

МІНІСТЕРСТВО ОСВІТИ ТА НАУКИ УКРАЇНИ
ДЕРЖАВНИЙ УНІВЕРСИТЕТ ТЕЛЕКОМУНІКАЦІЙ

Навчально-науковий інститут захисту інформації

На рецензію

Завідувач кафедри УІКБ

доктор економічних наук, професор

_____ С.В. Легомінова

«___» _____ 20__ р.

До захисту

Завідувач кафедри УІКБ

доктор економічних наук, професор

_____ С.В. Легомінова

«___» _____ 20__ р.

КВАЛІФІКАЦІЙНА РОБОТА

другого магістерського рівня вищої освіти

на тему:

**МЕТОДИКА ВИЯВЛЕННЯ ІНФОРМАЦІЙНОГО ВПЛИВУ ЗА
ДОПОМОГОЮ ЗАСТОСУВАННЯ СЕНТИМЕНТЕНТНОГО АНАЛІЗУ
ТЕКСТІВ**

СТУДЕНТ: Аркуша Даніїл Олегович

(підпис)

КЕРІВНИК: канд. техн. наук Щавінський Юрій Віталійович

(підпис)

НОРМОКОНТРОЛЕР: канд. військ. наук, доц.

Якименко Юрій Михайлович

(підпис)

Київ – 2022

ДЕРЖАВНИЙ УНІВЕРСИТЕТ ТЕЛЕКОМУНІКАЦІЙ
Навчально-науковий інститут захисту інформації
Кафедра управління інформаційною та кібернетичною безпекою
Спеціальність «Кібербезпека»
Спеціалізація «Управління інформаційною безпекою»
Другого магістерського рівня вищої освіти

«ЗАТВЕРДЖУЮ»

Завідувач кафедри УІКБ

д-р екон. наук, проф. _____ С.В. Легомінова
 (підпис)

«__» _____ 2022 р.

ЗАВДАННЯ
на кваліфікаційну роботу
 студенту Аркуші Даніїлу Олеговичу

1. Тема роботи «Методика виявлення інформаційного впливу за допомогою застосування сентиментного аналізу текстів», затверджена наказом по університету від “12” жовтня 2022 р. № 122.

2. Термін здачі студентом оформленої роботи «05» січня 2023 р.

3. Об’єкт дослідження: сентиментний аналіз текстів.

4. Предмет дослідження: методи сентиментного аналізу текстів в інформаційних мережах.

5. Мета дослідження: розробка методики виявлення інформаційного впливу, алгоритмів та програмного продукту для здійснення сентиментного аналізу тексту на основі нейронних мереж для застосування в управлінні кібербезпекою.

6. Перелік питань, які мають бути розроблені:

1. Провести аналіз основних положень і характеристик сентиментного аналізу, проаналізувати підходи і задачі сентиментного аналізу в кібербезпеці та методи визначення емоційного забарвлення інформації, розміщеної в соціальних мережах, визначити недоліки та переваги застосування сентиментного аналізу в управлінні кібербезпекою.

2. Дослідити принципи побудови моделі сентиментного аналізу, проаналізувати архітектуру штучних нейронних мереж, переваги їх застосування у кібербезпеці при організації кіберрозвідки.

3. Визначити технічні напрями побудови моделі сентиментного аналізу тексту із застосування мов програмування для контролю емоційного забарвлення змісту інформації в соціальних мережах з метою попередження потенційних загроз.

4. Розробити модель сентиментного аналізу та методику її застосування для автоматизованого визначення емоційного забарвлення інформації в соціальних мережах та визначити її валідність і ефективність.

7. Дата видачі завдання «15» вересня 2022 р.

КАЛЕНДАРНИЙ ПЛАН

Дата видачі завдання: «19» вересня 2022 р.

№ з/п	Етапи кваліфікаційної роботи	Термін виконання етапів	Примітка
1.	Визначення об'єкту, предмету, мети та завдань дослідження.	26.09.2022	
2.	Збір та аналіз літератури.	03.10.2022	
3.	Написання розділу 1.	17.10.2022	
4.	Написання розділу 2.	31.10.2022	
5.	Написання розділу 3.	14.11.2022	
6.	Написання розділу 4.	28.11.2022	
7.	Формулювання висновків за результатами дослідження.	01.12.2022	
8.	Оформлення роботи.	08.12.2022	
9.	Оформлення презентації.	15.12.2022	
10.	Отримання рецензії на роботу.	23.12.2022	
11.	Захист в ДЕК.	17.01.2023	

Студент

(підпис)

Д.О. Аркуша

Науковий керівник канд. техн. наук

(підпис)

Ю.В. Щавінський

РЕФЕРАТ

Кваліфікаційна робота присвячена дослідженню проблем виявлення потенційних загроз кібернетичній безпеці держави у контенті соціальних мереж в умовах інформаційного протиборства за допомогою сентиментного аналізу (скорочена назва – сентиментний, яка буде використовуватись в роботі). Робота складається зі вступу, чотирьох розділів, що містять 19 рисунків і одну таблицю, висновків та списку використаних джерел, що містить 65 найменувань і два додатки. Загальний обсяг роботи становить 88 аркушів, крім того, 7 аркушів займають список використаних джерел і 9 сторінок займають 2 додатки.

Об'єкт дослідження – сентиментний аналіз текстів.

Предмет дослідження – методи сентиментного аналізу текстів в інформаційних мережах.

Мета роботи – розробка методики виявлення інформаційного впливу, алгоритмів та програмного продукту для здійснення сентиментного аналізу тексту на основі нейронних мереж для застосування в управлінні кібербезпекою.

Методи дослідження – загально-наукові – аналіз і синтез, індуктивний і дедуктивний, аналогія, моделювання; метод системного аналізу, методи розробки алгоритмів – динамічне програмування, сходження, дерева розв'язків; одновимірні і багатовимірні статистичні методи, метод класифікації Байєса, методи навчання штучних нейронних мереж, контент-аналіз, методи розвідки кіберпростору.

У вступі проаналізований сучасний стан інформаційного захисту суспільства, визначена актуальність роботи кіберпідрозділів з попередження загрозам інформаційної безпеки особистості, суспільства, держави.

В першому розділі розкриті сутність, зміст та основні поняття сентиментного аналізу текстів, зроблений аналіз підходів і задач сентиментного аналізу в кібербезпеці, проаналізовані методи визначення емоційного забарвлення інформації, розміщеної в соціальних мережах, зроблений аналіз методів кіберрозвідки у кіберпросторі, визначені недоліки та переваги застосування сентиментного аналізу в управлінні кібербезпекою з метою організації превентивних заходів попередження загроз інформаційній безпеці.

Зроблені висновки про необхідність розроблення програмного забезпечення для ефективного застосування підрозділами кібербезпеки сентиментного аналізу у превентивній роботі виявлення загроз інформаційній безпеці суспільства.

В другому розділі із застосуванням системного аналізу визначені принципи побудови моделі сентиментного аналізу, проаналізована архітектура штучних нейронних мереж, визначені переваги їх застосування у кібербезпеці при організації кіберрозвідки. Зроблений висновок про актуальність розроблення моделей сентиментного аналізу на основі штучного інтелекту для застосування у кібербезпеці.

В третьому розділі визначена методика виявлення інформаційного впливу і розглянуті технічні аспекти побудови моделі сентиментного аналізу тексту із застосування мови програмування Python для визначення емоційного забарвлення змісту інформації в соціальних мережах, сформульована гіпотеза ефективності застосування моделі сентиментного аналізу підрозділами кібербезпеки при визначенні потенційних загроз інформації в соціальних мережах. Для підвищення ефективності управління кібербезпекою розглянуто переваги моделей, розроблених на основі штучних нейронних мереж та сформована методика виявлення інформаційного впливу за допомогою сентиментного аналізу.

У четвертому розділі для обґрунтування висновків і підтвердження гіпотез, визначених у третьому розділі, було розроблено програмне забезпечення на мові Python для здійснення сентиментного аналізу текстів. На базі безкоштовного віртуального середовища Google Colaboratory створено програмний продукт (нейронну мережу згорткового типу), що призначений для здійснення сентиментного аналізу текстів. Для навчання та тестування нейронної мережі було завантажено набір даних, а саме коментарі соціальної мережі Facebook, які були в загальному доступі на сайті <http://text-machine.cs.uml.edu>. Набір даних містив 31 185 російськомовних публічних постів (data.csv). Пости, включені в набір даних, мали довжину від 10 до 800 символів, принаймні 50% з яких були алфавітними, а принаймні 30% використовували російську кирилицю. Завдання складалося у визначенні позитивного, негативного, нейтрального забарвлення постів, у яких

виявлено відношення до подій в Україні у 2014 році з виключенням виразів подяки, привітань та поздоровлень, які будуть засмічувати результат.

Із 31 185 постів із статистикою Каппа 0,58, що за версією Флейса є задовільним результатом, для навчання мережі дані розділили на тренувальні (23370) та тестові (7790). З метою оптимізації ваг був обраний алгоритм RMSprop (Root Mean Square propagation – середньоквадратичне розповсюдження).

Навчання здійснювалося протягом 20 епох і в результаті було досягнуто точності близько 86%, що є дуже гарним результатом для нейронних мереж та підтверджує валідність отриманої моделі. Її застосування в управлінні кібербезпекою дозволить своєчасно в автоматизованому режимі визначати потенційні загрози на основі моніторингу інформації, розміщеної в соціальних мережах в кіберпросторі.

Наукова новизна одержаних результатів

1. В магістерській роботі вперше запропоновано застосування sentimentного аналізу в кібербезпеці для визначення емоційного забарвлення текстів в соціальних мережах з метою організації превентивних заходів в управлінні кібербезпекою.

2. Удосконалена модель sentimentного аналізу, яка розроблена на мові програмування Python запропонованими сумісним застосуванням словників та емотиконів.

3. Дістали подальшого розвитку застосування методів контент-аналізу при організації управління кібербезпекою з метою проведення превентивних заходів попередження негативного інформаційного впливу на особистість, суспільство, державу.

4. Проаналізований стан інформаційного захисту в Україні та визначена необхідність застосування штучного інтелекту для побудови моделей sentimentного аналізу в кібербезпеці при аналізі кіберзагроз.

Результати роботи оприлюднені на Всеукраїнської науково-практичної конференції (м. Київ, 27 жовтня 2022 року) та III науково-практичній конференції (м. Київ, 01 грудня 2022 року).

Практичне значення одержаних результатів.

Запропонована модель сентиментентного аналізу на основі штучних нейронних мереж дозволить своєчасно виявити негативний вплив контенту соціальної мережі та попередити загрози кіберпростору держави при застосуванні підрозділами кіберрозвідки.

Подальші дослідження застосування сентиментентного аналізу, як одного із превентивних способів кіберзахисту дозволять своєчасно виявити кіберзагрози, відстежувати настрої та моніторити підготовку терористичної діяльності в кіберпросторі.

Галузь застосування – Інформаційна та кібернетична безпека

Ключові слова: ІНФОРМАЦІЙНА БЕЗПЕКА, УПРАВЛІННЯ КІБЕРБЕЗПЕКОЮ, ЗАХИСТ ІНФОРМАЦІЇ, СЕНТИМЕНТЕНТНИЙ АНАЛІЗ, СЕНТИМЕНТНИЙ АНАЛІЗ, НЕЙРОН, НЕЙРОННІ МЕРЕЖІ, АЛГОРИТМ, PYTHON, GOOGLE COLABORATORY.

ЗМІСТ

ВСТУП.....	10
Розділ 1 ОСНОВИ СЕНТИМЕНТНОГО АНАЛІЗУ ТЕКСТУ	14
1.1 Основні поняття сентиментного аналізу тексту	14
1.2 Підходи та завдання сентиментного аналізу тексту в кібербезпеці	18
1.3 Аналіз методів розпізнавання емоційного забарвлення тексту	21
Висновки до розділу 1	44
Розділ 2 ОСНОВИ ЗАСТОСУВАННЯ ШТУЧНИХ НЕЙРОННИХ МЕРЕЖ В КІБЕРБЕЗПЕЦІ	45
2.1 Модель штучного нейрона та його функції активації.....	45
2.2 Архітектура штучної нейронної мережі	51
2.3 Переваги застосування штучних нейронних мереж в кібербезпеці	55
Висновки до розділу 2	60
Розділ 3 МЕТОДИКА ЗАСТОСУВАННЯ СЕНТИМЕНТНОГО АНАЛІЗУ ДЛЯ ВИЯВЛЕННЯ ІНФОРМАЦІЙНОГО ВПЛИВУ	61
3.1 Згорткові нейронні мережі як інструмент вирішення задач аналізу тональності тексту.....	61
3.2 Алгоритми і правила навчання нейронної мережі	65
3.3 Методика застосування сентиментного аналізу для виявлення інформаційного впливу соціальних мереж.....	71
Висновки до розділу 3	75
Розділ 4 РОЗРОБКА МОДЕЛІ НА ОСНОВІ НЕЙРОННОЇ МЕРЕЖІ ТА ЇЇ НАВЧАННЯ ДЛЯ ЗДІЙСНЕННЯ АНАЛІЗУ ТОНАЛЬНОСТІ ТЕКСТУ В СОЦІАЛЬНИХ МЕРЕЖАХ.....	77
4.1 Обґрунтування вибору мови програмування	77
4.2 Результат навчання нейронної мережі	78
Висновки до розділу 4.....	85
ВИСНОВКИ.....	86
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ	89
ДОДАТКИ	96

ВСТУП

Актуальність теми. Шлях, яким рухається Україна у розбудові власної кібербезпеки, потребує докорінних та невідкладних змін. Це не лише точка зору лідерів думок вітчизняного кіберзахисту. Глобальні системи масової комунікації сприяли утворенню кіберпростору, який разом з іншими фізичними просторами сьогодні визнається одним з можливих театрів воєнних дій. Унікальні можливості кіберпростору дозволяють використовувати різні інструменти інформаційно-психологічного впливу на індивідуальну свідомість людини, колективну та масову свідомість суспільства. Основним інструментом цього впливу є інформація в різних її проявах та емоційних забарвленнях в соціальних мережах.

