вузлів.

Другий член $d \sum_{i \in M(A)} \frac{PR(i)}{L(i)}$ – це частина формули, за якою обчислюється

PageRank вузла A на основі PageRank вузлів, що на нього посиляються. Для кожного вузла i у множині $M(A)$ його PageRank ділиться на кількість посилань $L(i)$, що він має, і це значення додається для всіх вузлів, що посиляються на A .

4.3. Візуалізація мереж

Візуалізація мереж дозволяє представити результати аналізу у зрозумілому вигляді, що значно полегшує виявлення взаємозв'язків та динаміки інформаційних війн.

До основних видів візуалізації можна віднести:

1. Графові діаграми, в яких використовують вузли та зв'язки для представлення інформаційної мережі. Розмір і колір вузлів можуть відображати їхню центральність або впливовість. Застосування силових алгоритмів (Force Atlas 2, Fruchterman-Reingold) дозволяє структурувати мережу природним чином.
2. Кластеризовані графи, які відображають групи взаємопов'язаних вузлів у різних кольорах, що допомагає визначити спільноти та інформаційні бульбашки.
3. Теплові карти (Heatmaps), які використовуються для візуалізації щільності інформаційних потоків у мережі та дозволяють оцінити рівень активності певних тем або наративів у різних сегментах інформаційного простору.

До найбільш доступних інструментів візуалізації можна віднести:

- Gephi – найпопулярніший інструмент для побудови графів та кластерного аналізу.
- NetworkX (Python) – використовується для обчислення метрик центральності та побудови графових структур.
- Graphviz – дозволяє створювати графічні зображення мережевих структур.

Наприклад, якщо потрібно проаналізувати інформаційний вплив на платформі соціальної мережі, можна побудувати граф взаємодій користувачів, де колір буде позначати їхню належність до певного інформаційного кластера, а розмір вузла – рівень впливу.

4.4. Практична частина. Завантаження, аналіз і візуалізація мережевих структур в Gephi

У цьому пункті розглядається, як на практиці проводити завантаження, аналіз та візуалізацію мережевих структур за допомогою програмного забезпечення Gephi. Цей інструмент є одним із найбільш популярних для аналізу графів, оскільки він надає зручний інтерфейс для візуалізації та обчислення метрик центральності, кластеризації та інших характеристик мереж.

Етап 1. Підготовка даних

Перед початком роботи у Gephi потрібно підготувати дані у форматі, придатному для завантаження. Найчастіше використовуються формати:

- CSV (Comma-Separated Values) – для таблиць зв'язків.
- GraphML – для збереження графів.

Приклад структури CSV-файлу для зв'язків:

Source,Target,Weight
NodeA,NodeB,1
NodeA,NodeC,2
NodeB,NodeD,3

Цей файл описує зв'язки між вузлами (Source та Target) та їхню вагу (Weight). Якщо вага не вказана, можна працювати лише зі списком зв'язків у форматі:

NodeA,NodeB
NodeA,NodeC
NodeB,NodeD

Етап 2. Завантаження даних у Gephi

1. Відкриття Gephi:

- запустіть програму Gephi;
- виберіть File > Open або File > Import Spreadsheet, щоб завантажити ваші дані.

2. Завантаження CSV-файлу:

- у вікні імпорту виберіть ваш CSV-файл;
- вкажіть, що перший рядок файлу містить заголовки (First row as header);
- натисніть Next, а потім Finish.

3. Перегляд графа:

- після завантаження ви побачите граф у вікні Overview. На цьому етапі граф може виглядати хаотично, оскільки вузли ще не організовані (рис. 4.1).

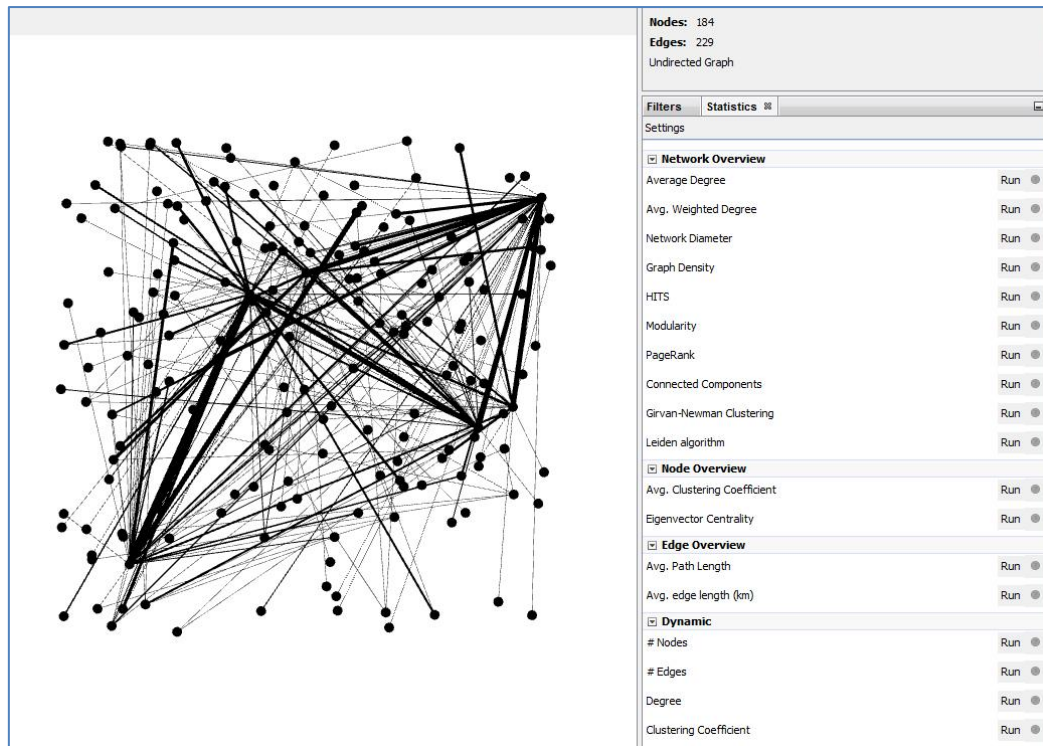


Рисунок 4.1: Первинне відображення графа після завантаження

Етап 3. Аналіз мережі

Gerhi надає потужні інструменти для аналізу графів. Ось основні кроки:

1. Обчислення метрик центральності

У вікні Statistics натисніть Run поруч із такими показниками:

- PageRank: для визначення впливових вузлів (рис. 4.2).
- Betweenness Centrality: для виявлення посередників.
- Modularity: для кластеризації мережі.

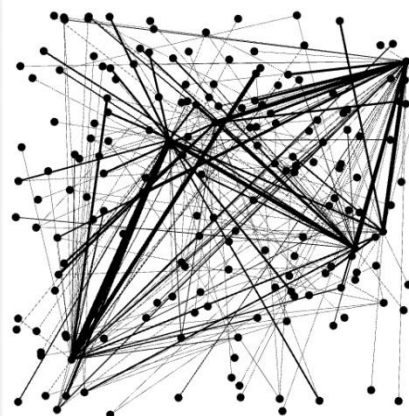
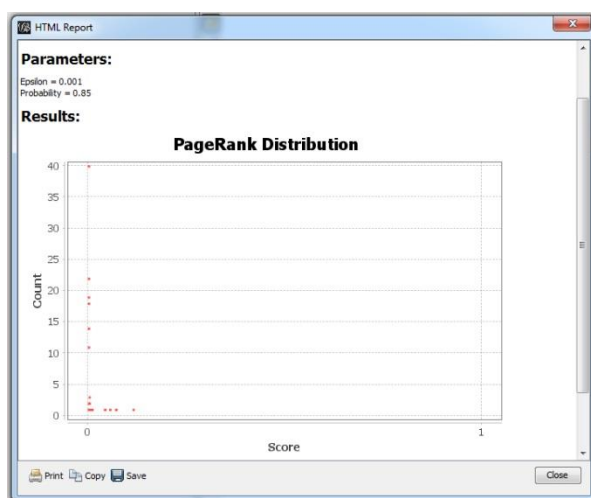


Рисунок 4.2: Розрахунок PageRank

2. Кластеризація мережі

- після обчислення модулярності (Modularity), вузли будуть автоматично розподілені по кластерах (рис. 4.3).
- у вікні Partition виберіть Modularity Class, щоб застосувати різні кольори до різних кластерів.

3. Аналіз щільності зв'язків:

- використовуйте інструмент Heatmap для візуалізації активності зв'язків у мережі.

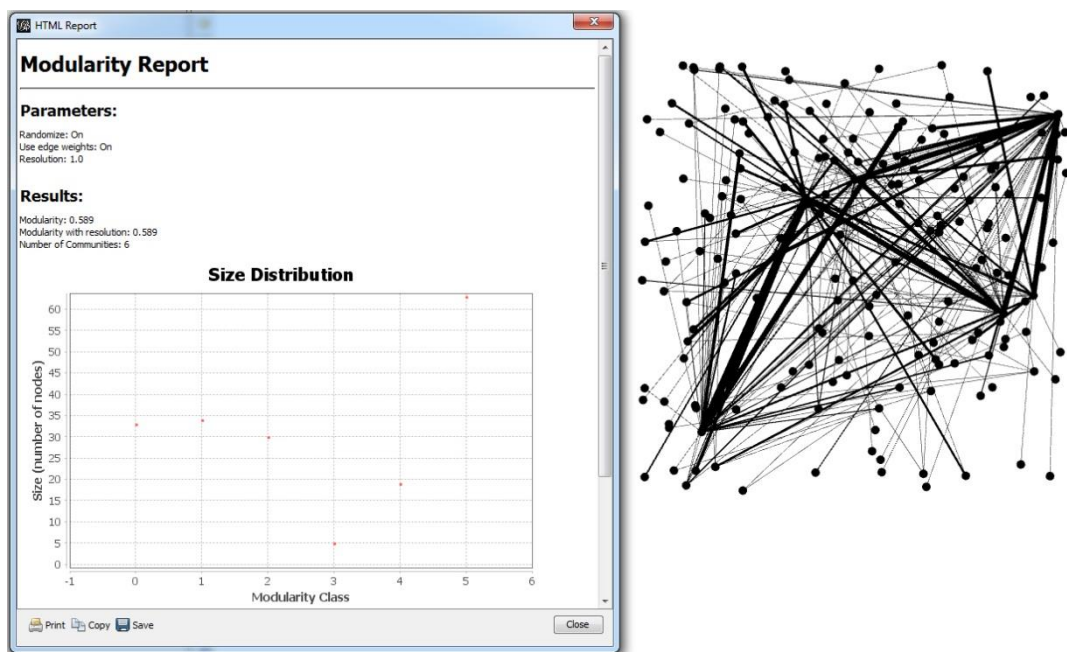


Рисунок 4.3: Розрахунок Modularity

Етап 4. Візуалізація графа

Для покращення читабельності графа використовуйте алгоритми розташування вузлів.

1. Застосування силового алгоритму

У вікні Layout виберіть алгоритм ForceAtlas 2:

- натисніть Run для організації вузлів у природному порядку;
- коли граф стане достатньо структурованим, натисніть Stop.

2. Налаштування візуалізації

У вікні Appearance налаштуйте розмір вузлів за значенням PageRank:

- виберіть Ranking > Node Size;
- встановіть параметр PageRank для зміни розміру вузлів;
- налаштуйте колір вузлів за кластерами. Виберіть Partition > Modularity Class.

3. Додавання міток

Приклад налаштованої кластерізованої мережі наведено на рис. 4.4.

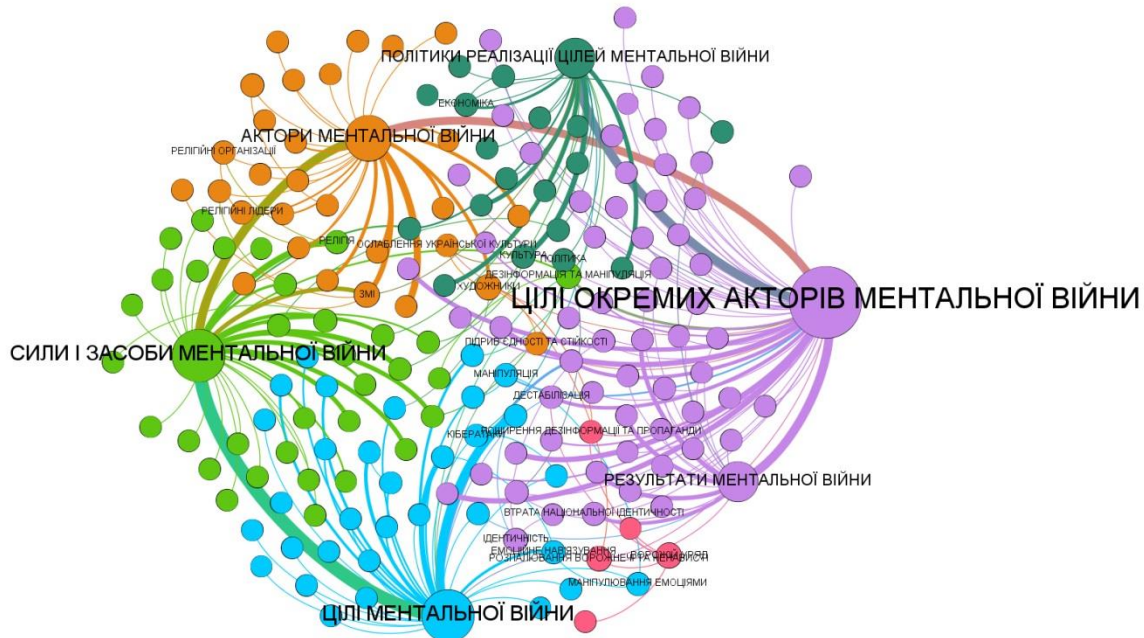


Рисунок 4.4: Приклад візуалізації кластерізованої мережі

Після завершення візуалізації графа можна проаналізувати результати:

- Ключові вузли. Великі вузли з високим PageRank є основними агентами впливу. Ранжирування вузлів побудованої мережі за PageRank у режимі «Лабораторія даних» наведено на рис. 4.5.
- Кластери. Різнокольорові групи вузлів представляють спільноти, які взаємодіють між собою.
- Посередники. Вузли з високою Betweenness Centrality є критичними точками для поширення інформації.

Для подальшого використання або презентації результатів ви можете експортувати граф:

- Зображення. Виберіть File > Export > PNG/SVG для збереження графа у графічному форматі.
- Дані. Виберіть File > Export > GraphML для збереження графа з усіма обчисленими метриками.

Label	PageRank	Modularity Class	Degree
ЦІЛІ ОКРЕМИХ АКТОРІВ МЕНТАЛЬНОЇ ВІЙНИ	0.115659	5	56
СИЛИ І ЗАСОБИ МЕНТАЛЬНОЇ ВІЙНИ	0.072316	1	35
ЦІЛІ МЕНТАЛЬНОЇ ВІЙНИ	0.071329	0	37
АКТОРИ МЕНТАЛЬНОЇ ВІЙНИ	0.056583	2	27
РЕЗУЛЬТАТИ МЕНТАЛЬНОЇ ВІЙНИ	0.043936	5	21
ПОЛІТИКИ РЕАЛІЗАЦІЇ ЦІЛЕЙ МЕНТАЛЬНОЇ ВІЙНИ	0.043138	4	22
ДЕЗІНФОРМАЦІЯ ТА МАНІПУЛЯЦІЯ	0.011431	1	6
ХУДОЖНИКИ	0.010753	2	6
ЗМІ	0.00998	2	5
ВОРОЖІЙ УРЯД	0.009956	3	5
РОЗПАЛЮВАННЯ ВОРОЖНЄЧІ ТА НЕНАВИСТІ	0.009198	3	5
ІДЕНТИФІКАЦІЯ	0.006163	5	3
ЕКОНОМІКА	0.006161	4	3
ВТРАТА НАЦІОНАЛЬНОЇ ІДЕНТИФІКАЦІЇ	0.006091	5	3
РОЗПАД ПОЛІТИЧНОЇ СИСТЕМИ	0.00605	5	3
ПОЛІТИКА	0.006015	4	3
РЕЛІГІЙНІ ОРГАНІЗАЦІЇ	0.00598	2	3
РЕЛІГІЙНІ ЛІДЕРИ	0.00598	2	3
МАНІПУЛЯЦІЯ	0.005976	0	3
РЕЛІГІЯ	0.005959	1	3
КУЛЬТУРА	0.005924	4	3
ДЕСТАБІЛІЗАЦІЯ	0.005899	5	3
КІБЕРАТАКИ	0.00587	0	3
ЕМОЦІЙНЕ НАВ'ЯЗУВАННЯ	0.005844	0	3
МАНІПУЛЮВАННЯ ЕМОЦІЯМИ	0.005803	0	3
ПОШИРЕННЯ ДЕЗІНФОРМАЦІЇ ТА ПРОПАГАНДИ	0.005783	3	3
ОСЛАБЛЕННЯ УКРАЇНСЬКОЇ КУЛЬТУРИ	0.005769	2	3
ПІДРИВ ЄДНОСТІ ТА СТИКОСТІ	0.00573	2	3
РОСІЙСЬКІ СПЕЦСЛУЖБИ	0.005151	2	2
ЕКОНОМІЧНІ ІНСТРУМЕНТИ	0.00511	1	2

Рисунок 4.5: Ранжирування вузлів побудованої мережі за PageRank

Висновки до розділу

У даному розділі були розглянуті методи аналізу та візуалізації даних, що є ключовими інструментами для дослідження інформаційно-психологічних ментальних війн. Аналіз мереж дозволяє ефективно виявляти основні інформаційні кластери, ключові вузли впливу та структуру інформаційного середовища. Візуалізація, у свою чергу, забезпечує наочне представлення взаємозв'язків між елементами інформаційної війни, спрощуючи ідентифікацію осередків дезінформації, механізмів маніпулятивного впливу та каналів поширення інформації.

Особливу увагу приділено методам кластеризації мереж, таким як алгоритми Louvain та Girvan-Newman, а також семантичній кластеризації на основі контенту. Ці методи дозволяють виділяти групи взаємопов'язаних вузлів, що допомагає аналізувати інформаційні «бульбашки», координовані групи акаунтів та інформаційні фронти. Крім того, важливим аспектом є ранжування вузлів за їхньою впливовістю, де використовуються такі метрики, як PageRank, Betweenness Centrality, Closeness Centrality та Eigenvector Centrality. Ці показники дозволяють виявляти ключових агентів впливу, посередників у поширенні інформації та найбільш активні вузли в мережі.

Візуалізація даних, зокрема за допомогою інструментів Gephi, NetworkX та Graphviz, забезпечує можливість структурувати граfi природним чином, відображати кластери та оцінювати щільність інформаційних потоків. Графові діаграми, теплові карти та кластеризовані граfi є основними підходами до представлення результатів аналізу.

Практична частина розділу демонструє, як застосовувати програмне забезпечення Gephi для завантаження, аналізу та візуалізації мережових структур. Описаний процес включає підготовку даних, обчислення метрик центральності, кластеризацію мережі, застосування силових алгоритмів для покращення читабельності графа та інтерпретацію результатів. Такий комплексний підхід дозволяє ефективно аналізувати соціальні мережі, ідентифікувати ключові акаунти та розуміти динаміку інформаційних війн.

Отже, методи аналізу та візуалізації даних є невід'ємною частиною дослідження інформаційно-психологічних ментальних війн, надаючи дослідникам потужні інструменти для розуміння складних взаємодій у сучасному інформаційному просторі.

Питання для самоконтролю

1. Що таке кластеризація мереж, і як вона допомагає аналізувати інформаційно-психологічні війни?
2. Які основні цілі кластеризації мереж у контексті інформаційних потоків?
3. Як можна виявити інформаційні "бульбашки" за допомогою методів кластеризації?
4. У чому полягає значення ідентифікації груп акаунтів або ресурсів, що діють узгоджено?
5. Які основні інформаційні фронти можуть бути виявлені за допомогою кластеризації?
6. Як працює алгоритм PageRank для ранжування вузлів у графі, і яке значення він має для аналізу інформаційних мереж?
7. Які типи кластеризації існують, і як вони застосовуються для аналізу інформаційних потоків?
8. Які переваги має модулярна кластеризація (Modularity-based Clustering) для виявлення взаємопов'язаних підмереж?
9. У чому полягає роль семантичної кластеризації (Semantic Clustering) у аналізі текстових даних?
10. Які інструменти використовуються для візуалізації інформаційних потоків, і які їх переваги?
11. Як візуалізація спрощує ідентифікацію осередків дезінформації та механізмів маніпулятивного впливу?
12. Як можна використовувати кластеризацію для аналізу груп ботів або пропагандистських каналів?
13. Як визначається оптимальний поріг для виявлення релевантних зв'язків у графі?

14. Як людська верифікація допомагає покращити точність результатів аналізу?

Література за розділом

1. Качинський А.Б. Структурно-функціональна модель системи забезпечення інформаційної та інформаційно-психологічної безпеки. Допов. Нац. акад. наук України. 2023. №1. – С. 16-23.
2. Singh, A., Rose Anish, P., Verma, A. et al. A data decomposition-based hierarchical classification method for multi-label classification of contractual obligations for the purpose of their governance. Sci Rep 14, 12755 (2024). <https://doi.org/10.1038/s41598-024-63648-x>
3. Douglas A. Luke A. User's Guide to Network Analysis in R. Springer; 1st ed. 2015 edition (December 21, 2015). 250 pp. ISBN: 978-3319238821
4. Mikhail Z. Zgurovsky, N.D. Pankratova. System analysis: Theory and applications. Springer Science & Business Media, 2007. – 447 p. ISBN 3540488804, 9783540488804
5. Семантичний нетворкінг на основі великих мовних моделей: монографія /Ланде Д.В., Страшной Л.Л. – Київ: ТОВ "Інжиніринг", 2025. – 274 с. ISBN 978-617-8180-01-0
6. Cherven Ken. (2015) Mastering Gephi Network Visualization. Packt Publishing, 378. ISBN 978-1783987344

5. Ознаки центральності у графах

Центральність у графах є важливим показником для аналізу інформаційно-психологічних війн, оскільки дозволяє визначати найвпливовіші вузли в мережі. У контексті інформаційних війн центральність допомагає виявити ключові акаунти, медіа або суб'єкти, що поширюють або контролюють інформаційні потоки.

Аналіз центральності дозволяє визначити ключових агентів впливу в інформаційних війнах. Методи ранжування та їх застосування, зокрема, у середовищі системи Gephi допомагають автоматизувати дослідження інформаційних потоків, ідентифікувати джерела впливу та прогнозувати розвиток інформаційних кампаній.

Цей розділ розглядає основні показники центральності, методи їх обчислення, а також способи ідентифікації стратегічних вузлів у мережі.

Існує кілька основних метрик центральності, які визначають впливовість вузла в мережі:

5.1. Degree Centrality (Ступенева центральність)

У мережевих моделях, які використовуються для аналізу інформаційно-психологічних війн (ІПВ), показник Degree Centrality (ступенева центральність) є одним із ключових метрик для оцінки важливості окремих понять (вузлів) у семантичній мережі. Цей показник допомагає зрозуміти, які поняття є найбільш значущими та впливовими, а також як вони взаємодіють з іншими елементами мережі.

Degree Centrality – це міра кількості прямих зв'язків (ребер), які має конкретний вузол у мережі. У контексті інформаційно-психологічних війн вузли з високою ступеневою центральністю є "ключовими" поняттями, які безпосередньо пов'язані з багатьма іншими поняттями. Такі вузли часто є центрами поширення інформації, ідей або стратегій.

Формула ступеневої центральності для вузла v має вигляд:

$$C_d(v) = \frac{\deg(v)}{N-1},$$

де $\deg(v)$ – кількість зв'язків вузла, N – загальна кількість вузлів у мережі. Для зваженої мережі враховується вага кожного ребра.

У мережах, побудованих для аналізу інформаційно-психологічних війн, показник Degree Centrality дозволяє виявити ключові поняття. Наприклад, якщо поняття "пропаганда" має високу Degree Centrality, це означає, що воно активно взаємодіє з іншими поняттями, такими як

«соціальні мережі», «дезінформація» чи «психологічний вплив». Вузли з високою Degree Centrality часто є джерелами або каналами поширення інформації. Крім того, вузли з високою Degree Centrality є потенційними мішенями для протидії, оскільки їх вилучення може значно зменшити ефективність мережі.

Наведемо деякі приклади застосування Degree Centrality у мережах ІПВ.

Наприклад, при аналізі ролі джерел поширення інформації, де вузли представляють поняття, пов'язані з інформаційними кампаніями, соціальні мережі часто мають високу Degree Centrality. Висока Degree Centrality соціальних мереж, пов'язаних з поняттями «генерація штучних мемів», «поширення дезінформації», «алгоритми рекомендацій», підкреслює їх роль як ключового каналу поширення інформації. Звідси впливає стратегія протидії – фокус на блокуванні або регулюванні алгоритмів рекомендацій у соцмережах.

Інший приклад – у мережі, де аналізуються методи впливу на суспільство, поняття «дезінформація» часто має високу Degree Centrality, причому дезінформація пов'язана з поняттями «маніпуляція інформацією», «психологічний вплив», «суспільна стабільність». У цьому випадку висока Degree Centrality показує, що дезінформація є центральним елементом у механізмах ІПВ. Звідси впливає можлива стратегія протидії – розробка систем для виявлення та запобігання поширенню дезінформації.

У випадку дослідження довгострокової кампанії високу Degree Centrality має поняття «стратегічна дезінформація», воно пов'язане з поняттями «наративи», «систематична кампанія», «суспільна поляризація». Висока Degree Centrality вказує на те, що стратегічна дезінформація є основним інструментом для досягнення довгострокових цілей. Звідси впливає необхідність аналізу та демонтажу наративів, які лежать в основі стратегічної дезінформації.

Обчислення Degree Centrality у мережах ІПВ включає наступні кроки:

Крок 1. Побудова мережі, а саме:

- Визначення вузлів (понять), наприклад, «пропаганда», «соціальні мережі», «дезінформація».
- Визначення ребер (зв'язків), наприклад, «пропаганда → поширюється через соціальні мережі».

Крок 2. Обчислення Degree Centrality:

- Для кожного вузла підраховується кількість ребер, які входять або виходять з нього.

- У зважених мережах враховується також вага ребер.

Крок 3. Візуалізація результатів:

- Вузли з високою Degree Centrality виділяються на графічному представленні мережі (наприклад, більшим розміром або іншим кольором).

До переваг Degree Centrality відносяться простота обчислення, чіткий фокус (цей показник чітко вказує на ключові вузли, які є центральними у мережі), практична користь (вузли з високою Degree Centrality можуть бути цільовими точками для протидії).

До обмежень можна віднести ігнорування глобальних зв'язків. Degree Centrality враховує лише безпосередні зв'язки, ігноруючи глобальну структуру мережі. Крім того, неспроможність врахувати напрямок, що впливає з того, що у спрямованих мережах Degree Centrality не завжди враховує напрямок зв'язків.

Показник Degree Centrality є важливим інструментом для аналізу мережових моделей у контексті інформаційно-психологічних війн. Він дозволяє виявити ключові поняття, які є центральними у процесах поширення інформації, маніпуляції свідомістю та психологічного впливу. Приклади із реальних ІПВ демонструють, як Degree Centrality може допомогти зрозуміти механізми ворожих дій та розробити ефективні стратегії протидії. Однак для більш глибокого аналізу слід комбінувати Degree Centrality з іншими метриками, такими як Betweenness Centrality або Eigenvector Centrality, щоб отримати повнішу картину.

5.2. Closeness Centrality (Близькість до інших вузлів)

Серед показників центральності особливо важливим є Closeness Centrality (Близькість до інших вузлів). Closeness Centrality – це метрика, яка вимірює, наскільки «близьким» є даний вузол до всіх інших вузлів у мережі. Вона базується на концепції найкоротших шляхів (shortest paths). Чим менше середня відстань від даного вузла до всіх інших вузлів, тим вища його Closeness Centrality.

Closeness Centrality визначає, наскільки швидко інформація може поширюватися від певного вузла по всій мережі.

Висока центральність за близькістю означає, що вузол знаходиться в центрі інформаційних потоків.

Використовується для виявлення акаунтів, які можуть швидко реагувати на інформаційні атаки.

Формула центральності за близькістю для вузла v :

$$C_c(v) = \frac{1}{\sum_{u \neq v} d(v,u)}$$

де $d(v,u)$ – довжина найкоротшого шляху між вузлом v і вузлом u , $\sum_{u \neq v} d(v,u)$ сума довжин найкоротших шляхів від вузла v до всіх інших вузлів у мережі.

Вузли з високою Closeness Centrality можуть швидко отримувати інформацію від інших вузлів і так само швидко її передавати. У контексті інформаційно-психологічних війн це можуть бути ключові медіа-канали, блогери або лідери думки, які мають великий вплив на аудиторію.

Важливо зазначити, що вузли з високою Closeness Centrality часто є "стратегічними точками" у мережі, які можуть контролювати потік інформації. Наприклад, якщо вузол має короткий шлях до багатьох інших вузлів, він може ефективно розповсюджувати пропаганду або нейтралізувати чужу інформацію.

У разі спроби дестабілізації мережі (наприклад, блокування певних каналів комунікації) вузли з високою Closeness Centrality є критичними цілями, оскільки їх вилучення може значно уповільнити поширення інформації.

У рамках аналізу мережевих моделей інформаційно-психологічних війн, Closeness Centrality може бути використана для ідентифікації ключових акторів, отримання відповідей на запитання «які медіа-ресурси або інформаційні платформи є найближчими до широкого кола користувачів?», «хто є "швидкими посередниками" у поширенні інформації?».

Цей показник при аналізі стратегій впливу дозволяє визначати, які вузли можуть ефективно розповсюджувати контрольовану інформацію або запобігати поширенню дезінформації, як можна оптимізувати маршрути передачі інформації, щоб вона досягала цільової аудиторії найшвидшим шляхом.

Вузли з найбільшим значенням Closeness Centrality є найбільш критичними для функціонування мережі. Їх вилучення може призвести до значного уповільнення поширення інформації. Такі вузли слід захистити вузли від атак або маніпуляцій.

При моделюванні поширення інформації можна визначати як швидко інформація може досягти конкретного сегменту аудиторії через певний вузол, які вузли можуть слугувати "брандмауерами", щоб зупинити поширення дезінформації.

У соціальних мережах, таких як Twitter або Facebook, користувачі з високою Closeness Centrality (наприклад, популярні блогери або журналісти) можуть швидко поширювати інформацію серед великої кількості людей.

У мережі новинних агенцій або телеканалів вузли з високою Closeness Centrality (наприклад, велика кількість партнерств або перехресних посилань) можуть ефективно розповсюджувати інформацію через різні канали.

У разі кібернетичних та інформаційних атак (наприклад, поширення фейків) вузли з високою Closeness Centrality є ключовими цілями для компрометації або контролю.

До обмежень Closeness Centrality можна віднести те, що у великих мережах обчислення Closeness Centrality може бути обчислювально дорогим, оскільки потрібно знаходити найкоротші шляхи між усіма парами вузлів. Якщо частина мережі недоступна для аналізу (наприклад, закриті групи в соціальних мережах), результати можуть бути неточними.. Якщо мережа є слабко зв'язаною (наприклад, існують окремі кластери без прямих зв'язків), Closeness Centrality може бути менш інформативною.

Closeness Centrality є важливим інструментом для аналізу мережевих моделей у контексті інформаційно-психологічних війн. Вона дозволяє ідентифікувати ключові вузли, які можуть швидко передавати інформацію, контролювати потік даних і впливати на структуру мережі. Однак її застосування потребує уважного врахування обмежень і контексту, у якому працює мережа.

5.3. Eigenvector Centrality (Центральність власного вектора)

Одним із ключових аспектів аналізу мережевих моделей є вивчення центральності (centrality), яка описує важливість окремих вузлів (наприклад, осіб, організацій, медіа-джерел) у мережі. Серед різних показників центральності особливо важливим є Eigenvector Centrality (Центральність власного вектора).

Цей показник дозволяє оцінювати впливовість вузла, враховуючи впливовість його сусідів. Він використовується для ідентифікації ключових лідерів думок або ресурсів з великим охопленням. Високе значення Eigenvector Centrality означає, що вузол пов'язаний з іншими впливовими вузлами.

Формула для розрахунку Eigenvector Centrality має вигляд:

$$A \cdot x = \lambda \cdot x,$$

де A – матриця суміжності мережі (матриця, де елемент A_{ij} дорівнює 1, якщо між вузлами i та j є зв'язок, і 0 – якщо немає), x – вектор центральності (власний вектор), λ – власне значення, яке визначає масштаб центральності.

Розв'язок цього рівняння дає вектор x , де кожен елемент x_i відповідає центральності вузла i .

Показник Eigenvector Centrality можна інтерпретувати наступним чином:

- *Якість зв'язків важливіша за кількість.* Вузол з невеликою кількістю зв'язків може бути більш центральним, якщо ці зв'язки ведуть до інших високорангових вузлів. Наприклад, у контексті ПІВ, медіа-джерело, яке має небагато контактів, але ці контакти – провідні новинні агентства, буде вважатися центральним.
- *Глобальна структура мережі.* Eigenvector Centrality враховує загальну структуру мережі, а не лише локальні зв'язки. Це робить її потужним інструментом для аналізу складних систем, таких як інформаційні мережі.
- *Кумулятивний вплив.* Центральність вузла зростає, коли він взаємодіє з іншими центральними вузлами. Це створює ефект "силового поля" у мережі, де популярні вузли стають ще популярнішими.

У рамках аналізу мережевих моделей інформаційно-психологічних війн, Eigenvector Centrality може бути використана для таких завдань, як ідентифікація ключових акторів, аналіз стратегій впливу, оцінка вразливості мережі, моделювання поширення інформації. Враховуючи цей показник можна визначати які медіа-ресурси або інформаційні платформи є найвпливовішими у мережі? Наприклад, блогери або журналісти, які взаємодіють з іншими впливовими особами. Можна відповісти на такі питання: «хто є «стратегічними посередниками» у поширенні інформації?», «які вузли можуть ефективно розповсюджувати контрольовану інформацію або запобігати поширенню дезінформації?», «як можна оптимізувати маршрути передачі інформації, щоб вона досягала цільової аудиторії через впливових посередників?», «які вузли є найбільш критичними для функціонування мережі?», «як захистити ключові вузли від атак або маніпуляцій?», «як швидко інформація може досягти конкретного сегменту аудиторії через певний вузол?», «які вузли можуть слугувати "брандмауерами", щоб зупинити поширення дезінформації?».

У соціальних мережах, таких як Twitter або Facebook, користувачі з високою Eigenvector Centrality (наприклад, популярні блогери або журналісти) можуть швидко поширювати інформацію серед великої кількості людей, оскільки вони взаємодіють з іншими впливовими особами.

У мережі новинних агенцій або телеканалів вузли з високою Eigenvector Centrality (наприклад, велика кількість партнерств або перехресних посилань) можуть ефективно розповсюджувати інформацію через різні канали.

У разі інформаційних атак (наприклад, поширення фейків) вузли з високою Eigenvector Centrality є ключовими цілями для компрометації або контролю.

До обмежень Eigenvector Centrality можна віднести розмір мережі, неповноту даних, зв'язність мережі.

У великих мережах обчислення Eigenvector Centrality може бути обчислювально дорогим, оскільки потрібно знаходити власні значення і вектори матриці суміжності.

Якщо частина мережі недоступна для аналізу (наприклад, закриті групи в соціальних мережах), результати можуть бути неточними.

Зв'язність мережі. Якщо мережа є слабо зв'язаною (наприклад, існують окремі кластери без прямих зв'язків), Eigenvector Centrality може бути менш інформативною.

Eigenvector Centrality дозволяє ідентифікувати ключові вузли, які можуть ефективно передавати інформацію, контролювати потік даних і впливати на структуру мережі. Однак її застосування потребує уважного врахування обмежень і контексту, у якому працює мережа.

5.4. Центральність посередництва (betweenness centrality)

Центральність посередництва (Betweenness Centrality) – це метрика, яка вимірює важливість вузла у мережі через його роль у з'єднанні інших вузлів. Вона показує, наскільки часто даний вузол знаходиться на найкоротшому шляху між іншими парами вузлів у мережі. У контексті інформаційно-психологічних війн ця метрика може допомогти ідентифікувати ключових посередників, які контролюють потік інформації або є критичними для передачі даних.

Центральність посередництва оцінює, наскільки часто вузол опиняється на найкоротших шляхах між іншими вузлами, та дозволяє вирішувати завдання:

1. Визначає вузли-мостики, які контролюють потоки інформації між різними групами.
2. Висока центральність посередництва означає, що вузол може блокувати або змінювати траєкторію поширення інформації.
3. Використовується для виявлення акаунтів, які можуть виступати посередниками між різними інформаційними кластерами.

Betweenness Centrality для вузла v обчислюється на основі того, скільки разів цей вузол знаходиться на найкоротшому шляху (Shortest Path) між будь-якими двома іншими вузлами у мережі. Чим більше таких шляхів проходить через вузол v , тим вища його Betweenness Centrality.

Математично, Betweenness Centrality для вузла v обчислюється за формулою:

$$C_b(v) = \sum_{s \neq v \neq t} \frac{\sigma_{st}(v)}{\sigma_{st}},$$

де σ_{st} – кількість найкоротших шляхів між вузлами s і t , а $\sigma_{st}(v)$ – кількість таких шляхів, що проходять через вузол v . Сума береться по всім парам вузлів s і t , де $s \neq v \neq t$.

Якщо вузол v не бере участі в жодному найкоротшому шляху між s і t , то його внесок у $C_b(v)$ дорівнює 0.

Наведемо такий приклад використання: якщо в мережі є кілька ізольованих кластерів, вузли з високою центральністю посередництва можуть відігравати роль зв'язкових агентів між ними. Наприклад, у соцмережах це можуть бути користувачі, що поширюють інформацію між різними спільнотами.

Вузли з високою Betweenness Centrality є «воротами» або «перехрестями», через які проходить велика частина інформації. У контексті інформаційно-психологічних війн це можуть бути медіа-канали, платформи або особи, які контролюють доступ до інформації. Вони є критичними для функціонування мережі. Якщо такий вузол вилучається або блокується, це може призвести до значного розриву комунікаційних каналів.

Крім того, вузли з високою Betweenness Centrality мають можливість маніпулювати інформацією, яка проходить через них. Наприклад, вони можуть спотворювати повідомлення або зупиняти їх поширення. Такі вузли можуть ефективно розповсюджувати інформацію серед різних груп або сегментів мережі, оскільки вони є "посередниками" між цими групами.

У рамках аналізу мережевих моделей інформаційно-психологічних війн Betweenness Centrality може бути використана для ідентифікації

ключових посередників, аналізу стратегій впливу, оцінки вразливості мережі, моделювання поширення інформації.

Цей показник дозволяє знаходити відповіді на питання:

- які медіа-ресурси, платформи або особи є основними каналами для передачі інформації між різними групами? Наприклад, соціальні мережі, новинні агенції або лідери думки;
- хто є "критичними точками" у мережі, через які проходить більшість інформаційних потоків?
- як можна оптимізувати маршрути передачі інформації, щоб вона досягала цільової аудиторії через ключових посередників?;
- які вузли можуть слугувати "брандмауерами", щоб зупинити поширення дезінформації?
- які вузли є найбільш критичними для функціонування мережі? Їх вилучення може призвести до значного зменшення впливу мережі;
- як захистити ключові вузли від атак або маніпуляцій?
- як швидко інформація може досягти конкретного сегменту аудиторії через певний вузол?;
- які вузли можуть слугувати "брандмауерами", щоб зупинити поширення дезінформації?

У соціальних мережах, таких як Twitter або Facebook, користувачі з високою Betweenness Centrality (наприклад, популярні блогери або журналісти) можуть швидко поширювати інформацію серед великої кількості людей, оскільки вони є «посередниками» між різними групами.