Для того, щоб мінімізувати збитки від потенційного негативного впливу соціальних мереж, важливо фокусуватися не лише на захисті, але й на побудові правильних процесів реагування на зміст розміщеної в соціальних мережах інформації.

Значну роль у налагодженні цих процесів відіграє навчання кіберпідрозділів визначення емоційного забарвлення та забезпечення дієвим інструментом їх превентивної розвідувальної роботи.

Аналіз інформації є основним напрямком діяльності розвідувальних органів кібербезпеки у кіберпросторі або з його використанням [1]. Останнім часом особливо актуальним став аналіз тонального забарвлення інформації в соціальних мережах з метою виявлення настроїв у суспільстві, підготовки кібератак та їх попередження. Ще одна суттєва проблема – в Україні все ще недостатньо ефективно працює система кіберрозвідки (Threat Intelligence). Відсутність дієвої системи інформаційно-аналітичного забезпечення та великий об'єм інформації в соціальних мережах потребує проведення наукових досліджень, розроблення і застосування ефективних засобів для аналізу.

Аналіз думок відіграє важливі роль у багатьох галузях людської діяльності (політиці, економіці та суспільному житті). Аналіз тональності використовується в

таких областях як сфера послуг, охорона здоров'я, фінансове обслуговування, політичні вибори та інформаційна безпека [2]. Наприклад, у маркетингу, якщо продавець знає про задоволення замовника певним товаром, він може краще оцінити попит на продукт. Також для політиків, вони будуть знати, чи підтримують їх люди чи ні. У роботі [3] модель аналізу використовується для передбачень торгових успіхів.

Інше дослідження тональності змісту [4] було пов'язане з опитуваннями громадської думки. У роботі [5] думки з Twitter використовувалися для передбачення результатів політичного вибору. У дослідженнях [6,7] дані з Twitter про фільми застосовувалися для передбачення доходів квиткових кас.

Різниця між людьми та машинами полягає в тому, що люди мають здатність сформулювати особисті думки, і мрія штучного інтелекту полягає в тому, щоб машина поводитися як людина. Область комп'ютерної лінгвістики, яка аналізує відношення автора тексту до певних подій, продуктів, послуг і т.д., називається сентиментентним (скорочена назва, яка використовується в роботі – сентиментний) аналізом (sentiment analysis). Ґрунтуючись на аналізі настроїв, що присутні в текстах, в першу чергу зосереджуються на думках, що вказують на позитивні або негативні емоції автора.

На сьогоднішній день джерел для сентиментного аналізу стає все більше, оскільки люди діляться поглядами на різні теми через соціальні мережі, такі як Twitter, Facebook, вони залишають коментарі та огляди продуктів на спеціалізованих сайтах. Мікроблогінг в соціальних мережах - надзвичайно популярний спосіб обміну думками, і він щодня виробляє величезну кількість повідомлень та може розглядатися як багате джерело надійних повідомлень, які можна збирати та використовувати для аналізу настроїв. Чинником, що утруднює виявлення тональності в мережі, є існування великої кількості засобів масової комунікації, в яких при текстовому спілкуванні дуже часто ігноруються правила граматики та правопису. Прикладом є мова текстових повідомлень на мобільних телефонах, в яких використовуються аббревіатури, емотикони та скорочені пропозиції, але очевидно, що такий стиль листа можна знайти і у формах комунікації через інші електронні пристрої (наприклад комп'ютери).

Задача класифікації настроїв - це не нова наукова сфера. Проте раніше основна увага досліджень полягала у аналізі великих документів, а не на мікроблогах. В даній роботі основна увага приділена сентиментному аналізу саме на повідомленнях мікроблогів. Для такого аналізу було використано апарат нейронних мереж згорткового типу, що дало можливість виявляти сентиментне забарвлення тексту. В свою чергу наявність такої інформації, отриманої з соціальних мереж дозволяє виявляти неявний негативний вплив для його подальшої нейтралізації. Тому тема магістерської роботи є своєчасною та актуальною.

Мета і завдання дослідження - розробка методики виявлення інформаційного впливу, алгоритмів та програмного продукту для здійснення сентиментного аналізу тексту на основі нейронних мереж для застосування в управлінні кібербезпекою.

Для досягнення цієї мети в роботі необхідно виконати наступні **завдання**:

1. Провести аналіз основних положень і характеристик сентиментного аналізу, проаналізувати підходи і задачі сентиментного аналізу в кібербезпеці та методи визначення емоційного забарвлення інформації, розміщеної в соціальних мережах, визначити недоліки та переваги застосування сентиментного аналізу в управлінні кібербезпекою.

2. Дослідити принципи побудови моделі сентиментного аналізу, проаналізувати архітектуру штучних нейронних мереж, переваги їх застосування у кібербезпеці при організації кіберрозвідки.

3. Визначити технічні напрями побудови моделі сентиментного аналізу тексту із застосування мов програмування для контролю емоційного забарвлення змісту інформації в соціальних мережах з метою попередження потенційних загроз.

4. Розробити модель сентиментного аналізу та методику її застосування для автоматизованого визначення емоційного забарвлення інформації в соціальних мережах та визначити її валідність і ефективність.

Об'єкт дослідження - сентиментний аналіз текстів.

Предмет дослідження — методи сентиментного аналізу текстів в інформаційних мережах.

Методи дослідження. Загально-наукові – аналіз і синтез, індуктивний і дедуктивний метод, аналогія, моделювання; метод системного аналізу, методи розробки алгоритмів – динамічне програмування, сходження, дерева розв’язків; одновимірні і багатовимірні статистичні методи, метод класифікації Байєса, методи навчання штучних нейронних мереж, контент-аналіз, методи розвідки кіберпростору.

Наукова новизна одержаних результатів.

1. В магістерській роботі вперше запропоновано застосування сентиментного аналізу в кібербезпеці для визначення емоційного забарвлення текстів в соціальних мережах з метою організації превентивних заходів в управлінні кібербезпекою.

2. Удосконалена модель сентиментного аналізу, яка розроблена на мові програмування Python запропонованими сумісним застосуванням словників та емотиконів.

3. Дістали подальшого розвитку застосування методів контент-аналізу при організації управління кібербезпекою з метою проведення превентивних заходів попередження негативного інформаційного впливу на особистість, суспільство, державу.

4. Проаналізований стан інформаційного захисту в Україні та визначена необхідність застосування штучного інтелекту для побудови моделей сентиментного аналізу в кібербезпеці при аналізі кіберзагроз.

Результати роботи оприлюднені на Всеукраїнській науково-практичній конференції (м. Київ, 27 жовтня 2022 року) [8] та III науково-практичній конференції (м. Київ, 01 грудня 2022 року [64].

Практичне значення одержаних результатів. Запропонована модель сентиментного аналізу на основі штучних нейронних мереж дозволить своєчасно виявити негативний вплив контенту соціальної мережі та попередити загрози кіберпростору держави.

Подальші дослідження застосування сентиментного аналізу, як одного із превентивних способів кіберзахисту дозволять своєчасно виявляти кіберзагрози, відстежувати настрої та моніторити підготовку терористичної діяльності в кіберпросторі.

РОЗДІЛ 1

ОСНОВИ СЕНТИМЕНТНОГО АНАЛІЗУ ТЕКСТУ

Важливим завданням організацій, які займаються аналізом інформації, завжди було визначення відношення особистості, суспільства до конкретної теми дослідження. Із зростаючою доступністю та популярністю соціальних інтернет-ресурсів, таких як сайти онлайн-оглядів і особисті блоги, виникають нові можливості та виклики, оскільки люди тепер можуть і використовують інформаційні технології для пошуку та розуміння думок інших. Тому раптовий спалах активності в області дослідження думок і аналізу настроїв стосується обчислювальної обробки думок, настроїв і їх суб'єктивності в тексті. Він став прямою відповіддю на підвищення інтересу до нових систем аналізу, які безпосередньо мають справу з думками як з об'єктом дослідження.

Робота охоплює методи та підходи, які обіцяють безпосередньо активувати системи пошуку інформації, орієнтовані на думки та їх емоційне забарвлення стосовно різних тем і напрямків, які необхідно дослідити. Увага зосереджена на методах, які спрямовані на вирішення нових проблем, пов'язаних із додатками в соціальних мережах, що оцінюють настрої, порівняно з тими, які вже присутні в більш традиційному аналізі на основі фактів. Це стосується узагальнення оцінювального тексту, його емоційного забарвлення та ширші питання щодо конфіденційності, маніпулювання та економічного впливу, які стимулюють розвиток служб доступу до інформації, її захисту, орієнтованих на пошук загроз і в цілому національної безпеки держави.

1.1 Основні поняття сентиментного аналізу тексту

Сентиментний аналіз інформаційних потоків має великий потенціал застосування для моніторингових, аналітичних та сигнальних систем, які застосовуються при організації інформаційного та кібернетичного захисту.

В офіційних джерелах історично склалося так, що традиційний підхід до сентиментного аналізу є завданням класифікації тексту (частини тексту) на дві-три категорії (негативний, позитивний, нейтральний або просто: негативний або позитивний) [9, 10].

Саме з такого завдання почав свій розвиток аналіз тональності: оцінити сентимент відгуків з будь-якої тематики (кіно, ресторани, електроніка, а для інформаційної та кібербезпеки – терор, зброя, підрив, захист та інше, залежно від напрямку діяльності). Крім цього завдання, яке вирішує сентиментний аналіз тексту, при проведенні моніторингу контенту в соціальних мережах головним є ставлення або відношення до теми обговорення, тобто ставлення суб'єкта висловлювання до об'єкта, що обговорюється. Таким об'єктом, якому дається емоційна оцінка, називається об'єкт тональності, а носієм вираженої в тексті емоційної оцінки є певна особа – автор тексту або суб'єкт тональності. Але, якщо автор посилається на чийсь думку, тоді суб'єктом тональності буде той, на чию думку посилаються.

Таким чином, тональність висловлювання визначається трьома компонентами: суб'єктом тональності (хто висловив оцінку), об'єктом тональності (про кого або про що висловлено оцінку) та власне тональною оцінкою (як оцінили). Крім самої тональності, текст можна оцінювати за суб'єктивністю/об'єктивністю судження (Opinion Mining). Якщо це думка автора висловлювання, що містить суб'єктивну оцінку описуваного, то текст вважається суб'єктивним. І навпаки, якщо це повідомлення ЗМІ чи думка, що за умовчанням поділяється учасниками діалогу, воно вважається об'єктивним.

Аналіз тональності тексту - це область комп'ютерної лінгвістики, присвячена автоматичному виявленню оцінок, емоцій людини щодо згадуваних в тексті сутностей. Таких як, наприклад, продукт, послуга, подія, організація, особистість і т.д. або виявлення загальної оцінки, коли просто аналізується емоційність висловлювання.

З точки зору аналізу тональності, інформація в тексті ділиться на два класи:

1. Думки.

2. Факти.

Метою аналізу тональності тексту, як правило, є виявлення думок і їх властивостей.

Думки, в свою чергу, також можна розділити на два типи:

1. Проста думка.
2. Порівняльна думка.

У простій думці міститься ставлення автора до одного об'єкту. Думка може бути виражено як явно (пряма думка), так і неявно (непряма думка). Прикладом прямої думки є речення «Якість зв'язку відмінна», прикладом думки, вираженої неявно - «Після того, як я змінив налаштування, якість покращилася».

Об'єктом аналізу тональності тексту є будь-яка сутність (продукт, послуга, проблема, особистість і т.д.), щодо якої висловлюється думка в тексті. Дуже часто, об'єкти можуть складатися з окремих частин (компонентів) і властивостей. Компоненти і властивості складають безліч аспектів. Припустимо, що об'єктом думки є мобільний телефон. Він володіє такими складовими частинами, як камера, батарея, дисплей і т.д., а також наступним набором атрибутів: розмір, вага, зовнішній вигляд і т.д. Думка може бути виражена як щодо самого об'єкта, так і щодо його складових. Наприклад, в пропозиції «Хороший телефон, добротний акумулятор (швидко заряджається та повільно сідає)» перша його частина висловлює позитивну думку про об'єкт «телефон», а друга - про його складову частину - «акумулятор».

Автор думки - це людина, що виражає свою думку. Замість поняття «автор думки» може також використовуватися поняття «джерело думки».