У мережі новинних агенцій або телеканалів вузли з високою Betweenness Centrality (наприклад, велика кількість партнерств або перехресних посилань) можуть ефективно розповсюджувати інформацію через різні канали.

У разі інформаційних атак (наприклад, поширення фейків) вузли з високою Betweenness Centrality є ключовими цілями для компрометації або контролю.

До обмежень застосування показника Betweenness Centrality можна віднести розмір мережі, неповноту даних і зв'язність мережі. У великих мережах обчислення Betweenness Centrality може бути обчислювально дорогим, оскільки потрібно знаходити найкоротші шляхи між усіма парами вузлів. Якщо частина мережі недоступна для аналізу (наприклад, закриті групи в соціальних мережах), результати можуть

бути неточними. Якщо мережа є слабо зв'язаною (наприклад, існують окремі кластери без прямих зв'язків), Betweenness Centrality може бути менш інформативною.

Betweenness Centrality є важливим інструментом для аналізу мережевих моделей у контексті інформаційно-психологічних війн. Вона дозволяє ідентифікувати ключових посередників, які контролюють потік інформації, оцінити вразливість мережі та розробити стратегії для захисту або атаки на інформаційні канали. Однак її застосування потребує уважного врахування обмежень і контексту, у якому працює мережа.

5.5. Виявлення впливів в мережах

У цьому пункті розглядається, як математичні моделі центральності застосовуються до графових моделей інформаційно-психологічних війн, де:

- вузли – поняття (наприклад, «ворожа пропаганда», «національна ідентичність»),
- зв'язки – впливи (наприклад, логічні зв'язки, причинно-наслідкові взаємодії, семантичні асоціації).

Мета аналізу таких мереж – ідентифікувати ключові поняття, які є стратегічно важливими для поширення ідеологічних наративів, контрольованого впливу на суспільство або захисту від маніпулятивних кампаній.

5.5.1. Семантична модель

Семантична модель у контексті вивчення впливу в мережах – це інструмент, який дозволяє аналізувати взаємодії між поняттями через їхні семантичні зв'язки. У разі дослідження інформаційно-психологічної ментальної війни така модель допомагає виявити ключові концепти, їх взаємодію та механізми поширення впливу через ментальні конструкції, які формуються у свідомості людей.

Семантична модель представляє собою граф, де вершини (вузли) відповідають окремим поняттям, а ребра між ними показують зв'язки між цими поняттями. Наприклад граф: «ворожа пропаганда» → «дезінформація» → «соціальна напруга» означає, що поняття «ворожа пропаганда» пов'язане з поняттям «дезінформація», а те, у свою чергу, з поняттям «соціальна напруга».

Зв'язки між поняттями можуть бути різними за типом і силою. Вони визначаються на основі текстового аналізу. Наприклад:

- 1) частота спільного вживання слів. Якщо два слова часто зустрічаються разом у текстах, то між ними можна припустити сильний зв'язок;
- 2) контекстуальна близькість. Якщо два поняття часто використовуються в одному контексті, то це може свідчити про їхню семантичну близькість;
- 3) логічні зв'язки. Деякі зв'язки можуть бути встановлені на основі логічних причинно-наслідкових відносин (наприклад, «пропаганда» → «маніпуляція свідомістю»).

Для побудови семантичної моделі використовуються методи обробки природної мови (NLP), такі як:

- TF-IDF (Term Frequency-Inverse Document Frequency) для визначення важливості слів у тексті;
- Word Embeddings (наприклад, Word2Vec, GloVe) для представлення слів у вигляді векторів, які враховують їхню семантичну близькість;
- Тематичне моделювання (наприклад, LDA) для виділення тем і взаємозв'язків між ними.

У мережі понять деякі вузли мають більший вплив, ніж інші. Це можна визначити за допомогою метрик центральності, які розглядалися вище, наприклад:

- Degree Centrality: кількість зв'язків у вузла;
- Betweenness Centrality: роль вузла у передачі інформації між іншими вузлами;
- Eigenvector Centrality: вплив вузла, враховуючи вплив його сусідів. Наприклад, у вашій мережі поняття "дезінформація" може мати високу центральність, оскільки воно є проміжним ланцюгом між поняттями «ворожа пропаганда» і «соціальна напруга».

Семантична модель дозволяє аналізувати, як змінюються зв'язки між поняттями з часом. Наприклад, якщо слово «ворожа пропаганда» починає частіше з'являтися в медіа, то це може посилювати зв'язок з поняттям «дезінформація».

Динамічні зміни у мережі можуть свідчити про зростання або зменшення впливу певних концептів.

Семантична модель допомагає зрозуміти, як ментальні конструкції формується у свідомості людей. Наприклад:

- Повторне використання певних слів або фраз у медіа може створювати «ментальні шаблони», які впливають на сприйняття реальності.
- Зв'язки між поняттями, такими як «ворожа пропаганда» → «дезінформація», можуть формувати негативні емоційні реакції, що призводить до «соціальної напруги».

Семантична модель може бути використана для:

- 1) виявлення ключових наративів. Визначення, які поняття є найбільш впливовими у формуванні інформаційних кампаній;
- 2) прогнозування наслідків. Аналіз, як зміна одного поняття може вплинути на інші;
- 3) контролю інформаційного поля. Виявлення слабких місць у мережі понять, які можуть бути використані для протидії ворожій пропаганді.

Розглянемо мережу понять, пов'язаних з інформаційно-психологічною війною.

1. Побудова мережі

Вершини: «ворожа пропаганда», «дезінформація», «соціальна напруга», «маніпуляція свідомістю», «психологічний стрес».

Ребра:

- «ворожа пропаганда» → «дезінформація»;
- «дезінформація» → «соціальна напруга»;
- «соціальна напруга» → «психологічний стрес».

2. Аналіз центральності:

- вузол «дезінформація» може мати високу центральність, оскільки він є проміжним ланцюгом між поняттями «ворожа пропаганда» та «соціальна напруга»;
- вузол «соціальна напруга» може мати високу Betweenness Centrality, оскільки він є ключовим для передачі впливу до вузла «психологічний стрес».

3. Динаміка

Якщо частота вживання слова «ворожа пропаганда» зростає, то це може посилити зв'язок з поняттям «дезінформація», що, у свою чергу, призведе до зростання ваги поняття «соціальна напруга».

Таким чином, семантична модель є потужним інструментом для аналізу взаємодії понять у контексті інформаційно-психологічної ментальної війни. Вона дозволяє виявляти ключові концепти та їх взаємозв'язки, прогнозувати наслідки поширення певних наративів, розробляти стратегії протидії ворожій пропаганді.

Цей підхід особливо корисний для аналізу складних інформаційних полів, де важливо зрозуміти, як формуються ментальні конструкції та як вони впливають на суспільство.

5.5.2. Модель перколяції

Окремо розглянемо модель перколяції як математичний інструмент, що дозволяє досліджувати, як зміни в структурі мережі (наприклад, вилучення ключових вузлів або зв'язків) впливають на її глобальну зв'язність.

Перколяція описує процес, коли частина вузлів або зв'язків вилучається з мережі (або, навпаки, додається). Цей процес дозволяє визначити, як змінюється зв'язність мережі.

У контексті інформаційно-психологічної війни ця модель може допомогти зрозуміти, як вилучення або посилення певних понять (вузлів) впливає на поширення ідей, наративів та впливу в інформаційному полі.

Особливої уваги при вивченні перколяції приділяється критичній пороговій точці (p_c), що відповідає стану, коли мережа переходить від зв'язного стану до розрідженого (або навпаки). Критична порогова точка визначається як:

$$p_c = \frac{\langle k \rangle}{\langle k^2 \rangle - \langle k \rangle \langle k \rangle},$$

де $\langle k \rangle$ – середня кількість зв'язків на вузол (середній ступінь вузлів), $\langle k^2 \rangle$ – другий момент розподілу зв'язків (середнє значення квадратів ступенів вузлів).

Ця формула показує, при якій частці збережених зв'язків (p_c) мережа залишається зв'язною.

У теоретико-графовій моделі ІПВ перколяційна модель застосовується наступним чином:

1. Визначається початкова мережа понять. Припустимо, у нас є мережа понять, пов'язаних з інформаційно-психологічною війною:

«Ворожа пропаганда» → «Дезінформація» → «Соціальна напруга» → «Психологічний стрес».

Кожне поняття має певну кількість зв'язків з іншими поняттями.

2. Реалізується процедура вилучення ключових понять. У разі вилучення ключових понять (наприклад, «Дезінформація»), мережа може втратити свою зв'язність. Це означає, що ідеї, пов'язані з

«Ворожою пропагандою», не зможуть ефективно поширюватися далі через поняття «Соціальна напруга».

3. Визначається критична порогова точка. Якщо частка збережених зв'язків (p) перевищує p_c , мережа залишається зв'язною, й ідеї продовжують поширюватися.

Якщо $p < p_c$, мережа розпадається на окремі компоненти, і поширення ідей стає неможливим.

У контексті аналізу ІПВ реалізується визначення ключових понять, вилучення яких призведе до руйнування мережі. Наприклад, якщо "Дезінформація" має високий ступінь зв'язності, її вилучення може суттєво послабити вплив ворожої пропаганди. При цьому, якщо частка збережених зв'язків падає нижче p_c , то це може свідчити про те, що інформаційна кампанія втрачає ефективність.

Все це дозволяє виявляти слабкі місця у мережі понять, які можуть бути використані для протидії ворожій пропаганді.

Як модельний приклад розглянемо мережу, де вузлами є поняття, пов'язані з ІПВ з вершинами:

- «Ворожа пропаганда»,
- «Дезінформація»,
- «Соціальна напруга»,
- «Психологічний стрес».

Як ребра визначимо:

- «Ворожа пропаганда» → «Дезінформація».
- «Дезінформація» → «Соціальна напруга».
- «Соціальна напруга» → «Психологічний стрес».

Обчислення параметрів:

- середня кількість зв'язків ($\langle k \rangle$): 2;
- другий момент розподілу зв'язків ($\langle k^2 \rangle$): 5;
- критична порогова точка:

$$p_c = \frac{\langle k \rangle}{\langle k^2 \rangle - \langle k \rangle \cdot \langle k \rangle} = \frac{2}{5 - 4} = 2.$$

Аналіз впливу показує:

- якщо частка збережених зв'язків (p) становить 0,8 ($p > p_c$), мережа залишається зв'язною, і ідеї продовжують поширюватися;

- якщо частка збережених зв'язків (p) становить 0,5 ($p < p_c$), мережа розпадається на окремі компоненти, і поширення ідей стає неможливим.

Вилучення ключового поняття:

Припустимо, що поняття «Дезінформація» вилучається з мережі. Це призводить до того, що зв'язки між "Ворожною пропагандою" і «Соціальною напругою» розриваються. Як наслідок, мережа стає менш зв'язною.

Практичне застосування моделі перколяції можливо для вирішення таких задач:

- Виявлення ключових наративів. Модель перколяції дозволяє виявити, які поняття є найбільш важливими для збереження зв'язності мережі. Наприклад, «Дезінформація» може бути ключовим поняттям, яке треба протидіяти.

- Прогнозування ефективності кампаній. Якщо частка збережених зв'язків (p) падає нижче p_c , це може свідчити про те, що інформаційна кампанія втрачає ефективність.

- Розробка стратегій протидії. Виявлення слабких місць у мережі понять, які можуть бути використані для протидії ворожій пропаганді. Наприклад, вилучення ключових понять або зменшення їхньої видимості у медіа.

Таким чином, модель перколяції є потужним інструментом для аналізу зв'язності мережі понять у контексті інформаційно-психологічної війни. Вона дозволяє виявляти ключові поняття, які є критичними для збереження зв'язності мережі, прогнозувати наслідки вилучення або посилення певних понять, а також, розробляти стратегії протидії ворожій пропаганді.

Цей підхід може бути особливо корисним для аналізу складних інформаційних полів, де важливо зрозуміти, як формуються та поширюються ідеї та наративи.

5.5.3. Модель розповсюдження інформації: SIR-модель

SIR-модель (Susceptible-Infected-Recovered) – це класична математична модель, яка спочатку була розроблена для опису поширення епідемій. Однак її можна адаптувати для аналізу поширення інформації або ідей у соціальних та інформаційних мережах. У контексті інформаційно-психологічної війни ця модель дозволяє досліджувати, як поняття (вузли графа) активуються, поширюються та втрачають актуальність.

До основних компонент SIR-моделі відносяться стан понять, параметри моделі і опис динаміки процесу.

1. Стан понять:

- S (Susceptible) – поняття, які ще не активовані, але можуть бути активовані через зв'язки із активними поняттями;
- I (Infected) – активовані поняття, які активно поширюються та впливають на інші поняття;
- R (Recovered) – поняття, які втратили актуальність або були "забуті" (більше не беруть участі у поширенні).

2. Параметри моделі:

- β (швидкість поширення інформації) – ймовірність того, що активоване поняття (I) передасть свій вплив сусідньому поняттю (S);
- γ (швидкість "вилікування") – ймовірність того, що активоване поняття (I) перейде у стан R (втратить актуальність).

3. Динаміка процесу:

- поняття з S переходять у I через взаємодію з активованими поняттями (I);
- поняття з I поступово перетворюються на R , коли їхній вплив слабшає.

SIR-модель описує поширення інформації через поняття:

$$\frac{dS}{dt} = -\beta SI, \quad \frac{dI}{dt} = \beta SI - \gamma I, \quad \frac{dR}{dt} = \gamma I.$$

Опишемо основні етапи застосування SIR-моделі:

1. Початкова активація

Деякі поняття в графі стають активними (I) через зовнішній вплив (наприклад, медіа-кампанії, ворожа пропаганда). Наприклад, поняття «ворожа пропаганда» може бути активовано через інформаційну атаку.

2. Поширення

Активовані поняття (I) передають свій вплив сусіднім поняттям (S) через зв'язки у графі.

Швидкість передачі визначається параметром β . Чим вища β , тим швидше поширюється інформація.

Наприклад: «Ворожа пропаганда» → «дезінформація».

«Дезінформація» → «соціальна напруга».

3. Забуття

Поняття, які довго залишаються активними (I), поступово втрачають актуальність і переходять у стан R .

Швидкість забуття визначається параметром γ . Чим вища γ , тим швидше поняття втрачають силу.

Наприклад: «Дезінформація» $\rightarrow$ «забуття» (перехід у R).

У графовій моделі ІПВ деякі поняття мають високу центральність (наприклад, «національна ідентичність»). Це означає, що вони мають багато зв'язків із іншими поняттями. Такі поняття є ключовими для поширення інформації, оскільки:

1. Активація ключових понять. Якщо активовано поняття з високою центральністю, воно швидко поширює вплив на багато інших понять. Наприклад, активація поняття «національна ідентичність» може активувати поняття «патріотизм», «єдність» тощо.
2. Сприйнятливості до впливу. Поняття з високою центральністю також є більш сприйнятливими до активації через зв'язки із іншими поняттями.
3. Прогнозування поширення. Виявлення понять з високою центральністю допомагає прогнозувати, як буде поширюватися інформація у мережі.

Розглянемо як приклад граф понять, пов'язаних з ІПВ:

1. Мережа понять

Вершини: «національна ідентичність», «ворожа пропаганда», «дезінформація», «соціальна напруга», «патріотизм».

Ребра: «національна ідентичність» $\rightarrow$ «патріотизм»; «ворожа пропаганда» $\rightarrow$ «дезінформація»; «дезінформація» $\rightarrow$ «соціальна напруга».

2. Параметри

Початковий стан:

- S : «національна ідентичність», «патріотизм», «соціальна напруга»;
- I : «ворожа пропаганда»;
- R : порожня множина;
- $\beta=0,7$ (висока ймовірність передачі впливу);
- $\gamma=0,3$ (помірна швидкість забуття).

3. Динаміка

Крок 1. «Ворожа пропаганда» активує вузол «дезінформація».

Крок 2. «Дезінформація» активує вузол «соціальна напруга».

Крок 3. Вузол «Ворожа пропаганда» переходить у R (втрачає актуальність).

4. Результат

На кінець процесу:

- S: «національна ідентичність», «патріотизм»;
- I: «соціальна напруга»;
- R: «ворожа пропаганда», «дезінформація».

SIR-модель у контексті ППВ дозволяє аналізувати поширення ідей та понять у мережі враховуючи динаміку активації, поширення та забуття. Поняття з високою центральністю є основними джерелами активації, а параметри β та γ визначають швидкість поширення та забуття.

Ця модель виявилась потужним інструментом для аналізу складних інформаційних полів, де важливо зрозуміти, як формуються та поширюються ідеї та наративи. Вона знаходить практичне застосування при вирішенні наступних питань:

- виявлення ключових понять для протидії ворожій пропаганді;
- прогнозування наслідків інформаційних кампаній;
- розробка стратегій захисту інформаційного поля.

5.6. Приклади застосування у реальних ситуаціях

У ході інформаційних кампаній проти вакцинації ключовими поняттями можуть бути:

- «дезінформація» (висока Betweenness Centrality);
- «здоров'я» (висока Closeness Centrality);
- «наука» (висока Eigenvector Centrality).

У політичних наративах ключовими поняттями можуть бути:

- «державний суверенітет» (висока Eigenvector Centrality);
- «економічна безпека» (висока Betweenness Centrality);
- «патріотизм» (висока Degree Centrality).

Математична формалізація та моделі дозволяють точно виявляти ключові поняття у графових моделях інформаційно-психологічних війн. Застосування метрик центральності (Degree, Betweenness, Closeness, Eigenvector, PageRank) разом з моделями перколяції, семантичного аналізу та SIR-моделлю забезпечує комплексний аналіз структури та функціонування мереж.

5.7. Практична частина. Розрахунок центральності

Цей пункт демонструє, як використовувати для аналізу Gephi центральності в інформаційних мережах. Він включає всі етапи – від створення графа до інтерпретації результатів. Такий підхід дозволяє автоматизувати дослідження інформаційних потоків, ідентифікувати ключові вузли та прогнозувати розвиток інформаційних кампаній.

Gephi надає широкий набір інструментів для аналізу графів. У цьому пункті розглядається, як за допомогою програмного забезпечення Gephi можна автоматизувати процес обчислення основних показників центральності для аналізу інформаційних мереж.

Етап 1. Формування інформаційної мережі

Інформаційну мережу ІПВ «ментальна війна» будемо формувати шляхом знаходження понять предметної області та зв'язків між ними, для чого будемо послідовно знаходити пари зв'язаних понять за допомогою LLM, які виводяться у вигляді нумерованого списку у форматі «поняття; поняття», що відповідає формату CSV.

Для завантаження отриманої мережі у систему Gephi можна скористатись режимом “Agency list”, вибрати кодування windows-1251 і у явному вигляді задати розподільувач – крапку з комою, semicolon.

Для отримання даних для мережі створимо первинний промпт і модифікуємо його у відповідності із вже розглянутим фреймворком безкодового програмування.

Запишемо первинний промпт і запит на його перетворення у вигляді:

Перетвори у відповідність з наведеною методикою промпт: "Знайди 20 зв'язаних між собою пар понять за тематикою «ментальна війна» і виведи їх у вигляді нумерованого списку у форматі «Поняття 1; Поняття2»."

Після застосування фреймворку безкодового програмування отримуємо новий структурований промпт.

1. Аналіз завдання

Промпт складається з таких елементів:

Потрібно знайти пари понять, пов'язаних між собою за тематикою «ментальна війна».

Результат має бути представлено у вигляді нумерованого списку.

Кількість пар: 20.

Це завдання можна реалізувати за допомогою композиції примітивів: "Цикл" для ітерації по парам та "Функція" для генерації кожної пари.

2. Формалізація за допомогою примітивів

Примітив "Цикл" (For-Loop)

S : Множина ітерацій – це кількість пар, які потрібно згенерувати. У нашому випадку $S=\{1,2,...,20\}$.

F : Операція, яка застосовується до кожного елемента множини. У нашому

випадку це генерація однієї пари понять за тематикою «ментальна війна».

Примітив "Функція" (Abstraction)

Fextract(x, параметр) : Функція, яка генерує одну пару понять.

x : Вхідний текст або контекст, що описує тематику «ментальна війна».

параметр : Додаткові інструкції, наприклад, тип зв'язку між поняттями (антоніми, синоніми, причинно-наслідкові зв'язки тощо).

3. Реалізація через композицію примітивів

3.1. Промпт для "Функції"

Промпт для генерації однієї пари понять:

"Знайди два поняття, пов'язані між собою за тематикою 'ментальна війна', і виведи їх у форматі 'Поняття 1;Поняття 2'."

3.2. Промпт для "Циклу"

Тепер об'єднаємо цей процес у цикл для генерації 20 пар:

"Для кожного номера від 1 до 20 виконай наступне: знайди два поняття, пов'язані між собою за тематикою 'ментальна війна', і виведи їх у форматі 'Поняття 1;Поняття 2'. Результати виведи у вигляді нумерованого списку."

4. Кінцевий промпт

Остаточний промпт, що враховує всі базові примітиви:

"Створи нумерований список із 20 пар понять, пов'язаних між собою за тематикою 'ментальна війна'. Для кожної пари виведи їх у форматі 'Поняття 1;Поняття 2'."

5. Синтаксичне представлення (AST)

У вигляді абстрактного синтаксичного дерева (AST):

Prompt ::= Цикл(20, Функція("Знайти пару понять за тематикою 'ментальна війна'"))

Виконання кінцевого промпту видає результати вигляду:

Результат:

Пропаганда;Дезінформація
Ідеологія;Ментальні стереотипи
Психологічний тиск;Емоційне вигорання
Медіа-вплив;Соціальна маніпуляція
Когнітивний дисонанс;Інформаційна перевантаженість
Культурний вплив;Ідеологічна експансія
Символічна боротьба;Конфлікт ідентичностей
Інформаційна атака;Цифрова оборона
Психологічна стійкість;Стресостійкість
Соціальна ізоляція;Психологічна резистентність
Інформаційна війна;Кібернетична боротьба
Меметика;Культурна колонізація
Емоційна маніпуляція;Психологічна залежність
Ідеологічна боротьба;Культурний конфлікт
Соціальна депресія;Колективна апатія
Інформаційний шум;Когнітивна блокада
Психологічна агресія;Емоційна реакція
Деструктивне мислення;Ментальна деформація
Інформаційна стратегія;Психологічна тактика
Соціальна тривожність;Колективний страх

Результат 5-разового виконання цього структурованого промпту у відповідності із концепцією «рою віртуальних агентів» зберігаємо у файлі mental.csv

Етап 2. Відображення графової структури

Для відображення графа, який буде моделювати інформаційну мережу, у середовищі системи Gephi виберемо режим File → Open, позначимо файл mental.csv для завантаження, після чого введемо параметри:

Separator: Semicolon;

Import as: Ajacency list;

Charset: windows-1251.

Після завантаження, визначенн кластерів, укладання (див. попередні розділи), мережа набуває вигляд, зображений на рис. 5.1.



Рисунок 5.1: Модель взаємозв'язку понять ментальної війни, що згенеровані шляхом застосування мовної моделі Qwen2.5-max

Етап 3. Обчислення показників центральності

Тепер в «Лабораторії даних» системи Gephi розглянемо показники центральності для вузлів побудованого графа.

1. Degree Centrality (Ступенева зважена центральність)

На рис. 5.2 наведено перелік вузлів з найбільшими ступенями. Це такі вузли-поняття, як «Соціальна тривога», «Психологічна атака», «Емоційна вразливість», «Культурний шок» тощо. На рис. 5.3 наведено перелік вузлів, що мають найбільшу Closeness Centrality, на рис. 5.4 – вузлів, що мають найбільшу центральність посередництва, на рис. 5.5 – вузлів, що мають найбільшу центральність власного вектора. Як можна побачити, що списки найбільш вагомих вузлів за різними критеріями не співпадають.

Label	Weighted Degree
Соціальна тривога	5.0
Психологічна атака	5.0
Емоційна вразливість	5.0
Культурний шок	5.0
Інформаційна блокада	5.0
Інформаційний хаос	5.0
Соціальна конформність	5.0
Ментальна колонізація	5.0
Психологічна асиметрія	5.0
Емоційна симетрія	5.0
Інформаційна диверсія	5.0
Когнітивна стабільність	5.0
Ідеологічна ескалація	5.0
Ментальна деескалація	5.0

Рисунок 5.2: Вузли, що мають найбільшу ступеневу центральність

2. Closeness Centrality (Близькість до інших вузлів, рис. 5.3)

Label	Closeness Centrality
Психологічна атака	1.0
Емоційна вразливість	1.0
Культурний шок	1.0
Інформаційна блокада	1.0
Інформаційний хаос	1.0
Соціальна конформність	1.0
Ментальна колонізація	1.0
Психологічна асиметрія	1.0
Емоційна симетрія	1.0
Інформаційна диверсія	1.0
Когнітивна стабільність	1.0
Ідеологічна ескалація	1.0
Ментальна деескалація	1.0
Культурна колонізація	1.0
Психологічна напруга	1.0
Емоційний вибух	1.0

Рисунок 5.3: Вузли, що мають найбільшу Closeness Centrality

3. Betweenness Centrality (Центральність посередництва, рис.5.4)

Label	Betweenness Centrality
Колективна надія	4.0
Соціальна конформність	3.0
Інформаційна стратегія	3.0
Соціальна тривога	3.0
Соціальна депресія	3.0
Культурна інвазія	2.0
Ідеологічна резистентність	2.0
Меметичний вплив	2.0
Соціальна адаптація	2.0

Рисунок 5.4: Вузли, що мають найбільшу центральність посередництва

4. Eigenvector Centrality (Центральність власного вектора, рис. 5.5)

Label	Eigenvector Centrality
Інформаційна стратегія	1.0
Соціальна конформність	1.0
Колективна надія	0.910662
Соціальна депресія	0.788664
Соціальна тривога	0.788664
Культурна інвазія	0.614797
Ідеологічна резистентність	0.614797
Меметичний вплив	0.614797
Соціальна адаптація	0.614797

Рисунок 5.5: Вузли, що мають найбільшу центральність власного вектора

Висновки до розділу

У даному розділі розглянуто методи аналізу центральності у графах, які є важливими інструментами для дослідження структури інформаційно-психологічних ментальних війн. Центральність у мережах дозволяє виявляти ключові вузли (поняття, медіа-ресурси, акаунти тощо), що відіграють стратегічну роль у поширенні інформації, дезінформації та ідеологічних наративів.

Основними показниками центральності, розглянутими у розділі, є:

1. Degree Centrality (Ступенева центральність). Визначає кількість прямих зв'язків вузла, що допомагає ідентифікувати ключові поняття, які активно взаємодіють з іншими елементами мережі. Наприклад, поняття "дезінформація" може мати високу Degree Centrality через свою взаємодію з іншими концепціями, такими як "маніпуляція" або "суспільна напруга". Однак цей показник не враховує глобальних зв'язків у мережі.

2. Closeness Centrality (Близькість до інших вузлів). Вимірює, наскільки швидко інформація може поширюватися від певного вузла до всіх інших. Вузли з високою Closeness Centrality є стратегічними точками, які можуть оперативно реагувати на інформаційні атаки.

Наприклад, медіа-канали або блогери з високою Closeness Centrality можуть швидко досягати великої аудиторії.

3. Eigenvector Centrality (Центральність власного вектора). Оцінює впливовість вузла, враховуючи якість його зв'язків. Вузли з високою Eigenvector Centrality пов'язані з іншими впливовими вузлами, що робить їх особливо значущими. Наприклад, поняття "державний суверенітет" може мати високу Eigenvector Centrality через зв'язки з іншими важливими концепціями, такими як "права людини" чи "міжнародне право".

4. Betweenness Centrality (Центральність посередництва). Визначає, наскільки часто вузол є посередником у передачі інформації між іншими вузлами. Вузли з високою Betweenness Centrality контролюють потоки інформації між різними групами або кластерами. Наприклад, соціальні мережі або новинні платформи можуть бути критичними точками, через які проходить більшість інформаційних потоків.

Також у розділі розглянуто комплексні моделі аналізу графових структур, такі як:

- семантична модель, яка аналізує семантичні зв'язки між поняттями;
- модель перколяції, що дозволяє оцінити вплив видалення ключових вузлів на зв'язність мережі;
- SIR-модель, яка описує динаміку поширення інформації через мережу.

Застосування метрик центральності та моделей аналізу дозволяє вирішувати практичні завдання, такі як:

- ідентифікація ключових акторів у інформаційних кампаніях (наприклад, аналіз проти-вакцинаційних наративів);
- аналіз політичних наративів (наприклад, виявлення ролі понять "державний суверенітет" або "патріотизм");
- оцінка вразливості мережі та розробка стратегій захисту від інформаційних атак.

Таким чином, використання математичних моделей та метрик центральності забезпечує комплексний підхід до аналізу інформаційно-психологічних війн, дозволяючи точно виявляти ключові поняття, прогнозувати поширення інформації та розробляти ефективні стратегії протидії маніпулятивним кампаніям. Однак слід враховувати обмеження кожного методу, зокрема обчислювальну складність для великих мереж та можливу неповноту даних.

Питання для самоконтролю

1. Що таке центральність у графах, і як вона використовується для аналізу інформаційних мереж?
2. Як визначається Degree Centrality, і які вузли вона виділяє у мережі?
3. У чому полягає значення Closeness Centrality для оцінки швидкості поширення інформації?
4. Як працює Betweenness Centrality, і яку роль вона відіграє у контролі комунікацій між групами?
5. Що таке Eigenvector Centrality, і чому вона важлива для виявлення впливових вузлів?
6. Які вузли є найбільш критичними для функціонування мережі, і як їх захистити від атак?
7. Як моделюється поширення інформації через певний вузол, і як це допомагає прогнозувати інформаційні атаки?
8. Які обмеження існують при обчисленні Eigenvector Centrality, і як їх можна подолати?
9. Як впливає розмір мережі на обчислення центральності, і які методи можуть бути застосовані для великих графів?
10. Як можна використовувати вузли з високою Betweenness Centrality для блокування комунікації між групами?
11. Як вузли з високою Eigenvector Centrality можуть бути використані для поширення контрнарративів?
12. Як візуалізуються результати аналізу центральності, і як це допомагає інтерпретувати дані?
13. Як можна застосувати аналіз центральності для виявлення вразливостей у соціальних мережах?
14. Які етапи включає процес аналізу центральності, і як вони взаємодіють між собою?
15. Які практичні застосування аналізу центральності у контексті інформаційно-психологічних війн?

Література за розділом

1. Основи теорії складних мереж: навч. пос. / Снарський А.О., Ланде Д.В., Субач І.Ю. ІСЗЗІ КПІ ім. Ігоря Сікорського. – Київ: ТОВ "Інжиніринг", 2023. – 225 с. ISBN 978-9662344950.
2. Візуалізація та аналіз мережових структур : навч. пос. / Ланде Д.В., Субач І.Ю. – Київ : КПІ ім. Ігоря Сікорського, Вид-во "Політехніка", 2021. – 80 с. ISBN 978-9662577143.

3. Mark Newman. Networks (2nd edn). Oxford University Press, 2018. – 800 p. ISBN 978-0198805090
4. Ланде Д., Сулема О. Порівняльна оцінка критеріїв центральності в ієрархічних мережах // Інформаційні технології та безпека, 2015. – Т. 3. – N 2(5). – С. 80-87.
5. Сулема О.К., Ланде Д.В. Знаходження оптимальної ієрархії у квазіієрархічному графі за критеріями центральності // Реєстрація, зберігання і обробка даних, 2015. – Т. 17. – N 4. – С. 3-10.
6. Douglas A. Luke A User's Guide to Network Analysis in R. Springer; 1st ed. 2015 edition (December 21, 2015). 250 pp. ISBN: 978-3319238821

6. Моделювання інформаційно-психологічних війн

6.1. Структура й особливості інформаційно-психологічної війни

Термін «інформаційна війна» першим використав Томас Рон у звіті «Системи зброї й інформаційна війна», підготовленому в 1976 р. для компанії Boeing. Він зазначав, що інформаційна інфраструктура набуває статусу ключового компоненту американської економіки: саме вона стає найбільш вразливою ціллю як у воєнний, так і мирний час.

Проблема зацікавила американських військових. Військово-повітряні сили США почали активно обговорювати цей предмет ще з 1980 р. Саме тоді було досягнуто єдиного розуміння, що інформація може виступати як ціллю, так і зброєю. Так, у директиві міністра оборони США (DoD Directive 3600 від 21 грудня 1992 року) цей термін був введений офіційно. Зазначений документ був доповнений «Доктриною боротьби з системами контролю та управління», прийнятою Міністерством оборони США в лютому 1996 року. Вважається, що саме тоді була сформульована доктрина інформаційної війни. В 1993 р. в директиві №30 Комітету начальників штабів Міністерства оборони США вже були викладені основні принципи ведення інформаційної війни.

В серпні 1995 р. національний Інститут Оборони США публікує роботу Мартіна Лібікі «Що таке інформаційна війна?», в якій автор визначає сім форм інформаційної війни: командно-управлінська; розвідувальна; психологічна; хакерська; економічна; електронна; кібервійна.

У 1997 р. було дано наступне визначення інформаційної війни – «дії, спрямовані на досягнення інформаційної переваги в інтересах національної стратегії і здійснювані шляхом впливу на інформацію й інформаційні системи противника при одночасному захисті власної інформації і своїх інформаційних систем».

В жовтні 1998 р. Міністерство оборони США вводить у дію «Об'єднану доктрину інформаційних операцій», в якій вперше офіційно підтверджується факт підготовки американців до проведення наступальних інформаційних операцій [1-3].

Розуміння категорії «інформаційна війна» пройшло непростий шлях. Якщо на початку 90-х під інформаційними війнами розуміли операції психологічного й ідеологічного впливу, то в останні 20-років стали розглядати наступні їх типи: ментальні війни, консцієнтальні війни, когнітивні війни, контентні війни, кібервійни, мережоцентричні війни,

інформаційно-алгоритмічні війни, поведінкові війни, цивільні мережеві війни, кібертерористичні війни тощо [4- 6]. Інформаційно-психологічна війна є найвищим ступенем інформаційного протиборства і спрямована на розв'язання суспільно-політичних, ідеологічних, а також національних, територіальних та інших конфліктів між державами, народами, націями, класами й соціальними групами у спосіб широкомасштабної реалізації способів і методів інформаційного насильства [7].

В структурі інформаційно-психологічної війни ключову роль відіграє *пропаганда*, яку досліджує низка дисциплін, від історичних наук до міждисциплінарного системного підходу. Загальним для цих досліджень є те, що вони розглядають пропаганду як різновид комунікації [8].

Найчастіше вживаними синонімами слова «пропаганда» є брехня, викривлення, обман, маніпуляція, контроль над розумом, психологічна війна, ідеологічна обробка, балаканина, сфабриковані матеріали. Нині популярними стали наступні терміни: підтасовка фактів, альтернативні факти, фейкові новини, організоване переконання. Загалом до пропаганди сформувалося негативне і підозріле ставлення.

На нашу думку, найбільш придатне визначення терміну «пропаганда» зробили Гарт С. Джоует і Вікторія О'Донел [8]. *Пропаганда* – умисна, систематична спроба сформувані сприйняття, маніпулювати пізнанням і керувати поведінкою для досягнення реакції, що сприяє бажаним намірам пропагандиста.

За допомогою цього визначення можна ідентифікувати й оцінити характеристики пропаганди, проаналізувати якість контексту, відправника, наміри, повідомлення, канал, аудиторію і реакцію.

Комунікативний підхід до дослідження пропаганди дає змогу розглядати комунікативний процес переконування як такий, що спрямований на змінення утверджених ідей або думок, які людина має про світ, себе чи інших, власне для зміни *переконань*. Пропагандист для зміни старих переконань або створення нових спирається на старі переконання, які уже існують в умах аудиторії. Переконування є інтерактивним процесом, у якому обидві сторони залежать одна від одної, намагаються задовольнити потреби як переконуючого, так і того, кого переконують, і це передбачає зміну комунікативних ролей. Окремої уваги заслуговує трансактивність цього процесу, коли взаємодіючі сторони формують усталену думку і створюють новий спільний погляд на що-небудь.

Зазначимо, що пропагандист намагається досягти певної реакції людини й суспільства, намагаючись змінити установки і поведінку

аудиторії. Під *установкою* розуміємо психологічний стан відносно стійкої схильності суб'єкта до певної програми поведінки в певній ситуації. Можливі три різні форми реакції :

1. *Формування реакції*. За допомогою прикладів пропагандист здійснює спроби викликати реакцію аудиторії, навчаючи її елементам бажаної поведінки.

2. *Посилення реакції*. Якщо в аудиторії вже сформоване певне ставлення до предмету, пропагандист апелює до відповідної установки та спонукає відчувати її ще сильніше, демонструючи свої переконання за допомогою певних форм поведінки.

3. *Зміна реакції*. Цей спосіб вимагає від аудиторії перейти від однієї установки до іншої, перейти від нейтральної позиції до позитивної або негативної, змінити поведінку або прийняти нову.

Важливим терміном практики пропаганди є «якір» – відправна точка для змін. *Якір* – це та інформація, якою вже володіє людина та довіряє їй, і що буде використовуватися пропагандистом для змін і прищеплення нових поглядів і поведінки аудиторії. Пропагандисту для того, щоби переконати людей необхідно зв'язати зміни й те, у що вони вже вірять. В центрі уваги цього процесу переконання знаходяться *цінності* – певні властивості того чи іншого предмета, явища задовольняти потреби, бажання, інтереси як окремих людей, так й суспільства загалом. Цінності, упорядковані відповідно до їхньої відносної важливості, виступають нормативними орієнтирами для поведінки людини.

В процесі формування у людей переконань щодо ідей, людей або об'єктів одночасно формуються установки до них. Тому установки є важливим об'єктом пропаганди.

Особливо важливим для пропагандиста є вплив на *групові норми* – керівні принципи та правила, які встановлюються, схвалюються і приймаються членами групи як зразок поведінки в групі. Пропагандисти часто використовують їх як якорі. Різновидом групових норм є норми референтно групи. *Референтна група* – це реальна або ідеальна група на яку орієнтована людина, цінності, ідеали та норми поведінки якої вона поділяє.

Пропагандист, нагадуючи аудиторії про її поведінку в минулому, що відповідала його бажанням, може закликати дотримуватися такої або подібної моделі поведінки в майбутньому. І навпаки, якщо певна модель поведінки має для нього негативні наслідки, пропагандист може агітувати членів аудиторії уникати негативних наслідків, не дотримуючись такої поведінки. Він також може доносити бажані

моделі поведінки. Отже, пропаганда – це продумана спроба донести ідеологію, що має відповідну ціль, до аудиторії, змінити або зберегти баланс сил, які вигідні пропагандисту.

Розглядають її наступні види: біла, сіра, чорна пропаганда, і як окремий випадок субпропаганда або фасилітативна комунікація. Відповідно до [8] ці види пропаганди мають наведені нижче визначення.

Біла пропаганда – це форма комунікації, за умови, що переважно інформація в повідомленні є точною, а її джерело правильно ідентифіковане. Метою білої пропаганди є завоювання довіри аудиторії задля реалізації майбутніх планів пропагандиста. Її прикладами є: наративи, які ми чуємо від «Голосу Америки» в мирний час; національні свята з їх очевидним патріотизмом і регіональним шовінізмом; патріотичні висвітлення журналістами результатів міжнародних спортивних змагань.

Сіра пропаганда – це форма комунікації, що знаходиться десь між білою та чорною пропагандою. Правдивість інформації в повідомленні є сумнівною, а її джерело може бути або не бути вірно ідентифіковано. Сіра пропаганда широко поширена: держави та приватні компанії давно запровадили практику вихваляння різних своїх історій; під час бойових дій уряди використовують вітчизняні ЗМІ для пропаганди, надаючи журналістам тексти для їх висвітлення; сіра пропаганда також використовується для компрометації ворогів і конкурентів. До вище зазначеного можна додати: викривлену інформацію, рекламу, що рекламує якість продукту, яку не може досягти; відомі благодійні організації, що збирають гроші на дослідження, але використовують їх з іншою метою.