Тональністю або сентиментом називають емоційне забарвлення, виражене в тексті. Можна виділити тональність трьох типів: позитивну, негативну і нейтральну. Однак у науковій літературі нейтральну тональність визначають по-різному, тому дуже часто беруть до уваги тільки перші два типи тональності.

Деякі дослідники визначають нейтральну тональність як деяке проміжне положення між позитивною і негативною, а деякі визначають її як відсутність будь-якого емоційного забарвлення. Таким чином, в аналізі емоційного забарвлення

тексту проста думка визначається через кортеж (впорядкована та скінченна сукупність), що складається з п'яти елементів (entity, aspect, sentiment value, holder, time), де entity (сутність) - це сутність аспекту (aspect – погляд, точка зору) про який автор (holder - власник) висловив думку (sentiment value) в певний момент часу (time) [11].

Другий тип думок - порівняльний - має три підтипи [12]:

1. Порівняння аспектів об'єктів на користь одного (Non-equal gradable comparison - нерівнозначне градуюче порівняння). Воно виражає відносини «щось краще / гірше чогось», які ранжують сутності відповідно до уподобань автора (найчастіше використовується порівняльна ступінь у порівнянні прикметників). прикладом є пропозиція «Windows краще, ніж Linux». Цей підтип включає в себе також переваги: «Я вважаю за краще використовувати Widows, ніж Linux».

2. Прирівнювання аспектів різних об'єктів (Equative comparison). Воно виражає відношення тотожності, яке свідчить, що дві або більше сутностей є рівними на підставі порівняння деяких загальних аспектів. Наприклад, «У своїх основних функціональних особливостях Widows та Linux майже однакові».

3. Перевага одного об'єкта над іншими (Superlative comparison). Цей підтип висловлює те ж відношення «щось краще / гірше чогось», що і перший, але в даному випадку одна сутність розташовується над іншими, тобто переважно використовується найвищий ступінь прикметників: «Widows – найкраще програмне забезпечення». В аналізі тональності тексту можна також зустріти поняття, тісно пов'язане з поняттям думки - суб'єктивність. Об'єктивне речення являє собою якусь фактичну інформацію про світ, тоді як суб'єктивне речення висловлює почуття, погляди, переконання людини. Прикладом об'єктивного речення є «iPhone є продуктом компанії Apple», а прикладом суб'єктивного - «Мені подобається iPhone». Таким чином, пояснивши необхідні терміни, можемо перейти до опису завдань сентиментного аналізу тексту.

1.2 Підходи та завдання сентиментного аналізу тексту в кібербезпеці

Визначення думки та сили думки при здійсненні аналізу необхідні для того, щоб допомогти зрозуміти, яка роль емоцій у неформальному спілкуванні, а також щоб визначити недоречні та аномальні емоційні висловлювання, які потенційно пов'язані з загрозливою поведінкою особистості стосовно оточуючих та суспільства. Існуючі алгоритми визначення думки здебільшого спрямовані на комерційне використання. Наприклад, створені для визначення думок про продукти, а не для вивчення людської поведінки чи її ставлення до об'єкту.

Більшість алгоритмів аналізу тональності не зовсім підходить для цього завдання, оскільки вони використовують неявні ознаки тональності, які можуть швидше відображати особливості жанру чи теми, а не емоційний аспект. Як результат, такі неявні ознаки можуть давати помилкові результати для дослідження соціальних мереж, виявляючи хибні моделі. Тому при організації заходів кіберрозвідки в соціальних мережах потрібні нові підходи, які будуть відрізнятися від стандартних, що застосовують у кібербезпеці.

Головне завдання організаційних заходів управління кібербезпекою - забезпечення захисту життєво важливих інтересів людини і громадянина, суспільства та держави, національних інтересів України у кіберпросторі, попередження кіберзагроз. Кіберзагроза - наявні та потенційно можливі явища і чинники, що створюють небезпеку життєво важливим національним інтересам України у кіберпросторі, справляють негативний вплив на стан кібербезпеки держави, кібербезпеку та кіберзахист її об'єктів [13].

Виявлення кіберзагроз з метою організації превентивних заходів є основною метою діяльності підрозділів кіберрозвідки.

Кіберпростір включає загальнодоступні соціальні мережі, де кожен користувач має право і можливість висловити своє відношення до інформації на різну тематику, яка може складати загрозу національним інтересам держави.

Для визначення відношення до такої інформації на основі емоційного забарвлення (тональності) тексту, що аналізується, підрозділи кібербезпеки мають застосовувати сентиментний аналіз.

Сентиментний аналіз є основою інтелектуального аналізу текстів і його розвинені держави використовують для цілей національної безпеки та розвідки. Крім того, програмне забезпечення інтелектуального аналізу тексту можна використовувати для створення великих досьє інформації про конкретних людей та події, як спосіб визначення характеристик повідомлень, які можуть бути небажаним матеріалом (актом розпалювання ворожнечі, закликами до повалення конституційного ладу, підготовкою терористичних актів тощо) [8].

Для вирішення цих завдань найчастіше використовують такі основні алгоритми:

Статистичний метод. Для його використання необхідні заздалегідь розмічені за тональністю колекції тексту. Вони служать для навчання моделі, за допомогою якої відбувається визначення тональності тексту або фрази.

Метод, заснований на словниках та правилах. Для цього потрібно скласти словники позитивних і негативних слів і виразів. Цей метод може використовувати як списки шаблонів, так і правила з'єднання тональної лексики всередині речення, засновані на граматичному та синтаксичному аналізі.

Для вирішення завдань кібербезпеки доцільно вибрати статистичний метод, що базується на моделі машинного навчання з учителем. Його перевага - можливість масштабувати завдання в рамках невеликого корпусу даних, що навчається. Потрібно домогтися того, щоб фахівець кібербезпеки міг додати до ІТ-системи посилання на новину або текст із соціальних мереж, а система визначила б емоційне забарвлення повідомлення.

Аналіз тональності зазвичай визначають як одну з задач комп'ютерної лінгвістики, тобто мається на увазі, що ми можемо знайти і класифікувати тональність, використовуючи інструменти обробки природної мови. З точки зору комп'ютерної лінгвістики текст на природній мові є неструктурованою

інформацією. З огляду на те, що думка - це кортеж, формулюють 6 основних завдань [11] аналізу тональності тексту (рис. 1.1):

Завдання 1 (витяг об'єктів та їх класифікація): Потрібно буде витягти всі слова, що виражають сутність в тексті, синонімічні слова згрупувати в класи (або кластери). Кожен кластер відповідає одній унікальній сутності.

Завдання 2 (витяг аспектів і їх класифікація): Потрібно буде витягти всі слова, що виражають аспекти, і синонімічні одиниці аналогічно розподілити по класах. Кожен такий «аспектний» кластер сутності відповідає одному унікальному аспекту.



Рис. 1.1. Основні завдання визначення тональності тексту

Завдання 3 (витяг автора і класифікація): З тексту або із структурних даних необхідно витягти слова, що виражають автора думки, і також класифікувати їх. Це завдання аналогічно першим двом.

Завдання 4 (витяг часу і його стандартизація): Необхідно також витягти час, коли було написано повідомлення, і привести різні його формати до одного. Це завдання також аналогічне тим, що описані вище.

Завдання 5 (визначення полярності): Необхідно визначити, чи є думка про аспект позитивною, негативною або нейтральною, або привласнити відповідне чисельне значення (яке обумовлюється заздалегідь).

Завдання 6 (вивід кортежу): Останнє завдання полягає в тому, щоб породити підсумковий кортеж думки, виражений в документі, на підставі результатів попередніх завдань.

1.3 Аналіз методів розпізнавання емоційного забарвлення тексту

Думка - це центральне поняття практично будь-якої діяльності людини. Думки є одним з ключових джерел впливу на нашу поведінку. Кожного разу, коли людина приймає рішення, для неї важливо знати, що про це думають інші люди. Підприємства та молодіжні організації завжди цікавляться думкою покупців і клієнтів про їх продукти та послуги. Окремим покупцям також цікаво знати думку користувачів товару перед його купівлею. Багато людей прагнуть дізнатися, що думає громадськість з приводу того чи іншого політичного кандидата перед голосуванням на виборах.

В минулому, коли людина потребувала чиеїсь думки, вона запитувала її у членів своєї сім'ї і друзів. Компанії проводили опитування та анкетування. Разом з величезним розвитком соціальних мереж (сайтів з відгуками, форумів, блогів, мікроблогів, Twitter та ін.) в Інтернеті, люди стали набагато частіше звертатися до подібних ресурсів, щоб приймати рішення з урахуванням думки громадськості.

Сьогодні якщо людина хоче придбати будь-який товар, у неї є можливість дізнатися думку не тільки члена сім'ї, але ще досвідчених користувачів по всьому світу. Великим організаціям більше не обов'язково проводити опитування і анкетування, тому що подібна інформація вже в надлишку знаходиться у вільному доступі. Однак у зв'язку з величезною кількістю сайтів, що містять відгуки, для звичайного користувача буде вельми проблематичним узагальнити безліч думок, виражених у текстовому форматі різними емоційними словами. Саме для цього й існують автоматичні системи аналізу тональності тексту.

Для розпізнавання емоційного забарвлення інформації в кіберпросторі існує деяка кількість алгоритмів і методів, створених вітчизняними та зарубіжними фахівцями.

У роботі [14] розглядаються методи аналізу тональності тексту (сентимент аналізу), які необхідні для автоматичного визначення ставлення автора до згаданої теми. Ці методи протестовані на базі, яка складається із 169 коментарів, що стосується банків. Один із них - метод, заснований на словниках, інший - метод з використанням бібліотеки сентимент аналізу Стенфорда та API-сервісу перекладів Яндекс.

У першому випадку при аналізі у коментарях підраховується кількість позитивних та негативних слів, також враховується частка «не» перед словом. Якщо частка «не» зустрічалася перед словом, то слово набувало протилежного емоційного забарвлення. При порівнянні слів використовувався стеммер Портера [15]. Стеммінг - відсікання від слова закінчень і суфіксів, щоб частина, що називається stem, була однаковою для всіх граматичних форм слова. Перед обробкою також забираються прийменники та словосполучення як неінформативні.

У другому використовувалась бібліотека обробки природної мови для англійської, що включала теоретико-графовий метод, розробленої Стенфордом і Yandex Translate API для попереднього перекладу тексту.

Як зазначає автор, в результаті перевірки були отримані наступні результати: тональність була визначена краще методом, заснованим на словниках - оцінка була визначена правильно для 71% відгуків. Другим методом тональність тексту було визначено для 56 % відгуків.

Неточності цих методів автор пояснює такими проблемами:

- численні орфографічні помилки у відгуках;
- немає зв'язку з об'єктів: тональність визначається для всього тексту;
- не завжди про відношення автора можна сказати за наявністю або відсутністю позитивних, негативних або нейтральних відгуків.

Поліпшити результати автоматичного визначення тональності тексту можна за допомогою використання методів автоматичного виправлення орфографічних помилок, вдосконалення словників (для методів, заснованих на словниках) та навчальної вибірки (для методів машинного навчання). Також можна підвищити точність роботи алгоритмів, застосовуючи розробки з інших проблем природної мови, таких як:

- автоматичне реферування (automatic summarization);
- виявлення кореферентності, референціональної множини (coreference resolution);
- аналіз порівнянь;
- вилучення об'єктів з текстів (Named-entity recognition, NER) та виявлення відносин між ними (Relationship extraction).

В роботі [16] представлений новий алгоритм SentiStrength, в якому використовується кілька нових методів для одночасного отримання позитивної та негативної тональності із коротких неформальних текстів в електронному вигляді. Програма використовує словник емоційно забарвлених слів з відповідними вагами, що відображають силу тональності, і використовує приклади ненормативного правопису, що широко вживаються, та інших загальних текстових способів вираження тональності.

Новий внесок цієї роботи полягає в наступному:

- використовується метод машинного навчання для оптимізації ваги емоцій;
- використовуються методи для отримання тональності з неформальних текстів, що містять ненормативні слова з літерами, що повторюються, використовується метод, що виправляє правопис.

На додаток до цього, система вводить п'ятибальну подвійну шкалу для позитивної та негативної тональності, корпус з 1041 коментарів з MySpace для цієї системи та абсолютно нову систему оцінки тональності текстів, яка поєднує нові та вже існуючі методи. Така система має вигляд:

- 5 - дуже сильна негативна тональність (напр., болісно);
- 4 - сильна негативна тональність (напр., ненавидіти);

- 3 - помірно негативна тональність (напр., ворожість);
- 2 - злегка негативна тональність (напр., незручність);
- 2 - злегка позитивна тональність (напр., задовольняти);
- 3 - помірно позитивна тональність (напр., щасливий);
- 4 - сильна позитивна тональність (напр., закоханий);
- 5 - дуже сильна позитивна тональність (напр., захоплений).

Як правило, як зазначають фахівці [17-19], визначення думки зазвичай включає три стадії, хоча деякі завдання можуть вимагати більше кроків.

По-перше, вхідний текст поділяється на частини, наприклад, пропозиції, і кожен відрізок перевіряється на наявність тональності - чи є пропозиція суб'єктивною чи об'єктивною.

По-друге, суб'єктивні пропозиції аналізуються на наявність тональності.