Чорна пропаганда – це форма комунікації, за умови, що інформація в повідомленні є хибною та неправдивою, а її джерело приховується або приписується фальшивому авторитету, що поширює брехню. Метою чорної пропаганди є здійснення активованої ідеології. В окремих випадках вона має агітаційний характер: пробудити аудиторію, спонукаючи її досягнути мету, що передбачена пропагандистом і призводить до значних змін. Інколи чорна пропаганда має інтегративний характер: вона ставить за мету зробити аудиторію пасивною, сприймаючою ідеологію пропагандиста як таку, що не кидає виклик його інтенціям.

Іншим виміром пропаганди є *субпропаганда* (фасилітативна комунікація). Тут основним завданням пропагандиста є поширення незнайомої доктрини, реалізація якої на практиці потребує значного періоду часу, в першу чергу для формування в аудиторії умонастроїв щодо її сприйняття. Сама по собі субпропаганда може й не бути

пропагандою, але це спілкування, покликане формувати позитивне ставлення до потенційного пропагандиста. Іншою формою фасилітативної пропаганди є надання допомоги країнам у відновленні своєї економіки й організацій після війни або конфлікту.

З погляду комунікації субпропаганда спрямована на підтримку відкритих ліній і контактів до моменту, коли вони будуть потрібні пропагандисту. Субпропаганда (фасилітативна комунікація) існує переважно у вигляді фінансової допомоги, випусків новин, теле- і інтернет-програм, прес-релізів, книг, брошур, періодичних видань, культурних програм, виставок, фільмів, семінарів, мовних курсів, довідкових служб тощо. Всі вони призначені для створення дружньої атмосфери для тих, хто може пригодитися згодом.

На саму пропаганду, як на форму комунікації, впливають технології для відправлення повідомлень, які наявні на даний момент часу. З їх розвитком пропагандисти отримують все більш сучасні і складні пристрої, саме тому кібертероризм можна віднести до ще однієї форми пропаганди [8-9].

Кібертероризм – це сукупність кіберзагроз та атак на комп'ютери, мережі й інформацію, що там зберігається, і використовуються для залякування людей або уряду з метою досягнення політичних і соціальних цілей. Кібертероризм є сучасною формою пропаганди. Це вид комунікації, що використовується в межах певної ідеологічної програми і підтримується централізованим фінансуванням. Хакери отримують гроші за здійснення кібернападів від імені терористів; хактивісти атакують сайти ворога із-за політичних мотивів тощо [10].

Дезінформація – це неправдива і неповна інформація, що вводить в оману та передається цільовій особі, групі або державі [8].

Основу дезінформації становлять новини, спеціально спрямовані на послаблення ворога, конкурента. Вона розміщується журналістами в ЗМІ. Ці новини подаються як правдиві та достовірні. Наразі існує все більше доказів того, що дезінформація практикується більшістю провідних країн світу.

Впродовж останніх років цифровізація змінює наш світ, наше життя та робочі процеси. Всі ми невідворотно стали частиною цієї цифрової трансформації у повсякденному житті – як професійному, так і особистому. Стало зрозуміло, що «цифрова революція» – або четверта промислова – не мала історичних прецедентів та розвивалася нелінійно. Вона створила специфічне середовище циркуляції інформаційних потоків, завдяки чому стали можливими нові, глобальні прояви маніпуляції масовою свідомістю.

У політичній сфері ця роль стала очевидною, коли Оксфордський словник обрав поняття «пост-правда» словом року. Словом року воно стало через те, що частота його використання у публічному дискурсі зросла на тисячі відсотків. Згідно з Оксфордським словником «пост-правда» означає, що інформація не обов'язково має бути правдивою, вона радше апелює до емоцій чи до суб'єктивних переконань. Політика «пост-правди» породила небезпечну медіатехнологію спрямованої дії «фейк-ньос». Необов'язковість фактичної основи стала характерною рисою даного явища.

Соціальні мережі і новітні засоби комунікації стали основним інструментом інформаційних війн і активним актором світової політики. Наразі вони претендують на звання «світової громадської думки». Хоча часто використовуються як зручний інструмент маніпуляції правдою в залежності від тієї чи іншої політичної цілі. Так працює механізм підміни правдивого факту псевдоправдою сумнівних, але у вищій мірі вирогідних емоційних доводів – гра на довірі аудиторії, емоційна апеляція до страху, фобії, інстинктів, масових стереотипів.

Таким чином, політика «пост-правди» як спосіб легалізації брехні затвердилась у сучасному медіа просторі, підмінив реальну дійсність віртуальною, ставши поживним середовищем для поширення сфальшованих, придуманих, брехливих новин – фейк-ньос.

Як зазначають сучасні дослідники [11; 12; 13], *фейк* – це подання фактів у спотвореному вигляді або подання свідомо неправдивої інформації. До того ж фейк – це спосіб маніпуляції свідомістю у спосіб надання неповної інформації, спотворення контексту, частини інформації з метою підготувати аудиторію до дій чи думок, які потрібні маніпулятору.

Фейкові новини – підробка чи імітація новин (маніпуляція, спотворення фактів, дезінформація), які створено з ігноруванням редакційних правок і процесів, прийнятих у ЗМІ для збереження відповідності та перевіреності, та які не витримують навіть поверхневих перевірок на відповідність та реальність.

Існують різні види класифікації фейків. Фейкові новини розрізняються по місцю, часу подій, персонажах, цілям і завданнях. До фейків відносять не тільки підроблені тексти, фальшиві фото-, відео- і аудіозаписи, але й віртуальні псевдоособистості (інтернет-боти) і фальшиві акаунти тощо.

Загалом, «фейк-ньос» можна характеризувати як сурогатне, паразитичне явище у світі медіа епохи постмодерна, пов'язане зі знижками статусу істинних цінностей (правдивої та точної інформації),

тріумф несистемного мислення, заміну духовних спрямувань особистості споживацькими інстинктами натовпу.

Фейки суттєво відрізняються від класичної дезінформації. Дезінформація формується на основі поєднання правди, напівправди та брехні. Для фейку це не потрібно. Чим більш малоймовірна фейкова інформація, тим більш дієво вона впливає на масову свідомість. В той час, коли дезінформація охоплює простір різних журналістських жанрів, фейк має цілком конкретний журналістський жанр – новину, особливого виду, фіктивний не існуючий в реальності інформаційний привід.

Дезінформація не обов'язково розрахована на миттєвий ефект, фейк – обов'язково є сенсаційним, шокуєчим. Він завжди є спотворенням реальної ситуації, свідомо сфабрикованою фальшивкою [10, 12]. Часто дезінформація існує приховано тривалий час, фейк навпаки спливає в результаті інформаційного вкиду у точно визначений момент часу.

Завдяки ефекту «емоційного зараження» у суспільстві фейк поширюється «вірусним» способом. Отже, резонансний характер фейкової інформації і вірусний механізм її поширення роблять фейк ефективним. Після того як Інтернет став глобальним явищем, різні дослідження його впливу на суспільство дали протилежні результати: Facebook, X (Twitter) тощо стали ефективними інструментами для пропаганди, поширення дезінформації й обміну підробленими новинами.

Останнім часом, особливо TikTok став важливим засобом комунікації соціальних груп, що поширюють підроблені відео і фото матеріали. Ставши важливою складовою технології «Глибоких фейків», ця платформа перетворилася на зброю масового ураження. Наразі під «Глибокими фейками» розуміють маніпулятивні відео або фото матеріали, створені із застосуванням штучного інтелекту. Сфабриковані цифрові матеріали створюють ілюзію справжніх зображень і звуків. Проблема в тому, що технологія «Глибоких фейків» стає дедалі досконалішою, викликаючи про цьому низку складних політичних, юридичних і етичних проблем.

Так TikTok – китайський соціомедійний застосунок для створення, перегляду та поширення відеофайлів, онлайн-трансляцій і коротких відео став однією з основних платформ, де фейкові відео збирають мільйони переглядів.

Це найпопулярніша платформа для коротких відео тривалістю до 3 хвилин. Новим користувачам можуть рекомендувати неправдивий контент одразу після приєднання до мережі. Тож сумнівний контент не потрібно навіть шукати. Згідно з дослідженнями вебсайту News Guard

щомісячна кількість активних користувачів TikTok сягнула мільярда, більша її частина – молодше 30 років. Платформа на цей час намагається зупинити потік фейків. Якщо Facebook, Instagram і X позначають фейкові відео про війну, то TikTok тільки починає впроваджувати цю функцію.

Еббі Річардс (Abby Richards), незалежна дослідниця, яка створила відео про загрозу дезінформації, зазначає, що багато вірусних жахливих стрімів зробили для того, щоб заробити грошей за допомогою системи подарунків TikTok (<https://www.bbc.com/ukrainian/features-61220423>).

Дослідження, здійсненні ВВС, розкрили деякі з найбільш поширених категорій фейкового контенту в TikTok: фейкові прямі трансляції, реальні відео, відео старих конфліктів, обмежена прозорість. Зокрема, фейкові прямі трансляції: користувач знаходить драматичне відео старого конфлікту чи якихось військових навчань, додає фейковий звук потужного вибуху чи інтенсивної перестрілки, запускає пряму трансляцію і, щойно збереться значна аудиторія, просить пожертвування на свій канал. В TikTok замість реальних відео користувачі регулярно використовують кадри військових відеоігор або комп'ютерні зображення (CGI): багато сцен з таких відеоігор, як Arma 3 та Call of Duty. Особливо популярними виявились фальшиві повітряні бої, зокрема кадри збиття винищувачів. Деякі акаунти задля більшої реалістичності використовують уривки з фільмів про війну, серіалів або реаліті-квестів, а також епізодів з відео матчу зі страйкболу.

ВВС аналізувала відео старих конфліктів і бойових дій у Лівії, Сирії та Чечні, які видавали за конфлікт в Україні. Щоб приховати справжнє джерело, на оригінальні кадри наклали фейкові позначки екстрених новин та субтитри нібито прямого ефіру. Деякі користувачі викладають відео бойових дій 2014 року, які легко сплутати з нинішніми, а також кадри військових навчань чи парадів у Росії чи в Україні.

Е. Річардс зазначає, що TikTok міг би вирішити проблему вірусних старих відео, якби вжив простих заходів, наприклад, зробив більш чіткою дату публікації відео.

Обмежена прозорість, пов'язана з тим, що попри те, що TikTok співпрацює з незалежними фахівцями з перевірки фактів, але якщо Facebook і Instagram позначають фейковий контент про Україну, то в TikTok таке маркування зустрічається рідко.

На відміну від деяких своїх конкурентів, TikTok не надає інструменти для аналітики й прозорості вченим, дослідникам і журналістам, тому незалежним фактчекерам доводиться постійно вручну аналізувати величезні обсяги контенту.

Ще одна проблема, за словами Е. Річардс, полягає у відсутності спільноти користувачів TikTok, які готові викривати фейки на платформі.

Війна росії проти України призвела до появи в інтернеті вірусних фейкових відео і фото, які активно розповсюджують користувачі соціальних мереж. Фахівець з дезінформації Маріанна Спрінг (Marianna Spring) (<https://www.bbc.com/ukrainian/features-60734162>) надала такі поради з того, як відрізнити фейки про війну в Україні від правди:

1. Перевірте, коли саме зняли те чи інше відео. Автори багатьох фейкових роликів намагаються видати за нові ті відео, що були зроблені під час інших конфліктів. Завжди перевіряйте, чи є ознаки того, що відео старе. Google Maps допоможе з'ясувати, чи дійсно відео зняли там, де каже автор. А за допомогою зворотного пошуку зображень, які є на низці вебсайтів, можна побачити, чи публікували це відео/фото раніше в інтернеті. Якщо так, це точна ознака того, що його переробили з попередньої події.

2. Хто публікує контент і чому. Завжди звертайте увагу на те, хто публікує той чи інший контент – і чи можна вважати це джерело авторитетним й таким, якому можна довіряти. Деякі фейкові відео розповсюджують ті, хто прагне отримати якомога більше «шерів» і «лайків». Інші намагаються у такий спосіб просунути свої власні політичні погляди або сіяти сумніви і тривогу. У багатьох з них мало або зовсім немає підписників, на сторінці стоїть не особисте фото, а взяте з інтернету, а ім'я – надто загальне, щоб точно встановити особу. Автори таких акаунтів почали викладати дописи лише після початку вторгнення. Важко сказати, хто взагалі веде ці облікові записи.

3. Думайте перед тим, як натиснуть на «шер». Відео та фото, що розповідають про війну, часто дуже страшні й викликають у людей емоційну реакцію, яка примушує їх одразу ділитися ними, навіть без додаткових перевірок. Дезінформація розповсюджується дуже швидко саме тому, що її автори намагаються скористатися емоціями та упередженнями. І йдеться не лише про негативні емоції – інколи обнадійливі історії також можуть стати вірусними. Вони підіймають бойовий дух, але ускладнюють розуміння того, що відбувається насправді, й сприяють хаосу. Тому треба розуміти – чому те чи інше джерело щось говорить та чи підтверджується це доказами.

Очевидно, що нині процес поширення фейків в інформаційному полі став одним з різновидів інформаційних операцій (ІО) зі всіма їх властивостями:

- 1) високою ступінню подібності текстів, не будучи їх точними копіями;

2) поширення текстів відбувається за короткий інтервал часу, що дозволяє розгляти їх як нову значну подію;

3) джерела, що публікують ці тексти, є маловідомими, регіональними або спеціалізованими ЗМІ. Їх поява у великих інформаційних агентствах можлива тільки під час пасивної фази ІО.

Наразі однією з популярних методик ідентифікації фейкових новин є графова модель поширення фейків [14], у відповідності з якою новинне повідомлення включає в себе повний текст, час опублікування, джерело ЗМІ і значення тематичного класифікатора. Оцінюється число повідомлень визначеного типу та змісту в рівні проміжки часу. Якщо їх кількість більше порогового значення, робиться висновок про проведення інформаційної операції.

Динаміка зміни кількості повідомлень під час проведення інформаційної операції відповідає закону зміни характеру інформаційних потоків за обраний період (рис. 6.1).

Під час активної фази публікуються схожі повідомлення в різних, як правило, маловідомих ЗМІ. Оскільки повідомлення з'являються впродовж короткого проміжку часу, системи моніторингу засобів масової інформації сприймають їх як повідомлення про нову, популярну подію/новину. Під час пасивної фази частина інших ЗМІ реагує на повідомлення активної фази інформаційної операції. Тим самим аудиторія збільшується, а ефективність ІО підвищується.

Для дослідження вищенаведеної моделі у роботі [15] було обрано новину-фейк про причетність українського заводу «Южмаш» до виробництва та поставки ядерних ракет для КНДР. Для чого використовувалися інструменти веб-аналітики, парсингу новин та RSS-даних новинної стрічки за період (з 01.08.2017 по 31.08.17) Основна ідея була наступною: будується графік зміни інтенсивності появи новини (за кількістю в одиницю часу) і порівнюється з заданим шаблоном. У разі їх збігу робиться висновок про факт проведення інформаційної операції.

В результаті дослідження були зібрані дані про кількість результатів щоденної пошукової видачі за запитами, що містили комбінації ключових слів: «КНДР», «Північна Корея», «Україна», «ракета(и)», «виготовлення» тощо. Був здійснений попередній синтаксичний аналіз та аналіз релевантності отриманих новин для виключення незадовільних або невідповідних за змістом результатів. Після опрацювання отриманих даних був побудований графік зміни кількості пошукових видач. Результат наведено на рис. 6.2.

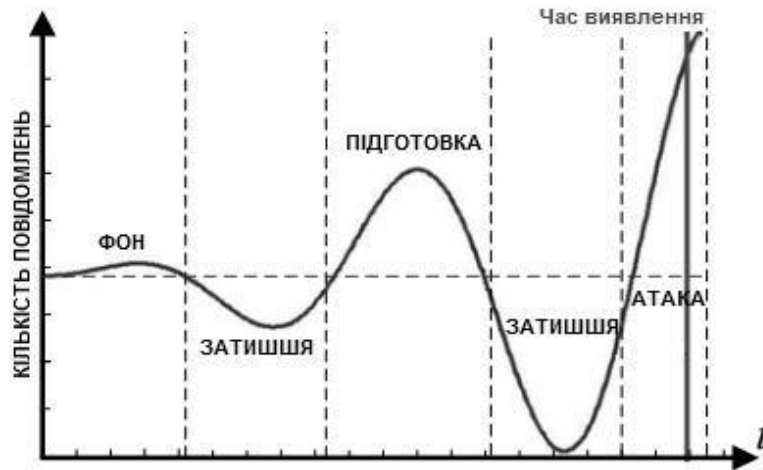


Рисунок 6.1: Динаміка кількості тематичних повідомлень

За допомогою отриманого графічного представлення вдалося виділити усі типові етапи, що властиві здійсненню інформаційних операцій: фон, етапи затишшя та підготовки, а також етап атаки, який передбачає форсування появи новин, що містять визначене фейкове повідомлення.

Отриманий результат підтверджує коректність графової моделі поширення фейків, а саме: фон (низький пошуковий попит за ключовими словами), попереднє затишшя, підготовка (сплеск популярності запиту), затишшя, атака (значний сплеск запитової популярності новини), релаксація.

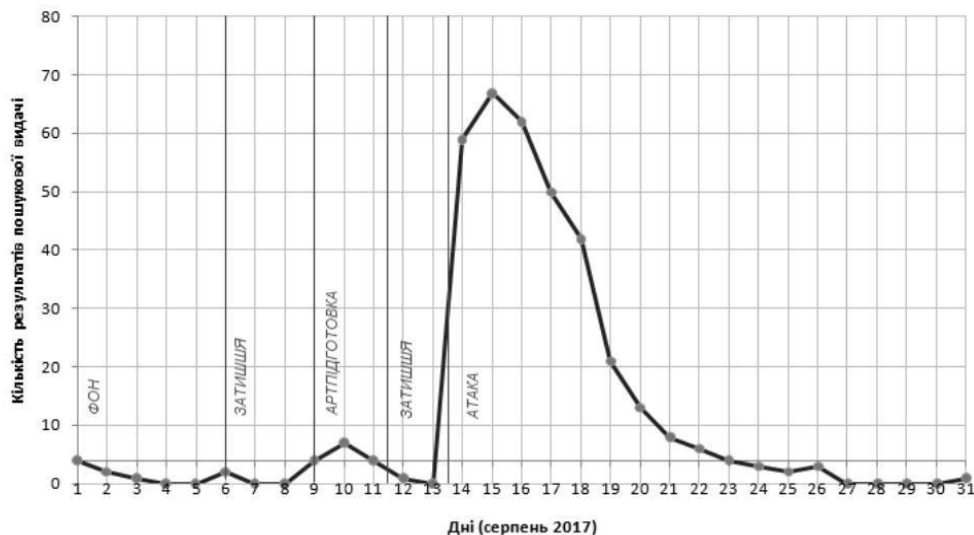


Рисунок 6.2: Динаміка кількості тематичних повідомлень під час інформаційної операції

Графічне представлення залежності кількості новин від часу проведення інформаційної операції, в якому чітко розпізнаються

відповідні етапи, підтверджує той факт, що наведений метод ідентифікації фейків дає гарні результати.

Основними суб'єктами проведення інформаційно-психологічної війни є [15]:

1. Керівництво іноземних держав;
2. Спецслужби іноземних держав та її агентура;
3. Засоби масової інформації (іноземні та підконтрольні вітчизняні);
4. Неурядові організації (іноземні й підконтрольні вітчизняні);
5. Інтернет-ресурси.

Таким чином, наразі тема інформаційно-психологічних війн є актуальною, адже інформація стала одним з найдорожчих та найважливіших інструментів, за допомогою яких можна впливати на безпеку особи, суспільства та держави. В умовах формування інформаційного суспільства інформаційно-психологічні операції є найбільш ефективним і універсальним засобом реалізації зовнішньої та внутрішньої політики. Нині інформаційно-психологічне протиборство пронизує всі форми боротьби: дипломатичної, економічної, воєнної тощо. Отже, інформаційно-психологічна війна сьогодні виступає всеосяжною, цілісною стратегією, яка показує всю значимість володіння інформацією в державному управлінні і є становим хребтом політики національної безпеки.

6.2. Практична реалізація результатів моделювання

Ментальна війна є найбільш масованим і небезпечним типом інформаційно-психологічних війн, що, зокрема, здійснюються тривалий час росією проти України. Мета інформаційно-психологічної війни – промивка мізків та реформування громадської думки у потрібному руслі. На відміну від традиційної війни, коли армія воює з військом ворога, інформаційно-психологічну війну [16] ведуть проти мільйонів цивільного населення. Вона здійснюється як у мирний, так і у воєнний час. Зокрема, метою російської агресії проти України є не просто загарбання територій, а й знищення нашої національної ідентичності [17].

Ми розглядаємо категорію «ментальна війна» як складну систему S , що породжує низку важливих проблем національної безпеки:

$$S = \langle A, R, Z \rangle,$$

де A – множина її підсистем і елементів, R – множина відношень між ними і Z – множина цілей системи «ментальна війна».

Розв'язання таких проблем є ключовим завданням системного аналізу. Методологія системного аналізу політики інформаційної безпеки, пов'язаної зі викликами та загрозами «ментальної війни», запозичує й інтегрує досвід багатьох дисциплін: політології, психології, соціології, культурології, журналістики тощо.

У роботі методологія системного аналізу політики інформаційної безпеки, пов'язаної з «ментальною війною», розглядається п'ять процедур, що утворюють інформаційно-аналітичний цикл [18]: *структурування проблем, прогнозування, рекомендація, моніторинг, оцінювання*. У своїй сукупності ці п'ять інформаційно-аналітичних процедур базуються на низці специфічних методів і методик системного аналізу.

Політико-аналітична процедура «*структурування проблеми*» є першим кроком системного дослідження складної системи «ментальна війна», що передбачає надання інформації про умови, що породжують проблеми і потребує визначення поняття «національний менталітет».

Під поняттям «*національний менталітет*», враховуючи результати попередніх досліджень, ми розглядаємо інтегральну характеристику народу, якому властива певна сукупність соціально-психологічних настанов та навичок свідомості, що є характерними для тієї або іншої культурної спільноти й чітко проявляються у типовій поведінці людей через стереотипи сприйняття та реакцій [19-21].

Структурний аналіз функціонування великої системи передбачає розуміння її складу, а саме підсистем й елементів. Для складної системи «ментальна війна» такими підсистемами є: *мова, культура, віра й історична пам'ять*. Елементами даної системи, що є найменшими й неподільними складовими зі своїми власними атрибутами: *вірування, цінності, установки, стереотипи, норми та традиції*.

За допомогою політико-аналітичної процедури структурування проблеми ми отримали ієрархічну графову модель складної системи «ментальна війна». Взаємодія і взаємозв'язок підсистем й елементів даної системи породжують її основну емерджентну властивість – *національний менталітет* (рис. 6.3).

Політико-аналітична процедура «*прогнозування*» надає інформацію про майбутні наслідки дій згідно зі запропонованими альтернативами (сценаріями) політики безпеки. Для даного кроку аналітичного циклу важливу роль відіграють ієрархії, які є основним засобом поділу людиною реальності на групи та класи [22].



Рисунок 6.3: Ієрархічна модель складної системи «Національний менталітет»

Дана ієрархічна модель – це певна абстракція системи, математичне дослідження якої, кількісна оцінка пріоритетів альтернатив політики безпеки здійснюється за допомогою методу аналізу ієрархій (МАІ).

Виклики та загрози, що виникають у сучасному інформаційному просторі та суспільстві, зумовлюють необхідність розробки на об'єктивних засадах хоча би приблизних сценаріїв того, що відбувається. Для цього можна скористатися процедурою МАІ. Її можна використати для дослідження та кількісної оцінки взаємодії елементів і підсистем, а також їх впливу на систему загалом.

У МАІ реалізовано рух із вершини до підсистем і елементів, коли елементи кожного нижчого рівня перебувають під впливом елементів іншої групи – вищого рівня. Здійснений у такий спосіб процес планування сценаріїв дає змогу підвищити якість розумних припущень, визначити якими можуть бути наслідки, якщо ви ризикнете діяти згідно з цими припущеннями, в яких випадках цей ризик може бути він виправданим [23]. МАІ може стати систематичною процедурою для ухвалення оптимальних рішень при сценарному плануванні механізмів протидії ментальним війнам, а також науково-методичним забезпеченням для управління інформаційними ризиками.

За допомогою структурного аналізу системи «ментальна війна» ми отримали знання про її склад і структуру. Однак, такий підхід не дає змоги розглянути наслідки реалізації різних сценаріїв державної політики інформаційної безпеки якщо складові системи розглядати окремо. Системи є цілісними сутностями, поведінку яких важко

передбачити. Їх функціонування породжує нові емерджентні властивості, що виходять за межі їхніх окремих складових [24]. Основним механізмом функціонування таких систем є загальносистемний механізм зворотного зв'язку. Згідно даного закону існує вплив результатів функціонування суспільства, включаючи елементи, зв'язки, властивості, відношення та зміни, на всіх його членів [25]. Тому, з погляду інформаційної безпеки, запропонована вище ієрархічна модель породжує таку важливу, для інформаційної безпеки, категорію суспільства як *«ідентичність»*.

Політико-аналітична процедура *«рекомендація»* надає інформацію про відносну цінність або ціну майбутніх наслідків у процесі розв'язання або послаблення проблем інформаційної безпеки, пов'язаної з викликами й загрозами *«ментальної війни»*.

Ментальність та ідентичність за значенням є близькими і взаємозалежними поняттями. Крос-факторний аналіз складної системи *«ментальна війна»* важливий тому, що робота експертної групи стимулює плідні дискусії щодо рекомендацій, які стосуються негативних наслідків ворожої пропаганди на ментальність народу. Ментальність інтегрує націю, робить її неповторною самобутньою, відрізняє від інших, складає основу національної ідентичності.

Національна ідентичність означає суб'єктивне відчуття людини належності до певної держави чи нації. Ідентичність виступає як об'єктивація ментальності [19]. Основні механізми протидії інформаційних загроз, пов'язані зі викликами й загрозами *«ментальної війни»*, мають бути спрямованими на запобігання кризи національної ідентичності.

За рахунок масованого використання росією маніпулятивних технологій, на разі існує реальна загроза формування альтернативної національної *«ідентичності»*. Рекомендації щодо основних механізмів відвертання кризи національної ідентичності, пов'язаної зі викликами й загрозами *«ментальної війни»*, мають бути заснованими на результатах крос-факторного аналізу, в першу чергу, взаємодії різних змінних *«чотирикутника»* (мова, культура, віра й історична пам'ять), а також виявлених трендів та тенденцій інформаційної й інформаційно-психологічної безпеки.

Політико-аналітична процедура рекомендації щодо кризи національної ідентичності вкрай важлива, оскільки криза викликає серйозні деструктивні процеси в сфері ціннісних настанов, культурних традицій світоглядних засад існування населення [20]. Слабка національна ідентичність призводить до серйозних проблем з національної безпеки.

Крайній випадок того, до чого може призвести відсутність національної ідентичності – це розпад держави та громадянська війна.

Основним завданням політико-аналітичної процедури «моніторинг», як системи постійного спостереження за явищами і процесами, що пов'язані з «ментальними війнами», є специфікація діяльності системи забезпечення інформаційної безпеки, спрямованої на обґрунтування управлінських рішень щодо безпеки особи, суспільства та держави. Для чого використовується крос-факторний аналіз для дослідження складних соціо-культурних систем.

Крос-факторний аналіз дозволяє аналізувати взаємозв'язки між змінними у «чотирикутнику» (мова, культура, віра й історична пам'ять), зрозуміти як різні тенденції або події впливають одна на одну, а також виявити нові тренди. Інформація, отримана за допомогою крос-факторного аналізу, надає додаткову картину взаємозв'язків трендів і змінних, що важливо, картину залежних і незалежних чинників ворожої пропаганди, її рухомих і рушійних сил. Крос-факторний аналіз – вкрай важливий метод визначення ключових змінних і тенденцій, які негативно впливають на ідентичність народу у результаті пропагандистських кампаній.

Для практичного застосування крос-факторний аналіз є доволі простим методом, що здійснюється невеликою групою експертів і включає наступні базові кроки [26]:

1. Скласти список трендів або базових змінних, що аналізуються. При чому часто корисно визначати змінні як тренди, що діють у певному напрямку.
2. Скласти схему (рис. 6.4) і занести у список тренди та змінні.
3. Заповнити кожну гратку оцінкою ступеня впливу тренду *A* на тренд *B* (позитивне, негативне, відсутній вплив) за шкалою від -2 до $+2$.
4. Порахувати суми абсолютних значень оцінок в рядах, які показують, наскільки значним чинником є дана змінна, тобто у якій мірі вона впливає на інші змінні. Суми абсолютних значень оцінок у стовпцях показують міру залежності кожної змінної.
5. Кількісні оцінки трендів наносяться на діаграму (рис. 6.5), однією з осей якої є ступінь залежності, іншою – сила чинників.

		ЗМІННІ						
	ТРЕНДИ	Вірування	Цінності	Настанови	Стереотипи	Норми	Традиції	СУМА
		1	2	3	4	5	6	7
1	Вірування		-1	2	1	-2	0	6
2	Цінності	1		-2	-1	2	1	7
3	Настанови	2	0		-2	1	-1	6
4	Стереотипи	0	-2	1		1	1	5
5	Норми	2	2	-1	1		0	6
6	Традиції	0	1	-2	2	0		5
7	СУМА	5	2	2	7	6	3	

↑
Сама залежна змінна

← Самий сильний чинник

Рисунок 6.4: Приклад матриці крос-факторного аналізу основних трендів впливу пропаганди на національну ідентичність

Якщо процедуру крос-факторного аналізу здійснювати періодично, то можна отримати чітку динамічну картину основних тенденцій і трендів впливу ворожої пропаганди на ідентичність народу. При цьому моніторинг може виконувати наступні функції:

- *узгодженості*: чи узгоджуються заходи та дії системи забезпечення інформаційної безпеки зі процедурами й стандартами, встановленими державними органами;
- *аудиту*: визначити, чи дійсно ресурси та механізми відвертання інформаційних загроз, орієнтовані на певні цільові групи і доходять до них;
- *пояснення*: надає інформацію, що допомагає пояснити, чому наслідки альтернатив інформаційної політики та програм безпеки відрізняються;
- *звітування*: надає інформацію, що допомагає скласти звіт про зміни у суспільстві, які настають після практичної реалізації широкого комплексу альтернатив інформаційної політики та програм безпеки після певного періоду.

Політико-аналітична процедура «оцінювання» надає інформацію про цінність або ціну наслідків діяльності системи забезпечення протидії інформаційним загрозам, засновану на аналізі контрольних індикаторів, у процесі розв’язання або послаблення гостроти проблеми інформаційної безпеки.

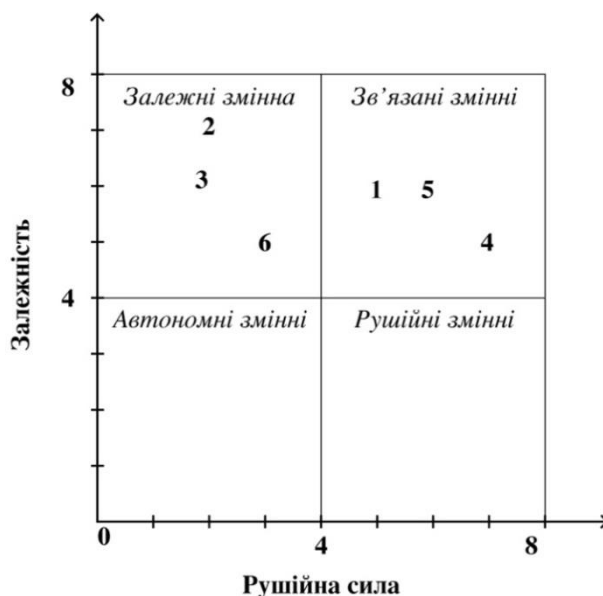


Рисунок. 6.5: Приклад крос-факторної діаграми

Щоби з'ясувати, чи насправді системі забезпечення вдалося досягти цілей, визначених політикою інформаційної безпеки, успішне її впровадження має містити підсумкове оцінювання. Метод згортання критеріїв має важливе значення для неструктурованих проблем, якими є інформаційні війни.

Коефіцієнти важливості у методі згортання критеріїв визначаються за допомогою принципу парного порівняння МАІ. Він реалізується за допомогою наступного алгоритму:

Крок 1. Критерії $F_1, F_2, \dots, F_m$ упорядковуються за ступенем їхньої важливості.

Крок 2. Порядкова шкала з k вербальними градаціями розглядається як єдина для кількісних і якісних критеріїв. Кожна градація повинна мати зрозумілий і докладний опис властивостей і характеристик критеріїв або їх оцінку (наприклад, відмінно, добре, погано тощо) і відповідно для кожної градації порядковий номер переваги або бал. Розташування градацій відповідає порядку, коли найкраща вербальна оцінка (наприклад, «чудово») отримує найбільший бал, а найгірша (наприклад, «жахливо») – найменший.

Крок 3. Кількісні і якісні критерії приводяться до однорідного вигляду і єдиної порядкової шкали. Для чого кожний критерій (як кількісний, так і якісний) розбивається на k градацій (рівнів) таким чином, щоби значення кількісних критеріїв і властивостей якісних критеріїв для кожної градації відповідали прийнятій вербальній шкалі і кількісним характеристикам градацій.

Крок 4. Здійснити попарне порівняння критеріїв між собою і побудувати матрицю попарних порівнянь критеріїв. Для матриці попарних порівнянь визначити максимальне власне значення та відповідний йому власний вектор і проаналізувати узгодженість матриці парних порівнянь для розрахованого відношення узгодженості (ВУ). Якщо $ВУ \leq 0,1$, то матриця попарних порівнянь може вважатися узгодженою, інакше, необхідно буде переглянути оцінки попарних порівнянь й уточнити їх таким чином, щоб матриця попарних порівнянь стала узгодженою.

Крок 5. Визначити коефіцієнти важливості критеріїв. Вони дорівнюють відповідним елементам власного вектору, розрахованого на попередньому кроці.

Крок 6. Слабоструктуровані проблеми ухвалення рішень характеризуються багатокритеріальністю, кількісними й якісними характеристиками, що суттєво ускладнює розробку адекватної математичної моделі. Ідея методу згортання кількох критеріїв $F_1, F_2, \dots, F_m$ в один глобальний критерій F передбачає використання формули

$$F = F_1 + F_2 + \dots + F_m,$$

за умови нормування ($\alpha_i \in [0, 1]$):

$$\alpha_1 + \alpha_2 + \dots + \alpha_m = 1.$$

Глобальний критерій F розраховується для кожного можливого варіанту рішень.

Крок 7. Визначити найліпше рішення. Воно є найбільшим значенням глобального критерія серед усіх можливих значень варіантів рішень.

Метод згортання критеріїв розглядає системи як єдине ціле, досліджує функціонування систем й їх емерджентні властивості. Він сприяє поліпшенню й критичному осмисленню цінностей, що лежать в основі вибору цілей і завдань, спрямованих проти зміни ідентичності народу. Результати, отримані за допомогою методу згортання критеріїв, мають ключове значення для вибору адекватних механізмів протидії ворожій пропаганді й дезінформації. Може сприяти визначенню нових або переглянутих альтернатив політики інформаційної безпеки, пов'язаної з «ментальними війнами».

Пропонована методологія синтезу системи «ментальна війна» надає можливість розв'язання фундаментального питання існування українського народу: «Як у «чотирикутнику» (мова, культура, віра й історична пам'ять), незалежно від рівня динамічної складності

інформаційних загроз, зберегти національну ідентичність, яка об'єднує суспільство в єдине ціле»? Наразі – це ключове питання інформаційної й інформаційно-психологічної безпеки.

6.3. Аналіз результатів

Мережа «ментальна війна» (далі МВ, див. П.3.5) – це загальний, але потужний спосіб представлення та вивчення простих і складних взаємовідносин. Її моделлю, розробленою за допомогою генеративного штучного інтелекту, є ненаправлений граф [26], першим кроком аналізу якого є оцінка *розміру*. Розмір мережі МВ є самою головною характеристикою моделі, що дорівнює кількості вузлів ($N = 378$) і ребер мережі ($L = 712$).

Ще однією базовою та важливою характеристикою соціальної мережі є *щільність* – доля наявних зв'язків щодо їх максимально можливого числа в мережі. Для ненаправленої мережі щільність ρ дорівнює [27]:

$$\rho = \frac{L}{k(k-1)/2},$$

де L – кількість наявних зв'язків (ребер) в мережі, k – кількість вузлів. Таким чином, щільність може змінюватися від одиниці до нуля: чим більша взаємозв'язаність мережі, тим ближче вона до 1 (повна мережа). Фактичне число зв'язків набагато менше одиниці – у нашому випадку $\rho = 0,01$ – на два порядки менше.

Хоча загальний розмір мережі, її щільність є цікавими характеристиками, але більш корисною може бути її компактність з урахуванням розміру та ступенем взаємозв'язаності. Важливим показником компактності є *діаметр мережі* – найбільша довжина будь якого самого короткого шляху між всіма парами вузлів. При цьому діаметр відповідає найбільшому маршруту або ланцюжку передачі інформації (або інших ресурсів) по мережі. Попри те, що загалом соціальні мережі можуть мати великий розмір, наша графова модель «ментальної війни» має малий діаметр – 6. Причиною чого є її щільність та кластеризація. *Кластеризація* є ще однією з фундаментальних характеристик мереж. Значення кластеризації C , або тенденція до створення кластерів – сильно зв'язаних груп вузлів також може змінюватися від 0 до 1. У побудованій у попередніх розділах мережі МВ, як уже зазначалося, $C = 0,447$, що є досить великим значенням, тобто мережа схильна до утворення кластерів.

Наведемо базові характеристики важливості вузлів мережі МВ. Вузол є важливим, якщо зв'язки роблять його прозорим для інших. Для вузлів ненаправлених мереж ми розглядаємо таку характеристику, як

центральність. Існують декілька способів оцінки цієї характеристики вузлів мережі.

Самою простою мірою центральності є уявлення про те, що вузол, з найбільшою кількістю прямих зв'язків, є більш важливим ніж вузли з меншою кількістю зв'язків. Таким чином, *центральність за ступенем* $C_D(v)$ – це ступінь вузла v (див. Розділ 5).

Для вузлів графової моделі ментальної війни міру центральності за ступенем можна умовно розділити на три рівні: 94 – 56 зв'язків, 20 – 11 зв'язків та 9 – 5 зв'язків. Вузли та ребра першого рівня разом відображають сили та засоби ментальної війни, сценарії її реалізації, політики акторів та їх цілі тощо. Другий рівень представлений механізмами реалізації МВ (медіа, культура, пропаганда тощо). Виявляється, що на третьому рівні знаходяться механізми протидії ментальної війни (освіта, мораль, релігія, національна єдність тощо).

Досліджуючи центральність окремого вузла, ми можемо розширити уявлення про мережу, визначаючи наскільки близько цей вузол знаходиться до інших вузлів. Це приводить нас до поняття *центральність за близькістю* $C_C(v)$ (див. Розділ 5). Центральність за близькістю розраховується як обернена величина до суми найкоротших відстаней від даного до всіх інших вузлів.

Наприклад, якщо значення центральності за близькістю вузла «*сили та засоби ментальної війни*» у побудованій вище мережі МВ дорівнює $C_C(v) = 0,5$, то це означає, що даний вузол знаходиться в середньому на відстані, яка дозволяє йому бути досить «близьким» до інших вузлів графу. Цей факт вказує на те, що вузол має хорошу доступність до інших частин графа. Значення центральності за близькістю для вузлів, які мншші, але близькі до 0,5 може вказують на те, що ці вузли мають гарну доступність до інших частин графу, але не є самими центральними.