По-третє, може виділятися об'єкт, щодо якого висловлюється думка.

Зазвичай метод отримання думки визначає лише позитивну чи негативну тональності а не окремі емоції (щастя, здивування), не визначає силу тональності.

Тим не менше, проведені в роботах [17-19] дослідження з визначення думки можуть допомогти при виділенні одночасно позитивних та негативних емоцій завдяки тому, що більшість методів можуть бути теоретично пристосовані під ценове завдання.

Загальноприйнятим підходом до аналізу тональності є вибір алгоритму машинного навчання та метод вилучення характеристик із тексту, де характеристиками зазвичай є слова, але також можуть бути основи слів або слова з частковою розміткою, які також можуть бути об'єднані в біграми (тобто два наступні слова) або триграми [20]. Були розроблені також складніші варіанти, що використовують, наприклад, автоматичне вилучення характеристик [21].

Альтернативним методом визначення тональності є виявлення найбільш ймовірної середньостатистичної тональності для слів у тексті, шляхом оцінки того, як часто вони зустрічаються разом із словами із заданого початкового списку емотивної лексики, на який було однозначно визначено тональність (наприклад,

«хороший», «жахливий»). Для цього зазвичай використовуються пошукові системи для оцінки відносної частоти спільної появи. Вони спираються на припущення, що позитивні слова частіше зустрічаються з позитивними словами, ніж з негативними і навпаки. Для цього методу необхідний відносно невеликий вхідний набір лексичних даних. Він може пристосовуватися до різних предметних областей, тому що набір початкових загальних ключових слів може використовувати породження нового словника кожної області застосування. Цей метод з використанням первинного списку слів, дає дуже хороші результати при обробці різних контекстів і розпізнає слова предметної області, пов'язані з тональністю, [22].

Однак подібні методи часто схильні виділяти слова, які асоціюються з тональністю, але явно її не виражають (такі як «відчувати», «Ірак» чи «пізно»). Такі поняття називають «неявно емоційними словами» (з англ. indirect affective words) і від «явно емоційних слів» (direct affective words) [23]. Використання неявно емоційних слів є недоліком для деяких типів дослідження виявлення тональності із соціальних мереж, а також для деяких комерційних застосувань, оскільки вони роблять методи залежними не тільки від предметної області, а й від часу (наприклад, 3G, ймовірно, вже не є надійним ознакою у позитивних відгуках на мобільні телефони, Windows-95 – вже не є надійною ознакою позитивного відгуку на операційні системи персональних комп'ютерів).

Лексичний підхід полягає у використанні вже існуючого набору понять із відомою тональністю, який створюють заздалегідь. Потім застосовується алгоритм для передбачення тональності тексту на підставі цих слів [23].

Словниковий метод може бути забезпечений іншою інформацією, наприклад, списком емотиконів та набором семантичних правил, які, наприклад, справлялися б із запереченням. Існують різні модифікації цього методу, які включають негативні частинки слова, які посилюють тональність інших слів («абсолютно ненавидіти») та загальні структури речень.

Більш складний підхід полягає у визначенні характеристик тексту, які можуть бути потенційно суб'єктивними в деяких контекстах і в подальшому використанні контекстної інформації для того, щоб вирішити, чи є суб'єктивними

характеристики в кожному новому контексті. Більше того, були розроблені різні методи поліпшення стандартних ресурсів, такі як виявлення складних слів. Як і описані вище методи з початковим списком слів, лексичний підхід також може знаходити слова, пов'язані з тональністю певної предметної області, наприклад, слово «маленький» є позитивним словом для оглядів на портативні електронні пристрої [24-26].

Ще одна програма SO-CAL використовує лексичний метод для класифікації текстів на позитивні чи негативні, а також використовує словник слів, яким надано оцінку за загальною для двох тональностей шкалою від -5 до +5 [27]. Словник у SO-CAL був створений на основі корпусу, розміченому людьми, а також із використанням словника General Inquirer. У результаті вийшло 2252 прикметників, 745 прислівників 1142 іменників та 903 дієслів, для всіх іменників та дієслів була проведена лематизація, яка передбачає трансформацію форми слів у початкову форму. Кожному слову було надано значення «пріоритетної тональності» – стандартної тональності даного слова на всій безлічі контекстів, а не в окремо взятому.

У SO-CAL також є список 187 емоційних виразів, що складаються з декількох слів, а також список підсилювальних виразів, які підвищують або знижують силу тональності та алгоритм, який справляється із запереченням. Підсумкова тональність речення складається із середнього арифметичного значення тональності всіх слів, знайдених у ньому після всіх перетворень. Ця програма показує хороші результати для збалансованого набору даних, що складаються в основному з текстів з різних сайтів та текстів новин [28].

Хоча аналіз тональності зазвичай пов'язаний з думками, Т. Вілсон у роботі [29] поширила його на психологічну площину з метою визначення з тексту прихованого внутрішнього стану автора. Як зазначає автор, в основі багатьох методів автоматичного аналізу модальності лежить великий лексикон, де у слів помічена їх апіорна полярність (яка також називається семантичною орієнтацією). Проте контекстуальна полярність словосполучення, у якому вживається окреме слово, може дуже відрізнятись від апіорної полярності цього слова. Позитивні

слова вживаються у фразях, які виражають негативні емоції, або навпаки. Також, досить часто слова, які є позитивними або негативними поза контекстом, є нейтральними у контексті, тобто їх вживають зовсім не для того, щоб виразити емоцію. Здійснене в роботі дослідження полягало у автоматичному розрізненні апріорної і контекстуальної полярності з акцентом на з'ясуванні важливих для вирішення цього завдання ознак.

Оскільки важливим аспектом проблеми є з'ясування, коли емоційно забарвлені слова вживаються у нейтральних контекстах, проаналізовано ознаки нейтрального і емоційно забарвленого значення, а також ознаки позитивної і негативної контекстуальної полярності. Аналіз включав оцінку продуктивності ознак у різних алгоритмах машинного навчання. В усіх алгоритмах машинного навчання, за винятком одного, найкращі результати досягаються шляхом комбінування усіх ознак. Іншим аспектом аналізу було з'ясування впливу нейтральних уживань на продуктивність ознак позитивної і негативної полярності. Ці експерименти свідчать, що присутність нейтральних уживань значно погіршує продуктивність цих ознак і що можливо найкращим способом підвищення результатів розпізнавання усіх видів полярності є удосконалення здатності системи розпізнавати нейтральні слововживання.

Завданням дослідження даних в соціальній мережі MySpace було визначити не тільки думки чи внутрішній стан автора, але й роль вираженої тональності для спілкування онлайн. Звідси дослідження було зосереджено на виявленні тональності, вираженої в кожному повідомленні, незалежно від того, чи відображає воно прихований внутрішній стан автора, інтерпретацію повідомлення або прихований внутрішній стан читача.

Для того, щоб отримати надійну оцінку експертів щодо випадкової вибірки коментарів з MySpace, було проведено два пробні анотування з використанням двох різних наборів даних (загальним обсягом 2600 коментарів). Анотування було проведено з метою виявлення відповідної шкали та головних проблем в оцінці. Для експертів було складено і потім удосконалено інструкцію, а також створено

онлайн-систему для випадкового вибору коментарів, які були представлені експертам.

Одним із головних результатів пробної розмітки виявилось те, що експерти відносили окличні пропозиції до позитивних, якщо в них не був явно виражений негативний контекст. Наприклад, фраза «Привіт!» інтерпретувалася більшістю експертів як позитивна за умовчанням, тому що вона виражає інтенсивність у контексті, який не дає жодних підказок щодо тональності емоції. І навпаки, фразу "невдаха!!!" експерти оцінили як негативнішу, ніж просто «невдаха», оскільки знаки вигуків асоціюються з негативним словом. Згодом інструкції були переглянуті, і в них було явно занесено твердження, що об'єднання, мабуть, нейтрального вигуків та позитивної тональності є допустимим [29].

Для остаточного рішення було вибрано більше тисячі коментарів з MySpace (в середньому з 20 словами та 101 символом на коментар) для їх оцінки за наступною п'ятибальною шкалою для позитивної та негативної тональності:

(немає позитивної емоції або інтенсивності) 1 2 3 4 5 (вкрай позитивна емоція);

(немає негативної емоції) 1 -2 - 3 - 4 - 5 (вкрай негативна емоція).

Важливу роль в опрацюванні природної мови відіграє людська оцінка, часто представлена у формі оціночних суджень. Незважаючи на деякі застосування класичних непараметричних і вузькоспеціалізованих підходів до оцінювання цих видів суджень, існує ціла низка їх досліджень у контексті оцінювання сенсорного розрізнення та людських суджень, які є ключовими в ньому, підкріплена строгою статистичною теорією і програмним забезпеченням у вільному доступі, які можна використати в опрацюванні природної мови. Фахівцями досліджено логарифмічні лінійні моделі Бредлі-Террі, як один з підходів, який застосовано до вибіркового даних для опрацювання природної мови [30].

Експертам було надано як усні інструкції з оцінки текстів, і брошура, пояснююча завдання (складена з урахуванням роботи [30]). Найзначніші інструкції були наведені у четвертому розділі цієї роботи. Початкова версія брошури містила приклади коментарів з відповідною позитивною чи негативною оцінкою, але на

практиці дана інформація ніяк не допомагала експертам під час пробної оцінки. З цієї причини набір коментарів не використовувався, тому ступінь згоди між експертами міг бути оцінений більш реалістично, оскільки виключалася можливість того, що деякі коментарі були надто схожі на ті, що зустрічалися в завданні.

Щодо застосування емотиконів. Різними людьми емотикони сприймалися по-різному частково через їхній життєвий досвід, через особисті розбіжності і навіть статі автора. Для розробки системи оцінка повинна давати послідовний погляд на тональність даних, а не середньостатистичне сприйняття людей.

До експерименту, описаному в роботі [31], було залучено експертів однієї статі (жіночого), пробна оцінка проводилася для того, щоб відібрати людей, які давали однорідні результати. Спочатку було обрано 5 експертів, проте результати двох із них згодом виявилися непридатними через аномальність: один експерт давав значно вищі оцінки коментарям із позитивною тональністю, інший давав загалом непослідовні результати. Для трьох експертів було обчислено та округлено середні значення за кожним коментарем. Із цього вийшов золотий стандарт для експериментів. Нижче наведено кілька прикладів текстів та оцінок [31]:

- hey witch wat cha been up too (оцінки: +: 2,3,1; -: 2,2,2);
- omg my son has the same b-day as you lol (оцінки: +: 4,3,1; -: 1,1,1);
- HEY U HAVE TWO FRIENDS!! (Оцінки: +: 2,3,2; -: 1,1,1);
- What's up with that boy Carson? (Оцінки: +: 1,1,1; -: 3,2,1).

Розробниками було пораховано ступінь згоди між трьома експертами.

Для цього було використано коефіцієнт, відомий під назвою альфа Криппендорфа [32]. Значення коефіцієнта для пропозицій із позитивною тональністю вийшло 0,5743, для негативних – 0,5634. Дані результати досить надійні, це показує, що оцінки експертів у багатьох випадках збігаються. Однак вони є недостатньо надійними, тому що нормою вважаються значення $>0,67$. Тим не менше, використання середньої оцінки експертів як золотий стандарт є цілком розумним методом отримання оцінок тональності.

Алгоритм даного інструменту ефективніше за інші методи справляється з аналізом позитивної тональності коротких неформальних текстів. Результат роботи програми оцінювався за допомогою точності та коефіцієнта кореляції між оцінками експертів та оцінками програми. Пізніше розробники зробили кілька спроб покращення роботи програми для негативної тональності шляхом розширення вихідних даних.

Розглянуті програми створювалися для аналізу коротких текстів англійською мовою і потребують для використання іншими мовами зміни файлів вихідних даних у вигляді тематичних словників.

На жаль, вдалося знайти лише кілька досліджень на тему аналізу тональності текстів українською мовою. Вони описані у роботах [33-35]. У них описано етапи створення тонального словника з використанням анотованого корпусу текстів, де предметом досліджень було аналіз емоційного стану тексту українською мовою на різних рівнях. Повний словник, однак, ніде не представлений, а ті, що складені, стосуються різних тематичних направленостей і в кожному із них відсутня направленість кібербезпеки.

Тому першочергове завдання для аналізу контенту соціальних мереж в українському сегменті – це створення тематичного словника.

В умовах агресії Російської Федерації проти України, анексії Криму, території Донецької та Луганської областей, однією з найважливіших проблем стає забезпечення інформаційної безпеки України. Важливість цієї проблеми, зокрема, визначається в нових редакціях Стратегії національної безпеки України та Воєнної доктрини України, нормативно-правових актах, які визначають напрями та завдання співпраці України з НАТО та іншими міжнародними безпековими організаціями. Відповідно до Стратегії національної безпеки України, затвердженої Указом президента України від 14 вересня 2020 року № 392/2020, основним пріоритетом забезпечення інформаційної безпеки нашої держави є “забезпечення наступальності заходів політики інформаційної безпеки на основі асиметричних дій проти всіх форм і проявів інформаційної агресії” [36].

У роботі [37] визначена важливість забезпечення національної безпеки держави у кібернетичному просторі, обґрунтована актуальність та необхідність проведення розвідувальних заходів у кібернетичному просторі противника, визначені етапи, складові та методи кібернетичної розвідки, а також критичні дані, які необхідно добути.