Центральність за проміжністю $C_B(v)$ (див. Розділ 5) є ще однією характеристикою центральності вузлів – це міра центральності вузла в графі, яка показує, наскільки часто цей вузол є проміжним на найкоротших шляхах між іншими вузлами. Вузол буде мати високе значення центральності за проміжністю, якщо він займає особливе положення в мережі. Для інших вузлів він є важливою «станцією» для комунікаційних шаблонів, що проходять через мережу. Важливості вузлів за проміжністю знаходяться у побудованій мережі МВ в межах від 16245 до 23000.

Ще однією характеристикою центральності вузла мережі є *центральність власного вектора* $C_E(v)$ (див. Розділ 5) – це міра важливості вузла в графі, яка враховує не лише кількість зв'язків (ступінь вузла), але й важливість цих зв'язків – вузол є важливим, якщо він з'єднується з іншими важливими вузлами. Даний показник допомагає визначити, які вузли є ключовими в графі, враховуючи не лише їх безпосередні зв'язки, але й зв'язки їхніх сусідів. Для самих ключових вузлів подубованої графової моделі МВ, розраховані значення власного вектору центральності знаходяться в межах від 1,0 до 0,478.

Засновники компанії Google запропонували ще один показник центральності – *PageRank Centrality* – $PR(v)$ міру важливості вузла v в графі, яка використовується для оцінки впливу або значущості вузла на основі його зв'язків з іншими вузлами. Для важливих вузлів графу, розраховані значення даного показника центральності знаходяться в межах від 0,053 до 0,033.

Підсумовуючи результати аналізу різних показників центральності вузлів графової моделі МВ необхідно зазначити, що, в першу чергу, соціальні мережі корисні своїми патернами взаємодії і тим, як вони впливають на її вузли. Можна виокремити сукупність ключових вузлів графу, а саме: сили та засоби ментальної війни, результати МВ, актори та цілі МВ, політики акторів, загальні цілі МВ. Всі вони близькі між собою: займають особливе положення в мережі для зв'язку між її різними частинами; відіграють ключову роль у передачі інформації або ресурсів між ними.

Соціальні системи, що представлені мережами, як правило, характеризуються складною структурою. Вони складаються із низки відносно щільних підгруп – сукупності вузлів, які мають відносно більшу кількість внутрішніх зв'язків і відносно невелику кількість зв'язків між даною групою та рештою мережі. Із чого випливає, що для нашої моделі МВ важливо знайти та визначити ці підгрупи. Для чого використовується поняття модулярності.

Модулярність – це різниця між числом зв'язків, внутрішніх для всіх кластерів і значенням цього числа для випадкової мережі. Чим вищий показник модулярності вузла $Q(v)$, який може приймати значення від -1 до $+1$, тим більший рівень кластеризації мережі щодо даного групування мережі. Отримане середнє по всіх вузлах мережі МВ значення показника модулярності $Q = 0,599$ означає, що графова модель МВ має досить чітко виражені кластери. Це вказує на те, що вузли всередині цих кластерів пов'язані між собою значно тісніше, ніж із вузлами з

інших кластерів. Таке значення вважається досить високим і свідчить про якісне розбиття графа на групи.

З одного боку, як показники центральності, так і модулярності роблять їх зручними для інтерпретації графової моделі, а з іншого боку, мережа МВ належить до класу штучних мереж, яка в результаті свого розвитку перетворюється у самоорганізовану складну систему. Її еволюція відбувається під впливом численних факторів: політичних, воєнних, культурних тощо. Є достатньо доказів, що мережа МВ має загальні статистичні властивості, які впливають на її функціонування (рис. 6.6). Вона є прикладом, для яких характерні ознаки цих систем: динаміка росту, оптимізація, вкладення у двовимірний простір.

Для степеневого розподілу функція щільності ймовірності має наступний вигляд:

$$p(x) = Cx^{-\alpha}, \quad x \geq x_{\min},$$

де $\alpha \approx 2,24$ – оцінка показника степеневого розподілу, $x_{\min} = 1$ – мінімальне значення, з якого починається степеневий розподіл, C – нормувальна константа.

Для перевірки точності степеневої моделі використано статистичний показник $K-S$ показує високу відповідність між емпіричним і теоретичним розподілом ($K-S = 0,04$), що підтверджує коректність апроксимації.

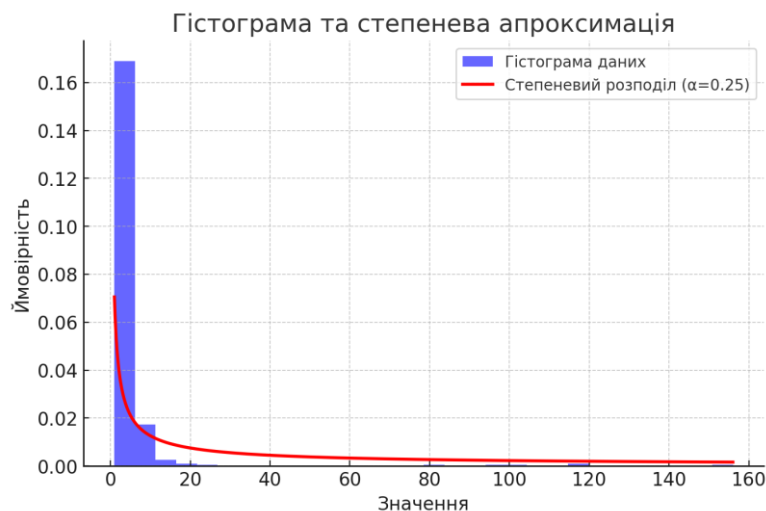


Рисунок 6.6: Гістограма мережі МВ й її степенева апроксимація

Аналіз статистичних властивостей мережі МВ підтвердив існування в ній вузлів з високими ступенями, які практично відсутні в мережах з пуасоновим чи експоненційним законами розподілу. Саме їх наявність пояснює поведінку моментів і спричиняє багато різних специфічних властивостей безмасштабних мереж.

Розмір околу вузла, який визначається його ступенем, вказує на його важливість, наприклад, в якості хаба в мережі. У великих мережах, створених випадковим приєднанням вузлів, вузли з високими ступенями є рідкісними, у той час як в реальних мережах вони часто зустрічаються з набагато вищою ймовірністю. Формально це визначається поведінкою хвоста розподілу ступенів вузлів [29]. Показник $\alpha = 2,24$ у подальшому можна використати для класифікації мережі MB: якщо $\alpha < 3$, тоді не існує кінцевого другого моменту, і мережа не має перколяційного порогу по відношенню до розростання її вузлів [30]. Її компонент зв'язаності залишається стійким до випадкових роз'єднань будь-якої кількості вузлів.

6.4. Практична частина. Аналіз достовірності повідомлень в інтернет через мережевий нетворкінг

У цьому розділі розглядається як аналізувати інформаційні повідомлення, визначати рівень їхньої достовірності шляхом побудови семантичної мережі та оцінки достовірності значимих зв'язків між окремими поняттями. Демонструється як використовувати мережевий нетворкінг (побудову і аналіз мереж на основі текстів) для оцінки достовірності повідомлень.

Етап 1. Побудова найпростішої семантичної мережі з повідомлення

На першому етапі формується найпростіша інформаційна мережа на основі змісту повідомлення, щоб надалі дослідити достовірність зв'язків в цій мережі. Цей етап включає процедури безпосереднього створення семантичної мережі на основі тексту за допомогою великої мовної моделі (LLM) і візуалізації результатів.

Такий підхід дозволяє автоматизувати дослідження інформаційних повідомлень, щоб надалі визначати достовірність зв'язків між ключовими вузлами мережі, що відповідають поняттям із повідомлення.

Інформаційну мережу на основі повідомлення, що досліджується, будемо формувати шляхом знаходження понять предметної області та виявлення зв'язків між ними, для чого будемо послідовно знаходити пари зв'язаних понять із повідомлення за допомогою LLM і виводити їх у вигляді нумерованого списку у форматі «Поняття 1;Поняття 2» (таке представлення відповідає формату CSV).

Для отримання даних для мережі створимо первинний промпт і модифікуємо його у відповідності із вже розглянутим у попередніх розділах фреймворком безкодового програмування. Для цього завантажимо у LLM фреймворк безкодового програмування як

зовнішній файл, після чого запишемо первинний промпт і запит на його перетворення у вигляді:

Перетвори відповідно до наведеної методики безкодового програмування промпт: "Обери пари ключових слів, пов'язаних причинно-наслідковими зв'язками з тексту, і представ їх у вигляді нумерованого списку рядків виду "Причина;Наслідок", наприклад, "дощ;калюжа"."

Після застосування фреймворку безкодового програмування отримуємо новий, структурований промпт, який і будемо безпосередньо застосовувати:

1. Аналіз завдання

Промпт складається з таких елементів:

Потрібно обирати пари ключових слів, пов'язаних причинно-наслідковими зв'язками, із тексту.

Результат має бути представлено у вигляді нумерованого списку.

Формат кожної пари: "Причина;Наслідок".

Це завдання можна реалізувати за допомогою композиції примітивів: "Цикл" для ітерації по парам та "Функція" для генерації кожної пари.

2. Формалізація за допомогою примітивів

Примітив "Цикл" (For-Loop)

S: Множина ітерацій – це кількість пар, які потрібно згенерувати. У нашому випадку множина не має фіксованої кількості, тому вона є динамічною.

F: Операція, яка застосовується до кожного елемента множини. У нашому випадку це генерація однієї пари ключових слів із тексту.

Примітив "Функція" (Abstraction)

Fextract(x, параметр): Функція, яка генерує одну пару ключових слів.

x : Вхідний текст, з якого потрібно виділити пари.

параметр: Тип зв'язку між словами (у нашому випадку – причинно-наслідковий).

3. Реалізація через композицію примітивів

3.1. Промпт для "Функції"

Промпт для генерації однієї пари ключових слів:

"Знайди два ключових слова, пов'язаних причинно-наслідковим зв'язком із тексту, і виведи їх у форматі 'Причина;Наслідок'."

3.2. Промпт для "Циклу"

Тепер об'єднаємо цей процес у цикл для генерації всіх пар:

"Для кожного причинно-наслідкового зв'язку у тексті знайди два ключових слова, пов'язаних цим зв'язком, і виведи їх у форматі 'Причина;Наслідок'. Результати виведи у вигляді нумерованого списку."

4. Кінцевий промпт

Остаточний промпт, що враховує всі базові примітиви:

"Створи нумерований список пар ключових слів, пов'язаних причинно-наслідковими зв'язками із тексту. Для кожної пари виведи їх у форматі 'Причина;Наслідок'."

5. Синтаксичне представлення (AST)

У вигляді абстрактного синтаксичного дерева (AST):

Prompt::= Цикл(Динамічна множина, Функція("Знайти пару ключових слів із тексту"))

Опрацюємо за допомогою кінцевого промπτу повідомлення з пропагандистського російського інформаційного агентства:

Лише за кілька днів позиції України в ЄС різко похитнулися. Лідер ХДС/ХСС Фрідріх Мерц заявив, що Київ не зможе перемогти Росію, а в Австрії вибори виграла партія, яка виступає проти антиросійських санкцій. З його [Мерца] точки зору, баланс сил виглядає „жахливим і дуже відрезвляючим”. Причому Київ опинився в такому становищі через союзників, які діяли „занадто мало і занадто пізно”. У свій час він підтримав ідею Макрона ввести війська НАТО на Україну. Крім того, був головним лобістом відправлення туди дальнобійних ракет Taurus. Також лідер ХДС/ХСС вважає, що Берлін зобов'язаний допомогти Києву вступити до ЄС. Вибори канцлера ФРГ відбудуться наступної осені. Мерц офіційно включився в боротьбу за цю посаду 17 вересня. Рейтинг ХДС/ХСС – 33%. У СДПГ (чинний глава держави Олаф Шольц – з цієї партії) – 14. Старший науковий співробітник сектору економіки європейських країн Центру європейських досліджень ІММО РАН Володимир Оленченко вважає, що Австрія і Німеччина переглядають ставлення до Києва. Правда, це відбувається не дуже швидко. „Австрійці і німці втомилися від української теми. Це одна з причин перемоги АПС і зміни риторики Мерца”, – зазначає він у розмові з РІА Новини.

Виконання кінцевого промπτу видає результати вигляду, що відповідає формату CSV:

Союзники діяли занадто мало і занадто пізно;Київ опинився у важкому становищі
Ідея Макрона ввести війська НАТО на Україну;Підтримка Фрідріха Мерца
Відправлення дальнобійних ракет Taurus;Зміцнення обороноздатності України
Баланс сил виглядає жахливим;Відрезвляючий ефект для України
Перемога партії проти антиросійських санкцій у Австрії;Зміна ставлення до Києва
Втома австрійців і німців від української теми;Перемога АПС на виборах
Зміна риторики Фрідріха Мерца;Нова політична стратегія
Позиції України в ЄС похитнулися;Зменшення підтримки серед європейських країн
Включення Фрідріха Мерца в боротьбу за посаду канцлера;Зростання рейтингу
ХДС/ХСС
Рейтинг ХДС/ХСС 33%;Зниження рейтингу СДПГ до 14%
Перегляд ставлення Німеччини та Австрії до Києва;Уповільнення інтеграції
України до ЄС
Відсутність перемоги над Росією;Сумніви у союзників щодо України
Допомога Берліна Україні;Прискорення вступу Києва до ЄС
Ідеологічна підтримка України;Політичні зміни в Європі
Протиросійські санкції;Економічний тиск на Росію
Економічна втома європейських країн;Зменшення підтримки санкцій проти Росії
Політична активність Фрідріха Мерца;Зміна політичного ландшафту Німеччини
Високий рейтинг ХДС/ХСС;Зниження впливу СДПГ
Уповільнення перегляду ставлення до Києва;Затримка допомоги Україні
Втома населення від української теми;Зменшення фінансування для України

Для візуалізації цієї мережі застосуємо сервіс Semantic Networking (рис. 6.7) для чого у поле введення завантажимо отриманий файл у форматі CSV.

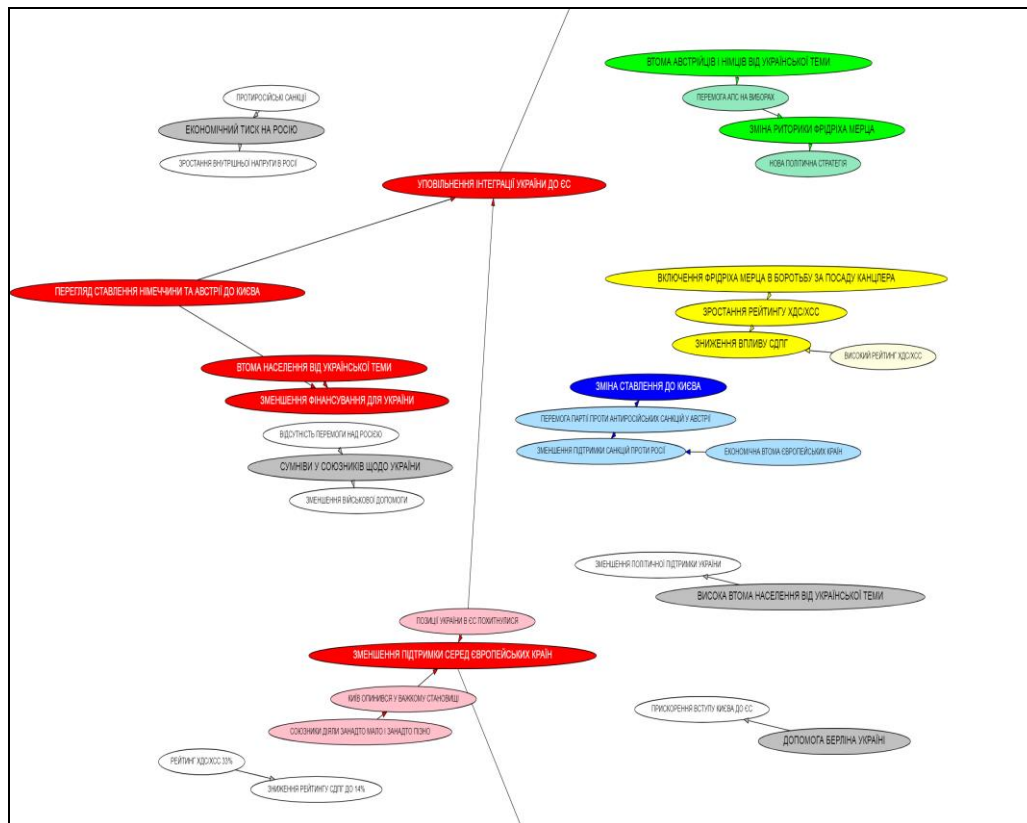


Рисунок 6.7: Фрагмент причинно-наслідкової мережі понять із повідомлення

Етап 2. Оцінка достовірності окремих зв'язків мережі і загальна оцінка достовірності

Для остаточної експертної оцінки достовірності тексту повідомлення через побудовану семантичну мережу модифікуємо промпт наступним чином:

За характером причинно-наслідкових зв'язків семантичної мережі оціни рівень достовірності всього тексту, джерела мережі. Кожний зв'язок оцінюється в діапазоні від 0 до 1. Можеш дати експертну оцінку достовірності всього тексту в діапазоні від 0 до 1. Ось мережа: ...

В результаті формалізації за методикою безкодового програмування створюється такий кінцевий промпт, який надається до виконання системі LLM:

Для кожного причинно-наслідкового зв'язку у семантичній мережі:

1. Оціни достовірність зв'язку в діапазоні від 0 до 1.
2. Збережи результати оцінок.

Після обробки всіх зв'язків:

1. Обчисли середнє значення достовірності.
2. Якщо середня достовірність перевищує 0.7, встанови загальну оцінку достовірності тексту як середнє значення; інакше – зменш її на 0.1.

Виведи загальну оцінку достовірності тексту в діапазоні від 0 до 1.

Ось мережа: ...

Виконання структурованого промпту видає результат вигляду:

Оцінка достовірності мережі повинна враховувати такі аспекти:		
– Суб'єктивність тверджень:		
– Багато зв'язків є інтерпретацією подій, а не об'єктивними фактами		
– Відсутність чітких доказів для багатьох причинно-наслідкових зв'язків		
– Можливі альтернативні пояснення:		
– Для кожного зв'язку існують потенційні альтернативні причини		
– Недостатня увага до контексту та супутніх факторів		
Оцінка окремих зв'язків :		
ЗВ'ЯЗОК	ОЦІНКА (0-1)	КОМЕНТАР
Союзнники діяли занадто мало і занадто пізно; Київ опинився у важкому становищі	0.6	Логічний зв'язок, але перехід може бути надмірно спрощеним
Ідея Макрона ввести війська НАТО на Україну; Підтримка Фрідріха Мерца	0.4	Слабка кореляція між цими подіями
Відправлення дальнобійних ракет Taugus; Зміцнення обороноздатності України	0.7	Прямий зв'язок, але ефективність сумнівна
...	...	...
Втома населення від української теми; Зменшення фінансування для України	0.5	Існують інші фактори, що впливають на фінансування
Загальні зауваження:		
Переважає більшість зв'язків має оцінку нижче 0.7		
Відсутність кількісних показників для багатьох тверджень		
Надмірна спрощеність багатьох причинно-наслідкових зв'язків		
Ігнорування можливих проміжних ланок у ланцюгах причинності		
Розрахунок загальної достовірності:		
– Середня оцінка: ~0.52		
– Медіана: ~0.5		
З урахуванням критичного підходу та високої планки вимог:		
Застосовуємо пониження на 0.1 (через недостатню обґрунтованість)		
Фінальна оцінка: 0.42		

Як можна бачити, загальна оцінка достовірності повідомлення не є високою.

Висновки до розділу

У цьому розділі було розглянуто комплексний аналіз інформаційно-психологічних ментальних війн, зосереджуючись на методах кластеризації мереж, ранжування вузлів за їхньою впливовістю та способах візуалізації інформаційних потоків. Цей підхід дозволяє ефективно ідентифікувати ключові агенти впливу, групи акаунтів або медіаресурсів, які беруть участь у поширенні інформації, а також оцінювати структуру і функціонування мережі.

Кластеризація мереж допомагає виявляти інформаційні фронти, що є важливими для розуміння механізмів поширення пропаганди, контрпропаганди та нейтральної інформації. Розподіл вузлів на кластери не лише полегшує аналіз великих даних, але й дозволяє

виділяти стратегічні точки, які мають значущий вплив на інформаційне середовище.

Ранжування вузлів за різними метриками центральності забезпечує точне виявлення найважливіших учасників інформаційних потоків. Це дозволяє ідентифікувати основних агентів впливу, посередників у передачі інформації, а також оцінювати їхню роль у розповсюдженні дезінформації або протидії їй. Зокрема, показники центральності за проміжністю допомагають визначати "ворота" для інформації, через які можна контролювати комунікаційні канали, а центральність за власним вектором дозволяє виявляти вузли, які переважно взаємодіють з іншими впливовими елементами мережі.

Візуалізація графів, особливо за допомогою інструментів типу Gephi, забезпечує наочне представлення даних, покращуючи їхню інтерпретацію. Використання силових алгоритмів розташування вузлів, кольорового кодування кластерів та виділення ключових вузлів робить аналіз доступним і зрозумілим. Це особливо важливо для практичного застосування результатів, наприклад, у контексті протидії інформаційним атакам або захисту національної ідентичності.

Окрім технічних аспектів, у розділі було розглянуто теоретичні моделі, такі як SIR-модель, перколяція та семантичний аналіз, які допомагають прогнозувати динаміку поширення інформації та оцінювати вразливість мережі. Особливу увагу приділено ролі ментальності народу як основного об'єкту захисту від інформаційно-психологічних загроз. У контексті сучасних конфліктів, таких як російська агресія проти України, це набуває стратегічного значення.

Таким чином, застосування математичних моделей, графових методів та інструментів візуалізації надає потужні інструменти для аналізу інформаційно-психологічних війн. Це дозволяє не лише розуміти механізми маніпуляції масовою свідомістю, а й розробляти ефективні стратегії протидії дезінформації, зберігаючи національну ідентичність та безпеку суспільства. Однак слід враховувати обмеження методів, зокрема обчислювальну складність для великих мереж та можливу неповноту даних, що вимагає постійного удосконалення аналітичних підходів.

Питання для самоконтролю

1. Які основні характеристики сучасних інформаційно-психологічних війн?
2. Які існують математичні моделі для прогнозування поширення інформації в інформаційному полі?

3. Як крос-факторний аналіз допомагає виявляти взаємозв'язки між ключовими чинниками національної ідентичності?
4. Як графові моделі допомагають у виявленні інформаційних операцій?
5. Які основні етапи інформаційної атаки можна виділити при аналізі даних?
6. Як визначити "ворожі вузли" у мережі поширення інформації?
7. Як метод згортання критеріїв допомагає у прийнятті стратегічних рішень щодо інформаційної безпеки?
8. Які основні типи пропаганди існують і як вони відрізняються за своєю ефективністю?
9. Як інформаційні загрози впливають на національну безпеку держави?
10. Які механізми психологічного впливу найчастіше використовуються у сучасних інформаційних війнах?
11. Як соціальні платформи можуть бути використані для протидії дезінформації?
12. Як оцінити ефективність заходів щодо протидії інформаційним загрозам?
13. Які основні виклики інформаційної безпеки в цифрову епоху?
14. Як інформаційні війни впливають на формування суспільної думки?
15. Які існують засоби візуалізації складних інформаційних потоків?
16. Як інформаційні технології змінюють характер сучасних конфліктів?

Література за розділом

1. Інформаційні операції та безпека суспільства: загрози, протидія, моделювання: монографія / В.П. Горбулін, О.Г. Додонов, Д.В. Ланде. – К.: Інтертехнологія, 2009. – 164 с. – Бібліогр.: С. 153-162. ISBN: 978-966-1648-12-7
2. Когут Ю.І. Кібервійни, кібертероризм, кіберзлочинність (концепції, стратегії, технології): практичний посібник. / Ю.І. Когут. – Київ: Консалтингова компанія «СІДКОН»; ВД «Дакор», 2022. – 284 с. ISBN: 9786179510069
3. Перлос Н. Ось таким, як мені кажуть, буде кінець світу: перегони кіберозброєнь. / Н. Перлос. – Харків: Фоліо, 2022. – 574 с.
4. Гаріс Ш. ВІЙН@: битви в кіберпросторі. / Ш. Гаріс. – Київ: «Ніка-Центр»; – Львів: Видавництво «Анетти Антоненко», 2019. – 296 с.

5. Померанцев П. Це не пропаганда. Подорож на війну проти реальності. / П. Померанцев. – Київ: Yakaboo, 2020. – 285 с.
6. Сінгер П., Бруклін Е. Війна лайків. Зброя в руках соціальних мереж. / П. Сінгер, Е. Бруклін. – Харків: Книжковий клуб “Клуб сімейного дозвілля”, 2019. – 319 с.
7. Качинський А.Б. Структурно-функціональна модель системи забезпечення інформаційної та інформаційно-психологічної безпеки. Допов. Нац. акад. наук України. 2023. №1. – С. 16-23.
8. Garth S. Jowett, Victoria O'Donnell. Propaganda and Persuasion. SAGE Publications, Inc; 1st edition (November 1, 1986). 232 pp. ISBN: 978-0803923980
9. Кулеба Д. Війни за реальність. Як перемагати у світі фейків, правд і спільнот. / Д. Кулеба. – К.: Книголав, 2019. – 384 с.
10. Сенгер Д.Е. Досконала зброя. Війна, саботажі страх у кіберепоху. / Д.Е. Сенгер. – Львів: Видавництво «Астролябія», 2022. – 496 с. ISBN: 978-617-664-237-4
11. Paul M.A. Linebarger. Psychological Warfare Paperback. CreateSpace Independent Publishing Platform (January 30, 1848). 234 p. ISBN: 978-1511545884
12. Патрикаракос Д. Війна у 140 знаках. Як соціальні медіа змінюють військовий конфлікт ХХІ століття./ Д. Патрикаракос. – Київ: Yakaboo, 2019. – 352 с. ISBN: 978-617-7544-18-9
13. Б'юкенен Б. Хакери і держави. Кібервійни як нові реалії сучасної геополітики. / Б. Б'юкенен. – Київ: «Наш формат», 2024. – 352 с.
14. Литвиненко О. Інформаційні впливи та операції. Теоретико-аналітичні нариси. / О. Литвиненко. – Київ: НІСД, 2003. – 240 с.
15. Ткаченко А.В., Качинський А.Б. Ідентифікація фейкових новин в соціальних ЗМІ / А.В. Ткаченко, А.Б. Качинський. // Всеукраїнська науково-технічна конференція студентів, аспірантів ті молодих вчених. – К.: ФТІ КПІ, 2019. – 49-50.
16. Smith, C.H., 1953. Psychological warfare. Naval War College Review, 5(6), pp.39-61.
17. Постанова Верховної Ради України «Про Заяву Верховної Ради України "Про використання політичним режимом російської федерації ідеології рашизму, засудження засад і практик рашизму як тоталітарних і людиноненависницьких"». Відомості Верховної Ради (ВВР), 2023, № 117, ст.533.
18. Качинський А.Б. Безпека складних систем. / А.Б. Качинський. – К.: ТОВ «Видавництво «Юстон», 2017. – 498 с.

19. Сміт Е.Д. Національна ідентичність. / Е.М. Сміт. Перлос Н. – К.: Основи, 1994. – 390 с.
20. Фукуяма Ф. Ідентичність. Потреба в гідності й політика скривдженості. / Ф. Фукуяма. – К.: Наш Формат, 2020. – 192 с.
21. Томпсон Ева М. Трубадури імперії: Російська література і колоніалізм. / Ева М. Томпсон. – К: Вид-во Соломії Павличко «Основи», 2006. – 368 с.
22. Thomas L. Saaty. The Analytic Hierarchy Process: Planning, Priority Setting, Resource Allocation (Decision Making Series). McGraw-Hill (January 1, 1980). 287 pp. ISBN: 978-0070543713
23. Thomas L. Saaty. Decision Making with Dependence and Feedback: The Analytic Network Process: the Organization and Prioritization of Complexity – RWS Publications, 1996. 370 pp. ISBN: 978-0962031793
24. Медовз Д. Мистецтво мислити системно. Розв'язання проблем від особистого до глобального масштабу. / Д. Медовз. – Київ: Віват, 2023. – 304 с.
25. Згуровский М.З., Панкратова Н.Д. Основы системного анализа. / М. Згуровский, Н.Д. Панкратова. – Київ: Видавнича група BHV, 2007. – 544 с.
26. Mats Lindgren, Hans Bandhold. Scenario planning. The link between future and strategy. Palgrave Macmillan (January 1, 2003). 240 pp. ISBN: 978-0333993170
27. Douglas A. Luke A User's Guide to Network Analysis in R. Springer; 1st ed. 2015 edition (December 21, 2015). 250 pp. ISBN: 978-3319238821
28. Menczer, Filippo, Santo Fortunato, and Clayton A. Davis. A first course in network science. Cambridge University Press, 2020. 300 pp. ISBN: 978-1108471138
29. B. Mandelbrot. The fractal geometry of nature. Times Books; 2nd prt. edition (January 1, 1982). 468 pp. ISBN: 978-0716711865
30. Провост Ф., Фоусот Т. Data science для бізнесу. Як збирати, аналізувати і використовувати дані. / Ф. Првост, Т. Фоусетт. – Київ: «Наш формат», 2019. – 400 с. ISBN: 978-617-7730-03-2

7. Динамічні аспекти інформаційно-психологічних війн

Інформаційно-психологічні війни на відміну від традиційних війн, не є статичним явищем – вони розвиваються у часі, адаптуючись до змін у соціальному, політичному та технологічному середовищі. Для ефективного аналізу та управління цими процесами критично важливо враховувати їхню динамічну природу.

Розглянемо динамічні аспекти інформаційних війн, зокрема на тому, як структура і функціонування інформаційних мереж змінюються у часі. Розглядаються графові моделі ІПВ, де вузли представляють поняття (наприклад, «ворожа пропаганда», «національна ідентичність», «патріотизм»), а зв'язки відображають взаємодії між цими поняттями (логічні, причинно-наслідкові, семантичні).

Такий підхід дозволяє аналізувати не лише статичну структуру мережі, але й її еволюцію під впливом зовнішніх факторів, таких як медійні кампанії, політичні події, соціальні тренди тощо.

Динаміка мереж є важливим аспектом для розуміння механізмів поширення інформації та формування суспільної думки. Це обумовлюється такими причинами, як зміна актуальності понять, адаптація до нових умов, необхідність прогнозування та управління, виявлення ключових моментів.

Певні поняття можуть набувати чи втрачати значущість залежно від контексту. Наприклад, під час глобальної пандемії поняття «коронавірус» стало центральним у інформаційному полі, тоді як інші теми могли відійти на другий план.

Мережі адаптуються до змін у зовнішньому середовищі. Наприклад, політична криза може призвести до появи нових понять (наприклад, «санкції», «гібридна війна») та до зміцнення зв'язків між існуючими поняттями.

Розуміння динаміки мереж дозволяє прогнозувати, які поняття стануть ключовими у майбутньому, і розробляти стратегії для захисту від маніпуляцій або активного впливу на інформаційне поле.

Аналіз еволюції мереж допомагає виявити критичні моменти, коли система може перейти у новий стан. Наприклад, різке зростання зв'язків навколо поняття «дезінформація» може сигналізувати про початок масштабної інформаційної атаки.

Розуміння динамічних аспектів ІПВ має практичне значення для медіа-стратегів при розробці ефективних медійних кампаній на основі аналізу актуальних понять та трендів, політиків з метою прогнозування

наслідків політичних рішень та управління інформаційним полем, для захисту від маніпуляцій шляхом виявлення ключових точок для блокування поширення дезінформації, науковців для розробки нових методів аналізу та моделювання інформаційних процесів.

Таким чином, динамічні аспекти інформаційних війн є критично важливими для розуміння механізмів поширення ідей та формування суспільної думки. Аналіз еволюції мереж у часі дозволяє не лише виявляти ключові поняття та зв'язки, але й прогнозувати їхню поведінку у майбутньому. У наступних підрозділах ми детально розглянемо механізми та моделі, які лежать в основі динаміки мереж ІПВ.

Розглядаються такі аспекти динамічних мереж ІПВ, як еволюція мереж у часі, механізми зміни ваги зв'язків, моделі динаміки мереж (Temporal Layers Model, Growth and Decay Model, SIR тощо).

При дослідженні еволюції мереж особлива увага приділяється питанням «як з'являються нові поняття та зв'язки?», «як змінюється структура мережі під впливом зовнішніх факторів?», «як вилучаються поняття та зв'язки, які втратили актуальність?». Також досліджуються питання зміни ваги зв'язків, а саме, «як сила впливів між поняттями змінюється з часом?», «які фактори впливають на посилення або ослаблення зв'язків?».

7.1. Основні характеристики динамічних мереж ІПВ

Динамічні мережі ІПВ мають декілька характерних ознак, які визначають їхню поведінку, а саме: нелінійність; складність взаємодій; чутливість до зовнішніх факторів; емерджентність.

Зміни у мережі часто відбуваються нелінійно. Наприклад, невелика подія (наприклад, поширення однієї фейкової новини) може призвести до лавиноподібного ефекту, який змінює всю структуру мережі.

Зв'язки між поняттями можуть бути багатшаровими. Наприклад, зв'язок між поняттями «економічні проблеми» і «політична нестабільність» може бути посилено через додаткові шари, такі як «соціальна напруга» або «медійний вплив».

Динаміка мережі суттєво залежить від зовнішніх подій. Наприклад, політичні рішення, природні катаклізми або технологічні інновації можуть кардинально змінити структуру взаємодій.

У мережах можуть виникати нові властивості, які не можна передбачити на основі аналізу окремих вузлів або зв'язків. Наприклад, поява нового наративу може призвести до формування цілого кластера пов'язаних понять.

7.2. Методологія дослідження динаміки

Для аналізу динамічних аспектів ІПВ використовуються різні методи та інструменти, зокрема: графові моделі; моделі центральності; моделі розповсюдження інформації; аналіз часових рядів; семантичний аналіз.

Мережі представляються у вигляді графів, де вузли – поняття, а зв'язки – впливи. Еволюція графа у часі вивчається за допомогою часових шарів (temporal layers) або послідовностей графів.

Метрики центральності, такі як Degree Centrality, Betweenness Centrality, Closeness Centrality та Eigenvector Centrality, застосовуються для оцінки важливості вузлів у різні моменти часу.

Моделі розповсюдження інформації. SIR-модель та її модифікації використовуються для аналізу того, як поняття активуються, поширюються та втрачають актуальність.

Зміни у структурі мережі (кількість вузлів, зв'язків, ваги зв'язків) аналізуються як часові ряди, що дозволяє виявити тренди та сезонність.

Семантичний аналіз тестів (наприклад, аналіз частоти спільного вживання слів у документах) використовується для виявлення нових понять та зв'язків між ними.

7.3. Еволюція мереж у часі

У графових моделях ІПВ, де вузли – це поняття (наприклад, «ворожа пропаганда», «національна ідентичність»), а зв'язки – це впливи (логічні, причинно-наслідкові, семантичні), аналіз еволюції мереж у часі допомагає відповісти на такі питання:

- Які поняття стають більш впливовими або навпаки втрачають актуальність?
- Як змінюється структура мережі під впливом зовнішніх факторів (наприклад, політичних подій, медійних кампаній)?
- Як нові поняття (вузли) інтегруються в існуючу мережу?

Розглянемо математичну формалізацію, моделі та приклади еволюції мереж у часі. Для опису еволюції мереж у часі використовуються динамічні графи, які представляють собою послідовність графів $G_t = (V_t, E_t)$, де V_t – множина вузлів (понять) у момент часу t , $E_t \subseteq V_t \times V_t$ – множина зв'язків (впливів) у момент часу t .

Матриця суміжності A_t для графа G_t описує взаємодії між вузлами в момент часу t :

$$A_{ij}(t) = \begin{cases} w_{ij}(t), & \text{якщо існує вплив від } i \text{ до } j \text{ у момент } t; \\ 0, & \text{якщо вплив відсутній.} \end{cases}$$

Тут $w_{ij}(t)$ – вага зв'язку (сила впливу) у момент часу t . Ваги можуть змінюватися з часом, що відображає динаміку взаємодій.

До основних аспектів еволюції мереж відносять додавання/вилучення вузлів або зв'язків і зміну ваги зв'язків.

Нові вузли-поняття можуть з'являтися у мережі через події, медійні кампанії або зміни в суспільному дискурсі. Наприклад, у ході війни можуть з'явитися нові поняття, такі як «гібридна війна» або «інформаційний фронт». Також поняття можуть втрачати актуальність або зникати з дискусій. Наприклад, термін «холодна війна» могли використовувати менше після кінця Холодної війни, але він знову став актуальним у сучасному контексті.

Таким чином, кількість вузлів $|V_t|$ у графі може змінюватися з часом:

$$|V_t| = |V_{t-1}| + \Delta V_t,$$

де ΔV_t – зміна кількості вузлів (додавання або вилучення).

Нові логічні або причинно-наслідкові зв'язки між поняттями можуть виникати через зміни в суспільстві. Наприклад, після соціальної кризи може з'явитися зв'язок між поняттями «економічні проблеми» і «політична нестабільність». Разом з цим деякі зв'язки можуть зникати, якщо поняття перестають бути пов'язаними. Наприклад, зв'язок між поняттями «технологічний прогрес» і «екологічні загрози» може змінюватися залежно від контексту.

Кількість зв'язків $|E_t|$ у графі також може змінюватися за формулою:

$$|E_t| = |E_{t-1}| + \Delta E_t,$$

де ΔE_t – зміна кількості зв'язків.

Сила впливів між поняттями може змінюватися з часом. Наприклад, зв'язок між поняттями «ворожа пропаганда» і «дезінформація» може посилюватися під час інформаційної кампанії.

Вагові значення зв'язків $w_{ij}(t)$ можуть змінюватися з часом:

$$w_{ij}(t) = w_{ij}(t-1) + \Delta w_{ij}(t),$$

де $\Delta w_{ij}(t)$ – зміна ваги зв'язку.

Метрики центральності (наприклад, Degree, Betweenness, Closeness) також змінюються з часом. Наприклад, Degree Centrality для вузла v у момент часу t :

$$C_{D(t)}(v) = \sum_{u \in V_t} A_{vu}(t).$$

7.4. Моделі еволюції мереж

Для опису та аналізу змін структури мереж у часі використовують різноманітні моделі еволюції мереж, які дозволяють врахувати динаміку додавання або видалення вузлів і зв'язків. Серед них особливе місце посідають модель часових шарів, що розглядає мережу як послідовність графів у певні моменти часу, та модель росту та спаду, яка формалізує процеси появи й зникнення елементів мережі. Ці підходи забезпечують основу для глибшого розуміння еволюційних механізмів складних систем, таких як соціальні мережі, транспортні системи чи біологічні мережі.

В рамках *моделі часових шарів* (Temporal Layers Model) граф розглядається як послідовність шарів, кожен з яких відповідає певному моменту часу. Наприклад:

- шар 1: граф у момент t_1 ;
- шар 2: граф у момент t_2 .

Це дозволяє аналізувати, як змінюється структура мережі між шарами.

Модель росту та спаду (Growth and Decay Model) описує як нові вузли та зв'язки додаються до графа, а старі вузли та зв'язки видалюються. Рівняння для розрахунку зміни кількості вузлів:

$$\frac{d|V_t|}{dt} = r_v - d_v,$$

де r_v – швидкість появи нових вузлів, d_v – швидкість видалення вузлів.

Аналогічно для зв'язків:

$$\frac{d|E_t|}{dt} = r_e - d_e,$$

де r_e – швидкість появи нових зв'язків, d_e – швидкість видалення зв'язків.