Одним із важливих засобів розвідки у кіберпросторі є інструменти віддаленого аналізу та ідентифікації досліджуваних об'єктів противника. Проте, незважаючи на той факт, що у теперішній час питання побудови інструментів добування розвідувальних даних приділена велика увага, головним питанням залишається – ефективне застосування інструментів розвідки кіберпростору.

У теперішній час, методи розвідки кібернетичного простору класифікують, як активні та пасивні, кожен з яких має свої переваги та недоліки. При використанні пасивного методу збору розвідувальної інформації контакту з досліджуваним об'єктом не відбувається. При безпосередньому добуванні даних не генерується трафік, не реєструється з'єднання з хостом або сервером у системному журналі подій, а також скорочується загальна завантаженість на досліджуваний сегмент мережі при скануванні.

Не дивлячись на всі переваги пасивного методу добування інформації у нього є і недоліки. Для проведення пасивного аналізу завжди потрібно інтегруватися у сегмент досліджуваного об'єкта мережі для виконання ролі датчика, через який буде проходити та аналізуватися мережевий трафік, який циркулює в даному сегменті між об'єктами розвідки. Або, як варіант, потребує віддаленого з'єднання з інформаційним ресурсом досліджуваного об'єкта для генерації мережевого трафіку та його подальшого аналізу для ідентифікації віддаленого об'єкта, що в свою чергу втрачає актуальність прихованого віддаленого аналізу. Також великим недоліком є велика ймовірність помилкової ідентифікації досліджуваного об'єкта, що впливає на якісні показники добутої інформації у ході проведення розвідувальних заходів.

Метод активного добування розвідувальних даних полягає у безпосередньому контактуванні з досліджуваним об'єктом розвідки з

використанням методів прихованого сканування, що дає можливість бути непомітним при здійсненні впливу на противника. Також на відміну від пасивного методу добування даних при активному добуванні ймовірність помилкової ідентифікації віддаленого об'єкта зменшується вразі, що підвищує точність розвідувальної інформації та ефективність подальшого формування кібернетичного впливу на основі добутих розвідувальних даних про об'єкт дослідження.

У свою чергу, для порівняльного аналізу засобів розвідки можна застосувати наступні критерії побудови програмного забезпечення [38]:

- масштабованість (можливість додавання нових ресурсів, а також можливість керування єдиною розподіленою системою кібернетичної розвідки);
- відкритість (можливість інтеграції в систему додаткових розроблених компонентів);
- кроссплатформеність (можливість перенесення додатку на різну платформу сімейства операційних систем Windows, MacOS X, Unix);
- методи добування розвідувальних даних досліджуваного об'єкта (TCP, UDP та приховане сканування віддаленого об'єкта);
- час виконання розвідувальних заходів (зменшення використання часу на добування інформації);
- якість добутих розвідувальних даних (зведення до мінімуму помилкової ідентифікації об'єкта).

Проте, аналіз досліджень та публікацій [16, 18-20, 24, 28-29, 37], а також досвід експлуатації засобів кібернетичної розвідки показує, що жодний з них у повній мірі не відповідає наведеним критеріям.

Враховуючи переваги зазначених засобів добування розвідувальних даних над іншими, фахівці відмічають у своїх роботах [37-39], що в умовах обмеження часу наявні у підрозділах кіберрозвідки інструменти не здатні у короткі проміжки часу добути бажану розвідувальну інформацію про противника і потребують розроблення моделей, здатних в автоматизованому режимі збирати та аналізувати інформацію з різних джерел, включаючи соціальні мережі. Крім того,

враховуючи величезний об'єм інформації та необхідність постійного моніторингу з причини її мінливості, при побудові моделей необхідно застосовувати штучний інтелект.

Враховуючи активне застосування комп'ютерних технологій в усіх сферах життя, а також перехід значної частини злочинності з реального світу у кіберпростір, який став як місцем і засобом вчинення кримінальних правопорушень, так і джерелом інформації, яка становить оперативний інтерес, фахівці виділяють кіберрозвідку, як ще один вид кримінальної розвідки.

Авторами [37-40] у своїх дослідженнях досить часто використовуються такі поняття, як:

- комп'ютерна розвідка – оперативно-розшуковий захід, який полягає у цілеспрямованому пошуку та отриманні інформації з комп'ютерних систем та мереж, доступ до яких не обмежується їх власником, володільцем чи розпорядником або не пов'язаний з подоланням системи логічного захисту, що здійснюється працівниками оперативних та оперативно-технічних підрозділів з метою виявлення відомостей криміногенного та кримінального характеру;

- кримінологічна розвідка, яка полягає у легальному зборі, аналізі, оцінці та інтерпретації інформації про криміналізацію суспільних відносин, загрози та ризику кримінального та криміногенного змісту, реальний стан злочинності, криміналітету, його статус і можливості протидіяти правопорядку, з метою забезпечення національної, громадської, групової та особистої безпеки;

- аналітична розвідка – розвідувальний пошук, технічна розвідка, комплексне вивчення матеріалів прихованого спостереження та оперативної установки, а також аналіз повідомлень, публікацій та виступів у засобах масової інформації, статистичних даних, відомостей автоматизованих банків даних;

- «віртуальна» розвідка – комплекс заходів щодо здобуття, обробки й аналізу розвідувальної інформації в кібернетичному (телекомунікаційному, віртуальному) просторі за допомогою різних видів програмно-математичного

впливу. При цьому «віртуальна розвідка» включає в себе кіберрозвідку, радіорозвідку, розвідку систем супутникового зв'язку тощо.

При проведенні кіберрозвідки під час оперативно-розшукової діяльності та в ході розслідування злочинів використовуються наступні засоби та методи [38]:

- аналіз відкритих джерел інформації (OSINT –Open source intelligence), в тому числі – соціальних мереж;
- соціальна інженерія, обмін думок;
- зняття інформації з транспортних телекомунікаційних мереж та з електронних інформаційних систем;
- отримання та аналіз інформації, яка знаходиться в операторів та провайдерів телекомунікацій, про зв'язок, абонента, надання телекомунікаційних послуг, у тому числі отримання послуг, їх тривалості, змісту, маршрутів передавання тощо.

У висновках автор [38] зазначив, що кіберрозвідка є перспективним напрямом сучасної оперативно-розшукової діяльності і потребує подальшого вивчення та дослідження, поєднаного з розробкою новітніх методів і засобів пошуку, збору та аналізу інформації, розміщеної в кіберпросторі, з урахуванням досягнень науково-технічного прогресу.

За підрахунками, 80% даних у світі неструктуровані та не організовані заздалегідь визначеним чином. Найбільше це відбувається з текстовими даними, такими як електронні листи, висловлювання в чатах та соціальних мережах, опитування, статті та документи. Ці тексти зазвичай важкі, трудомісткі і дорогі для аналізу та сортування. Системи аналізу почуттів дозволяють значно спростити цей процес.

Деякі з переваг аналізу настроїв включають наступне:

1. Масштабованість. Існує занадто багато даних для обробку вручну. Аналіз тональності дозволяє ефективно та економічно обробляти дані в масштабі.

2. Аналіз у реальному часі. Ми можемо використовувати аналіз настроїв для виявлення критичної інформації, яка дозволяє усвідомлювати ситуацію під час конкретних подій у режимі реального часу.

3. Відповідні критерії. Люди не дотримуються чітких критеріїв оцінювання настроїв тексту. За підрахунками, різні люди погоджуються лише в 60-65% випадків, коли судять про настрої до певного тексту. Це суб'єктивне завдання, на яке сильно впливають особисті переживання, думки та переконання. Використовуючи централізовану систему аналізу настроїв, компанії можуть застосовувати однакові критерії до всіх своїх даних. Це допомагає зменшити помилки та покращити узгодженість даних. І як наслідок мати більш точний результат.

Основні методи визначення тональності тексту фахівці [41-42] розділяють на наступні групи (рис. 1.2):



Рис. 1.2. Групи методів визначення тональності тексту

1. Методи, засновані на правилах (rule-based approach). Вони складаються з певного набору правил, на підставі яких система робить висновки щодо тональності тексту. Прикладом такого правила для речення «Я купив приголомшливий телефон» є наступне твердження: Якщо прикметник («приголомшливий») входить в множину позитивних прикметників («приголомшливий», «чудовий», «хороший», ...) і не входить в множину негативних прикметників («жахливий», «огидний», «поганий», ...), то класифікувати тональність як позитивну.

Переваги: даний підхід може давати хороші результати при великому наборі правил;

Недоліки: складання великого набору правил - дуже трудомісткий процес; дуже часто правила прив'язуються до певної тематичної області; цей підхід не дуже підходить для аналізу мікроблогів через «зашумленість» даних, що обумовлена наявністю помилок.

2. Методи, засновані на використанні словників емоційної і оціночної лексики (sentiment lexicon). В даному випадку кожному окремому слову в словнику присвоюється значення тональності (шкали обговорюються заздалегідь). Для отримання підсумкового значення тональності часто використовують простий спосіб: беруть середнє арифметичне або обчислюють суму значень тональності всіх слів з документа. Більш складний спосіб - навчання класифікатора (наприклад, нейронної мережі).

Переваги: даний метод простий у застосуванні;

Недоліки: метод не універсальний, для нової предметної області потрібно складання нового словника.

3. Методи, засновані на машинному навчанні з учителем (supervised learning). Це найбільш поширений метод. Його суть полягає в тому, щоб навчити алгоритм класифікації на основі колекції документів (вибірки), класи яких відомі заздалегідь.

Переваги: висока точність при визначенні тональності; проблема залежності від конкретної предметної області вирішується шляхом навчання класифікатора на основі вибірки з даної області, так як класифікатор сам виділяє ознаки, які

впливають на тональність; проводиться безліч досліджень з метою поліпшення точності.

Недоліки: необхідна розмічена колекція текстів (розмітка є вельми трудомістким процесом).

4. Методи, засновані на машинному навчанні без учителя (unsupervised learning). Відмінність від попереднього методу полягає в тому, що приховані закономірності і взаємозв'язки між об'єктами виявляються з нерозміченої вибірки даних (або беруться розмічені дані, але така інформація при навчанні алгоритму не використовується).

Переваги: не потрібна розмічена колекція документів;

Недоліки: низька точність, в порівнянні навчання з учителем

Далі розглянемо більш детально методи навчання з учителем. Алгоритм реалізації даного підходу (рис. 1.3) можна коротко описати таким чином:

Алгоритм реалізації методу навчання з учителем



Рис. 1.3. Схема алгоритму методу навчання з учителем

1) спочатку необхідно зібрати колекцію тверджень, на основі якої навчатиметься класифікатор;

2) кожне твердження необхідно представити у вигляді вектору ознак (аспектів);

3) далі кожному твердженню потрібно присвоїти правильний тип тональності;

4) необхідно вибрати алгоритм класифікації і метод для навчання класифікатора;

5) застосування отриманої моделі.

Перш ніж переходити до вищеописаного алгоритму, необхідно вирішити, яка кількість класів і який тип класифікації будуть використані. При використанні плоскої класифікації досить складно досягти високих результатів. Вже відомо, що набагато кращі результати дає ієрархічна класифікація. Всі документи з навчальної вибірки повинні представляти собою n-мірні вектори аспектів. Від того, який набір характеристик буде використаний, безпосередньо залежить якість результатів.

Найбільш поширеними способами подання документів - це або у вигляді так званого «мішка слів» (bag-of-words), або у вигляді n-грам [42-44].

«Мішок слів». Припустимо, є 2 твердження:

(1) Я люблю читати книги. Мій друг любить читати книги теж.

(2) Я також люблю читати газети.

На їх основі виділяється список зустрічаються слів (часто слова призводять до їх початкову форму за допомогою стемінгу):

["Я", "люблю", "читати", "книги", "також", "газети", "друг", "теж"]

Відповідно до цього списком документи можна уявити наступним чином:

(1) (1, 2, 1, 1, 2, 0, 0, 0, 1, 1)

(2) (1, 1, 1, 1, 0, 1, 1, 1, 0, 0)

Розмірність векторів відповідає розмірності списку. Числа показують те, скільки разів слово зустрічається в даному документі.

N-грами. Припустимо, є твердження «Я купив приголомшливий ноутбук».

Набором уніграм в даному випадку буде список послідовностей («Я», «купив», «приголомшливий», «ноутбук»), набором біграм - («Я купив», «купив приголомшливий», «приголомшливий ноутбук»). Найчастіше в аналізі використовуються уніграми, біграми, або їх комбінація: («Я», «купив», «приголомшливий», «ноутбук», «я купив», «купив приголомшливий», «приголомшливий ноутбук»). У завданнях текстової класифікації в більшості випадків використовуються n-грами з $1 < n < 3$. Далі, щоб скласти вектор, необхідно присвоїти кожному його значенню вагу. В інформаційному пошуку популярним методом оцінки ваги є міра TF-IDF (від англ. TF - term frequency, IDF - inverse document frequency) - статистичний показник, що використовується для оцінки важливості слів у контексті документа, який є частиною колекції документів чи корпусу. Вага (значимість) слова пропорційна кількості вживань цього слова у документі, і обернено пропорційна частоті вживання слова у інших документах колекції. Проте цей метод не дуже ефективний в аналізі тональності. Є ще його модифікація - дельта TF-IDF, суть якої полягає в тому, щоб привласнити більшу вагу для слів, які мають позитивну або негативну тональність.