7.5. Моделі розповсюдження інформації

У Розділі 5 вже було розглянуто базову SIR-модель, яка успішно застосовується для опису динаміки поширення інфекцій та ідей у мережах. У контексті аналізу динамічних процесів у соціальних та когнітивних системах, зокрема в інформаційно-психологічних війнах, ця модель послужила основою для створення більш складних і реалістичних підходів. Розповсюдження інформації у таких системах має специфічні особливості – включаючи часові затримки, повторне активування понять, порогові ефекти, нерівномірність впливу між елементами та психологічні чинники. Тому для моделювання цих процесів використовують різноманітні модифікації класичних моделей, що враховують природу когнітивного простору.

Одними з найпоширеніших є SEIR-та SIS-моделі, які розширюють можливості SIR-підходу, додаючи латентні стани або передбачаючи циклічне повернення понять до стану сприйнятливості. Разом з тим, у сучасних дослідженнях все частіше застосовуються дискретні моделі, засновані на порогових механізмах, каскадному поширенні та агентному моделюванні. Ці підходи дають змогу врахувати не лише структуру мережі, а й складні взаємодії між окремими елементами, їхню вразливість до впливу, силу зв'язків між ними та реакції на зовнішні стимули. Такі моделі забезпечують гнучкий інструментарій для аналізу динаміки інформаційних кампаній, прогнозування поширення ідей та оцінювання ефективності контрпропагандистських стратегій.

Нагалаємо, що SIR-модель описує, як активність понять змінюється з часом:

$$\frac{dS}{dt} = -\beta SI, \quad \frac{dI}{dt} = \beta SI - \gamma I, \quad \frac{dR}{dt} = \gamma I.$$

У контексті аналізу динамічних мереж для інформаційно-психологічної війни, зокрема ментальних війн існує кілька моделей, які є розширеннями або модифікаціями класичної SIR-моделі. Ці моделі враховують специфіку динамічних процесів, таких як зміна структури мережі, взаємодія між поняттями, часові затримки та психологічні фактори. Розглянемо основні з них.

7.5.1. SEIR-модель (Susceptible-Exposed-Infected-Recovered)

Однією з найвідоміших моделей поширення інформації є SIR-модель, яка поділяє множину понять або індивідів на три стани: S (сприйнятливі – *Susceptible*), I (інфіковані, тобто активно поширюють інформацію – *Infected*) і R (відновлені, тобто втратили інтерес або припинили поширення – *Recovered*). У контексті аналізу когнітивного

простору це дозволяє моделювати, як ідеї набирають обертів і зникають з суспільної свідомості. Одним із важливих розширень цієї моделі є SEIR-модель, яка додає проміжний стан E (*Exposed* – експонований, або латентний), що відображає період, коли поняття або особа уже піддалася впливу, але ще не почала активно його поширювати. Це робить модель більш реалістичною для опису процесів, де перед поширенням інформації потрібен час на осмислення або внутрішню активацію.

В рамках моделі досліджуються такі етапи:

- $S \rightarrow E$: поняття переходить у латентний стан через контакт з активними поняттями (I);
- $E \rightarrow I$: поняття стає активним після деякої затримки (наприклад, часу для обробки інформації);
- $I \rightarrow R$: активне поняття втрачає актуальність.

Динаміка станів моделі описується системою диференціальних рівнянь:

$$\begin{aligned}\frac{dS}{dt} &= -\beta \cdot S \cdot I; \\ \frac{dE}{dt} &= \beta \cdot S \cdot I - \sigma \cdot E; \\ \frac{dI}{dt} &= \sigma \cdot E - \gamma \cdot I; \\ \frac{dR}{dt} &= \gamma \cdot I,\end{aligned}$$

де S – сприйнятливі поняття, E – латентні (експоновані) поняття, I – активовані поняття, R – поняття, які втратили актуальність, β – швидкість передачі впливу ($S \rightarrow E$), σ – швидкість переходу з $E \rightarrow I$, γ – швидкість забуття ($I \rightarrow R$).

У ході функціонування моделі поняття з S переходить у E через контакт з активованими поняттями (I); поняття з E переходить у I після затримки, яка визначається параметром σ ; поняття з I втрачає актуальність і переходить у R зі швидкістю γ .

Ця модель враховує часовий лаг між активацією ідеї та її поширенням. Наприклад, поняття "дезінформація" може бути спочатку прийнятим (E), а потім стати активним (I).

7.5.2. SIS-модель (Susceptible-Infected-Susceptible)

На відміну від SIR-моделі, у якій поняття або особа після втрати активності переходить у стан стійкості (R), у SIS-моделі (від

Susceptible-Infected-Susceptible) відсутній остаточний стан «відновленого». У цій моделі поняття, що перестало бути активним ($I \rightarrow S$), повертається до стану S – тобто знову стає сприйнятливим до повторної активації. Це означає, що одна й та сама ідея може знову набирати популярність, навіть якщо на певному етапі її поширення припинилося.

Така динаміка добре відображає реальні процеси в інформаційних системах, де окремі теми чи повідомлення регулярно повертаються в поле громадської уваги – наприклад, через нові події, медійні кампанії або повторне посилення пропагандистського тиску. Модель особливо корисна для аналізу ситуацій, коли суспільство не формує тривалої "імунітетної" реакції на певну ідею, і вона може знову активуватися під впливом зовнішніх або внутрішніх факторів.

В рамках моделі розглядаються такі етапи:

- $S \rightarrow I$: поняття активується;
- $I \rightarrow S$: поняття втрачає актуальність, але залишається доступним для повторної активації.

Динаміка станів описується такими рівняннями:

$$\begin{aligned}\frac{dS}{dt} &= -\beta \cdot S \cdot I + \gamma \cdot I; \\ \frac{dI}{dt} &= \beta \cdot S \cdot I - \gamma \cdot I,\end{aligned}$$

де S – сприйнятливі поняття, I – активовані поняття, β – швидкість передачі впливу ($S \rightarrow I$), γ – швидкість переходу з $I \rightarrow S$.

У ході функціонування моделі поняття із стану S переходить у I через контакт з активованими поняттями (I), а поняття з I повертається до S , коли воно втрачає актуальність.

Моделі корисна для аналізу циклічного характеру інформаційних кампаній. Наприклад, поняття "соціальна напруга" може бути активованим, а потім повернутися до стану S , коли воно знову стає чутливим до нових інформаційних атак.

7.5.3. Threshold Models (Моделі порогового типу)

У моделях порогового типу активація понять залежить від порогової кількості активних сусідів. Поняття переходить у стан I , лише якщо достатньо багато його сусідів уже активовані.

В моделі передбачено такі етапи функціонування:

- кожне поняття має свій пороговий параметр (θ), який визначає, скільки активних сусідів потрібно для активації;

– якщо частка активних сусідів перевищує θ , поняття активується.

Для кожного поняття i :

$$x_i(t+1) = \begin{cases} 1, & \text{якщо } \sum_{j \in N(i)} x_j(t) \geq \theta_i, \\ 0, & \text{інакше,} \end{cases}$$

де $x_j(t)$ – стан поняття j на момент часу t (1 – активоване, 0 – неактивоване), $N(i)$ – множина сусідів поняття i , θ_i – порогове значення для поняття i .

У цій моделі активація поняття i залежить лише від кількості активних сусідів ($\sum_{j \in N(i)} x_j(t)$). Якщо кількість активних сусідів перевищує або дорівнює пороговому значенню θ_i , поняття i активується.

Модель добре описує колективну поведінку у соціальних мережах. Наприклад, поняття «патріотизм» може бути активоване лише тоді, коли багато сусідніх понять («національна ідентичність», «єдність») вже активовані.

7.5.4. Independent Cascade Model (Модель незалежних каскадів)

Independent Cascade Model (ICM) – це ймовірнісна модель поширення інформації, яка описує процес активування понять або вузлів мережі через ланцюгову реакцію. У цій моделі кожне активоване поняття отримує одну можливість з певною ймовірністю p передати активність своїм сусідам. Якщо сусід було успішно активовано, він у свою чергу намагається активувати своїх сусідів на наступному кроці, продовжуючи таким чином каскадне поширення ідеї. Процес триває доти, доки всі нові активації стануть неможливими. ICM особливо корисна для моделювання швидкого та хаотичного розповсюдження ідей в інформаційному просторі, наприклад, в аналізі медійних кампаній або вивченні ефективності пропагандистських повідомлень.

Функціонування моделі включає етапи, що визначаються алгоритмом:

1. Процес починається з одного або кількох активованих понять.
2. На кожному кроці активовані поняття намагаються активувати своїх сусідів з ймовірністю p_{ij} .
3. Процес триває, поки всі можливі активації не завершаться.

Для кожного зв'язку $i \rightarrow j$:

$$P(i \rightarrow j) = p_{ij},$$

де p_{ij} – ймовірність того, що поняття i активує поняття j .

Модель корисна для аналізу поширення ідей через медіа. Наприклад, поняття «ворожа пропаганда» може активувати поняття «дезінформацію», яка, у свою чергу, активує «соціальну напругу».

7.5.5. Linear Threshold Model (Лінійна модель порогового типу)

Лінійна модель порогового типу – це розширення моделі порогового типу, де активація понять залежить від вагових коефіцієнтів зв'язків. Кожен сусід має вагу, і поняття активується, якщо сума ваг активних сусідів перевищує порогове значення.

Передбачено такі етапи функціонування:

- ваги зв'язків (w_{ij}) визначають силу впливу одного поняття на інше;
- порогове значення (θ_i) визначає, наскільки сильно поняття реагує на вплив.

Для кожного поняття i :

$$x_i(t+1) = \begin{cases} 1, & \text{якщо } \sum_{j \in N(i)} w_{ij} x_j(t) \geq \theta_i, \\ 0, & \text{інакше.} \end{cases}$$

де $x_i(t)$ стан поняття i на момент часу t , w_{ij} : вага зв'язку між поняттями i та j , θ_i : порогове значення для поняття i .

У цій моделі активація поняття i залежить від вагованих впливів активних сусідів ($\sum_{j \in N(i)} w_{ij} x_j(t)$). Поняття i активується, якщо сума ваг активних сусідів перевищує порогове значення θ_i .

Модель враховує нерівномірність впливу між поняттями. Наприклад, поняття «національна ідентичність» може мати великий вплив на «патріотизм», але менший – на поняття «соціальна напруга».

7.5.6. Agent-Based Models (Агентні моделі)

У агентних моделях кожне поняття (або ідея) представлено як агент, який має власні правила поведінки та взаємодіє з іншими агентами.

Агенти можуть мати різні характеристики:

- сприйнятливості до активації;
- силу впливу на інших агентів;
- тривалість активності.

Кожен агент (поняття) має свій стан $x_i(t)$ та правила взаємодії:

$$x_i(t+1) = f\left(x_i(t), \{x_j(t)\}_{j \in N(i)}\right),$$

де f – функція, яка визначає новий стан агента i на основі його поточного стану та станів сусідів.

Взаємодії між агентами можуть бути детермінованими або стохастичними.

Агенти взаємодіють згідно з певними правилами, які можуть бути детермінованими або стохастичними. Наприклад, правило може базуватися на порогових ефектах або ймовірностях.

Модель дозволяє симулювати складні взаємодії між поняттями. Наприклад, можна моделювати, як група понять («патріотизм», «єдність») реагує на зовнішній вплив («ворожа пропаганда»).

7.6. Приклади еволюції мереж у реальних ситуаціях

Еволюція когнітивних та соціальних мереж – це невід’ємна частина динаміки суспільних процесів, особливо в умовах інформаційно-психологічних війн, де поняття, ідеї та взаємозв’язки між ними постійно змінюються під впливом подій, медійних кампаній та політичних чинників. У реальних сценаріях можна спостерігати, як окремі терміни набирають ваги, як формуються нові асоціації між поняттями, а також як змінюється структура взаємодії між суб’єктами в соціальних і медійних мережах. Такі зміни можуть бути зафіксовані через аналіз метрик центральності, виявлення ключових вузлів, а також моделювання динаміки активації понять у часі. Розглянемо декілька прикладів, що демонструють механізми еволюції мереж у різних контекстах.

1. Еволюція понять у медіа. У медійному просторі поняття, такі як «пандемія», «вакцинація», «дезінформація», стали значно більш впливовими під час пандемії COVID-19. Це відображається у зростанні Degree Centrality та Betweenness Centrality цих понять.

2. Еволюція політичних наративів. У ході політичних конфліктів поняття, такі як «державний суверенітет», «економічна безпека», «патріотизм», стають більш важливими. Наприклад, зв’язок між «економічними санкціями» і «політичною стабільністю» може посилитися.

3. Еволюція соціальних мереж. У соціальних мережах нові тренди (наприклад, «кліматична криза») можуть створювати нові поняття та зв’язки, що призводить до зміни структури мережі.

Таким чином, еволюція мереж у часі є важливим аспектом аналізу інформаційно-психологічних війн. Застосування математичних моделей (часових шарів, росту та спаду, SIR-моделі тощо) дозволяє

виявити динаміку змін у структурі мережі, ідентифікувати ключові поняття та прогнозувати їхню поведінку.

7.7. Виявлення атак

У контексті інформаційно-психологічних війн (ІПВ), виявлення атак та розробка стратегій контрзаходів є критично важливими завданнями для захисту інформаційного поля від маніпулятивних кампаній, дезінформації та інших форм ворожого впливу.

Атаки в інформаційно-психологічних війнах можуть бути спрямовані на руйнування структури мережі, поширення дезінформації або зміну поведінки ключових понять. Нижче наведено основні типи атак, такі як атаки на ключові вузли, атаки на зв'язки, атаки на структуру мережі, атаки на динаміку мережі.

Атаки на ключові вузли спрямовані на вилучення або компрометацію вузлів з високою центральністю (Degree, Betweenness, Closeness, Eigenvector), наприклад:

- атака на поняття «державний суверенітет» може призвести до розриву зв'язків з іншими важливими поняттями, такими як «патріотизм»;
- компрометація поняття «національна ідентичність» може послабити національні наративи.

Атаки на зв'язки спрямовані на ослаблення або вилучення зв'язків між поняттями. Наведемо приклади:

- руйнування зв'язку між «економічними проблемами» і «політичною нестабільністю» може призвести до роз'єднання економічних та політичних наративів;
- ослаблення зв'язку між «дезінформацією» і «медіа» може зменшити можливості поширення фейків.

Атаки на структуру мережі спрямовані на руйнування глобальної структури мережі, наприклад:

- знищення зв'язності – вилучення ключових вузлів з високою Betweenness Centrality може призвести до розпаду мережі на окремі кластери;
- збільшення потоку інформації через певні зв'язки може призвести до їх перегрузки та зниження ефективності передачі ідей.

Атаки на динаміку мережі спрямовані на зміну динаміки мережі, наприклад за рахунок маніпуляції часовими трендами або введенням нових вузлів і зв'язків.

У результаті швидке посилення зв'язків навколо нового поняття (наприклад, «гібридна війна») може призвести до зміщення уваги суспільства, а додавання нових вузлів (наприклад, «кліматична криза») може змінити структуру мережі.

7.8. Математична формалізація атак

Для виявлення вузлів, які є цілями атак, використовуються метрики центральності. Наприклад:

- Degree Centrality: $C_D(v) = \sum_{u \in V} A_{vu}$.
- Betweenness Centrality: $C_B(v) = \sum_{s \neq v \neq t} \frac{\sigma_{st}(v)}{\sigma_{st}}$.
- Closeness Centrality: $C_C(v) = \sum_{u \neq v} \frac{1}{d(v, u)}$.

Атака на вузол v може бути формалізована як вилучення цього вузла з графа:

$$G' = G \setminus \{v\},$$

де G' – новий граф без вузла v .

Атака на зв'язок (i, j) може бути формалізована як вилучення ребра з графа:

$$E' = E \setminus \{(i, j)\},$$

де E' – нова множина зв'язків без ребра (i, j) .

Руйнування зв'язності мережі може бути формалізоване через перколяційну модель (див. Розділ 5).

Динаміка мережі може бути описана через часові зміни ваг зв'язків:

$$w_{ij}(t) = w_{ij}(t-1) + \Delta w_{ij}(t),$$

де $\Delta w_{ij}(t)$ – зміна ваги зв'язку. Атака може бути направлена на значне збільшення чи зменшення $\Delta w_{ij}(t)$.

7.9. Стратегії контрзаходів

У контексті інформаційно-психологічних війн стратегії контрзаходів мають на меті не лише протидіяти активним атакам на когнітивне та медійне поле, а й забезпечити стійкість, адаптивність та швидке відновлення структури мережі після зовнішніх впливів.

Оскільки поняття, зв'язки та взаємодії між ними формують основу інформаційного простору, ефективна стратегія повинна враховувати як локальні удари по окремим вузлам, так і системні атаки, спрямовані на руйнування загальної структури мережі. Для цього використовують комплекс підходів, що охоплює захист ключових понять, відновлення порушених зв'язків, підвищення загальної зв'язності мережі та швидке реагування на динамічні зміни.

Одним із базових принципів є захист ключових вузлів, які відіграють центральну роль у передачі ідей – наприклад, понять типу «національна ідентичність» чи «патріотизм». Виявлення таких вузлів здійснюється за допомогою метрик центральності, таких як *degree centrality* або *betweenness centrality*. Після їх ідентифікації проводяться профілактичні заходи – від освітніх програм до медійних кампаній, спрямованих на зміцнення довіри до цих понять.

Також важливим є відновлення зв'язків, особливо після деструктивних впливів, що порушують комунікаційні шляхи. Це може включати активацію альтернативних каналів поширення інформації або створення нових зв'язків для підвищення загальної зв'язності мережі.

Додатково застосовується підвищення стійкості структури мережі через кластеризацію та додавання нових вузлів, що зменшує вразливість до перехресних атак. Нарешті, прогнозування та швидке реагування дозволяють оперативно виявляти зміни в динаміці мережі та запускати автоматизовані механізми протидії.

Важливішою умовою стійкості когнітивної або медійної мережі є захист її найважливіших елементів – тих понять і зв'язків, які відіграють центральну роль у передачі ідей. Для цього аналізують структуру мережі за допомогою метрик центральності, щоб визначити, які вузли є критичними для її функціонування. Після їх виявлення розробляють стратегії підвищення їхньої стійкості, серед яких медійні кампанії, освітні ініціативи та інші форми активного протидії спробам дискредитації чи заміни ключових понять.

Не менш важливим є відновлення зв'язків, порушених внаслідок інформаційних атак. Це може передбачати активацію альтернативних шляхів поширення даних або навіть створення нових зв'язків між існуючими вузлами, щоб забезпечити безперебійну комунікацію в межах мережі. Такий підхід особливо ефективний, коли основні канали були уражені або втратили довіру серед аудиторії.

Щоб зробити мережу стійкою до системних атак, важливо не лише окремо захищати її частини, а й підвищувати загальну зв'язність структури. Це досягається через додавання нових вузлів або зміцнення внутрішніх зв'язків усередині груп схожих понять. Створення міцних

локальних кластерів дозволяє зберігати цілісність мережі навіть у разі масштабних впливів.

Окрім захисних заходів, необхідно мати механізми прогнозування та оперативного реагування. Моніторинг змін у структурі мережі дає змогу вчасно помічати ознаки атак, а автоматизовані системи реагування дозволяють швидко запускати контрзаходи. Такий комплексний підхід забезпечує не тільки відновлення мережі після атак, а й її адаптацію до нових видів загроз.

Одним із характерних прикладів інформаційної атаки є спроба маніпулювати сприйняттям поняття «дезінформація» з метою дискредитації медіа та поширення недовіри до джерел інформації. У цьому випадку ефективним контрзаходом є формування сильних зв'язків між поняттями «дезінформація» і «критичне мислення», що дозволяє не просто протистояти фейковим новинам, а й активно навчати суспільство розпізнавати маніпуляції на рівні освідомленості та аналізу.

Ще одним прикладом є спроба послабити взаємозв'язок між поняттями «економічні проблеми» та «політична нестабільність», щоб приховати реальний вплив економічного стану країни на її політичну ситуацію. Для протидії такій атаці можна використовувати поняття «соціальна напруга», через яке відновлюються ланки між економічним станом суспільства і його політичними настроями, забезпечуючи більш прозоре розуміння причин суспільних процесів.

Не менш серйозною загрозою є системна атака на структуру самої медійної мережі, спрямована на порушення зв'язності та блокування поширення достовірної інформації. Такі дії можуть призводити до інформаційної ізоляції окремих груп аудиторії чи масового замовчування важливих подій. Відповіддю на це може бути впровадження децентралізованих платформ, які забезпечують альтернативні канали поширення інформації, підвищуючи її доступність і стійкість до цензури чи маніпуляцій.

Таким чином, кожен тип атаки вимагає не лише точкового реагування, а й системного переосмислення структури когнітивної чи медійної мережі, щоб зберегти її цілісність і функціональність у умовах інформаційно-психологічної війни.

Виявлення атак та розробка стратегій контрзаходів у графових моделях інформаційно-психологічних війн дозволяє ефективно захищати інформаційне поле від маніпулятивних кампаній. Застосування математичних моделей (метрик центральності, перколяційних моделей, SIR-моделі) забезпечує комплексний аналіз структури та динаміки мережі, що є ключовим для успішного протидії атакам.

7.10. Прогнозування інформаційних атак

У контексті інформаційно-психологічних війн, прогнозування інформаційних атак є важливим завданням для забезпечення захисту інформаційного поля та запобігання поширенню дезінформації, маніпулятивних наративів та інших форм ворожого впливу.

Нижче детально розглянемо методи прогнозування, акцентуючи увагу на сценаріях атак, інструментах прогнозування (включаючи ARIMA, LSTM тощо) та прикладах.

7.10.1. Основні принципи прогнозування інформаційних атак

Прогнозування інформаційних атак є ключовою умовою ефективного захисту когнітивного простору, оскільки дозволяє не лише вчасно виявити потенційні загрози, а й сформуванати стратегію проактивної протидії. У цьому процесі важливо враховувати як структурні особливості мережі, так і її динамічні зміни, а також аналізувати минулі приклади інформаційних кампаній для виявлення типових сценаріїв агресивного впливу. Прогнозування базується на чотирьох основних принципах: аналіз уразливостей мережі, моніторинг її динаміки, моделювання гіпотетичних атак та використання історичних даних для встановлення повторюваних патернів. Разом ці підходи забезпечують комплексне розуміння можливих напрямів агресії та дають змогу протидіяти їм системно та заздалегідь.

Процес прогнозування починається з детального аналізу структури мережі з метою виявлення її найбільш уразливих місць – вузлів і зв'язків, які через свою центральність або функціональну роль можуть стати цілями для атаки.

Особлива увага приділяється оцінці того, як втрата окремих елементів може вплинути на загальну зв'язність і стабільність мережі. Далі здійснюється постійний моніторинг динаміки – відстежуються зміни в кількості вузлів і зв'язків, аналізуються часові тренди, щоб вчасно виявити аномалії, які можуть свідчити про підготовку до інформаційного нападу.

Для глибшого розуміння можливих сценаріїв агресії розробляються гіпотетичні моделі атак, які враховують різні способи впливу – від заміщення ключових понять до системного порушення комунікаційних шляхів. Ці сценарії тестуються за допомогою імітаційного моделювання, що дозволяє спрогнозувати реакцію мережі на різні типи втручань і вчасно підготувати адекватні контрзаходи. Не менш важливим є аналіз минулих інформаційних кампаній – вивчення тактики, засобів і послідовності дій агресора дає змогу встановити

повторювані шаблони поведінки, що значно підвищує точність прогнозів та ефективність протидії новим загрозам.

Таким чином, прогнозування інформаційних атак стає не просто аналітичним інструментом, а стратегічно важливим механізмом, що забезпечує стійкість когнітивного простору в умовах постійно змінної інформаційно-психологічної взаємодії.

7.10.2. Методи прогнозування інформаційних атак

Прогнозування інформаційних атак передбачає не лише аналіз структури когнітивної чи медійної мережі, а й глибоке розуміння того, які типи агресивного впливу найчастіше використовуються для її дестабілізації. Важливо вміти передбачати, які елементи мережі можуть стати цілями, які наслідки може мати їхнє ураження та як такі атаки можуть вплинути на загальну динаміку простору. У рамках цього аналізу виділяють кілька ключових сценаріїв, кожен з яких має специфічні механізми реалізації та потенційні наслідки.

Одним із найпоширеніших напрямів атак є спроба дискредитації або заміщення ключових вузлів – понять, які займають центральне місце в мережі через високе значення метрик центральності. Прикладом можуть бути такі поняття, як «дезінформація» або «національна ідентичність», які виступають важливими ланками у поширенні ідей або формуванні суспільної повестки. Якщо такий вузол буде вилучений або перетворений, це може призвести до значного порушення зв'язності мережі, фрагментації її структури та втрати цілісності інформаційного поля.

Не менш небезпечним є напрям атак, спрямованих на руйнування критичних зв'язків між стратегічно важливими поняттями. Наприклад, взаємозв'язок між «економічними проблемами» та «політичною нестабільністю» часто формує основу для розуміння причин соціального напруження. Якщо цей канал комунікації буде ослаблений або повністю зруйнований, це може призвести до втрати логічної цілісності між подіями та їхніми наслідками, що, у свою чергу, здатна викривити суспільне сприйняття реальності.

Ще одним рівнем агресивного втручання є системні атаки на саму структуру мережі, коли метою стає не окремий вузол чи зв'язок, а глобальна зв'язність усього інформаційного простору. Такі атаки можуть передбачати одночасне вилучення великої кількості вузлів або зв'язків, що призводить до розпаду мережі на ізольовані кластери, які втрачають змогу ефективно взаємодіяти між собою. Це особливо небезпечно у разі масштабних медійних кампаній, спрямованих на дестабілізацію довіри до джерел інформації або формування паралельних інформаційних реальностей.

Крім того, слід враховувати атаки, спрямовані не на структуру, а на динаміку мережі – зокрема, на часові зміни в активності понять та зв'язків. Наприклад, швидке формування нових зв'язків навколо певного поняття може свідчити про активну пропагандистську кампанію, яка має на меті переключити суспільну увагу або змінити ієрархію актуальних тем. Такі атаки вимагають не тільки оперативного виявлення, але й швидкого реагування, щоб запобігти необоротним змінам в інформаційному середовищі.

Таким чином, методи прогнозування інформаційних атак мають охоплювати широкий спектр можливих сценаріїв – від точкового втручання в окремі елементи до системних впливів на структуру та динаміку мережі в цілому. Лише комплексний аналіз цих факторів дозволяє ефективно протидіяти сучасним інформаційним загрозам та забезпечити стійкість когнітивного простору в умовах інформаційно-психологічних конфліктів.

7.10.3. Інструменти прогнозування

ARIMA (Autoregressive Integrated Moving Average)

ARIMA є одним із класичних методів для прогнозування часових рядів. Він використовується для аналізу змін у структурі мережі (кількість вузлів, зв'язків, ваги зв'язків):

$$X_t = \phi_1 X_{t-1} + \phi_2 X_{t-2} + \dots + \theta_1 \varepsilon_{t-1} + \theta_2 \varepsilon_{t-2} + \dots,$$

де X_t – це стан мережі у момент часу t . У контексті аналізу графів це може бути, наприклад, кількість вузлів, зв'язків або ваги зв'язків у графі; ϕ_i – параметри авторегресії. Вони показують, як значення X_t залежить від попередніх значень ($X_{t-1}, X_{t-2}, \dots$) (наприклад, якщо кількість нових зв'язків у графі залежить від минулих днів, то ϕ_i будуть описувати цю залежність); θ_i – параметри ковзного середнього. Вони враховують вплив попередніх помилок ($\varepsilon_{t-1}, \varepsilon_{t-2}, \dots$) на поточний стан. Помилки – це різниця між реальним значенням і прогнозованим; ε_t – шум (випадковий компонент), який моделює непередбачувані зміни у часовому ряді.

ARIMA використовується для прогнозування часових рядів, таких як кількість нових вузлів або зв'язків у графі, зміна ваги зв'язків у часі.

Цей метод добре підходить для короточасних прогнозів, коли зміни у структурі мережі мають лінійну або близьку до лінійної природу.

LSTM (Long Short-Term Memory)

LSTM є типом рекурентної нейронної мережі (RNN), яка добре підходить для аналізу часових рядів з довготривалими залежностями. LSTM може використовуватися для прогнозування динаміки мережі:

$$h_t = f(W_h h_{t-1} + W_x x_t + b),$$

Де h_t – приховані стани (hidden states) LSTM. Вони представляють внутрішній стан моделі на момент часу t . У контексті аналізу графів це може бути, наприклад, внутрішнє представлення структури мережі; h_{t-1} – приховані стани з попереднього моменту часу. LSTM враховує історичні дані, що робить його ефективним для моделювання довготривалих залежностей; x_t – вхідні дані на момент часу t . У контексті аналізу графів це може бути, наприклад, кількість нових вузлів або зміна ваг зв'язків; W_h, W_x – матриці вагових значень. Вони визначають, як приховані стани (h_{t-1}) та вхідні дані (x_t) впливають на поточний стан (h_t); b – зсув (bias) (це константа, яка додається до результату перед застосуванням активаційної функції); f – активаційна функція (наприклад, $\tanh$ або ReLU). Вона додає нелінійність до моделі, що дозволяє вловлювати складні закономірності у даних.

LSTM використовується для аналізу часових рядів з довготривалими залежностями, таких як прогнозування появи нових понять у графі (наприклад, коли поняття "кліматична криза" стане актуальним), аналіз динаміки зв'язків (наприклад, зміна ваг зв'язків у часі), виявлення сезонних трендів (наприклад, підвищення активності у соціальних мережах під час політичних подій).

LSTM особливо ефективний, коли зміни у структурі мережі мають складну, нелінійну природу, і коли важливо враховувати довготривалі залежності.

Машинне навчання

Одним із найбільш перспективних напрямів прогнозування інформаційних атак є застосування методів машинного навчання, які дозволяють не лише аналізувати великий обсяг минулих подій, а й виявляти приховані закономірності, на основі яких можна будувати точні прогнози щодо майбутнього розвитку подій. У цьому підході використовуються різноманітні алгоритми – від класичних дерев рішень до глибоких нейронних мереж, здатних моделювати складні взаємозв'язки в інформаційному просторі.

Основною перевагою машинного навчання є те, що моделі навчаються на реальних даних про минулі інформаційні кампанії: вони аналізують, які поняття ставали цілями, які типи атак були використані, як

відбувалося поширення ідей у мережі та якими були наслідки таких втручань. На основі цього алгоритми формують узагальнені шаблони поведінки агресора, що дає змогу вчасно виявити ознаки нової атаки навіть на етапі її зародження. Наприклад, модель може впізнати характерну активність навколо певного поняття, швидке формування нових зв'язків чи незвичайне зростання частоти вживання окремих термінів – все це може свідчити про початок інформаційної операції.

Крім того, сучасні методи машинного навчання дозволяють враховувати не тільки структуру когнітивної мережі, а й її динаміку в часі. Це особливо важливо для прогнозування атак, спрямованих на зміну суспільної повестки або маніпулювання громадською думкою через активізацію певних тем. Такі моделі можуть працювати в режимі моніторингу, постійно аналізуючи зміни в медійному та соціальному просторі та видаючи тривожні сигнали при виявленні потенційно небезпечних тенденцій.

Впровадження таких технологій дозволяє значно підвищити рівень готовності до протидії інформаційним впливам та зробити процес захисту когнітивного простору більш проактивним і науково обґрунтованим.

Семантичний аналіз

Семантичний аналіз виступає важливим інструментом у прогнозуванні інформаційних атак, особливо тих, що спрямовані на формування нових наративів або зміну акцентів у суспільній повестці. Цей підхід ґрунтується на глибокому вивченні текстових даних – матеріалів ЗМІ, соціальних мереж, блогів, офіційних повідомлень – з метою виявлення прихованих зв'язків між поняттями та визначення тематичних трендів, які поступово набирають обертів.

Одним із ключових етапів семантичного аналізу є вивчення частоти спільного вживання слів чи фраз у різних контекстах. Це дозволяє не лише встановити, які поняття починають асоціюватися одне з одним, а й виявити нові, ще не до кінця усвідомлені суспільством, зв'язки, які можуть стати основою для нових медійних кампаній або пропагандистських маніпуляцій. Наприклад, раптове зростання кількості текстів, у яких поряд вживаються слова «економічна криза» та «політична некомпетентність», може свідчити про підготовку до дискредитації влади через економічний фактор.

Крім того, семантичний аналіз дає змогу прогнозувати, які теми можуть стати центральними в найближчому майбутньому. Шляхом аналізу тематичних зміщень у масових медіа та соціальних мережах можна вчасно помітити зародження нових наративів, що формується навколо певних подій або понять. Це особливо важливо для виявлення атак,

спрямованих не на ураження вже існуючих структур, а на активне формування нового інформаційного середовища, яке згодом може замінити старі реалії новими, часто маніпулятивними.

Семантичний аналіз не лише допомагає розуміти сучасний стан когнітивного простору, а й забезпечує стратегічну перевагу через здатність передбачати майбутні зміни в інформаційному полі. Він стає потужним механізмом раннього виявлення потенційних загроз і дозволяє протидіяти інформаційним атакам на етапі їхнього формування, ще до того, як вони отримають широкий резонанс.

7.10.4. Додаткові підходи

Прогнозування інформаційних атак вимагає не лише аналізу структури мережі та її динаміки, а й застосування більш специфічних методів, які дозволяють оцінювати ризики на глибшому рівні, моделювати масштабні сценарії та вчасно виявляти перші ознаки потенційного нападу. Серед таких підходів особливе місце посідають оцінка ймовірності атаки, моделювання їх масштабних наслідків та аналіз соціальних сигналів.

Один із ключових елементів прогнозування – це *оцінка ймовірності атаки* на окремі вузли мережі. Цей процес ґрунтується на комбінації кількох факторів, серед яких – значення метрик центральності (що вказують на важливість вузла в структурі мережі) та аналіз змін у його оточенні з плином часу. Наприклад, якщо вузол має високу степеневу чи посередницьку центральність і водночас спостерігається раптове зростання активності навколо нього, це може свідчити про підвищену загрозу. Формалізована модель передбачає обчислення ймовірності атаки за допомогою формули, яка враховує ваги різних параметрів та темпи зміни зв'язків навколо вузла. Такий підхід дозволяє кількісно оцінити ступінь ризику для кожного елемента когнітивної мережі.

Ще одним важливим інструментом є *моделювання масштабних атак*, яке дає змогу зрозуміти, як мережа поводитиметься у разі одночасного ураження кількох вузлів або зв'язків. У такому сценарії використовують імітаційне моделювання: на першому етапі вилучають певну множину елементів, що можуть бути цілями агресора, а на другому – аналізують, як це вплине на загальну зв'язність і функціональність мережі. Для цього порівнюють стани мережі до і після "атаки", використовуючи спеціальні метрики, що кількісно вимірюють зміни. Це дозволяє не лише передбачити наслідки реальних атак, а й тестувати різні стратегії протидії, щоб обрати найбільш ефективні.

Не менш важливим є прогнозування на основі *аналізу соціальних сигналів* особливо в умовах швидкоплинної медійної дійсності. Цей

підхід полягає у постійному моніторингу соціальних мереж, таких як X (Twitter), Facebook або Telegram, з метою виявлення нових понять, хештегів, ключових слів чи тематичних трендів, які починають набирати популярності. Аналіз таких сигналів дозволяє вчасно помітити зародження нових наративів, які згодом можуть стати центральними у суспільній повестці або ж використовуватися як інструмент пропаганди. Це особливо корисно для виявлення атак, спрямованих на формування нового інформаційного середовища, що поступово замінює старі реалії.

Таким чином, додаткові підходи до прогнозування інформаційних атак забезпечують глибше розуміння ризиків, дозволяють тестувати можливі сценарії та вчасно реагувати на нові загрози. Вони доповнюють базові принципи прогнозування і формують комплексний механізм захисту когнітивного простору в умовах інформаційно-психологічної взаємодії.

7.10.5. Приклади прогнозування інформаційних атак

Прогнозування інформаційних атак є важливою складовою аналізу когнітивного простору в умовах інформаційно-психологічних війн. Наприклад, поняття «дезінформація» характеризується високим значенням метрики Betweenness Centrality, оскільки воно часто виступає посередником між ключовими темами, такими як «медіа» та «суспільство». Це робить його особливо привабливим об'єктом для атаки, спрямованої на порушення комунікаційних зв'язків між іншими поняттями. У разі успішного ураження цього вузла може відбутися фрагментація мережі та втрата зв'язності між важливими сегментами інформаційного поля.

Особливу увагу також заслуговує зв'язок між поняттями «економічні проблеми» та «політична нестабільність», який має високу вагу, що свідчить про його значущість у передачі соціального напруження в політичну площину. Атака на цей зв'язок може призвести до формування хибного сприйняття причинно-наслідкових залежностей у суспільстві, послаблення реакції на економічні виклики та зниження громадянської активності. Такий сценарій загрожує стабільності суспільства в цілому.

Ще однією потенційною ціллю є медійна мережа, яка містить велику кількість вузлів з високою Closeness Centrality – тобто таких, що швидко взаємодіють з іншими частинами системи. Якщо атака буде спрямована на ці елементи, це може суттєво уповільнити поширення достовірної інформації, знизити її доступність серед широкого загалу та підірвати довіру до незалежних джерел.

Застосування методів прогнозування, зокрема моделей типу ARIMA, LSTM та SIR, дозволяє не лише виявити вразливі місця в структурі мережі, але й моделювати розвиток подій у разі їхнього ураження. Такий комплексний підхід забезпечує основу для стратегічного планування контрзаходів, що дозволяє протидіяти інформаційним атакам більш ефективно та заздалегідь готуватися до можливих викликів.

7.11. Практична частина. Моделювання динаміки інформаційних графів

У цьому пункті показано, як використовувати промпт-інжиніринг для моделювання та аналізу динаміки інформаційних графів. Такий підхід дозволяє, з одного боку, швидко формувати моделі предметних областей за допомогою великих мовних моделей, а з іншого боку оцінювати глибину цих областей, її насиченість ключовими концептами.

Методологія семантичного приєднання

Мережі концептів є потужним інструментом для моделювання знань у різних галузях науки. У роботах Альберт-Ласло Барабаші та Реці Альберти було запропоновано модель переважного приєднання, яка пояснює формування безмасштабних мереж зі степеневим розподілом ступенів вузлів. У цьому пункті розглядається адаптація моделі Барабаші-Альберт для аналізу мережі концептів, яка формується ітеративно за допомогою великої мовної моделі. На відміну від класичного переважного приєднання, нові вузли приєднуються не лише на основі ступенів існуючих вузлів, але й враховуючи їх семантичну близькість до нових термінів. Цей процес отримав назву переважного семантичного приєднання.

Модель реалізується наступним чином. Спочатку розглядається семантична мережа, яка формується за певним алгоритмом з понять, що входять до документів з вибраної тематики, зокрема «ментальної війни». Тобто на початковому етапі мережа складається з фіксованої кількості зв'язків між поняттями, які екстраговані з конкретної статті. Після цього протягом кожного такту роботи цієї динамічної моделі до цієї мережі додаються нові вузли-поняття з вибраної предметної області. Після деякої великої кількості тактів, яка обумовлюється змістом нових концептів, що оцінюється експертом, генерація припиняється.

Генерація промту для реалізації моделі

Для реалізації моделі переважного семантичного приєднання скористуємось методологією безкодового формування систем, де суть запропонованого підходу для створення систем полягає у представленні промптів як аналогів програмних конструкцій (умовні оператори, цикли, функції) шляхом математичної формалізації їх логіки та взаємодій. Для формування промптів застосовуються такі основні примітиви («Циглинка»), такі як «Умова», «Цикл» і «Функція», а також методи композиції цих примітивів для побудови складних систем, зокрема семантичних мереж.