Формула розрахунку наступна [45]

$$V_{t,d} = C_{t,d} \times \log \left(\frac{|N| \times P_t}{|P| \times N_t} \right), \quad (1.1)$$

де $V_{t,d}$ - вага слова в документі;

$C_{t,d}$ кількість разів, скільки слово t зустрічається в документі d ;

$|P|$ - кількість тверджень з позитивною тональністю;

$|N|$ - кількість тверджень з негативною тональністю;

P_t - кількість позитивних тверджень, де зустрічається слово t .

Наступним кроком є вибір алгоритму класифікації. Розглянемо два найпоширеніших методи.

1. Наївний Байєсовський класифікатор (naïve Bayes classifier) [46] – один з найпростіших методів класифікації. Нехай у нас є рядок O , множина класів C , до одного з яких необхідно віднести рядок. Ми вибираємо клас так, щоб ймовірність приналежності об'єкта $P(C|O)$ до нього була максимальна

$$c = \arg \max P(C|O) \quad (1.2)$$

Обчислити $P(C|O)$ складно, але можна використовувати теорему Байєса [47]

$$P(C|O) = \frac{P(O|C) P(C)}{P(O)} \quad (1.3)$$

де $P(C)$ - апіорна ймовірність гіпотези C ;

$P(C|O)$ - ймовірність гіпотези C при настанні події O (апостеріорна ймовірність);

$P(O|C)$ - ймовірність настання події O при істинності гіпотези C ;

$P(O)$ - повна ймовірність настання події O .

Якщо класифікатор працює не з усім рядком, а з вектором ознак, можна представити формулу наступним чином

$$P(C|o_1 o_2 \dots o_n) = P(o_1 o_2 \dots o_n|C)P(C) \quad (1.4)$$

Тут робиться «наївне» припущення про те, що змінні O залежать тільки від класу C і не залежать один від одного.

Права частина рівняння приймає вид

$$P(o_1 o_2 \dots o_n|C)P(C) = P(C)P(o_1|C)P(o_2|C) \dots P(o_n|C) = P(C) \prod_i (o_i|C) \quad (1.5)$$

Кінцевий вид формули наступний [47]

$$c = \arg \max_{c \in C} P(c|o_1 o_2 \dots o_n) = \arg \max_{c \in C} P(c) \prod_i (o_i|C) \quad (1.6)$$

Код на Python розміщений в загальному доступі на сервері Хабр (Додаток Б).

Отже все, що потрібно зробити - це обчислити ймовірності $P(C)$ і $P(O|C)$. Обчислення цих параметрів і називається тренуванням класифікатора. Метод опорних векторів (SVM, support vector machine). Метою такої класифікації є пошук гіперплощини в просторі аспектів, котра розділяє всі об'єкти на два класи. Ідею методу зручно представити на наступному прикладі: припустимо, дано точки на площині, розділені на два класи (рис.1.4).

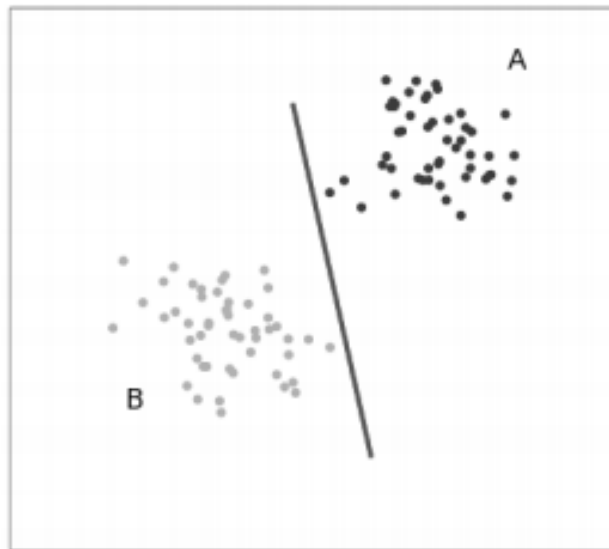


Рис. 1.4. Розбиття точок на класи

Проведемо лінію між класами. Потім всі нові точки (не з навчальної вибірки) будуть автоматично класифікуватися відповідно до умов: точка вище прямої потрапляє в клас А, а точка нижче прямої - в В.

Таку пряму можна назвати розділяючою прямою. Але в просторах високих розмірностей пряма вже не буде розділяти класи, тому замість неї розглядають

гіперплощину - простір, розмірність якого на одиницю менше, ніж розмірність початкового простору.

В даному прикладі існує кілька прямих, які поділяють точки на класи (рис. 1.5), однак необхідно вибрати пряму, відстань від якої до кожного класу максимальна.

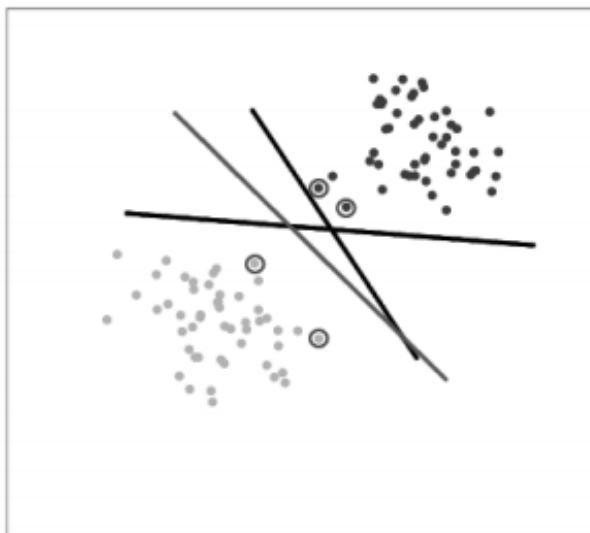


Рис.1.5. Приклад вибору необхідної гіперплощини

Вектори, що лежать ближче всіх до розділяючої площини, називаються опорними векторами (на рисунку вони обведені в кружки). Однак на практиці найчастіше зустрічаються випадки, які є лінійно нероздільними (рис. 1.6).

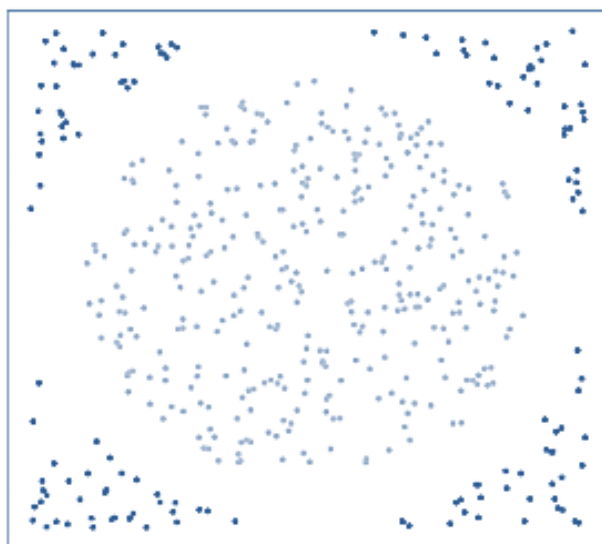


Рис.1.6. Приклад лінійно нероздільної вибірки

У таких випадках всі елементи навчальної вибірки вкладаються в простір більш високої розмірності, щоб в ньому вибірка була лінійно роздільною (рис. 1.7).



Рис. 1.7. Приклад вирішення проблеми з лінійно нероздільною вибіркою

Після вибору алгоритму класифікації та навчання класифікатора проводиться оцінка результатів за допомогою повноти та точності, або за допомогою ковзного контролю чи перехресної перевірки (cross-validation).

Повнота буде розрахована за формулою

$$R_{\Pi} = \frac{N_{correctly\ extracted\ opinions}}{N_{total\ number\ of\ opinions}} \quad (1.7)$$

де *N_{correctly extracted opinions}* – кількість правильно визначених думок;

N_{total number of opinions} - загальна кількість думок (як знайдених системою, так і не знайдених).

Точність буде розрахована за формулою

$$R_T = \frac{N_{correctly\ extracted\ opinions}}{N_{total\ number\ of\ opinions\ found\ by\ system}} \quad (1.8)$$

де *N_{total number of opinions found by system}* - загальна кількість думок, знайдених системою.

У випадку з перехресною перевіркою, дані розбиваються на k частин, потім на $k-1$ частинах даних проводиться тренування моделі, а частину, що залишилася використовують для тестування. І так k разів (рис. 1.8) [48].

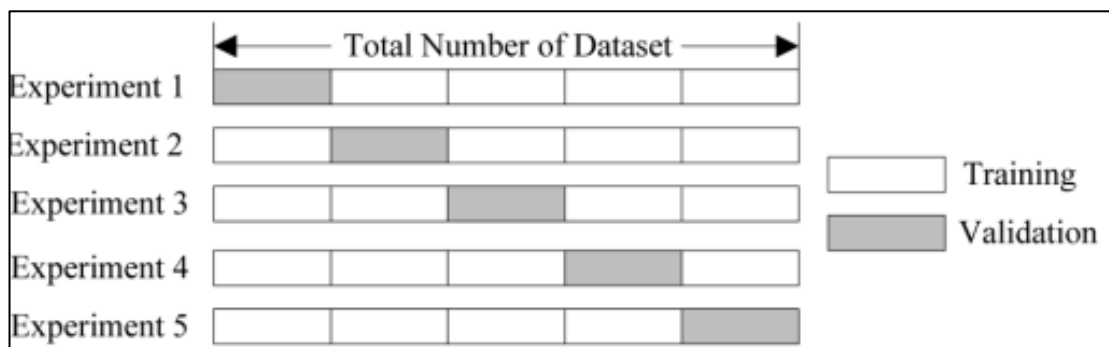


Рис.1.8. Представлення ковзаючого контролю

Висновки до 1 розділу

Унікальні можливості кіберпростору дозволяють використовувати різні інструменти інформаційно-психологічного впливу на індивідуальну свідомість людини, колективну та масову свідомість суспільства. Основним інструментом цього впливу є інформація у кіберпросторі в різних її проявах та емоційних забарвленнях, особливо у загальнодоступних соціальних мережах.

У розділі за допомогою аналізу визначена недостатня ефективність заходів кіберрозвідки та малоефективні засоби, які вони застосовують. Це не дозволяє своєчасно не тільки попереджувати кіберзагрози але й своєчасно на них реагувати. В умовах обмеження часу наявні у підрозділах кіберрозвідки інструменти не здатні у короткі проміжки часу добути бажану розвідувальну інформацію про противника і існує нагальна потреба розроблення моделей, здатних в автоматизованому режимі збирати та аналізувати інформацію з різних джерел, включаючи соціальні мережі.

Враховуючи величезний об'єм інформації та необхідність постійного моніторингу з причини її мінливості, при побудові моделей необхідно застосовувати штучний інтелект, застосовуючи методи і алгоритми, розглянуті у розділі.

РОЗДІЛ 2

ОСНОВИ ЗАСТОСУВАННЯ ШТУЧНИХ НЕЙРОННИХ МЕРЕЖ В КІБЕРБЕЗПЕЦІ

2.1 Модель штучного нейрона та його функції активації

Для побудови моделі на основі штучних нейронних мереж (ШНМ) необхідно визначитись з основами і принципами побудови.

Елементарною структурною одиницею штучної нейронної мережі є штучний нейрон, який представляє собою спрощену модель біологічного нейрона. На (рис. 2.1) зображено нейрон, входами якого можуть бути або вхідні дані, або виходи від такого ж нейрона [49].

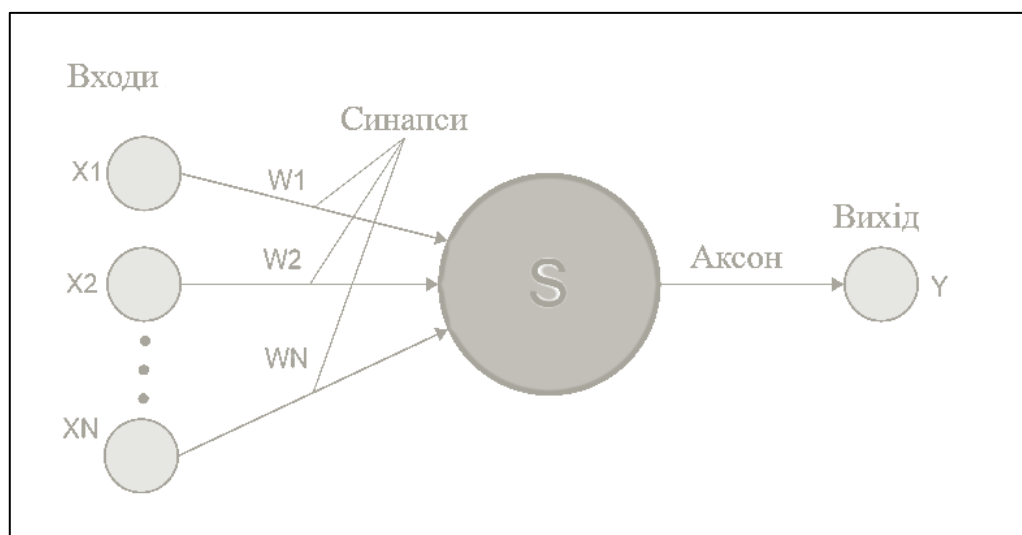


Рис. 2.1. Візуалізація штучного нейрона

Входи з'єднані з осередком нейрона S за допомогою синаптичних зв'язків.