Для отримання чіткого формального промпту формування мережі переважного семантичного приєднання, наведений у попередніх розділах фреймворк перетворюється на документ у форматі PDF і завантажується в систему LLM (наприклад, ChaGPT, Qwen або DeepSeek). Після чого завантажуюмо промпт-завдання для формування формалізованого промпту створення мережі переважного семантичного приєднання наступного вигляду:

У відповідності з наведеною методикою прошу сформувати загальний промпт для вирішення задачі: Нехай у нас є деякий вхідний текст [вхідний текст] за тематикою «Ментальна війна». З тексту виділяється 15 пар найбільш зв'язаних між собою понять за цією тематикою, які для утворення мережі виводяться у вигляді списку у форматі "Поняття 1;Поняття 2". Мережа буде неспрямованою. Після формування цієї мережі 20 разів здійснюється така процедура: Вибирається якесь нове поняття з заданої тематики і воно пов'язується з трьома поняттями вже до того сформованої мережі, які найсильніше з ним зв'язані і виводяться у вигляді "нове поняття; поняття". На цьому цикл завершається. Остаточно сформована мережа виводиться у вигляді списку "Поняття 1;Поняття 2". Надай сформований промпт

Нижче наведено вигляд формалізованого промпту, завдання для формування мережі переважного семантичного приєднання, сформованого у результаті виконання наведеного первинного промпту:

Ви отримали вхідний текст [вхідний текст], який належить до тематики «Ментальна війна». Ваше завдання – створити мережу понять, що описує взаємозв'язки між ключовими ідеями цього тексту. Для цього дотримуйтесь наведеної нижче методики:

Початковий аналіз тексту:

- Проаналізуйте вхідний текст і виділіть 15 пар найбільш сильно пов'язаних між собою понять за заданою тематикою.
- Запишіть ці пари у форматі "поняття;поняття". Мережа на цьому етапі є неспрямованою.

Розширення мережі:

Проведіть 20 ітерацій, де на кожній ітерації:

- Виберіть нове поняття з заданої тематики [тематика].
- Знайдіть три поняття з уже сформованої мережі, які найсильніше пов'язані з новим поняттям.
- Додайте до мережі нові зв'язки у форматі "нове поняття;поняття".

Формування результату:

Після завершення всіх ітерацій виведіть остаточну мережу у вигляді списку зв'язків у форматі "Поняття 1; Поняття 2".

Також наводяться деталі формалізації, формальний опис промпту через примітиви.

1. Примітив «Функція» для виділення початкових пар

`Fextract(текст, параметр="виділення 15 пар понять") → InitialNetwork`

Вхідні дані: текст за тематикою [тематика].

Операція: аналіз тексту для виділення 15 пар понять з найбільшим зв'язком.

Вихідні дані: список пар у форматі "поняття;поняття".

2. Примітив «Цикл» для розширення мережі

`For (i = 1 to 20):`

`Fextend (нове поняття, параметр="пошук 3 найближчих понять у мережі") → ExtendedNetwork`

Вхідні дані:

- нове поняття – поняття, обране з тематики [тематика];
- `InitialNetwork + ExtendedNetwork` – поточна мережа понять.

Операція: для кожного нового поняття знаходяться три найближчих зв'язаних поняття з поточної мережі.

Вихідні дані: нові зв'язки у форматі "нове поняття;поняття".

3. Примітив «Функція» для об'єднання результатів

`Fmerge (InitialNetwork, ExtendedNetwork) → FinalNetwork`

Вхідні дані :

- `InitialNetwork` – початкова мережа з 15 пар;
- `ExtendedNetwork` – мережа, розширена після 20 ітерацій.

Операція: об'єднання всіх зв'язків у єдиний список.

Вихідні дані: остаточна мережа у форматі "поняття;поняття".

Результати експерименту

Як документ для генерації початкової мережі розглядається наступне повідомлення Академії Української Преси:

Шукати шляхи нової пропаганди

20.11.2024

Пропаганда була, є і буде. Це "мрія", яку продають суспільству. Але вона завжди

спирається на емоції, лише мімікуючи під серйозну розмову. Пропаганда сильна лише тоді, коли за нею стоять реальні дії.

Сьогоднішня війна багато в чому "знищила" статус старих інструментів управління, але коли вона закінчиться, вони знову придуть. З цієї причини треба бути готовим до них. І одним із таких інструментів була і буде пропаганда. Масова свідомість відрізняється від індивідуальної, вона потребує свого підходу. І пропаганда якраз є розмовою з масовою свідомістю.

...

Академія Української Преси 2024.11.20 12:12

<https://www.aup.com.ua/news/shukati-shlyakhi-novoi-propagandi/>

У результаті опрацювання кінцевого промпту велика мовна модель Qwen2.5-max видає:

Попередній аналіз тексту та виділення 15 пар понять

На основі вхідного тексту, який стосується тематики «Ментальна війна», виділяються наступні 15 пар найбільш сильно пов'язаних між собою понять:

Пропаганда;емоції

Мрія;масова свідомість

Інформаційна війна;когнітивна війна

...

Розширення мережі через 20 ітерацій

Для кожної ітерації додаємо нове поняття з тематики «Ментальна війна» та знаходимо три поняття з уже сформованої мережі, які найсильніше пов'язані з ним.

Ітерація 1

Нове поняття: соціальний протест

Зв'язки:

соціальний протест;мрія

соціальний протест;кумулятивний ефект

соціальний протест;ментальна війна

Ітерація 2

Нове поняття: емоційна маніпуляція

Зв'язки:

емоційна маніпуляція;пропаганда

емоційна маніпуляція;емоції

емоційна маніпуляція;мрія

Ітерація 3

Нове поняття: віртуальна реальність

Зв'язки:

віртуальна реальність;візуальність

віртуальна реальність;мрія

віртуальна реальність;радянська манія контролю

(Продовження аналогічно для решти ітерацій.)

Остаточна мережа у форматі "Поняття 1; Поняття 2"

Пропаганда;емоції

Мрія;масова свідомість

Інформаційна війна;когнітивна війна

Пропагандистська боротьба;візуальність

... (і так далі для всіх ітерацій).

Ця мережа описує взаємозв'язки між ключовими поняттями тексту, що стосуються тематики «Ментальна війна».

Таким чином, при початковій мережі: $m_0 = 15$ вузлів (первинні поняття, екстраговані з статті) шляхом проведення 20 ітерацій сформовано мережу з 58 вузлів (рис. 7.1).

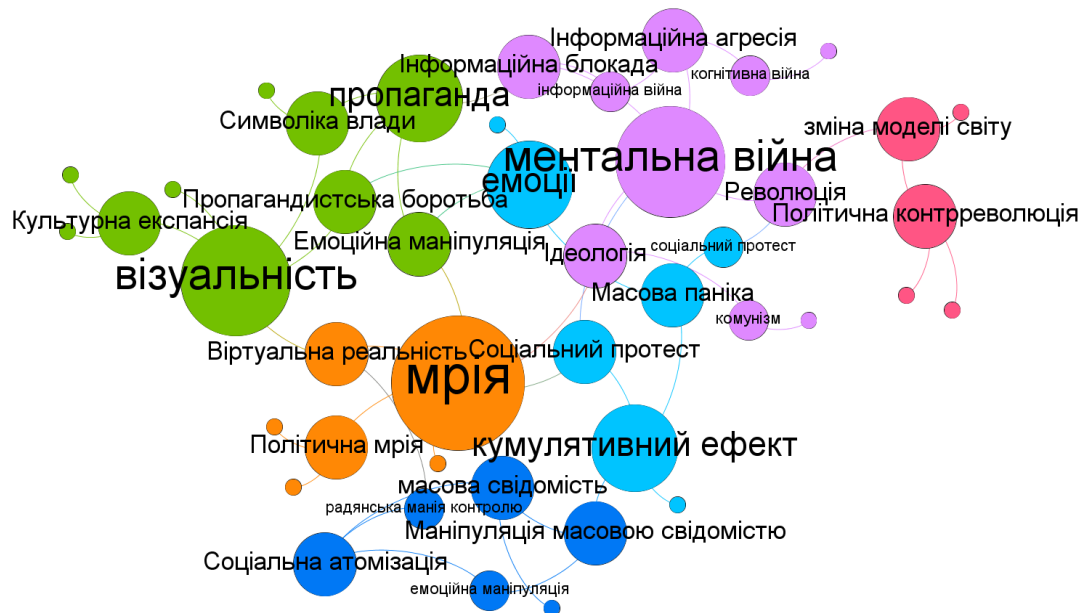


Рисунок 7.1: Вигляд згенерованої за допомогою системи Gephi мережі

Висновки до розділу

У цьому розділі було детально розглянуто динамічні аспекти інформаційно-психологічних війн, зокрема механізми еволюції мережових структур, моделі атак на інформаційні мережі та методи прогнозування та захисту від таких атак. Цей аналіз має практичне значення для розуміння того, як інформаційні потоки формуються, трансформуються та захищаються у сучасному суспільстві.

Одним із ключових аспектів розділу є аналіз динаміки мережових структур, а саме того, як інформаційні мережі змінюються з часом через додавання/вилучення вузлів та зв'язків, а також через зміну ваг зв'язків. Наприклад:

- поява нових понять (наприклад, «гібридна війна») може значно змінювати структуру мережі, створюючи нові комунікаційні канали;
- вилучення понять або ослаблення зв'язків між ними (наприклад, між поняттями «економічні проблеми» і «політичною нестабільністю») може призвести до розриву наративів або зниження їхнього впливу;

- зміна ваг зв'язків відображає еволюцію суспільних пріоритетів та інтересів, що є важливим для прогнозування інформаційних трендів.

Ці зміни демонструють, що інформаційні мережі є динамічними системами, чутливими до подій, медійних кампаній та соціальних змін. Саме тому аналіз часових шарів та динаміки мереж є критичним для розуміння механізмів поширення ідей.

В розділі також детально описано різні типи атак на інформаційні мережі, спрямовані на підрив їхньої структури та функціональності:

- атаки на ключові вузли. Спричиняють втрату важливих понять або концепцій, що може призвести до розриву комунікаційних каналів;
- атаки на зв'язки. Ослаблюють причинно-наслідкові зв'язки між поняттями, наприклад, між "економічними проблемами" і "політичною нестабільністю";
- атаки на динаміку мережі. Маніпуляція часовими трендами або введення нових понять може зміщувати увагу суспільства та впливати на його переконання.

Ці атаки часто спрямовані на ключові аспекти національної ідентичності, таких як «державний суверенітет» або «національна безпека», що робить їх особливо небезпечними у контексті інформаційно-психологічних війн.

Захист інформаційних мереж передбачає комплексний підхід:

- захист ключових вузлів, спрямований на підсилення зв'язків навколо важливих понять (наприклад, «критичне мислення» або «медіа грамотність») для зменшення вразливості до атак;
- кластеризація та підвищення зв'язності для формування сильних внутрішньокластерних зв'язків для зменшення впливу атак на окремі частини мережі;
- прогнозування та реагування з використанням інструментів, таких як ARIMA, машинне навчання та семантичний аналіз, для виявлення потенційних атак та швидкого реагування на них.

Ці методи дозволяють не лише захищати існуючі структури, а й прогнозувати майбутні загрози, що є важливим для забезпечення інформаційної безпеки.

Розуміння динамічних аспектів ІПВ має широке практичне застосування:

- медіа-стратегами при розробці ефективних медійних кампаній на основі актуальних трендів;

- політиками для прогнозування наслідків політичних рішень та управління інформаційним полем;
- для захисту від маніпуляцій, що передбачає виявлення ключових точок для блокування поширення дезінформації;
- науковцями для розробки нових методів аналізу та моделювання інформаційних процесів.

Особливу увагу приділено складності взаємодій у інформаційних мережах. Зв'язки між поняттями часто є багатошаровими, що ускладнює їх аналіз й крім того, нелінійність змін у мережі (наприклад, лавиноподібний ефект від однієї фейкової новини) підкреслює важливість моніторингу та прогнозування.

Таким чином, динамічний аналіз інформаційних мереж є потужним інструментом для розуміння механізмів поширення ідей, прогнозування інформаційних атак та розробки стратегій захисту. У контексті сучасних інформаційно-психологічних війн цей підхід допомагає захищати національну ідентичність та забезпечувати інформаційну безпеку суспільства. Однак слід враховувати обмеження методів, зокрема складність моделювання багатошарових взаємодій та чутливість до зовнішніх факторів, що вимагає постійного удосконалення аналітичних підходів.

Питання для самоконтролю

1. Як складність взаємодій у мережах впливає на аналіз інформаційних впливів?
2. Чому нелінійність є однією з ключових особливостей динаміки мереж ІПВ?
3. У чому полягає емерджентність у контексті інформаційних мереж?
4. Як графові моделі допомагають у дослідженні змін у структурі мереж?
5. Які метрики центральності найчастіше використовуються для аналізу динаміки мереж?
6. Як часові шари (temporal layers) застосовуються для вивчення еволюції графів?
7. Як SIR-модель допомагає у прогнозуванні активності понять у динамічних мережах?
8. Як описується зміна кількості зв'язків у графі з плином часу?
9. Як зміна ваги зв'язків впливає на структуру мережі?
10. Які основні типи атак на графові моделі ІПВ виділяються?
11. Як атаки на ключові вузли впливають на структуру мережі?

12. Як можна формалізувати атаку на зв'язки у графі?
13. Як перколяційна модель допомагає прогнозувати руйнування зв'язності мережі?
14. Які основні стратегії захисту ключових вузлів?
15. Як можна відновити зв'язки після атаки на мережу?
16. Як підвищення зв'язності мережі допомагає протистояти атакам?
17. Які основні етапи прогнозування інформаційних атак виділяються?
18. Як машинне навчання застосовується для прогнозування атак на мережі?
19. Як семантичний аналіз допомагає у прогнозуванні нових інформаційних трендів?
20. Як оцінюється ймовірність атаки на конкретний вузол у мережі?
21. Які приклади реальних ситуацій, де спостерігалася еволюція мереж, наведені?
22. Як можна використовувати історичні дані для прогнозування майбутніх інформаційних атак?

Література за розділом

1. Інформаційні операції та безпека суспільства: загрози, протидія, моделювання: монографія / В. П. Горбулін, О. Г. Додонов, Д. В. Ланде. – К.: Інтертехнологія, 2009. – 164 с. – Бібліогр.: С. 153-162. ISBN 978-966-1648-12-7
2. Information Operations Recognition. From Nonlinear Analysis to Decision-Making / A. Dodonov, D. Lande, V. Tsyganok, O. Andriichuk, S. Kadenko, A. Graivoronskaya – LAP Lambert Academic Publishing, 2019. – 292 p. ISBN 978-620-0-27697-1
3. Mehran Mesbahi and Magnus Egerstedt. Graph Theoretic Methods in Multiagent Networks. Princeton University Press, 2010. – 434 p. ISBN 9780691140612

8. Автоматизація аналізу та прогнозування

Інформаційно-психологічні ментальні війни характеризуються великою кількістю взаємопов'язаних даних, складними інформаційними потоками та високою динамічністю мережевих структур. Ручний аналіз таких процесів є трудомістким і недостатньо ефективним.

Автоматизація дослідження інформаційних впливів за допомогою машинного навчання (ML) та штучного інтелекту (AI) дозволяє прискорити аналіз великих обсягів даних, виявляти приховані закономірності у мережах, прогнозувати зміни у структурах інформаційного простору, автоматизувати моніторинг та виявлення загроз.

Автоматизація аналізу інформаційних війн дозволяє покращити ефективність досліджень та реагування на загрози. Використання ML та AI допомагає аналізувати великі обсяги даних, прогнозувати динаміку інформаційних атак та створювати аналітичні інструменти для моніторингу інформаційного середовища.

У цьому розділі розглядаються методи ML та AI для аналізу графів, прогнозування змін у мережах, розробка аналітичних інструментів та їх практичне застосування.

8.1. Використання ML та AI для аналізу графів

У контексті інформаційно-психологічних війн, використання машинного навчання та штучного інтелекту для аналізу графів є інструментом для автоматизації процесів ідентифікації важливих понять, прогнозування атак та розробки стратегій захисту. У графових моделях ІПВ, де вузли – поняття (наприклад, «ворожа пропаганда», «національна ідентичність»), зв'язки – впливи (логічні, причинно-наслідкові, семантичні), методи машинного навчання та штучного інтелекту дозволяють ефективно аналізувати структуру та динаміку мережі, прогнозувати зміни та ідентифікувати потенційні загрози.

Нижче детально розглянемо методи, математичну формалізацію та приклади.

8.1.1. Основні завдання ML та AI у аналізі графів

Аналіз когнітивних, медійних та соціальних мереж у сучасних умовах неможливий без застосування потужних інструментів машинного навчання та штучного інтелекту, які дозволяють не лише автоматизувати обробку великих обсягів даних, а й виявляти складні залежності, передбачати майбутню поведінку системи та формувати стратегії ефективної протидії зовнішнім загрозам. Методи ML та AI забезпечують глибоке розуміння структури графів, їхньої динаміки та

функціональної ролі окремих елементів, що має критичне значення для аналізу інформаційно-психологічних війн. Ці технології використовуються для вирішення широкого кола завдань – від класифікації понять до прогнозування атак і генерації рекомендацій щодо підвищення стійкості мереж.

Одним із важливих напрямів є *класифікація вузлів*, яка дозволяє автоматично визначати типи понять чи їхні функціональні ролі в мережі. Наприклад, алгоритми можуть ідентифікувати, чи є певне поняття «позитивним» чи «негативним» наративом, а також встановлювати, чи виступає вузол як «джерело інформації», «посередник» чи «споживач». Такий аналіз допомагає краще зрозуміти логіку поширення ідей і виділити елементи, що відіграють особливу роль у формуванні суспільної поведінки.

Ще однією важливою задачею є *прогнозування зв'язків*, спрямоване на виявлення потенційних нових асоціацій між поняттями та оцінку зміни їхньої сили в часі. Це особливо корисно для аналізу формування нових наративів, що можуть стати основою для масових медійних кампаній або пропагандистських маніпуляцій. Алгоритми аналізують історичні дані про взаємодії між елементами та виявляють приховані закономірності, які свідчать про наближення таких змін.

Важливим напрямком є також *виявлення аномалій* – відстеження нехарактерних змін у структурі мережі, які можуть бути ознакою інформаційної атаки. Наприклад, раптове зростання ступеня центральності окремого вузла чи різке формування великої кількості нових зв'язків може свідчити про активну кампанію з маніпуляції думкою. Вчасне виявлення таких сигналів дозволяє оперативно реагувати на загрозу.

Для аналізу локальних процесів у мережі використовується *сегментація мережі*, що передбачає виявлення кластерів – груп тісно пов'язаних понять або суб'єктів. Це дає змогу детально вивчати окремі наративи, що формуються в різних частинах мережі, а також виявляти «комунікативні бар'єри» між кластерами, які можуть свідчити про розкол у суспільній думці або дискредитацію певних джерел інформації.

Крім того, методи ML та AI широко застосовуються для *прогнозування атак* – виявлення потенційних цілей через аналіз метрик центральності та моделювання сценаріїв втрати ключових елементів. Цей підхід дозволяє не лише визначити найбільш вразливі місця в структурі мережі, але й оцінити, які зміни можуть відбутися в разі їхнього ураження.

Нарешті, *генерація рекомендацій* є практичним застосуванням отриманих аналітичних результатів. На основі отриманих даних алгоритми можуть запропонувати оптимальні стратегії захисту мережі, включаючи підсилення ключових вузлів, створення нових зв'язків для підвищення зв'язності чи впровадження механізмів швидкого реагування. Це перетворює аналіз даних у конкретний інструмент управління інформаційним простором.

Використання ML та AI у графовому аналізі відкриває нові можливості для розуміння складних когнітивних структур, передбачення їхньої поведінки та формування ефективних стратегій захисту від інформаційних загроз. Ці технології стають невід'ємною частиною сучасних методів аналізу інформаційно-психологічних війн, забезпечуючи науково обґрунтований підхід до дослідження динаміки когнітивного простору.

8.2. Graph Neural Networks

Аналіз когнітивного простору, його структури та динаміки вимагає застосування сучасних методів обробки даних, які здатні враховувати не лише окремі елементи інформаційної системи, але й їхні взаємодії. Одним із найбільш перспективних напрямів у цьому напрямку є використання Graph Neural Networks (GNN) – класу алгоритмів машинного навчання, спеціально розроблених для аналізу графових структур. Ці моделі забезпечують глибоке розуміння складних мереж понять, що має особливе значення в умовах інформаційно-психологічних війн, де ключову роль відіграють наративи, їхні взаємозв'язки та механізми поширення.

Graph Neural Networks відрізняються від традиційних нейронних мереж тим, що вони прямо оперують з графами, враховуючи не лише ознаки окремих вузлів, а й топологію зв'язків між ними. Це дозволяє моделям «бачити» контекст кожного елемента в межах усієї мережі, що є критично важливим при аналізі медійних полів, соціальних мереж чи когнітивних моделей суспільної свідомості. У контексті інформаційно-психологічних війн GNN можуть бути використані для широкого кола завдань: від прогнозування поширення ідей до виявлення ключових нарративів і формування стратегій протидії пропаганді.

Основна перевага GNN полягає в їхній здатності обробляти складні взаємодії між поняттями, враховуючи як безпосередні, так і опосередковані зв'язки. Наприклад, модель може не лише визначити, що певне поняття набуло популярності, а й зрозуміти, через які шляхи воно поширюється, хто є основними посередниками його активації та які інші наративи починають з ним асоціюватися. Такий підхід

дозволяє не просто фіксувати зміни, а й передбачати їхній подальший розвиток.

Один із ключових механізмів роботи GNN – агрегація інформації, коли кожен вузол збирає дані про стан своїх сусідів. На основі цієї інформації відбувається оновлення стану вузла, що моделює його зміну під впливом оточення. Після кількох таких ітерацій, аналогічних шарам нейронної мережі, система переходить до прогнозування : вона може класифікувати вузли, передбачати появу нових зв'язків, виявляти аномалії чи виділяти групи понять з однаковою семантикою.

Завдяки цьому GNN можуть виконувати широкий спектр аналітичних завдань. Зокрема, вони використовуються для класифікації вузлів – визначення типу поняття (наприклад, «позитивний», «негативний», «маніпулятивний») або його ролі в мережі («джерело», «посередник», «рецептор»). Важливою функцією є прогнозування поширення зв'язків та ідей, що дозволяє вчасно виявити зародження нового наративу чи активізацію старого. Крім того, GNN добре справляються з виявленням ключових наративів – тих понять і зв'язків, які формують основу поточної інформаційної повестки.

Особливе значення GNN мають для протидії ворожій пропаганді, оскільки вони дають змогу аналізувати не лише окремі факти, а й структуру інформаційного простору в цілому. На основі отриманих даних можна розробляти стратегії захисту – від посилення довіри до ключових понять до формування альтернативних комунікаційних шляхів, що компенсують втрату зв'язків через дезінформаційні кампанії.

Таким чином, Graph Neural Networks стають потужним інструментом для аналізу інформаційних полів, забезпечуючи глибоке розуміння динаміки когнітивних структур та їхньої реакції на зовнішні втручання. Впровадження цих технологій дозволяє не лише ефективно аналізувати інформаційно-психологічні процеси, а й формувати науково обґрунтовані стратегії протидії маніпулятивним впливам.

Коротко наведемо математичну формалізацію. Граф $G = (V, E)$ складається з множини вузлів (понять) V і множини зв'язків (взаємодій між поняттями) E . Кожен вузол $v_i \in V$ має:

- атрибути, де $x_i \in \check{Y}^d$ (наприклад, частота вживання слова, важливість поняття);
- стан, де $h_i^{(l)} \in \check{Y}^h$ на шарі l .

Кожен зв'язок $e_{ij} \in E$ може мати вагу – w_{ij} (сила взаємодії між поняттями).

Процес передачі повідомлень виглядає наступним чином.

Для кожного вузла V_i :

1) Збір інформації від сусідів

$$m_i^{(l+1)} = \text{AGGREGATE}\left(\{h_j^{(l)}, \forall j \in N(i)\}\right),$$

де $N(i)$ – множина сусідів вузла i , AGGREGATE – функція агрегації (наприклад, середнє значення, сума).

2) Оновлення стану вузла

$$h_i^{(l+1)} = \sigma\left(W \cdot \text{COMBINE}\left(h_i^{(l)}, m_i^{(l+1)}\right) + b\right),$$

де W – матриця вагових значень, b – зсув, σ – активаційна функція (наприклад, ReLU), COMBINE – функція комбінування стану вузла та інформації від сусідів.

3) Прогнозування

Після кількох шарів (L) GNN генерує вектори станів для всіх вузлів:

$$H^{(L)} = \left[h_1^{(L)}, h_2^{(L)}, \dots, h_n^{(L)}\right].$$

Ці вектори можуть бути використані для класифікації вузлів, прогнозування зв'язків, визначення глобальних характеристик графа.

Застосування GNN у ІПВ

Graph Neural Networks – потужний інструмент для моделювання складних систем понять, прогнозування їхньої поведінки та формування стратегій протидії пропаганді.

GNN дозволяють не просто фіксувати окремі елементи інформаційного поля, але й розуміти їхню роль у структурі мережі, враховуючи як безпосередні, так і опосередковані зв'язки між ними. Це надає можливість виявляти ключові поняття, передбачати поширення нових ідей, аналізувати динаміку наративів та розробляти проактивні стратегії захисту інформаційного простору. У рамках інформаційно-психологічних війн GNN можуть бути використані для чотирьох основних завдань: виявлення ключових понять, прогнозування поширення ідей, виявлення ворожих наративів та формування стратегій протидії пропаганді.

Одним із перших напрямів є виявлення ключових понять, тобто тих елементів мережі, які мають найвищу впливовість через свою позицію

та зв'язки. GNN може оцінювати центральність вузлів не лише на основі структури мережі, а й на основі стану кожного елемента після тривалого навчання, що дає глибше розуміння їхньої ролі в системі. Наприклад, у графі «ворожа пропаганда» → «дезінформація» → «соціальна напруга» модель може встановити, що «дезінформація» є центральним вузлом, через який поширюється вплив.

Ще однією важливою задачею є прогнозування поширення ідей. GNN здатні моделювати динаміку активності вузлів, враховуючи їхні взаємодії з оточенням. Це дозволяє аналізувати, як активація одного поняття може спричинити реакцію в інших частинах мережі. Наприклад, якщо активовано поняття «маніпуляція свідомістю», GNN може передбачити, що це призведе до подальшої активації таких тем, як «соціальна напруга» або «політична невдоволеність».

Не менш важливим є виявлення ворожих наративів, коли метою є визначення зловмисних інформаційних конструкцій, які формуються через специфічну комбінацію понять. GNN може аналізувати великі масиви текстових даних – соціальні мережі, медіа-матеріали, блоги – і виділяти з них семантичні зв'язки, характерні для пропагандистських кампаній. Наприклад, модель може виявити, що у текстах ворожої пропаганди часто разом зустрічаються поняття «дезінформація» і «національна безпека», що свідчить про спробу дискредитації державних інститутів.

На основі отриманих даних GNN може бути використана і для розробки стратегій протидії пропаганді. Модель виявляє слабкі місця в мережі – наприклад, поняття з низькою центральністю або зламані зв'язки – і пропонує способи їхнього підсилення. Так, GNN може рекомендувати активацію таких понять, як «патріотизм» або «єдність», щоб компенсувати негативний вплив ворожих наративів.

Серед переваг GNN у контексті інформаційно-психологічних війн варто виокремити:

- обробка складних взаємодій – модель враховує не лише окремі поняття, а й їхні взаємозв'язки, що забезпечує реалістичне представлення когнітивного простору;
- масштабованість – GNN ефективно працюють з великими графами, що особливо важливо при аналізі масових інформаційних полів, де кількість понять і зв'язків може бути надзвичайно великою;
- гнучкість – GNN можна адаптувати для різноманітних завдань, включаючи класифікацію, прогнозування, кластеризацію та аналіз часових змін;

- інтерпретованість – стани вузлів після навчання можуть бути проаналізовані людиною, що дозволяє отримувати цінні інсайти щодо структури інформаційного поля.

Наприклад, при аналізі графа понять «ворожа пропаганда» → «дезінформація» → «соціальна напруга» виконується така послідовність дій:

1. Побудова графа: вузли: окремі поняття; зв'язки: семантичні залежності між ними, побудовані на основі частоти спільного вживання або TF-IDF-аналізу.
2. Навчання GNN: модель аналізує взаємодії між поняттями, встановлює їхню важливість та прогнозує, як вони можуть активуватися в майбутньому.
3. Прогнозування: на основі поточної активності модель передбачає, які поняття будуть найбільш вразливими або найбільш впливовими.
4. Результати: виявлено, що "дезінформація" є ключовим поняттям для поширення ворожої пропаганди; – запропоновано стратегію активації контрпонять ("патріотизм", "єдність") для протидії.

Застосування GNN у інформаційно-психологічних війнах дозволяє не лише аналізувати структуру когнітивного простору, а й формувати ефективні стратегії протидії маніпулятивним впливам. Ці моделі забезпечують науково обґрунтований підхід до аналізу інформаційних полів, що робить їх важливим інструментом в інформаційній безпеці.

8.3. Node Embedding

Node Embedding – це метод представлення вузлів у вигляді векторів низької розмірності, які зберігають структурну інформацію про граф.

Node Embedding є важливим інструментом для аналізу графів понять у контексті ІПВ. Він дозволяє виявляти подібність між поняттями, класифікувати вузли за їхніми характеристиками, представляти вузли у вигляді векторів низької розмірності, зберігати структурну інформацію про граф, виявляти ключові поняття та прогнозувати взаємодії.

Цей метод особливо корисний для аналізу інформаційних полів, де важливо зрозуміти, як формуються та поширюються ідеї та наративи.

Node Embedding переводить кожен вузол V_i у графі $G = (V, E)$ у вектор $z_i \in R^d$, де $d = |V|$. Ці вектори зберігають структурну інформацію (наприклад, сусідство вузла), атрибути вузла (якщо вони є), зв'язки між вузлами. Вони можуть бути використані для машинного навчання, зокрема для класифікації вузлів, прогнозування зв'язків, кластеризації понять.

Основа математичної формалізації методу Node Embedding – це граф $G = (V, E)$, який складається з множини вузлів (понять) V і множини зв'язків (взаємодій між поняттями) E .

Мета Node Embedding – знайти відображення $f: V \rightarrow R^d$, яке переводить кожен вузол v_i у вектор $z_i \in \mathbb{R}^d$, де d – розмірність вектора ($d = |V|$), z_i – вектор, який зберігає інформацію про вузол v_i .

Для знаходження векторів z_i використовується функція втрат, яка максимізує збереження структурної інформації. Наприклад:

$$L = \sum_{(v_i, v_j) \in E} \log P(v_j | v_i),$$

де $P(v_j | v_i)$ – ймовірність того, що v_j є сусідом v_i , E – множина зв'язків у графі.

Ймовірність $P(v_j | v_i)$ часто моделюється за допомогою softmax:

$$P(v_j | v_i) = \frac{\exp(z_i, z_j)}{\sum_{k \in N(i)} \exp(z_i, z_k)},$$

де (z_i, z_j) – скалярний добуток векторів z_i та z_j , $N(i)$ – множина сусідів вузла v_i .

Розглянемо деякі алгоритми Node Embedding.

DeepWalk

DeepWalk – це один із перших алгоритмів, який запозичив підхід до побудови текстових вкладень (Word2Vec) і застосував його до графів. Основна ідея полягає в тому, щоб розглядати вузли графа як "слова", а шляхи по графу – як "речення", що дозволяє використовувати методи обробки природної мови для аналізу мереж.

DeepWalk використовує випадкові блукання по графу для генерації послідовностей вузлів, які потім обробляються за допомогою Word2Vec, в рамках якого генерується послідовність вузлів $(v_1, v_2, \dots, v_n)$ за допомогою випадкових блукань. Вектори z_i навчаються максимізувати ймовірність спостереження сусідніх вузлів:

$$L = \sum_{i=1}^n \log P(v_{i+k} | v_i),$$

де k – довжина контексту.

До переваг алгоритму можна віднести здатність зберігати локальну структуру графа та простоту реалізації, завдяки використанню готових методів Word2Vec.

Node2Vec

Node2Vec є розширенням DeepWalk, яке дає більший контроль над тим, яку частину графа досліджують випадкові блукання. У цьому алгоритмі вводяться два параметри, які регулюють характер блукань: p і q .

Випадкові блукання тепер не є повністю випадковими – вони формуються з урахуванням двох параметрів:

- p (return parameter) – впливає на ймовірність повернення до попереднього вузла. Велике значення p зменшує шанс того, що блукання повернеться назад, тобто «виштовхує» його вперед.
- q (in-out parameter) – визначає, наскільки ймовірно продовження блукання у напрямку вже відвіданих вузлів (внутрішній пошук) або на віддалені вузли (зовнішній пошук). Наприклад, мале q змушує блукання досліджувати більш віддалені області графа.

Ці параметри дозволяють гнучко регулювати баланс між локальними (ближні сусіди) та глобальними (структура всього графа) властивостями графа.

Алгоритм Node2Vec реалізується наступними кроками:

1. Генеруються випадкові блукання, в яких вибір наступного вузла залежить від параметрів p і q .
2. Отримані послідовності вузлів передаються в Word2Vec (звичайно, skip-gram).
3. Формула функції втрат аналогічна DeepWalk:

$$\max P(v_i | v_j; p, q).$$

Але тепер ця ймовірність формується з урахуванням параметрів, що контролюють тип дослідження графа.

До переваг алгоритму слід віднести більшу гнучкість через параметри p і q та можливість адаптації до різних задач: виявлення кластерів (якщо більше уваги на локальну структуру) або моделювання глобальної топології графа.

Graph Factorization

На відміну від попередніх алгоритмів, що базуються на випадкових блуканнях, Graph Factorization – це математично строгий підхід, який намагається наблизити матрицю суміжності графа за допомогою факторизації матриці.

Мета алгоритму – знайти представлення вузлів у вигляді векторів, які добре відображають структуру графа. Це досягається через апроксимацію матриці суміжності A як добутку двох менших матриць:

$$A \approx XX^T,$$

де X – це матриця, рядки якої є векторними представленнями вузлів.

Реалізація алгоритму полягає у тому, що:

- будується матриця суміжності графа A , де $A_{ij} = 1$, якщо між вузлами i та j є зв'язок, інакше 0;
- знаходиться матриця X , яка мінімізує функцію втрат:

$$L = \sum_{(v_i, v_j) \in E} \|A_{ij} - x_i^T x_j\|^2 + \lambda \|x_i\|^2,$$

де x_i – векторне представлення вузла i , E – множина зв'язків графа,

λ – параметр регуляризації, що запобігає перефїтінгу.

Таким чином, модель намагається навчити вектори x_i так, щоб скалярний добуток $x_i^T x_j$ був максимально близьким до наявності зв'язку між вузлами i та j .

Перевага алгоритму полягає у тому, що ця модель є дуже точною, особливо для графів з чітко вираженою структурою. Також алгоритм підходить для невеликих і середніх графів, де важлива точність, а не швидкість.

LINE (Large-scale Information Network Embedding)

LINE – це алгоритм node embedding, призначений для ефективного моделювання масштабних графів, особливо тих, що містять мільйони вузлів і зв'язків. На відміну від методів на кшталт DeepWalk чи Node2Vec, які використовують випадкові блукання, LINE прямо моделює структуру графа, зберігаючи перший та другий порядок взаємодій між вузлами. Це дозволяє більш точно передати як локальну, так і глобальну структуру мережі.

Мета LINE – навчити числові вектори для кожного вузла таким чином, щоб:

- зберегти перший порядок подібності – тобто безпосередні зв'язки між вузлами;
- зберегти другий порядок подібності – тобто схожість вузлів за їхніми оточеннями (сусідами).

Це робить LINE особливо корисним у задачах аналізу соціальних, медійних чи когнітивних мереж, де важливо врахувати не лише прямі зв'язки, а й їхню роль у ширшій структурі.

Перший порядок подібності характеризує безпосередню взаємодію між двома вузлами – наприклад, наявність зв'язку між ними. Для цього LINE максимізує ймовірність спостереження зв'язку між вузлами i та j , враховуючи вагу цього зв'язку w_{ij} . Формально це можна записати як:

$$L_1 = \sum_{(v_i, v_j) \in E} w_{ij} \log P(v_j | v_i),$$

де E – множина зв'язків графа, w_{ij} – вага зв'язку між вузлами i та j (наприклад, частота взаємодії чи сила семантичного зв'язку), $P(v_j | v_i)$ – умовна ймовірність того, що вузол j «спостерігається» при наявності вузла i .

Іншими словами, модель намагається навчити такі вектори v_i , щоб вузли, які мають сильний зв'язок, мали близькі представлення у векторному просторі.

Другий порядок описує структуру оточення вузла – тобто, наскільки схожі два вузли за своїми сусідами. Це особливо важливо в когнітивних мережах, де поняття можуть мати однакове значення через спільні асоціації, навіть якщо між ними немає прямого зв'язку.

У LINE це реалізується шляхом максимізації ймовірності того, що два вузли мають подібне оточення:

$$L_2 = \sum_{(v_i, v_j) \in E} \log P(N(v_j) | v_i),$$

де $N(v_j)$ – множина сусідів вузла j . Таким чином, модель навчається не лише впізнавати безпосередні зв'язки, але й формувати подібні вектори для вузлів, які мають аналогічні ролі в структурі графа, навіть якщо вони фізично не пов'язані.

LINE може опціонально використовувати як перший, так і другий порядок подібності. У більшості випадків обидва підходи поєднуються, щоб забезпечити комплексне представлення графа:

1. спочатку навчається модель першого порядку, щоб зберегти безпосередні взаємодії;
2. потім навчається модель другого порядку, щоб зберегти схожість за оточенням;
3. обидва набори векторів можуть бути об'єднані, щоб отримати загальне представлення вузлів.

Цей процес забезпечує добре збалансоване представлення графа, що поєднує локальну і глобальну структуру.

До переваг LINE слід віднести підтримку зважених графів – LINE враховує силу зв'язків (ваги), що особливо важливо в когнітивних мережах, де зв'язки мають різну силу (наприклад, частота спільного вживання понять), масштабованість – алгоритм ефективно працює навіть на дуже великих графах, що робить його придатним для аналізу великих медійних полів чи соціальних мереж, збереження локальної та глобальної структури – завдяки обробці першого і другого порядку, LINE здатний передавати як окремі взаємодії, так і загальну роль вузлів у мережі, гнучкість – можна використовувати лише один із порядків, якщо потрібно акцентувати увагу на конкретному типі взаємодій.

LINE особливо корисний у аналізі когнітивних мереж, де важливо зрозуміти не лише зв'язки між поняттями, а й їхню роль у системі. Наведемо приклади:

- Модель може виявити, що поняття «патріотизм» і «єдність» мають подібні оточення, хоча між ними немає прямого зв'язку.
- Також може бути використана для прогнозування появи нових зв'язків – наприклад, яке поняття найімовірніше почне асоціюватися з темою «дезінформація» в результаті активної медійної кампанії.
- LINE допомагає класифікувати поняття за їхньою структурною роллю – наприклад, виявляти ключові наративи, посередників або ізольовані групи понять.

Благодаря здатності моделювати перший та другий порядок взаємодій, LINE забезпечує точне і багатовимірне представлення графа, що може використовуватися для виявлення ключових наративів, прогнозування поширення ідей або формування стратегій протидії пропаганді.

Застосування Node Embedding у ІПВ

У сучасних інформаційно-психологічних війнах ключову роль відіграє аналіз когнітивного простору – тобто понять, їхньої структури, зв'язків між ними та механізмів поширення ідей. Для цього все частіше застосовують методи Node Embedding, які дозволяють представити поняття у вигляді числових векторів, зберігаючи при цьому важливі властивості графа. Цей підхід забезпечує не лише глибоке розуміння інформаційного поля, а й ефективне прогнозування, класифікацію та протидію пропагандистським кампаніям.

Одним із найважливіших завдань є виявлення ключових понять, тобто тих, що мають найвищу впливовість через свою структурну роль чи зв'язки з іншими елементами мережі. Для цього використовують

вектори, отримані за допомогою Node Embedding: на основі їхньої близькості в векторному просторі можна класифікувати вузли або об'єднувати їх у кластери. Наприклад, у графі «ворожа пропаганда» → «дезінформація» → «соціальна напруга» модель може встановити, що «дезінформація» є центральним ланцюгом, через який передається вплив. Такий аналіз дає змогу точково впливати на найвразливіші місця інформаційного поля.

Ще однією важливою задачею є прогнозування нових зв'язків, що дозволяє передбачати, як будуть формуватися нові наративи або активізуватися старі теми. У цьому випадку вектори окремих понять порівнюють між собою за допомогою спеціальної функції, яка оцінює ймовірність появи зв'язку:

$$p(i, j) = \sigma(z_i^T z_j),$$

де z_i та z_j – векторні представлення понять, а σ – сигмоїдна функція, що перетворює скалярний добуток у ймовірність. Це особливо корисно для виявлення потенційних нових асоціацій, які можуть бути частиною планової медійної кампанії.

Крім того, Node Embedding ефективно застосовується для аналізу семантичної близькості між поняттями. Використовуючи таку метрику, як косинусна відстань:

$$\text{similarity}(i, j) = \frac{z_i \cdot z_j}{\|z_i\| \cdot \|z_j\|}$$

можна виявити, які поняття починають мати спільний контекст, навіть якщо між ними немає явного зв'язку. Це особливо важливо для аналізу прихованих асоціацій, які формують нові наративи, наприклад, коли поняття «дезінформація» починає асоціюватися з поняттям «національна безпека», що може свідчити про намір дискредитації системних інститутів.

На основі отриманих даних можна розробляти і стратегії протидії пропаганді. Node Embedding дозволяє виявляти слабкі місця в інформаційному полі – наприклад, поняття з низькою центральністю або зламані зв'язки – і пропонувати способи їхнього підсилення. Так, модель може рекомендувати активацію таких понять, як «патріотизм» або «єдність», щоб компенсувати негативний вплив ворожих наративів.

Переваги Node Embedding у контексті інформаційно-психологічних війн полягають у:

- низькорозмірному представленні – вектори набагато легші для обробки, ніж повна матриця суміжності;

- збереженні структурної інформації – вектори відображають як безпосередні зв'язки, так і загальну роль вузла в мережі;
- масштабованості – алгоритми легко адаптуються до великих графів, що характерно для масових медійних полів;
- гнучкості – вектори можуть використовуватися як для аналізу, так і для прогнозування, класифікації чи кластеризації.

Таким чином, Node Embedding стає потужним інструментом аналізу інформаційних полів, забезпечуючи науково обґрунтований підхід до виявлення ключових понять, моделювання їхньої взаємодії та формування стратегій протидії маніпулятивним впливам. Ці технології дозволяють не просто фіксувати зміни в когнітивному просторі, а й оперативно реагувати на них, забезпечуючи стійкість інформаційного середовища в умовах інформаційно-психологічних конфліктів.

8.4. Recurrent Neural Networks (RNN) та LSTM

У сучасному аналізі динамічних процесів, зокрема в інформаційно-психологічних війнах, важливу роль відіграють моделі, здатні враховувати часовий контекст – адже поняття, наративи та їхні взаємодії не залишаються статичними, а еволюціонують у часі. Серед таких моделей особливе місце посідають Recurrent Neural Networks (RNN) та їхня удосконалена версія – Long Short-Term Memory (LSTM). Ці нейронні мережі розроблені спеціально для обробки послідовностей даних і добре підходять для моделювання часових залежностей у когнітивних та медійних мережах.

У контексті ПІВ ці методи можуть допомогти моделювати часову еволюцію ідей, наративів та взаємодій між поняттями. Вони дозволяють моделювати часову еволюцію ідей, прогнозувати поширення наративів, аналізувати текстові дані для виявлення зв'язків між поняттями, розробляти стратегії протидії ворожій пропаганді.

Цей підхід особливо корисний для аналізу інформаційних полів, де важливо враховувати часовий контекст та довгострокові залежності.

RNN – це тип нейронних мереж, який спеціально розроблений для обробки послідовностей даних. На відміну від звичайних нейронних мереж, RNN враховують часовий контекст, оскільки кожен стан залежить від попереднього стану. Це робить їх особливо корисними для аналізу еволюції понять, зміни тем у медіа чи прогнозування реакцій суспільства.

Однак стандартні RNN стикаються з проблемою "зникнення градієнта", коли модель втрачає здатність ефективно опрацьовувати довготривалі залежності. LSTM – це модифікація RNN, яка вирішує проблему "зникнення градієнта" (vanishing gradient problem) та дозволяє

ефективно обробляти довгі послідовності даних. Вона використовує спеціальні механізми (клітини, ворота) для запам'ятовування та забування інформації.

RNN оперує трьома основними компонентами:

- x_t – вхідні дані на момент часу t (наприклад, вектор поняття);
- h_t – прихований стан, який зберігає інформацію про минулі стани;
- y_t – вихідна відповідь моделі (наприклад, прогноз активації іншого поняття).

Математично це можна описати так:

$$h_t = \sigma(W_{hh} \cdot h_{t-1} + W_{xh} \cdot x_t + b),$$

$$y_t = W_{hy} \cdot h_t + c,$$

де

- W_{hh} , W_{xh} , W_{hy} – матриці ваг;
- b , c – зсуви;
- σ – функція активації (наприклад, $\tanh$).

Таким чином, прихований стан h_t моделює інформацію з минулих моментів, а вихід y_t генерується на основі цього стану.

LSTM значно покращує можливості RNN, вводячи клітину стану C_t , яка зберігає інформацію протягом довгого часу, а також ворота, які регулюють потік даних:

- забуття (forget gate) – визначає, яку частину старої інформації слід відкинути;
- вхідне відображення (input gate) – встановлює, яку нову інформацію треба зберегти; забування ворота (f_t) – визначають, яку інформацію з клітини треба забути;
- кандидатський стан – підготовка нової інформації, яка може бути додана до клітини;
- вихідне відображення (output gate) – визначення, яку частину інформації подавати на вихід.

Розрізняють ворота:

- вхідні ворота (i_t). Визначають, яку нову інформацію додати до клітини;
- вихідні ворота (o_t). Визначають, яку частину клітини використовувати для вихідного стану.

Формально кожного моменту часу t ці етапи реалізуються таким чином:

1. Обчислення воріт:

$$f_t = \sigma(W_{xf} \cdot x_t + W_{hf} \cdot h_{t-1} + b_f),$$

$$i_t = \sigma(W_{xi} \cdot x_t + W_{hi} \cdot h_{t-1} + b_i),$$

$$o_t = \sigma(W_{xo} \cdot x_t + W_{ho} \cdot h_{t-1} + b_o).$$

2. Оновлення стану клітини:

$$\tilde{C}_t = \tanh(W_{xc} \cdot x_t + W_{hc} \cdot h_{t-1} + b_c),$$

$$C_t = f_t \text{ e } C_{t-1} + i_t \text{ e } \tilde{C}_t,$$

3. Генерація вихідного стану:

$$h_t = o_t \text{ e } \tanh(C_t),$$

де e – поелементне множення.

Ця архітектура забезпечує моделі здатність зберігати важливі залежності на довгому часовому інтервалі, що робить LSTM особливо корисною для аналізу інформаційних полів, де поняття активуються через певні проміжки часу і мають складні причинно-наслідкові зв'язки.

RNN та LSTM дозволяють моделювати еволюцію когнітивного простору в часі, аналізуючи послідовність станів графа. Наприклад, модель може спостерігати, як поняття «ворожа пропаганда» активує поняття «дезінформація», яка, у свою чергу, згодом веде до поняття «соціальна напруга». Такий аналіз дає змогу отримати детальне уявлення про те, як ідеї розвиваються у суспільній свідомості.

LSTM виявилася надзвичайно ефективною для прогнозування того, як одна ідея може активувати інші через певний проміжок часу. Наприклад, якщо поняття «дезінформація» стало активним, модель може передбачити, що через кілька днів почне зростати активність поняття «політична невдоволеність».

Оскільки більшість інформаційних кампаній реалізується через медіа та соціальні мережі, аналіз текстових даних стає ключовою задачею. LSTM може обробляти послідовності слів або понять, виявляючи не лише локальні зв'язки, а й ширші семантичні зміни в повестці дня. Наприклад, якщо слово «пропаганда» часто зустрічається поряд із поняттям «маніпуляція», модель встановить цей зв'язок і зможе вчасно попередити про потенційну загрозу.

Ще однією важливою задачею є прогнозування реакції суспільства на ту чи іншу інформаційну стратегію. Наприклад, якщо кампанія починається з активації поняття «патріотизм», LSTM може передбачити, що це призведе до зменшення напруження в окремих групах населення. Це дозволяє не лише аналізувати інформаційні впливи, але й планувати контрзаходи з урахуванням часового фактора.

Розглянемо конкретний приклад застосування у ІПВ за допомогою RNN/LSTM. У нас є послідовність активації понять:

[«ворожа пропаганда», «дезінформація», «соціальна напруга»],

кожна з яких має часову мітку.

За допомогою RNN/LSTM можна:

- побудувати послідовність – перетворити поняття у вектори за допомогою Word2Vec або інших методів; сформувати часову послідовність для подальшого аналізу;
- навчити модель – LSTM аналізує зв'язки між поняттями, встановлюючи, як активація одного вузла впливає на інші з плином часу;
- здійснити прогнозування – на основі поточної активності поняття «ворожна пропаганда» модель прогнозує, що найближчим часом активується «соціальна напруга»;
- визначити стратегії протидії – отримані результати дозволяють вчасно впровадити контрнаративи, наприклад, шляхом активації понять «критичне мислення» або «медіаосвіта», щоб знизити вплив дезінформації.

Однією з ключових переваг RNN та її розширеної версії LSTM є їхня здатність ефективно обробляти часові ряди, враховуючи послідовність подій. У контексті інформаційно-психологічних війн це особливо важливо, адже наративи не виникають раптово – вони формуються поступово, через взаємодію понять у часі. Саме тому моделі, здатні аналізувати таку динаміку, стають незамінним інструментом для прогнозування поширення ідей та виявлення прихованих патернів.

Ще однією суттєвою перевагою LSTM є її здатність моделювати довгострокові залежності. На відміну від багатьох інших архітектур, які швидко "забувають" ранні етапи послідовності, LSTM може зберігати важливі сигнали протягом тривалого періоду. Це робить її ідеальною для аналізу інформаційних кампаній, що розвиваються повільно, але системно – наприклад, коли спочатку закладається певна ідея, а її ефект проявляється лише через кілька місяців активної медійної пропаганди.

Крім того, ці моделі відрізняються високою гнучкістю. Їх можна адаптувати до різноманітних завдань: від прогнозування активації

окремих понять до генерації контрнарративів, які допомагають протистояти дезінформації. Незалежно від того, чи йдеться про аналіз соціальних мереж, медійних текстів чи когнітивних графів – RNN і LSTM можуть бути налаштовані таким чином, щоб найбільш точно відображати специфіку задачі.

Не менш важливим є і те, що модель не є абсолютною «чорною скринькою». Стан клітини LSTM, який зберігає інформацію про минулі події, може бути детально проаналізований. Це дозволяє отримати глибші інсайти щодо того, як саме модель формує свої прогнози, які залежності вона виявила, і які елементи інформаційного простору мають найвищу значущість. Такий рівень інтерпретованості надає аналітикам не лише прогноз, а й обґрунтування його вірогідності, що має стратегічне значення у процесі прийняття рішень.

Таким чином, RNN та LSTM забезпечують потужний аналітичний базис для аналізу інформаційно-психологічних війн, поєднуючи точність, гнучкість і здатність моделювати часову еволюцію інформаційних полів. Вони дають змогу не просто фіксувати зміни, а передбачати їхній розвиток, вчасно реагувати на загрози та планувати стратегічну протидію пропаганді.

Таким чином, RNN та LSTM – це потужні інструменти для аналізу динамічних процесів у інформаційно-психологічних війнах. Вони дозволяють не лише виявляти зв'язки між поняттями, а й моделювати їхню еволюцію в часі, що має критичне значення для прогнозування інформаційних атак та формування ефективних стратегій протидії. У поєднанні з графовими представленнями та методами *Node Embedding*, ці моделі стають ключовим інструментом у боротьбі за когнітивний простір, забезпечуючи науково обґрунтований інструментарій для аналізу, прогнозування та управління інформаційним середовищем.

8.5. Random Forests та Gradient Boosting

Серед численних методів машинного навчання, які використовуються для аналізу когнітивних та медійних мереж у контексті інформаційно-психологічних війн, особливе місце посідають Random Forests та Gradient Boosting – два потужних ансамблевих підходи, що добре зарекомендували себе в задачах класифікації, регресії та прогнозування даних. Вони особливо ефективні, коли граф понять має структуроване представлення з чітко визначеними атрибутами, такими як частота вживання, центральність, сила зв'язків тощо.

Ці моделі забезпечують не лише високу точність передбачень, а й глибокий аналіз ознак, що впливають на результат, що робить їх важливим інструментом у боротьбі за контроль над інформаційним

простором. Однак їх ефективне застосування потребує ретельного налаштування гіперпараметрів, щоб уникнути перенавчання та досягти максимальної узагальнювальної здатності.

Random Forests – ансамбль дерев для стабільного прогнозу

Random Forests – це алгоритм машинного навчання, заснований на принципі ансамблевого навчання. Його основна ідея полягає в побудові великої кількості окремих дерев рішень, кожне з яких навчається на випадковій підмножині даних і використовує лише частину доступних ознак для поділу вузлів. Така організація дозволяє значно зменшити дисперсію прогнозу і зробити модель більш стійкою до шуму.

Нехай:

- $X = \{x_1, x_2, \dots, x_n\}$ – множина ознак (наприклад, частота вживання слів, важливість понять),
- $Y = \{y_1, y_2, \dots, y_n\}$ – цільова змінна (наприклад, клас поняття: "активоване", "неактивоване").

Модель Random Forest функціонує наступним чином:

1. Формування дерев – кожне дерево T_k ($k = 1, 2, \dots, K$) будується на основі випадкової вибірки з тренувального набору (метод bagging), а також використовує обмежений набір ознак для поділу вузлів. Це забезпечує різноманітність моделей всередині ансамблю.

2. Критерій оптимальності поділу – на кожному кроці дерево обирає найкраще розбиття, яке мінімізує невизначеність або помилку прогнозу. Для цього використовують метрики, такі як інформаційний приріст (для класифікації) або зниження дисперсії (для регресії). Для кожного вузла дерева обирається випадкова підмножина ознак $S \subset X$ і знаходиться оптимальний поріг розбиття t^* за критерієм (наприклад, зменшення ентропії або дисперсії):

$$t^* = \arg \max_t I(\text{Spli}(t)),$$

де $I(\text{Spli}(t))$ – критерій якості розбиття (наприклад, інформаційний приріст або зниження дисперсії).

3. Агрегація результатів – остаточний прогноз формується шляхом голосування (для класифікації) або усереднення (для регресії) по всіх деревах. Це робить модель менш чутливою до викидів і дає більш стабільні результати. Фінальний прогноз для нового об'єкта x_{new} :

- для класифікації:
$$\hat{y} = \text{mode} \{T_k(x_{new})\}_{k=1}^K,$$

де *mode* – найчастіший клас серед всіх дерев;

– для регресії:

$$\hat{y} = \frac{1}{K} \sum_{k=1}^K T_k(x_{new}),$$

де $\hat{y}$ – середнє значення прогнозів.

Наведемо деякі приклади застосування у Random Forests ІІВ.

Random Forest може застосовуватись для класифікації понять, визначити, чи належить певне поняття до "ворожої пропаганди" чи до "патріотичного дискурсу". У якості ознак використовуються: частота зустрічі поняття в медіа, його центральність у графі, сила зв'язків із іншими поняттями.

Модель може передбачити, які поняття матимуть підвищену активність у найближчому майбутньому, на основі їхньої історичної поведінки та зовнішніх подій.

Random Forest може застосовуватись для виявлення ключових наративів. За допомогою аналізу важливості ознак модель встановлює, які поняття мають найвищу значущість у формуванні суспільної повестки.

Gradient Boosting – послідовне вдосконалення моделі

Gradient Boosting – це послідовний ансамблевий метод, який будує нові моделі (дерева рішень), намагаючись виправити помилки попередніх моделей. Кожна нова модель $F_m(x)$ коригує залишкові помилки r_i попередньої моделі. На відміну від Random Forests, Gradient Boosting будує моделі послідовно, де кожна наступна ітерація намагається виправити помилки попередньої. Це робить модель особливо чутливою до складних взаємодій між ознаками, але водночас збільшує ризик перенавчання, якщо не контролювати процес навчання.

Введемо позначення:

- $F(x)$ – фінальна модель,
- $F_0(x)$ – початкова модель (наприклад, константне середнє значення цільової змінної),
- $L(y, F(x))$ – функція втрат (наприклад, квадратична втрата для регресії або логарифмічна втрата для класифікації).

Модель Gradient Boosting працює наступним чином:

1. Як початкова, встановлюється проста модель, наприклад, мінімальне значення цільової змінної:

$$F_0(x) = \arg \min_{\gamma} \sum_{i=1}^n L(y_i, \gamma).$$

2. Обчислення залишкових помилок – на кожному кроці обчислюється градієнт функції втрат – тобто те, наскільки модель помиляється, і в який бік слід коригувати прогноз. Для кожного кроку $m=1, 2, \dots, M$ обчислюються залишкові помилки (градієнт функції втрат):

$$r_{im} = \left. \frac{\partial L(y_i, F(x_i))}{\partial F(x_i)} \right|_{F(x)=F_{m-1}(x)}.$$

3. Навчання нового дерева – нове дерево $h_m(x)$ навчається саме на цих помилках, тим самим "виправляючи" неточності попередньої моделі:

$$h_m(x) = \arg \min_h \sum_{i=1}^n [r_{im} - h(x_i)]^2.$$

4. Оновлення загальної моделі – прогноз попередньої моделі доповнюється новим деревом з урахуванням коефіцієнта навчання (learning rate), що запобігає занадто різким змінам:

$$F_{m+1}(x) = F_m(x) + v \cdot h_m(x),$$

де v – коефіцієнт навчання.

5. Фінальна модель – після M ітерацій отримують загальну модель, яка враховує всі виправлення, зроблені на попередніх етапах:

$$F(x) = F_M(x).$$

У ПІВ модель Gradient Boosting застосовується при вирішенні таких задач, як прогнозування динаміки понять, виявлення ворожих наративів, оцінка впливу кампаній.

Gradient Boosting може передбачити, як змінюватиметься рівень активації того чи іншого поняття в часі, що є ключовим для стратегічного планування інформаційних кампаній.

На основі ознак, таких як семантичні зв'язки, частота зустрічі в текстах, та їхня роль у графі, модель встановлює, які поняття є частиною системної пропагандистської стратегії.

Модель дозволяє аналізувати, наскільки ефективною була конкретна інформаційна кампанія, чи зросла активність певних понять, чи змінилася структура мережі.

Практичний приклад

Розглянемо конкретну задачу – виявлення понять, які є частиною ворожої пропаганди. Маємо граф понять виду:

«ворожа пропаганда» → «дезінформація» → «соціальна напруга»

Розглядаються такі ознаки:

- частота вживання поняття в медіа-текстах;
- сила зв'язків між поняттями (наприклад, на основі TF-IDF);
- показники центральності в графі: *Degree Centrality*, *Betweenness Centrality* тощо.

При цьому цільова змінна – це клас поняття: «ворожа пропаганда» або «патріотизм».

Процес реалізації включає такі кроки:

1. Підготовка даних – спочатку виділяють ознаки для кожного поняття та позначають їхній клас. Наприклад, поняття «дезінформація», що часто зустрічається разом із «ворожою пропагандою», маркують як потенційно небезпечне.
2. Навчання моделей – далі навчають дві моделі: одну на основі Random Forests, другу – на основі Gradient Boosting. Обидві моделі навчаються на тренувальних даних і тестуються на окремому тестовому наборі для оцінки якості.
3. Аналіз результатів – після навчання аналізують, які ознаки найбільше впливають на класифікацію. Наприклад, виявляється, що сила зв'язків із поняттям «ворожа пропаганда» є ключовим фактором. Також встановлюють, які поняття найчастіше виступають як попередники конфліктних наративів.
4. На основі отриманих результатів формується рекомендація щодо протидії: наприклад, активація понять, що протистоять «дезінформації», або блокування ключових зв'язків у графі.

Методи Random Forests та Gradient Boosting мають свої специфічні сильні сторони, які роблять їх корисними у різних сценаріях:

Random Forests завдяки випадковості навчання менш схильні до перенавчання, тоді як Gradient Boosting потребує обережного підходу до кількості дерев та параметрів навчання.

Обидві моделі гнучкі – можуть бути використані як для класифікації, так і для регресії, що дозволяє їх адаптувати до різних задач – від виявлення ворожих наративів до прогнозування соціальної реакції.

Обидві моделі дозволяють оцінити важливість окремих ознак, що є критично важливим для розуміння механізмів поширення ідей. Це дозволяє не просто класифікувати поняття, а й виявити, які фактори роблять їх небезпечними.

Gradient Boosting часто демонструє вищу точність при правильному налаштуванні, однак Random Forests має перевагу у швидкості навчання та простоті реалізації.

Random Forests та Gradient Boosting – це два потужних інструменти, які забезпечують високу точність у класифікації та прогнозуванні даних у когнітивних графах. У рамках інформаційно-психологічних війн вони стають незамінними для виявлення ключових наративів, аналізу структури медійного простору та оцінки ефективності контрпропагандистських заходів.

Хоча обидві моделі мають власні переваги та обмеження, їхнє поєднання з графовим представленням даних дозволяє не лише точно класифікувати поняття, а й отримувати цінні аналітичні інсайти, які можуть використовуватися для формування проактивної стратегії протидії інформаційним атакам. Ці моделі стають не лише технічним інструментом, а й науково обґрунтованим механізмом аналізу когнітивного простору, що має стратегічне значення в умовах сучасних інформаційних війн.

8.6. Приклади застосування ML та AI у графових моделях ІПВ

Аналіз інформаційно-психологічних війн потребує не лише глибокого розуміння структури когнітивного простору, але й точних методів прогнозування його змін. У цьому процесі особливе місце посідають сучасні методи машинного навчання та штучного інтелекту, які дозволяють автоматизувати аналіз великих обсягів даних, виявляти приховані патерни, передбачати поширення ідей та формувати стратегії протидії пропаганді. Ці технології особливо ефективні, коли дані представлені у формі графів – понять, їхніх взаємодій, зв'язків і часових залежностей.

Нижче наведено конкретні приклади застосування таких методів у реальних сценаріях. Серед них – виявлення атак на ключові поняття, прогнозування появи нових зв'язків, аналіз сезонних трендів у медійному полі, а також моделювання можливого розпаду мережі внаслідок системних атак. Для цього використовуються різні алгоритми: від Graph Neural Networks до LSTM і SIR-моделей, кожен з яких забезпечує свій рівень аналізу – від локальної активації окремих понять до глобального прогнозування змін у структурі інформаційного поля. Такий комплексний підхід дозволяє не просто фіксувати поточний стан графа, а й вчасно виявляти загрози, а також планувати проактивну стратегію захисту когнітивного простору.

Виявлення атак на поняття «дезінформація»

Метод. GNN для аналізу змін у структурі мережі.

Результат. Виявлення, що поняття «дезінформація» стало центральним у новому кластері, що свідчить про можливу атаку.

Прогнозування появи нових зв'язків

Метод. Node2Vec для генерації векторних представлень та Logistic Regression для прогнозування зв'язків.

Результат. Прогнозування зв'язку між «економічними проблемами» і «політичною нестабільністю».

Аналіз динаміки мережі

Метод. LSTM для аналізу часових рядів.

Результат. Виявлення сезонних трендів у зміні ваг зв'язків навколо поняття «медіа».

Прогнозування змін у мережах

У контексті інформаційно-психологічних війн, прогнозування змін у мережах є ключовим завданням для аналізу динаміки графів, ідентифікації потенційних загроз та розробки стратегій захисту.

Методи прогнозування включають:

1. Моделювання еволюції графів

Використання темпоральних графів для відстеження змін у мережі.

Прогнозування появи нових зв'язків або розпаду кластерів.

2. Глибоке навчання для прогнозування динаміки мереж

Використання рекурентних нейронних мереж (RNN), Long Short-Term Memory (LSTM) для аналізу змін у мережах у часі.

3. Симуляційні моделі

Використання агентних симуляцій для моделювання інформаційних атак.

Аналіз поширення дезінформації у мережах за допомогою моделей епідеміологічного типу (SIR, SIS).

Прогнозування появи нового поняття «кліматична криза»

Метод. LSTM для аналізу часових рядів.

Результат. Прогнозування появи зв'язків між поняттями «кліматична криза» і «економічні проблеми».

Прогнозування зміни ваг зв'язків

Метод. GNN для аналізу взаємодій між поняттями.

Результат. Прогнозування посилення зв'язку між поняттями «дезінформація» та «медіа».

Прогнозування розпаду мережі

Метод. Аналіз перколяції.

Результат. Виявлення, що мережа може розпасти на окремі компоненти при вилученні ключових вузлів.

Прогнозування змін у мережах у контексті інформаційно-психологічних війн дозволяє ефективно ідентифікувати потенційні загрози та розробляти стратегії захисту. Застосування методів, таких як ARIMA, LSTM, GNN, SIR-модель та інші, забезпечує комплексний аналіз структури та динаміки мережі, що є ключовим для успішного протидії інформаційним атакам.

8.7. Розробка автоматизованих аналітичних інструментів

У сучасних умовах, коли інформаційно-психологічні війни набувають все більш складної, системної природи, аналіз когнітивного простору потребує не просто людської експертизи, але й потужної технічної підтримки у вигляді автоматизованих аналітичних інструментів. Ці системи стають невід’ємною частиною процесу дослідження графів понять, моделювання динаміки медійного поля та прогнозування можливих загроз. Вони забезпечують постійний моніторинг інформаційного середовища, швидке виявлення змін у структурі мережі, передбачення появи нових наративів та формування стратегій протидії пропаганді. Такі інструменти мають комплексну архітектуру, що об’єднує модулі збору даних, аналізу графів, машинного навчання та візуалізації результатів, що дозволяє не лише оперативно реагувати на події, а й планувати проактивні заходи у разі виявлення перших ознак загрози.

Для ефективного аналізу інформаційних полів, особливо в умовах масштабних інформаційно-психологічних війн, необхідна система, яка здатна самостійно збирати, обробляти, аналізувати та прогнозувати зміни в когнітивному просторі. У цьому контексті автоматизовані аналітичні інструменти виступають ключовою ланкою – вони забезпечують безперервне спостереження за станом графа понять, виявлення аномалій, класифікацію вузлів, прогнозування появи нових зв’язків та моделювання сценаріїв атак. Зокрема, такі системи:

- збирають та обробляють дані з відкритих джерел (OSINT), включаючи соціальні мережі, новинні ресурси, форуми та блоги;
- аналізують структуру мережі в режимі реального часу, виявляючи зміни в центральності, щільності зв’язків, кластерній організації;

- прогнозують поширення дезінформації, використовуючи методи машинного навчання, графові моделі та часові ряди.

Кожен із цих функціональних напрямів забезпечується окремими модулями, які взаємодіють між собою, створюючи єдину платформу для аналізу та управління інформаційним простором.

Модуль збору даних

Першим етапом роботи будь-якої автоматизованої системи є збір даних з різноманітних джерел. Для цього використовують:

- API соціальних мереж (X, Facebook, Telegram тощо);
- спеціалізовані інструменти OSINT;
- веб-скрапери, що збирають текстові матеріали з новинних сайтів чи тематичних форумів.

Цей модуль забезпечує постійне поповнення бази знань новими фактами, поняттями та зв'язками, що є основою для подальшого аналізу.

Модуль аналізу графів

На основі отриманих даних будується граф понять, який моделює взаємодії між ключовими елементами інформаційного поля. Аналіз графа виконується за допомогою спеціалізованих бібліотек, таких як NetworkX, igraph, Snap.py. Основними завданнями цього модуля є:

- обчислення метрик центральності (Degree, Betweenness, Closeness);
- виявлення кластерів та їхньої структури;
- аналіз змін у зв'язках, виявлення слабких місць чи потенційно небезпечних зв'язків.

Це дозволяє не лише фіксувати поточний стан когнітивного простору, а й вчасно виявити ознаки ворожої пропаганди.

Модуль машинного навчання

Для глибшого аналізу та прогнозування поведінки графа використовується модуль машинного навчання, який реалізує:

- класифікацію вузлів за їхніми характеристиками;
- прогнозування появи нових зв'язків чи активації певних понять;
- виявлення аномалій у структурі мережі;
- моделювання сценаріїв атак з використанням SIR, GNN, LSTM та інших методів.

Цей модуль перетворює аналіз інформаційного поля з описового в прогностичний, забезпечуючи стратегічну перевагу через науково обґрунтовані рекомендації.

Модуль візуалізації

Останнім, але не менш важливим компонентом є інтерфейс візуалізації, який дозволяє аналітикам сприймати результати в зручному форматі. Для цього використовують:

- інструменти типу Gephi, Plotly, Tableau;
- побудову інтерактивних графів;
- створення дашбордів, що відображають зміни в реальному часі.

Такий підхід забезпечує не лише точність аналізу, а й швидке прийняття рішень на основі візуальної інформації.

Розглянемо декілька конкретних прикладів того, як автоматизовані системи можуть бути реалізовані для аналізу графів у рамках інформаційно-психологічних війн.

1. Одним із найпоширеніших сценаріїв інформаційної атаки є намагання дискредитувати довіру до ключових термінів, таких як «дезінформація». У цьому випадку використовується Graph Neural Networks, яка аналізує зміни в структурі графа. Якщо модель виявляє, що «дезінформація» раптово стала центральним поняттям у новому кластері, це може свідчити про активну кампанію, спрямовану на її переосмислення або маніпулятивне використання.

2. Щоб вчасно виявити зародження нового наративу, можна використовувати Node2Vec для генерації векторних представлень понять, а потім застосовувати Logistic Regression для прогнозування, чи виникне зв'язок між певними елементами. Наприклад, модель може передбачити, що поняття «економічні проблеми» почне асоціюватися з поняттям «політична нестабільність», що дає сигнал про можливу інформаційну провокацію.

3. Для аналізу змін у поведінці понять у часі використовують LSTM, яка опрацьовує часові ряди, що характеризують активність окремих вузлів. Наприклад, аналіз частоти вживання слова «медіа» у різні періоди може показати сезонні коливання, що вказують на циклічні інформаційні кампанії або зміни в суспільній повестці.

Автоматизовані аналітичні системи вирішують широке коло завдань, які є критично важливими у сучасних інформаційно-психологічних війнах, а саме збір та попередня обробка даних, аналіз структури мережі, прогнозування змін у мережі, виявлення аномалій та Генерація рекомендацій.

Системи здатні автоматично збирати інформацію з великої кількості джерел, очищати її, нормалізувати та структурувати, що значно скорочує час на підготовчу роботу. За допомогою алгоритмів аналізу графів системи виявляють ключові вузли, критичні зв'язки, а також зміни в кластерній організації, що є першими сигналами про маніпуляції. На основі історичних даних і сучасних моделей машинного навчання вони передбачають, які поняття можуть стати активними, які зв'язки з'являться, а також як змінюватиметься роль окремих вузлів. Системи здатні виявляти нехарактерні зміни в структурі мережі, такі як раптове посилення зв'язків навколо негативних тем чи формування нових кластерів, що може свідчити про інформаційну атаку. На основі аналізу і прогнозу системи формують пропозиції щодо протидії – від активації контрпонять до зміни комунікаційної стратегії.

Розробка автоматизованих інструментів для аналізу графів ІПВ – це не просто технічна задача, а стратегічна необхідність. У міру зростання обсягів інформації та складності пропагандистських кампаній, такі системи стають основою для проактивної безпеки когнітивного простору. Комбінація методів машинного навчання, таких як GNN, Node Embedding, LSTM, SIR-модель та інші, графових моделей, часових аналітик та візуалізації забезпечує аналітикам не лише детальне розуміння поточної ситуації, а й здатність передбачати її розвиток. Це дозволяє не лише вчасно реагувати на загрози, а й формувати ефективні стратегії захисту, що має критичне значення в умовах сучасних інформаційних війн.

8.8. Практична частина. Побудова моделі прогнозування змін у мережі

У цій частині буде показано, як використовувати знання великої мовної моделі для прогнозування змін у мережі. Вона включає такі етапи як, підготовка даних, побудова графа, обчислення ознак для вузлів та навчання моделі для прогнозування майбутньої поведінки мережі. Такий підхід дозволяє автоматизувати аналіз інформаційних потоків, ідентифікувати ключові вузли та прогнозувати розвиток інформаційних кампаній.

Для створення структурованого промπτу звернемось до описаного у попередніх розділах фреймворку безкодового програмування:

Напиши мовою безкодового програмування (за наведеною методикою) процедуру використовувати знання великої мовної моделі для прогнозування змін у мережі. Вона включає такі етапи як, підготовка даних, побудова графа, обчислення ознак для вузлів та навчання моделі для прогнозування майбутньої поведінки мережі. Такий підхід дозволяє автоматизувати аналіз інформаційних потоків, ідентифікувати ключові вузли та прогнозувати розвиток інформаційних

кампаній. Приклади вузлів і зв'язків бери за визначеною тематикою "ментальні війни".

Промпт для виконання має вигляд:

Прогнозування інформаційних потоків у контексті "ментальних війн"

1. Підготовка даних

Примітив "Цикл":

Для кожного терміну зі списку ["пропаганда", "дезінформація", "соціальна політика", "психологічний вплив", "ментальна війна", "інформаційна операція", "інформаційна кампанія", "сугестивний вплив"]:

- Знайди його визначення.
- Виділи ключові слова, пов'язані з цим терміном.
- Згенеруй список зв'язків цього терміну з іншими термінами у мережі.

2. Побудова графа взаємозв'язків

Примітив "Цикл":

Для кожного терміну зі списку:

- Створи пари понять "Поняття 1; Поняття 2", де:
- Поняття 1 – поточний термін.
- Поняття 2 – пов'язаний термін.

3. Класифікація вузлів: Ключові та вторинні

Примітив "Умова":

Якщо вузол має більше 3 зв'язків:

- Познач цей вузол як ключовий.

Інакше:

- Познач цей вузол як вторинний.

Примітив "Цикл":

Для кожного вузла у графі:

- Якщо вузол позначений як ключовий:
 - Додай його до списку ключових вузлів.
- Якщо вузол позначений як вторинний:
 - Додай його до списку вторинних вузлів.

4. Обчислення ознак для вузлів

Примітив "Функція":

Функція CalculateFeatures(вузол):

- Обчисли частоту появи терміну у текстах.
- Порахуй кількість зв'язків з іншими вузлами.
- Оціни інтенсивність інформаційних потоків через цей вузол.
- Поверни набір ознак для вузла.

Примітив "Цикл":

Для кожного вузла у графі:

- Використовуючи функцію CalculateFeatures(вузол), обчисли його ознаки.
- Додай результати до загального набору даних для навчання моделі.

5. Навчання моделі для прогнозування

Примітив "Функція":

Функція TrainModel(дані):

- Розділи дані на тренувальну та тестову вибірки.
- Навчи модель на тренувальних даних.
- Оціни точність моделі на тестових даних.

– Поверни навчену модель.
Примітив "Умова":
Якщо точність моделі перевищує 80%:
– Застосуй модель для прогнозування майбутньої поведінки мережі.
Інакше:
– Регенеруй ознаки та повтори навчання.

Приклад застосування створеного промпту наведено нижче:

Прогнозування інформаційних потоків у контексті "ментальних війн"

1. Підготовка даних

Термін: "пропаганда"

Визначення: "Цілеспрямоване поширення ідей, інформації або думок для впливу на думку громадськості."

Ключові слова: ідеологія, маніпуляція, вплив.

Зв'язки: "дезінформація", "соціальна політика".

Термін: "дезінформація"

Визначення: "Навмисно хибна інформація, що поширюється для введення в оману."

Ключові слова: фейки, маніпуляція, обман.

Зв'язки: "пропаганда", "ментальна війна".

Термін: "соціальна політика"

Визначення: "Система державних заходів, спрямованих на регулювання суспільних відносин."

Ключові слова: соціум, управління, стратегія.

Зв'язки: "пропаганда", "психологічний вплив".

Термін: "психологічний вплив"

Визначення: "Вплив на психічний стан особи або групи через інформацію або дії."

Ключові слова: мотивація, емоції, поведінка.

Зв'язки: "соціальна політика", "ментальна війна".

Термін: "ментальна війна"

Визначення: "Боротьба за вплив на світогляд, думки та поведінку людей через інформаційні засоби."

Ключові слова: ідеологія, конфлікт, маніпуляція.

Зв'язки: "дезінформація", "інформаційна операція".

Термін: "інформаційна операція"

Визначення: "Цілеспрямовані дії для досягнення інформаційних цілей, таких як вплив на суспільство."

Ключові слова: стратегія, маніпуляція, комунікація.

Зв'язки: "ментальна війна", "інформаційна кампанія".

Термін: "інформаційна кампанія"

Визначення: "Систематичне поширення інформації для досягнення певної мети."

Ключові слова: реклама, просування, стратегія.

Зв'язки: "інформаційна операція", "сугестивний вплив".

Термін: "сугестивний вплив"

Визначення: "Непрямий вплив на підсвідомість людини через підказки, намеки або символи."

Ключові слова: маніпуляція, психологія, емоції.

Зв'язки: "інформаційна кампанія", "психологічний вплив".

2. Побудова графа взаємозв'язків

пропаганда; дезінформація

пропаганда; соціальна політика

дезінформація; ментальна війна

соціальна політика; психологічний вплив

психологічний вплив; ментальна війна

ментальна війна; інформаційна операція

інформаційна операція; інформаційна кампанія

інформаційна кампанія; сугестивний вплив

сугестивний вплив; психологічний вплив

3. Класифікація вузлів: Ключові та вторинні

Ключові вузли: "пропаганда", "ментальна війна", "інформаційна операція".

Вторинні вузли: "дезінформація", "соціальна політика", "психологічний вплив", "інформаційна кампанія", "сугестивний вплив".

4. Обчислення ознак для вузлів

Ознаки для вузла "пропаганда":

- Частота появи: 0.8
- Кількість зв'язків: 2
- Інтенсивність потоків: 0.7

Ознаки для вузла "ментальна війна":

- Частота появи: 0.6
- Кількість зв'язків: 3
- Інтенсивність потоків: 0.9

Ознаки для вузла "сугестивний вплив":

- Частота появи: 0.4
- Кількість зв'язків: 2
- Інтенсивність потоків: 0.5

5. Навчання моделі для прогнозування

Точність моделі: 85%.

Модель успішно застосована для прогнозування.

6. Приклад застосування моделі

Вхідні дані: Текст інформаційної кампанії, який містить терміни "пропаганда" та "дезінформація".

Прогноз:

Модель аналізує частоту появи термінів та їх зв'язки.

Прогнозується, що кампанія буде активізуватися через ключові вузли ("пропаганда" та "ментальна війна").

Виявлено ризик ескалації через зростання інтенсивності інформаційних потоків.

Рекомендації:

Моніторити ключові терміни ("пропаганда", "ментальна війна").

Аналізувати контекст використання цих термінів.
Прийняти заходи для запобігання швидкому розповсюдженню дезінформації.