Кожен синапс має свою вагу, при передачі в осередок нейрона вхідного параметра, він відповідно збільшується на вагу, тобто $x_i * w_i$. Стан нейрона S , визначається у вигляді формули

$$S = \sum_{i=1}^N x_i * w_i \quad (2.1)$$

де x – входи, w – синапси, N – кількість.

Вихід нейрона є функція його стану $y = f(S)$.

В основному, нейрони класифікують на основі їх положення в топології мережі. При класифікації розділяють [50]:

- вхідні нейрони - отримують вектор, що кодує вхідний сигнал. Як правило, ці нейрони не виконують обчислювальних операцій, а просто передають отриманий вхідний сигнал на вихід, можливо, посиливши або послабивши його;
- вихідні нейрони - являють собою виходи мережі. У вихідних нейронах можуть проводитися будь-які обчислювальні операції;
- проміжні нейрони - виконують основні обчислювальні операції.

Також, розрізняють декілька типів передавальних функцій.

Лінійна передавальна функція. Сигнал на виході нейрона лінійно пов'язаний із ваговою сумою сигналів на його вході [50]

$$f(x) = tx, \quad (2.2)$$

де t – параметр функції.

У ШНМ із шаруватою структурою, нейрони з передавальними функціями такого типу, як правило, складають вхідний шар. Крім простої лінійної функції можуть використовуватись також її модифікації. Наприклад напівлінійна функція (якщо її аргумент менше нуля, то вона дорівнює нулю, а в інших випадках, поводить себе як лінійна) або крокова (лінійна функція з насиченням), яку виражають формулою [51]

$$f(x) = \begin{cases} 0 & \text{if } x \leq 0 \\ 1 & \text{if } x \geq 1 \\ x & \text{else} \end{cases} \quad (2.3)$$

При цьому можливий зсув функції по обох осях. Недоліками крокової і напівлінійних активаційних функцій порівняно із лінійною можна назвати те, що вони не є диференційованими на всій числової осі, а отже не можуть бути використані при навчанні за деякими алгоритмами.

Порогова передавальна функція. Інша назва - функція Хевісайда (функція одиничного стрибка або ступенчата функція, дорівнює нулю при від'ємних значеннях аргумента та більше нуля при додатніх). Являє собою перепад. До тих пір поки зважений сигнал на вході нейрона не досягає певного рівня T -сигнал на виході дорівнює нулю. Як тільки сигнал на вході нейрона перевищує зазначений рівень - вихідний сигнал стрибкоподібно змінюється на одиницю.

Найперший представник шаруватих штучних нейронних мереж - перцептрон складався виключно з нейронів такого типу. Математичний запис цієї функції виглядає так [51]

$$f(x) = \begin{cases} 1 & \text{if } x \geq T \\ 0 & \text{else if } x < T \end{cases} \quad (2.3)$$

де $T = -\omega_0 x_0$ - зсув функції активації щодо горизонтальної осі.

Відповідно під x слід розуміти зважену суму сигналів на входах нейрона без урахування цього доданку. З огляду на те, що ця функція не є диференційованою на всій осі абсцис, її не можна використовувати в мережах, що навчаються за алгоритмом зворотного поширення помилки та іншими алгоритмами, що вимагають диференційованої передавальної функції.

Сигмоїдальна передавальна функція. Один із найчастіше використовуваних на даний час типів передавальних функцій. Введення функцій сигмоїдального типу було обумовлене обмеженістю нейронних мереж із пороговою функцією активації нейронів - за такої функції активації будь-який із виходів мережі дорівнює або нулю, або одиниці, що обмежує використання мереж не в задачах класифікації.

Вид сигмоїдальної функції на рис. 2.2.

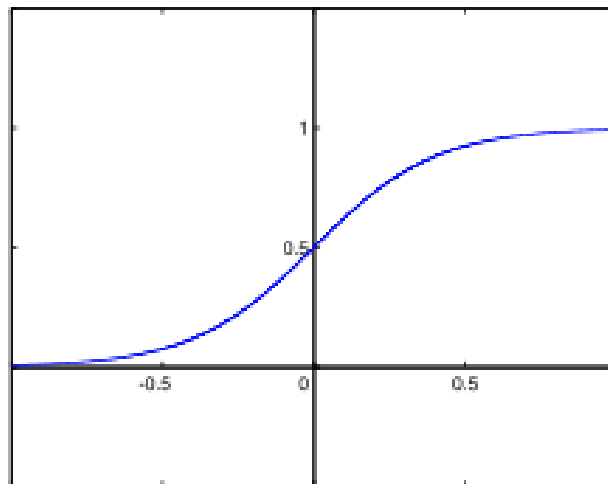


Рис. 2.2. Сигмоїдальна функція

Використання сигмоїдальних функцій дозволило перейти від бінарних виходів нейрона до аналогових. Функції передачі такого типу, як правило, властиві нейронам, що знаходяться у внутрішніх шарах нейронної мережі.

Логістична функція. Математично цю функцію можна виразити так [52]

$$\sigma(x) = \frac{1}{(1+\exp(-tx))} \quad (2.4)$$

де t - це параметр функції, що визначає її крутизну.

Коли t прямує до нескінченності, функція вироджується в порогову. Сигмоїда вироджується в постійну функцію із значенням 0,5. Область значень даної функції знаходиться в інтервалі (0,1). Важливою перевагою цієї функції є простота її похідної [52]

$$\frac{d\sigma(x)}{dx} = tf(x)(1 - f(x)) \quad (2.5)$$

Те, що похідна цієї функції може бути виражена через її значення, полегшує використання цієї функції при навчанні мережі за алгоритмом зворотного поширення. Особливістю нейронів з такою передавальною характеристикою є те, що вони посилюють сильні сигнали істотно менше, ніж слабкі, оскільки області сильних сигналів відповідають пологим ділянках характеристики. Це дозволяє запобігти насиченню від великих сигналів.

Гіперболічний тангенс. Використання функції гіперболічного тангенса [52]

$$th(A_x) = \frac{\exp(A_x) - \exp(-A_x)}{\exp(A_x) + \exp(-A_x)} \quad (2.6)$$

відрізняється від розглянутої вище логістичної кривої тим, що його область значень лежить в інтервалі $(-1; 1)$. Оскільки є вірним наступне співвідношення [52]

$$th\left(\frac{A}{2}x\right) = 2\sigma(x) - 1, \quad (2.7)$$

то обидва графіки відрізняються лише масштабом осей. Похідна гіперболічного тангенса, зрозуміло, теж виражається квадратичною функцією значення; властивість протистояти насиченню також має місце.

Таким чином, вибір функції активації визначається:

- специфікою завдання;
- зручністю реалізації на ЕОМ, у вигляді електричної схеми або іншим способом;
- алгоритмом навчання: деякі алгоритми накладають обмеження на вид функції активації, їх потрібно враховувати.

Найчастіше вид нелінійності не завдає принципового впливу на рішення задачі. Однак, вдалий вибір може скоротити час навчання в кілька разів.

Також, немає чітко визначеної процедури для вибору кількості нейронів і кількості шарів в мережі. Чим більше кількість нейронів і шарів, тим ширше

можливості мережі, тим повільніше вона навчається і працює і тим більше нелінійною може бути залежність вхід - вихід.

Кількість нейронів і шарів пов'язують [50]:

- із складністю завдання;
- з кількістю даних для навчання;
- з необхідною кількістю входів і виходів мережі;
- з наявними ресурсами: пам'яттю і швидкодією машини, на якій моделюється мережа.

Були спроби записати емпіричні формули для числа шарів і нейронів, але застосовність формул виявилася дуже обмеженою.

Якщо в мережі занадто мало нейронів або шарів:

- мережа не навчиться і помилка при роботі мережі залишиться великою;
- на виході мережі не будуть передаватися різкі коливання функції, що апроксимується у (х).

Перевищення необхідної кількості нейронів теж заважає роботі мережі. Якщо нейронів або шарів занадто багато тоді, як зазначено в роботі [52]:

- швидкодія буде низькою, а пам'яті буде потрібно багато;
- мережа перенавчиться: вихідний вектор буде передавати незначні і несуттєві деталі в досліджуваній залежності у (х), наприклад, шум або помилкові дані;
- залежність виходу від входу виявиться різко нелінійною: вихідний вектор буде суттєво і непередбачувано змінюватися при малій зміні вхідного вектора х;
- мережа буде нездатна до узагальнення: в області, де немає або мало відомих точок функції у (х) вихідний вектор буде випадковим і непередбачуваним, що не буде адекватним для вирішення.

2.2 Архітектура штучної нейронної мережі

Нейронні мережі – одна з найкращих парадигм програмування.

Під час традиційного підходу до програмування, комп'ютеру вказують «що робити», розбиваючи складні задачі на декілька простих, котрі комп'ютер може з легкістю виконати. В нейронних мережах, навпаки, не повідомляють як вирішити задачу. Натомість, комп'ютер вирішує сам, на основі даних спостережень та навченої математичної моделі.

Виділяють кілька основних типів нейронних мереж [52]:

Багат шарові мережі. У багат шарових мережах один або кілька нейронів об'єднуються в шари (рис. 2.3).

Шар - сукупність нейронів, на вхід яких подається один і той же загальний сигнал. У мережах такого типу, зовнішні вхідні дані подаються на входи нейронів першого шару, а вихідні дані є результатом останнього вихідного шару. Крім вхідного і вихідного шарів, в багат шарових мережах так само присутній один або кілька прихованих шарів. Зв'язки виходів нейронів від деякого шару і до деякого шару $i + 1$ називають послідовним [53].

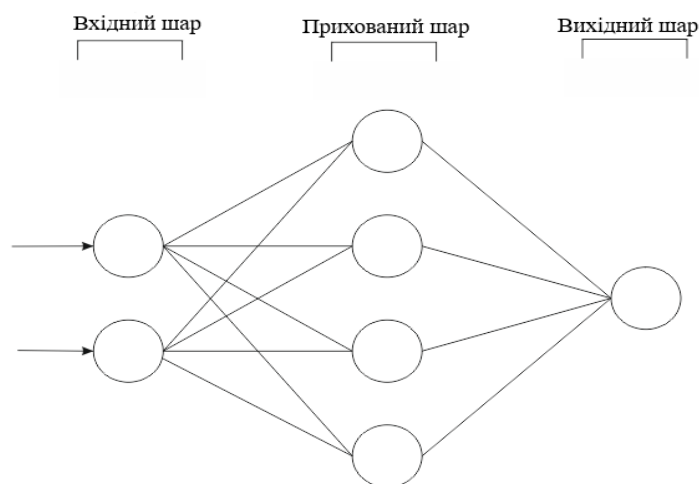


Рис.2.3. Багат шарова нейронна мережа

Повнозв'язні нейронні мережі. У повнозв'язній нейронній мережі кожен нейрон передає свій сигнал іншим нейронам (рис. 2.4). Вихідними сигналами,

можуть бути всі або деякі сигнали нейронів після кількох тактів функціонування. Всі вхідні сигнали, подаються на вхід всіх нейронів.

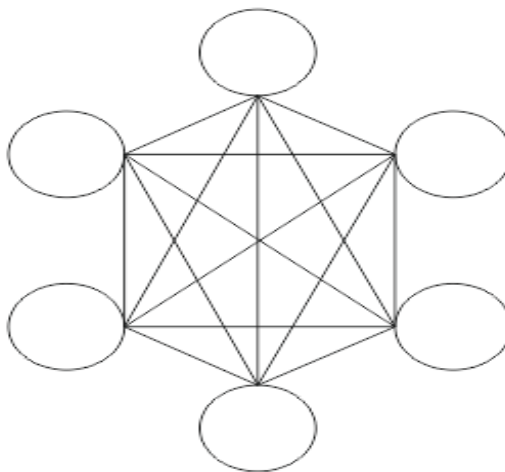


Рис.2.4. Повнозв'язна нейронна мережа

Кожна нейронна мережа має не тільки свою архітектуру, а й тип, який краще підходить для вирішення конкретного завдання. Наприклад, згорткові нейронні мережі (convolutional neural network CNN) набагато краще справляється з розпізнаванням образів, проблемами комп'ютерного зору та аналізом тональності тексту.

Їх відмінність від інших типів ШНМ полягає в тому, що кожен фрагмент зображення множиться на матрицю (ядро) згортки поелементно, а результат підсумовується і записується в аналогічну позицію вихідного зображення (рис. 2.5).