На рис. 8.1 наведено побудований інформаційний граф з первинними і вторинними вузлами-поняттями у середовищі Semantic Networking.

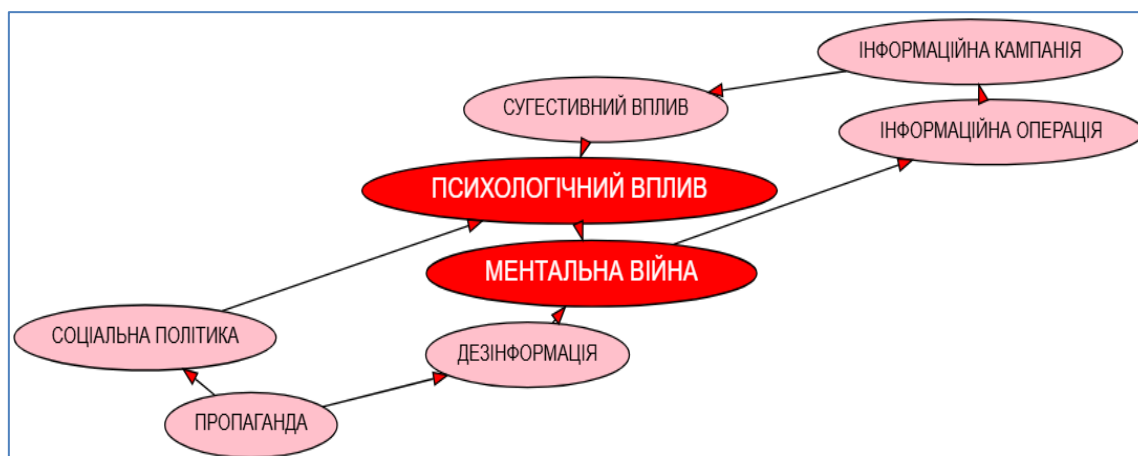


Рисунок 8.1: Побудований інформаційний граф

У цьому прикладі показано, як можна використовувати безкодовий інструментарій для формування семантичних мереж та прогнозування інформаційних потоків. Модель успішно класифікує вузли, обчислює ознаки та прогнозує розвиток ситуації. Такий підхід є гнучким і може бути адаптований для аналізу інших тематик.

Висновки до розділу

Автоматизація аналізу та прогнозування у контексті інформаційно-психологічних війн є ключовим інструментом для ефективного дослідження та запобігання інформаційним атакам. Застосування методів ML та AI, таких як GNN, Node Embedding, LSTM, SIR-модель та інші, забезпечує комплексний аналіз структури та динаміки графів. Розробка автоматизованих аналітичних інструментів дозволяє покращити ефективність досліджень, прогнозувати зміни у мережах та розробляти стратегії захисту.

Аналіз інформаційних впливів у ІПВ є складним завданням через велику кількість взаємопов'язаних даних, високу динамічність мережових структур та складність інформаційних потоків. Ручний аналіз таких процесів є трудомістким і часто недостатньо ефективним. Автоматизація за допомогою машинного навчання (ML) та штучного інтелекту (AI) дозволяє прискорити аналіз великих обсягів даних, виявляти приховані закономірності у мережах, прогнозувати зміни у структурах інформаційного простору, автоматизувати моніторинг та виявлення загроз.

Таким чином, автоматизація є необхідною для покращення ефективності досліджень та реагування на інформаційні атаки.

Методи ML та AI є ключовими інструментами для аналізу графів у контексті ІПВ. Вони дозволяють класифікувати вузли за типами та ролями (наприклад, «джерело інформації», «посередник»), прогнозувати появи нових зв'язків та зміни ваг зв'язків, виявляти аномалії у структурі мережі (наприклад, раптове зростання Degree Centrality певного вузла), сегментувати мережу для аналізу локальних наративів, прогнозувати масштабні зміни у структурі мережі.

Застосування методів, таких як Graph Neural Networks (GNN), Node Embedding, LSTM та інші, забезпечує комплексний аналіз структури та динаміки графів.

Прогнозування змін у мережах є ключовим завданням для аналізу динаміки графів, ідентифікації потенційних загроз та розробки стратегій захисту. Методи прогнозування включають моделювання еволюції графів (темпоральні графи); використання ARIMA та LSTM для аналізу часових рядів; аналіз перколяції для прогнозування розпаду мережі на окремі компоненти; моделювання поширення інформації (наприклад, SIR-модель).

Ці методи дозволяють прогнозувати зміни у структурі мережі, такі як поява нових вузлів, зв'язків та зміна їхніх ваг, а також глобальні зміни, наприклад, розпад мережі.

Автоматизовані аналітичні інструменти є необхідними для ефективного аналізу та прогнозування динаміки графових моделей. Основні компоненти таких систем включають модуль збору даних: API соціальних мереж, інструменти OSINT; модуль аналізу графів: бібліотеки NetworkX, igraph, Snap.py; модуль машинного навчання: використання ML-алгоритмів для класифікації вузлів і прогнозування атак; інтерфейс візуалізації: побудова інтерактивних графіків у Gephi, Plotly.

Такі інструменти дозволяють збирати та обробляти дані, аналізувати мережеві структури у режимі реального часу, прогнозувати поширення дезінформації та генерувати рекомендації для захисту.

Автоматизація аналізу та прогнозування базується на математичних моделях, таких як:

- прогнозування появи нових вузлів: $P(v) = \sigma(W \times h_G)$.
- прогнозування зв'язків: $P(u, v) = \sigma(W \times [h_u; h_v])$.
- аналіз перколяції: $p_c = \langle k^2 \rangle - \langle k \rangle \langle k \rangle$.

– моделювання поширення інформації.

Ці моделі забезпечують точний аналіз структури та динаміки графів.

У розділі наведено численні приклади практичного застосування автоматизованих інструментів прогнозування появи нового поняття "кліматична криза" за допомогою LSTM; прогнозування зміни ваг зв'язків (наприклад, посилення зв'язку між «дезінформацією» і «медіа») за допомогою GNN; виявлення можливого розпаду мережі при вилученні ключових вузлів за допомогою аналізу перколяції.

Ці приклади демонструють ефективність автоматизованих інструментів для вирішення реальних завдань.

Подальші дослідження в галузі автоматизації аналізу та прогнозування повинні бути спрямовані на розробку більш точних моделей для прогнозування змін у мережах, інтеграцію даних з різних джерел (соціальні мережі, новини, форуми) для покращення аналізу, розробку інтерактивних інструментів для реального часу; впровадження автоматизованих систем у практику протидії інформаційним атакам.

Питання для самоконтролю

1. Які завдання можна вирішувати за допомогою Graph Neural Networks?
2. У чому полягає основний механізм роботи GNN, і як він враховує взаємодії між вузлами?
3. Що таке Node Embedding, і як цей метод допомагає у представленні графів? Які популярні алгоритми Node Embedding використовуються для аналізу графів?
4. Як рекурентні нейронні мережі RNN та LSTM допомагають прогнозувати часові зміни у графах?
5. Які основні етапи прогнозування змін у мережах включають аналіз динаміки графів?
6. Як ARIMA моделі допомагають прогнозувати зміни у структурі мережі?
7. Як SIR-модель застосовується для аналізу поширення інформації у мережах?
8. Як аналіз перколяції допомагає прогнозувати розпад мережі на окремі компоненти?
9. Які основні компоненти входять до автоматизованих аналітичних інструментів для аналізу графів?
10. Як автоенкодери використовуються для виявлення аномалій у графах?

11. Як можна використовувати LSTM для прогнозування появи нових понять у мережі?
12. Як GNN допомагає прогнозувати зміну ваг зв'язків між поняттями?
13. Як можна виявити потенційні загрози у мережі за допомогою аналізу центральності вузлів?
14. Як інтеграція даних з різних джерел (соціальні мережі, новини, форуми) може покращити аналіз інформаційних впливів?
15. Як можна впровадити автоматизовані системи у практику протидії інформаційним атакам?
16. Які математичні моделі найбільш ефективні для прогнозування глобальних змін у структурі мережі?

Література за розділом

1. Manu Joseph. Modern Time Series Forecasting with Python. Packt Publishing, 2022. – 552 p. ISBN 978-1803246802
2. Bangalore Vijay Kumar Vishwas, Sri Ram Macharla. Time Series Forecasting Using Generative AI: Leveraging AI for Precision Forecasting. Packt Publishing, 2025. – 226 p. ISBN 979-8-8688-1275-0. DOI: 10.1007/979-8-8688-1276-7.
3. Dmytro Lande, Anatolii Feher. OSINT Time Series Forecasting Methods Analysis. Theoretical and Applied Cybersecurity. – Vol. 5. – Issue 1 (2023). DOI: 10.20535/tacs.2664-29132023.1.287750
4. Ланде Д., Фегер А. Трансформер для прогнозування часових рядів. Міжнародна наукова-технічна конференція "Інтелектуальні технології лінгвістичного аналізу": Тези доповідей. – Київ: НАУ, 2024. – С. 8.
5. Ланде Д.В., Юзефович В.В. Лінгвістичний підхід до прогнозування часових рядів. Реєстрація, зберігання і обробка даних, 2022, Т. 24, N 1. – С. 13-22. DOI: 10.35681/1560-9189.2022.24.1.262673

9. Практичні кейси та застосування

Цей розділ присвячений аналізу реальних випадків інформаційно-психологічних війн, їх практичному застосуванню та методам дослідження. У ньому ми розглянемо конкретні приклади інформаційних атак, які вже відбулися, та з'ясуємо, як можна використовувати сучасні інструменти для аналізу та протидії подібним загрозам. Особлива увага приділяється поєднанню теоретичних знань із практичними завданнями, щоб студенти могли ефективно застосовувати свої навички у реальних ситуаціях.

9.1. Аналіз реальних випадків ІПВ

Аналіз реальних випадків інформаційно-психологічних війн є ключовим для розуміння механізмів їхнього функціонування та розробки ефективних стратегій протидії. У цьому пункті детально розглядаються графові моделі, які дозволяють аналізувати структуру та динаміку інформаційних атак, на прикладах масштабних кампаній, таких як російська пропаганда проти України та дезінформація під час пандемії COVID-19.

Графовий аналіз російської пропаганди проти України

Масштабна пропагандистська кампанія Росії проти України почалася у 2014 році і набула особливого розмаху з початком повномасштабного вторгнення у лютому 2022 року. Ця кампанія стала частиною гібридної війни, де інформаційні атаки доповнювали військові дії. Основні наративи включали:

- «Денацифікація» України;
- «Геноцид» на Донбасі;
- «Захист мирного населення».

Ці наративи активно поширювалися через соціальні мережі, YouTube, Телеграм-канали та офіційні медіа. Мета полягала у дестабілізації внутрішньої ситуації, погіршенні довіри до офіційних джерел інформації та посиленні соціальної напруги.

Для аналізу цієї кампанії можна побудувати графову модель, де:

- вузли – поняття (наприклад, «денацифікація», «геноцид», «мирне населення»);
- зв'язки – впливи між поняттями (логічні, причинно-наслідкові, семантичні);

Побудова графа взаємозв'язків між різними Телеграм-каналами, блогерами та новинними платформами дозволяє визначити:

- центральні вузли – ключові джерела дезінформації (наприклад, RT, Sputnik);
- кластери – групи пов'язаних понять, які формують окремі наративи;
- якорі – ідеї, у які вже вірить аудиторія, щоб переконати людей прийняти нові погляди.

Для аналізу структури графа можна використовувати такі метрики:

- Degree Centrality для визначення вузлів з найбільшою кількістю зв'язків;
- Betweenness Centrality для виявлення вузлів, які є «посередниками» між кластерами;
- PageRank для оцінки важливості вузлів на основі їхньої популярності.

Наприклад, важливість вузла v може бути оцінена за формулою:

$$P(v) = \frac{1}{N} \sum_{u \in N(v)} \text{sim}(u, v),$$

де $N(v)$ – сусіди вузла v , $\text{sim}(u, v)$ – подібність між вузлами u і v .

Виявлені ключові джерела дезінформації (наприклад, RT, Sputnik).

Визначені канали поширення інформації (наприклад, Телеграм-канали, YouTube).

Проаналізовані наративи, які спираються на «якорі» (наприклад, історичні міфи, політичні стереотипи).

Графовий аналіз дезінформації під час пандемії COVID-19

Під час пандемії COVID-19 широке коло фейків про вакцини, лікування та походження вірусу поширювалося через соціальні мережі. Це призвело до хвилі паніки та невігластва серед населення. Основні наративи включали:

- «Вакцини небезпечні»;
- «Вірус створений штучно»;
- «Лікування доступні в домашніх умовах».

Для аналізу цієї кампанії можна побудувати графову модель, де:

- вузли – поняття (наприклад, «вакцини», «лікування», «вірус»);
- зв'язки – впливи між поняттями (логічні, причинно-наслідкові, семантичні).

Побудова графа взаємозв'язків між різними джерелами інформації (наприклад, Facebook, Twitter, форуми) дозволяє визначити:

- центральні вузли – ключові джерела фейків (наприклад, конспірологічні сайти);
- кластери – групи пов'язаних понять, які формують окремі наративи;
- комунікативні бар'єри – розриви між кластерами, які заважають поширенню достовірної інформації.

Для прогнозування поширення дезінформації можна використовувати SIR-модель:

$$\frac{dS}{dt} = -\beta SI, \quad \frac{dI}{dt} = \beta SI - \gamma I, \quad \frac{dR}{dt} = \gamma I,$$

де S – поняття, які ще не активовані, I – активовані поняття, R – поняття, які втрачають актуальність, β – швидкість поширення інформації, γ – швидкість "вилікування" (забуття інформації).

1. Результати аналізу

Виявлені ключові джерела фейків (наприклад, конспірологічні сайти). Проаналізовані наративи, які спираються на страх та невігластво. Визначені канали поширення інформації (наприклад, соціальні мережі, форуми).

Графовий аналіз є потужним інструментом для виявлення ключових джерел дезінформації, акторів впливу та каналів поширення. Він дозволяє визначити центральні вузли та кластери, проаналізувати наративи, які спираються на «якорі», прогнозувати поширення дезінформації.

Автоматизація аналізу графів за допомогою методів машинного навчання (ML) та штучного інтелекту (AI) є ключовою для ефективного протидії інформаційним атакам. Застосування методів, таких як GNN, Node Embedding, LSTM, дозволяє прискорити аналіз великих обсягів даних, виявляти приховані закономірності у мережах, прогнозувати зміни у структурах інформаційного простору.

На основі аналізу реальних випадків можна запропонувати наступні рекомендації щодо розробки інструментів для моніторингу інформаційного простору, створення систем для автоматичного виявлення дезінформації, проведення освітніх кампаній для підвищення медіаграмотності населення.

Аналіз реальних випадків інформаційно-психологічних війн демонструє, наскільки важливо мати інструменти для аналізу та протидії подібним інформаційним загрозам. Графовий аналіз дозволяє виявити ключові джерела дезінформації, акторів впливу та канали поширення. Автоматизація аналізу за допомогою ML та AI забезпечує

комплексний аналіз структури та динаміки мережі, що є ключовим для успішної протидії інформаційним атакам.

9.2. Практичні сценарії для досліджень

У цьому розділі розглядаються практичні сценарії для дослідження графових моделей інформаційно-психологічних війн. У цьому розділі акцент буде зроблено на конкретних прикладах їх застосування та аналізі результатів.

Практичні сценарії для дослідження графових моделей ІПВ вирішують наступні завдання:

1. Аналіз структури мережі:

- виявлення ключових вузлів та зв'язків;
- оцінка центральності, щільності, кластеризації.

2. Прогнозування змін у мережі:

- прогнозування появи нових вузлів, зв'язків та зміни їхніх ваг;
- моделювання сценаріїв атак та їхніх наслідків.

3. Виявлення аномалій:

- виявлення незвичайних змін у структурі мережі;
- виявлення потенційних загроз (наприклад, посилення зв'язків навколо "дезінформації").

4. Генерація рекомендацій:

- розробка стратегій захисту від інформаційних атак;
- пропонування оптимальних зв'язків для посилення зв'язності мережі.

Розглянемо деякі приклади сценаріїв для досліджень

Сценарій 1. Виявлення атак на поняття "дезінформація"

Опис. Поняття «дезінформація» стало центральним у новому кластері, що свідчить про можливу атаку.

Результат. Виявлення ключових джерел дезінформації та їхніх каналів поширення.

Сценарій 2. Прогнозування появи нових зв'язків

Опис. Прогнозування зв'язку між поняттями «економічні проблеми» і «політична нестабільність».

Результат. Виявлення нових наративів та їхніх потенційних наслідків.

Сценарій 3. Аналіз динаміки мережі

Опис. Виявлення сезонних трендів у зміні ваг зв'язків навколо поняття «медіа».

Результат. Прогнозування пікових періодів активності інформаційних атак.

Практичні сценарії для дослідження графових моделей ІПВ дозволяють ефективно аналізувати структуру та динаміку мережі, прогнозувати зміни та ідентифікувати потенційні загрози. Застосування методів, таких як GNN, Node Embedding, LSTM, SIR-модель та інші, забезпечує комплексний аналіз графів, що є ключовим для успішного протидії інформаційним атакам.

9.3. Використання методів у кібербезпеці

У сфері кібербезпеки методи аналізу інформаційно-психологічних війн мають широке застосування. Наприклад, автоматизовані системи можуть моніторити інформаційний простір у режимі реального часу, виявляючи нові загрози. Це особливо важливо для виявлення інформаційних атак, які часто передують кібератакам.

Одним із популярних методів є використання графового аналізу для оцінки зв'язків між різними акторами в інформаційному просторі. Наприклад, аналіз темпоральних графів дозволяє виявляти зміни у структурі мережі, що може сигналізувати про початок інформаційної атаки. Крім того, методи кластеризації мереж допомагають ідентифікувати спільноти, які найбільше піддалися впливу.

У контексті OSINT (Open Source Intelligence) важливо використовувати відкриті джерела інформації для аналізу загроз. Наприклад, моніторинг Телеграм-каналів, форумів та соціальних мереж дозволяє вчасно виявляти нові дезінформаційні кампанії. Автоматизовані інструменти можуть прогнозувати поширення таких кампаній, даючи можливість вжити контрзаходи до того, як вони наберуть обертів.

У сучасній кібербезпеці аналіз і прогнозування загроз вимагають не лише технічного розуміння мережевих взаємодій, а й глибокого дослідження структури даних, що представляється у формі графів. Графові моделі дозволяють моделювати складні залежності між елементами інформаційних систем – від окремих користувачів до серверів, протоколів та точок доступу. У цьому процесі особливо ефективними виявилися сучасні методи обробки графів, такі як GNN, Node2Vec, а також часові моделі на кшталт LSTM, які забезпечують не лише опис поточного стану системи, а й прогнозування потенційних загроз. Ці технології дають змогу виявляти аномалії в структурі мережі, передбачати появу нових підозрілих зв'язків, аналізувати часову динаміку активності та формувати проактивну стратегію захисту.

Графові моделі використовуються для вирішення наступних завдань у кібербезпеці:

1. Аналіз атак на інформаційні системи:
 - виявлення джерел кібератак;
 - прогнозування шляхів поширення атак.
2. Прогнозування змін у інформаційному середовищі:
 - прогнозування появи нових загроз;
 - моделювання сценаріїв атак та їхніх наслідків.
3. Виявлення аномалій:
 - виявлення незвичайних змін у структурі мережі;
 - виявлення потенційних загроз (наприклад, посилення зв'язків навколо поняття «дезінформація»).
4. Розробка стратегій захисту:
 - пропонування оптимальних зв'язків для посилення зв'язності мережі;
 - розробка рекомендацій для запобігання кібератакам.

Аналіз інформаційних систем через призму графових моделей стає дедалі актуальнішим у сфері кібербезпеки, особливо з огляду на зростання складності атак, що вимагають глибокого розуміння структури мережі, її динаміки та взаємодії окремих елементів. Графові моделі дозволяють представити інформаційну систему як сукупність вузлів (користувачів, серверів, термінальних пристроїв) та зв'язків (комунікаційних шляхів, запитів, передачі даних), забезпечуючи зручний формат для застосування сучасних методів машинного навчання, аналізу даних і прогнозування.

Нижче розглянемо основні методи та технології, які використовуються для аналізу графів у контексті кібербезпеки, щоб виявляти загрози, прогнозувати поведінку мережі та реагувати на потенційні атаки в режимі майже реального часу.

Graph Neural Networks – моделювання взаємодій у мережі

У кібербезпеці GNN використовують для аналізу структури інформаційної мережі, виявлення підозрілих вузлів, визначення аномальних зв'язків та моделювання поширення шкідливого коду.

Можливі такі застосування GNN, як виявлення внутрішніх загроз – GNN може встановити, чи користувач почав взаємодіяти з новими серверами або нехарактерними ресурсами; поширення шкідливого ПЗ – модель аналізує, як вразливості чи шкідливий код можуть поширюватися через мережу; класифікація вузлів – система може автоматично класифікувати пристрої на нормальні, підозрілі чи заражені.

До переваги застосування GNN можна віднести обробку складних та нерегулярних структур; врахування не лише характеристик вузла, а й його оточення; можливість інтеграції з часовими даними для динамічного аналізу.

Node Embedding – представлення вузлів у числовому просторі

Node Embedding – це процес перетворення вузлів графа у вектори у багатовимірному просторі, які зберігають структурні та семантичні властивості мережі. Такі представлення можуть бути використані для подальшого аналізу, наприклад, для класифікації, кластеризації чи прогнозування.

Наведемо деякі приклади застосування Node Embedding у кібербезпеці виявлення підозрілих вузлів, які почали вести себе інакше, ніж їхні «сусіди»; прогнозування появи нових зв'язків, що може свідчити про спробу підключення зловмисника до внутрішньої мережі; аналіз структури мережі для виявлення слабких місць.

LSTM і RNN – аналіз динаміки мережі

Моделі типу Long Short-Term Memory та Recurrent Neural Networks у кібербезпеці використовують для виявлення аномалій у поведінці мережі, аналізу послідовностей запитів, а також прогнозування активності зловмисників.

LSTM аналізує послідовність станів мережі: звернення до файлів, зміни в правах доступу, зв'язки між пристроями. На основі минулих даних модель передбачає, які дії можуть відбутися далі.

Приклади застосування LSTM і RNN у кібербезпеці включають виявлення несанкціонованого доступу через аналіз послідовності запитів до ресурсів; прогнозування появи нових точок доступу, що може вказувати на впровадження шкідливого коду; виявлення патернів поведінки зловмисника в мережі (наприклад, він може спочатку сканувати мережу, потім збирати дані, і нарешті – експлуатувати вразливості).

Одним із найпростіших, але ефективних способів аналізу графа є обчислення метрік центральності – показників, які визначають, які вузли мають найвищу значущість у структурі мережі.

У попередніх розділах вже розглядалися такі метрики центральності:

- Degree Centrality – кількість зв'язків вузла;
- Betweenness Centrality – кількість найкоротших шляхів, що проходять через вузол;
- Closeness Centrality – наскільки близький вузол до інших вузлів;

- Eigenvector Centrality – враховує не лише кількість зв'язків, а й впливовість їхніх сусідів.

Обчислення метрик застосовується в таких задачах як виявлення ключових серверів, які можуть бути мішенями атак, виявлення перехресних зв'язків, що порушують політики безпеки, визначення вузлів, які могли бути компрометовані, через раптове зростання їхньої активності.

Агентні моделі дозволяють симулювати поведінку різних суб'єктів у мережі, включаючи зловмисників. Це корисно для тестування стійкості мережі до різних типів атак.

При цьому враховуються такі особливості агентних моделей, а саме, що кожен агент (вузол) має свої правила поведінки, а також те, що модель відстежує, як зловмисники можуть рухатися по мережі, шукаючи вразливі точки.

Агентні моделі застосовуються для тестування стратегій захисту і прогнозування маршруту атаки, даючи відповідь на запитання, «як зловмисник може досягти ключових вузлів?».

SIR-модель, відома з епідеміології, успішно адаптована для аналізу поширення шкідливого програмного забезпечення або внутрішніх загроз в інформаційних системах. Такі моделі і їх розширення (SEIR, SIS тощо) можуть застосовуватись для таких задач, як прогнозування, як швидко вірус може поширитися по мережі; оцінка ефективності антивірусних заходів; стратегії ізоляції вузлів для мінімізації поширення.

Аналіз графів дозволяє виявляти аномальні зміни у структурі мережі, що часто є ознакою кібератаки. При цьому модель аналізує звичайні стани мережі. Якщо виявляється раптове зростання степені зв'язності, поява нових кластерів чи зміна ролі вузла – це сигнал тривоги.

Можливі такі приклади застосування як виявлення незвичайної активності одного з пристроїв, що відправляє великі обсяги даних; виявлення нових кластерів, що можуть бути частиною внутрішньої загрози або APT (Advanced Persistent Threat); виявлення вилучення вузлів – наприклад, при атаці типу DDoS.

У деяких випадках графи будуються не лише на основі фізичних зв'язків, а й на основі семантичних залежностей – наприклад, які запити виконуються, які файли відкриваються, які права доступу змінюються. При цьому використовуються Knowledge Graphs – графи знань, що відображають зв'язки між подіями, користувачами, запитами тощо, а також Ontology-based models – моделі на основі онтологій, що встановлюють логічні залежності між поняттями.

Семантичні графи та онтології застосовуються для виявлення нелогічних переходів – коли користувач, який раніше не мав доступу до бази даних, раптово починає її аналізувати, а також для прогнозування логічних атак, таких як використання SQL-ін'єкцій чи витоку даних.

Ефективна система кібербезпеки не може обходитися без візуалізації графів. Для цього використовують такі програмні комплекси, як Gephi, Cytoscape, Kibana тощо, а також інтерактивні дашборди, що відображають зміни в реальному часі. Ці засоби забезпечують візуалізація змін у структурі мережі, виявлення вузлів із аномальною активністю, відстеження маршруту поширення загрози.

Аналіз графів у кібербезпеці вимагає комплексного підходу, що поєднує структурний аналіз, машинне навчання, динамічне моделювання та візуалізацію результатів. За допомогою таких технологій, як GNN, Node Embedding, LSTM, SIR-моделі, агентне моделювання та графова аналітика, можна не лише виявити поточну атаку, а й передбачити її розвиток, визначити ключові вузли, які потребують захисту, та протидіяти внутрішнім та зовнішнім загрозам на рівні структури мережі.

Ці методи забезпечують не просто реактивну, а й проактивну безпеку, даючи змогу аналізувати інформаційні системи як структуровані графи, а не як набір дискретних подій. Такий підхід особливо важливий у сучасних умовах, коли кібератаки стають все більш автономними, системними та заснованими на аналізі мережі.

Таким чином, графові моделі стають не лише інструментом аналізу, а й стратегічним механізмом для забезпечення кібербезпеки, що дозволяє системам швидко реагувати на нові види загроз, аналізувати їхню структуру та еволюцію, а також планувати заходи протидії на науково обґрунтованій основі.

9.4. Приклади застосування графових моделей у кібербезпеці

У сучасній кібербезпеці аналіз і прогнозування загроз вимагають не лише технічного розуміння мережевих взаємодій, а й глибокого дослідження структури даних, що представляється у формі графів. Графові моделі дозволяють моделювати складні залежності між елементами інформаційних систем – від окремих користувачів до серверів, протоколів та точок доступу. У цьому процесі особливо ефективними виявилися сучасні методи обробки графів, такі як GNN, Node2Vec, а також часові моделі на кшталт LSTM, які забезпечують не лише опис поточного стану системи, а й прогнозування потенційних

загроз. Ці технології дають змогу виявляти аномалії в структурі мережі, передбачати появу нових підозрілих зв'язків, аналізувати часову динаміку активності та формувати проактивну стратегію захисту. Нижче наведено конкретні приклади їхнього застосування для виявлення кібератак, прогнозування їхнього розвитку та аналізу змін у поведінці мережі.

Виявлення кібератак на інформаційні системи

Технологія. GNN для аналізу змін у структурі мережі.

Результат. Виявлення, що поняття «дезінформація» стало центральним у новому кластері, що свідчить про можливу атаку.

Прогнозування появи нових зв'язків

Технологія. Node2Vec для генерації векторних представлень та Logistic Regression для прогнозування зв'язків.

Результат. Прогнозування зв'язку між «економічними проблемами» та «політичною нестабільністю».

Аналіз динаміки мережі

Технологія. LSTM для аналізу часових рядів.

Результат. Виявлення сезонних трендів у зміні ваг зв'язків навколо поняття «медіа».

Використання графових моделей у кібербезпеці дозволяє ефективно аналізувати структуру та динаміку мережі, прогнозувати зміни та ідентифікувати потенційні загрози. Застосування методів, таких як GNN, Node Embedding, LSTM, забезпечує комплексний аналіз графів, що є ключовим для успішного протидії кібератакам.

9.5. Практична частина. Проведення самостійного дослідження

Практична частина передбачає проведення самостійного дослідження інформаційно-психологічних впливів в рамках ментальної війни.

Першим кроком є вибір конкретної інформаційної кампанії або загрози для аналізу. Наступний етап – збір даних з відкритих джерел, таких як соціальні мережі, новинні сайти та блоги. Ці дані можна завантажити за допомогою звернення до глобальних пошукових сервісів, таких як Google, Bing, або корпоративних OSINT-систем.

Рекомендовані кейси для самостійних досліджень:

1. Пропагандистські наративи у соціальних мережах.
2. Меми як інструмент ментальної війни.
3. Фейкові новини та їх вплив.
4. Культурні стереотипи.

5. Радянська пропаганда.
6. Роль релігії у формуванні ментальних моделей.
7. Медіа та національна ідентичність.
8. Військова пропаганда.
9. Концепція «ворогів народу» в російській пропаганді.
10. Політичні «мрії» лідерів різних країн.
11. Емоції у прийнятті політичних рішень.
12. Символи та їх роль у формуванні ідеологій, зокрема Z-символіка у росії.

Після збору даних слід побудувати графову модель інформаційного простору. Це можна зробити за допомогою системи Gephi або сервісу Semantic Networking. Далі проводиться аналіз графа, включаючи кластеризацію та ранжирування вузлів, які застосовувались у попередніх розділах.

Останнім кроком є прогнозування розвитку інформаційної атаки. Після цього на основі отриманих даних можна розробити стратегії для нейтралізації виявлених загроз.

Цей комплексний підхід дозволяє студентам не лише зрозуміти механізми інформаційно-психологічних війн, а й набути практичних навичок їх аналізу та протидії.

Висновки до розділу

У цьому розділі були детально розглянуті реальні випадки інформаційних атак, методи їх аналізу та прогнозування, а також практичні сценарії для дослідження. Особливу увагу приділено використанню сучасних інструментів машинного навчання та штучного інтелекту для автоматизації процесів аналізу та протидії інформаційним загрозам.

Аналіз реальних випадків, таких як російська пропаганда проти України та дезінформація під час пандемії COVID-19, продемонстрував важливість графового аналізу для виявлення ключових джерел дезінформації, акторів впливу та каналів поширення. Застосування метрик центральності, кластеризації та моделей SIR дозволило не лише описати структуру мережі, але й прогнозувати поширення інформаційних атак. Автоматизація аналізу за допомогою ML та AI забезпечила комплексний підхід до вивчення динаміки мережі, що є ключовим для успішного протидії інформаційним загрозам.

У розділі були запропоновані практичні сценарії для дослідження графових моделей ІПВ. Ці сценарії охоплюють широкий спектр завдань:

- аналіз структури мережі (визначення ключових вузлів, кластерів та «комунікативних бар'єрів»);
- прогнозування змін у мережі (поява нових зв'язків, сезонні тренди);
- виявлення аномалій (наприклад, посилення зв'язків навколо поняття «дезінформація»);
- генерація рекомендацій (розробка стратегій захисту від інформаційних атак).

Застосування методів, таких як GNN, Node Embedding та LSTM, дозволило ефективно вирішувати ці завдання. Наприклад, GNN використовувалися для класифікації вузлів та прогнозування зв'язків, а Node Embedding – для виявлення аномалій та сегментації мережі.

Графові моделі ІПВ мають широке застосування у сфері кібербезпеки. Вони дозволяють виявляти джерела кібератак та прогнозувати їхні шляхи поширення, моделювати сценарії атак та їхні наслідки, розробляти стратегії захисту, враховуючи оптимальні зв'язки для посилення зв'язності мережі.

Автоматизовані системи моніторингу інформаційного простору, засновані на OSINT, допомагають вчасно виявляти нові дезінформаційні кампанії та прогнозувати їхнє поширення. Це особливо важливо для запобігання кібератакам, які часто передують інформаційним атакам.

Практична частина розділу передбачає проведення самостійного дослідження інформаційно-психологічних впливів. Для самостійного дослідження пропонується:

1. Обрати конкретну інформаційну кампанію для аналізу.
2. Зібрати дані з відкритих джерел (соціальні мережі, новинні сайти, блоги).
3. Побудувати графову модель інформаційного простору за допомогою інструментів, таких як Gephi або NetworkX.
4. Провести аналіз графа, включаючи кластеризацію та ранжирування вузлів.
5. Прогнозувати розвиток інформаційної атаки за допомогою машинного навчання.

Цей комплексний підхід дозволить не лише зрозуміти механізми інформаційно-психологічних війн, але й набути практичних навичок їх аналізу та протидії.

Таким чином, у цьому розділі демонструється, що графові моделі є потужним інструментом для аналізу та прогнозування інформаційно-психологічних війн. Використання сучасних методів ML та AI

забезпечує можливість швидко аналізувати великі обсяги даних, виявляти приховані закономірності у мережах, прогнозувати зміни у структурах інформаційного простору.

Практичні сценарії та приклади реальних випадків інформаційних атак підтверджують ефективність графового аналізу для виявлення загроз та розробки стратегій захисту. Застосування цих методів у кібербезпеці дозволяє покращити захищеність інформаційних систем та запобігти масштабним інформаційним атакам.

Питання для самоконтролю

1. Які ключові наративи використовувались у російській пропаганді проти України?
2. Як будується графова модель для аналізу інформаційних атак?
3. Які основні метрики використовуються для аналізу структури графа?
4. Як формулюється PageRank для оцінки важливості вузла?
5. Які методи використовуються для прогнозування поширення дезінформації?
6. Які основні наративи поширювалися під час пандемії COVID-19?
7. Як SIR-модель застосовується для аналізу інформаційних атак?
8. Які переваги дає автоматизація аналізу графів за допомогою ML та AI?
9. Які основні завдання вирішують практичні сценарії дослідження графових моделей?
10. Як оцінюється ймовірність появи нового зв'язку між вузлами?
11. Як виявляються аномалії у графах за допомогою автоенкодерів?
12. Як GNN застосовуються для аналізу інформаційних атак?
13. Які основні завдання вирішують графові моделі у кібербезпеці?
14. Як Node Embedding допомагає у виявленні аномалій?
15. Як LSTM використовується для аналізу динаміки зв'язків?
16. Як визначається важливість вузлів у графі?
17. Які методи використовуються для прогнозування змін у інформаційному середовищі?
18. Як OSINT допомагає у моніторингу інформаційного простору?
19. Які основні кроки при проведенні аналізу інформаційної кампанії?
20. Як можна використовувати результати аналізу для розробки стратегій захисту?

Література за розділом

1. Структурно-функціональна модель системи забезпечення інформаційної та інформаційно-психологічної безпеки. Допов. Нац. акад.наук України. 2023. №1. – С. 16-23.
2. Качинський А.Б., Ланде Д.В. Графова модель ментальної війни. Інформаційні технології та безпека. Матеріали XXIV Міжнародної науково-практичної конференції ІТБ-2024. – Київ: Інжиніринг. – С. 123-128. ISBN: 978-617-8180-00-3
3. Інформаційні операції та безпека суспільства: загрози, протидія, моделювання: монографія / В.П. Горбулін, О.Г. Додонов, – К.: Інтертехнологія, 2009. – 164 с. ISBN 978-966-1648-12-7
4. Dmitry Lande, Leonard Strashnoy. Directed Correlation Networks, Determined by the Dynamics of COVID-19 Distribution in Various Countries. Available at SSRN: <http://ssrn.com/abstract=3674041>, DOI: 10.2139/ssrn.3674041 (Posted: 28 Aug 2020). – 7 p.

Висновки

У цьому посібнику були розглянуті ключові аспекти аналізу, моделювання та протидії інформаційно-психологічним війнам, зосереджуючись на застосуванні сучасних технологій, зокрема великих мовних моделей та семантичного нетворкінгу. Цей підхід дозволяє ефективно виявляти, аналізувати та протистояти загрозам у сфері інформаційної безпеки, що є особливо важливим у контексті сучасних глобальних викликів.

Великі мовні моделі стали потужним інструментом для обробки та аналізу великих обсягів текстових даних. Вони дозволяють виявляти ключові концепти та зв'язки. LLM можуть автоматично виділяти важливі теми, наративи та взаємозв'язки між ними, що допомагає розуміти структуру інформаційних потоків; генерувати високоякісний контент. Моделі можуть створювати тексти, які відповідають контексту та меті, наприклад, для поширення достовірної інформації або демонтажу фейків; прогнозувати тренди. LLM допомагають прогнозувати динаміку поширення інформації та виявляти потенційні загрози.

Однак застосування LLM має свої обмеження, зокрема проблеми «галюцинацій» (генерація неправдивих даних) та залежність від якості вхідних даних. Для подолання цих обмежень запропоновано комбінувати автоматизовані методи з людською експертизою.

Семантичний нетворкінг на основі великих мовних моделей та роїв віртуальних експертів є інноваційним підходом до аналізу інформаційних потоків. Він дозволяє створювати графи взаємозв'язків. Вершинами графа є поняття, а ребрами – зв'язки між ними. Такий підхід полегшує ідентифікацію ключових вузлів та кластерів. Крім того в рамках семантичного нетворкінгу здійснюється аналіз динаміки мереж і прогнозування наслідків атак. Зміни у структурі графа (наприклад, зниження зв'язності між поняттями) допомагають розуміти механізми поширення інформації та виявляти атаки. З іншого боку, руйнування зв'язків між поняттями «економічні проблеми» та «політична нестабільність» може призвести до роз'єднання наративів.

Захист інформаційного простору передбачає комплексний підхід, який включає захист ключових вузлів, відновлення зв'язків шляхом активації альтернативних шляхів передачі інформації після атак, кластеризацію, тобто формування сильних внутрішньокластерних зв'язків для зменшення впливу атак на окремі групи.

Результати аналізу мають широке практичне застосування медіа-стратегами при розробці ефективних кампаній для протидії

дезінформації, політиками для прогнозування наслідків політичних рішень та управління інформаційним полем, науковцями для розробки нових методів аналізу та моделювання інформаційних процесів.

Незважаючи на значні досягнення, існують обмеження, а саме, складність взаємодій, тому що зв'язки між поняттями часто є багатoshаровими, що ускладнює їх аналіз; чутливість до зовнішніх факторів, коли невеликі події можуть призвести до лавиноподібних змін у структурі мережі; необхідність людської участі, оскільки автоматизовані методи потребують валідації експертами для забезпечення точності та надійності результатів.

Подальші дослідження можуть бути спрямовані на удосконалення методів аналізу динамічних мереж, розробку нових алгоритмів для виявлення дезінформації та інтеграцію LLM з іншими технологіями, такими як машинне навчання та штучний інтелект.

Таким чином, цей науково-методичний посібник надає комплексний огляд сучасних методів аналізу та протидії інформаційно-психологічним війнам. Застосування великих мовних моделей, семантичного нетворкінгу та аналізу графів є потужними інструментами для захисту інформаційного простору. Однак для досягнення максимальної ефективності необхідно поєднувати автоматизовані методи з людською експертизою, враховуючи складність та динамічність інформаційних потоків. Цей підхід допомагає не лише захищати суспільство від маніпуляцій, а й зберігати національну ідентичність та інформаційну безпеку у сучасному світі.