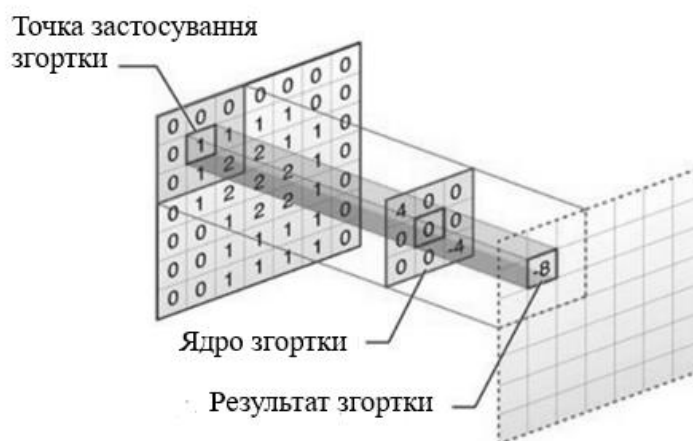


Рис.2.5. Згорткова нейронна мережа

Якщо кожне значення в реченні розглядалося б окремо, це призвело б мережу до швидкого перенавчання та її здатність розпізнавання тональності тексту була б точна тільки на навчальній вибірці.

Існують так само розгортаючі нейронні мережі (deconvolutional networks DN), так само звані зворотними графічними мережами і є зворотними до CNN (рис. 2.5) [52].

Їх завдання генерувати зображення по заданих ознаках. Наприклад, при передачі мережею слова "кіт" вона повинна буде згенерувати зображення схожі на котів.

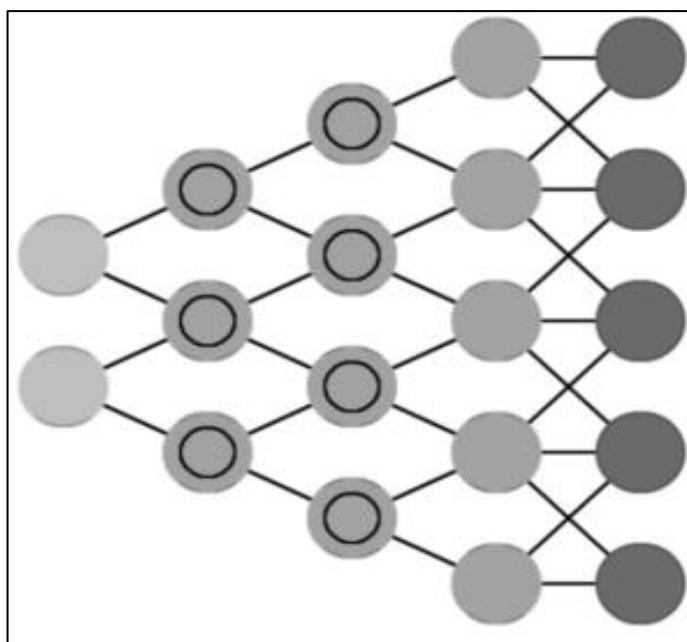


Рис.2.5. Розгортаюча нейронна мережа

Для передбачення слів у реченні, передбаченні наступного числа в послідовності, виділення головної думки тексту, генерації нової інформації схожою на дану використовуються рекурентні нейронні мережі (recurrent neural network RNN).

У рекурентних нейронних мережах нейрони обмінюються інформацією між собою, наприклад, в добавок до нового шматочку входять даних нейрон так само отримує інформацію про попередньому стані мережі (рис. 2.6). Таким чином, в

мережі реалізується так звана «пам'ять», що принципово змінює характер її роботи і дозволяє аналізувати будь-які послідовності даних [53].

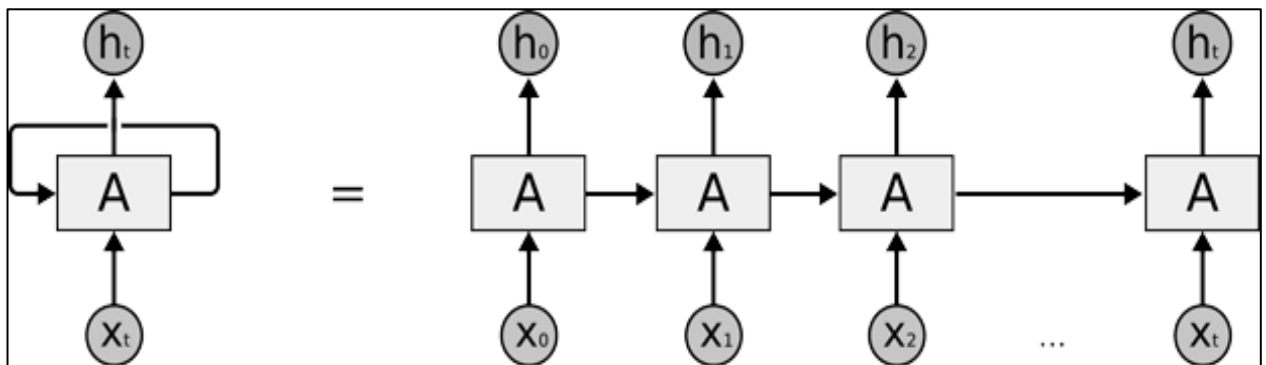


Рис.2.6. Рекурентна нейронна мережа

Найпростішими нейронними мережами, є мережі прямого поширення (feedforward neural network FNN), які взяті за основу для створення багатьох інших мереж [54]. Дані мережі прямолінійні і передають інформацію від входу до виходу. Нейрони кожного шару не пов'язані між собою, а сусідні шари зазвичай повністю пов'язані. Вид FNN можна розглядати як багатозарову мережу (рис. 2.7).

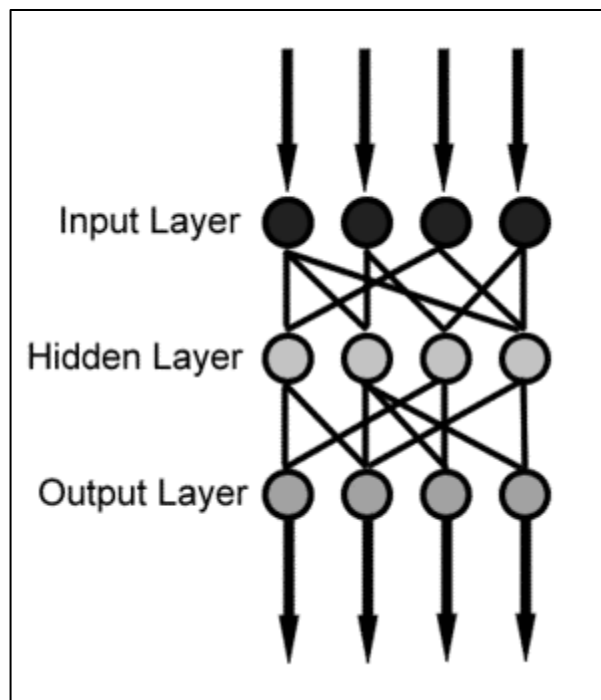


Рис.2.7. Мережа прямого поширення (feedforward neural network FNN)

Нейронні мережі Feedforward складаються з наступного:

- вхідний шар: цей шар складається з нейронів, які отримують входи і передають їх іншим шарам. Кількість нейронів у вхідному шарі має дорівнювати атрибутам або особливостям в наборі даних;
- вихідний шар є прогнозованою функцією і залежить від типу моделі, яка створюється;
- прихований шар: між вхідним і вихідним шаром існують приховані шари, засновані на типі моделі. Приховані шари містять величезну кількість нейронів, які застосовують перетворення до входів перед їх проходженням. У міру навчання мережі ваги оновлюються, щоб бути більш прогнозованими;
- ваги нейронів: відносяться до сили або амплітуди зв'язку між нейронами (від 0 до 1);

В мережах такого виду немає зворотних зв'язків. Прикладом нейронної мережі прямого поширення є перцептрон Розенблатта, від якого і беруть свій початок нейромережі прямого розповсюдження. В літературі часто термін перцептрон, багат шаровий перцептрон та нейромережа прямого поширення застосовуються синонімічно

2.3 Переваги застосування штучних нейронних мереж у кібербезпеці

Очевидно, що свою силу ШНМ черпають, по-перше, з розпаралелювання обробки інформації, по-друге, зі здатності до самонавчання. Ці властивості дозволяють їм вирішувати масштабні задачі, котрі на сьогоднішній день вважаються важкими для розв'язання.

У сучасному світі кібератак, які набувають все більшої ваги у суспільстві і державі, застосування ШНМ значно покращує інформаційний захист кіберпростору. ШНМ мають кілька переваг у сфері кібербезпеки. Деякі з них:

1. Виявлення невідомих кіберзагроз: Всім знаємо, що неможливо виявити всі кіберзагрози з боку людей. Існують мільйони кібератак з різними намірами.

Невідомі загрози можуть завдати мережі величезної шкоди. Вони можуть мати гірший вплив на мережу, перш ніж їх вдасться ідентифікувати, виявити та запобігти. ШНМ є однією з найкращих технологій для виявлення та припинення таких атак невідомих загроз.

2. Обробка великих обсягів даних: Велика активність компаній середнього розміру має величезний трафік, що свідчить про те, що в мережі що призводить до щоденної передачі великої кількості даних між клієнтами та компанією. Тому виникає необхідність захисту даних від зловмисників і програмного забезпечення. ШНМ роблять це, створюючи алгоритми для автоматичного виявлення загроз безпеці. Вони охоплюють широкий спектр елементів ІТ-мережі, включаючи спільні файли, електронні листи або відвідувані сайти. ШНМ обробляють дані більш ретельно за короткий час. Для цього використовується програмне забезпечення безпеки зі штучним інтелектом із більш надійним процесором. Такі процесори швидко сканують величезні дані і за лічені хвилини визначають аномалії, загрози та пропонують рішення.

3. Рівень помилок значно знижений в порівнянні з людськими зусиллями, оскільки ШНМ не втомлюються і не відволікаються від виконання вказаного набору завдань. Вони швидко виявляють будь-які кіберзагрозливі фактори, допомагаючи співробітникам більше зосередитися на інших завданнях, а не на загрозах.

4. Забезпечення кращого керування уразливими місцями: Управління вразливістю є головною функцією захисту мережі компанії. ШНМ швидко оцінює систему, тим самим збільшуючи можливості вирішення проблем у кілька разів. Вони визначають слабкі місця в мережі безпеки та допомагають підприємствам зосередитися на них.

5. Загрози, з якими стикаються бізнес-мережі, час від часу змінюються. З часом хакери змінюють свою тактику, щоб зламати мережу безпеки. Це ускладнює визначення пріоритетів завдань безпеки. Інколи доводиться мати справу з різними атаками одночасно: фішингові атаки, програми-вимагачі разом із атакою відмови в

обслуговуванні. Цьому можна запобігти, розгортаючи ШНМ із завданням комплексного виявлення усіх атак.

6. Особлива перевага ШНМ у покращенні часу виявлення та відповіді.

7. Покращення виявлення та автоматизації кіберзагроз: Автоматизація виявлення кіберзагроз означає, що можна швидко знайти зв'язки між потенційними ризиками та швидко діяти. ШНМ використовують аргументи, які дають змогу ідентифікувати підозрілі посилання, файли або загрози в реальному масштабі часу.

8. Захист аутентифікації. Доступні різні бізнес-сайти, які вимагають входу користувача, заповнення форм та здійснення онлайн-платежів. Такі підприємства потребують додаткового рівня безпеки. ШНМ роблять процес аутентифікації більш безпечним, використовуючи техніку фізичного розпізнавання образів, контурів обличчя, сітківки ока, сканування відбитків пальців і т.д..

Однак на практиці при автономній роботі нейронні мережі не можуть забезпечити готові рішення. Їх необхідно інтегрувати в складні системи.

Використання нейронних мереж забезпечує наступні корисні властивості систем [55]:

- нелінійність (nonlinearity). ШНМ можуть бути лінійними та нелінійними. Нейронні мережі, які побудовані із з'єднань нелінійних нейронів, являються нелінійними. Нелінійність є дуже важливою властивістю, особливо якщо сам фізичний механізм, що відповідає за формування вхідного сигналу, також являється нелінійним (наприклад, аудіо-дані);

- відображення вхідної інформації у вихідну (input-output mapping). Одною із популярних парадигм навчання є навчання із вчителем (supervised learning). Під цим розуміється зміна синаптичних ваг на основі маркованих навчальних прикладів (training sample). Кожен приклад складається із вхідного сигналу і відповідного йому бажаного відклику (desired response). З цієї множини випадковим чином вибирається приклад, а нейронна мережа модифікує синаптичні ваги для мінімізації розбіжностей бажаного вихідного сигналу і формованого

мережею згідно вибраного статистичного критерію. Це відбувається до тих пір, поки зміни синаптичних ваг не стануть незначними;

- адаптивність (adaptivity). Нейронні мережі володіють здатністю адаптувати свої синаптичні ваги до змін оточуючого середовища. Чим вищі адаптивні властивості системи, тим стійкішою буде її робота в нестаціонарному середовищі;

- очевидність відповіді (evidential response). В контексті задачі класифікації зображень можна розробити нейронну мережу, яка збирає інформацію не тільки для визначення конкретного класу, але і для збільшення достовірності прийнятого рішення. В результаті ця інформація може використовуватися для виключення сумнівних рішень, що підвищить продуктивність нейронної мережі.

- контекстна інформація (contextual information). Знання представляються в самій структурі нейронної мережі за допомогою її стану активації. Кожен нейрон мережі потенційно може бути схильний до впливу всіх інших її нейронів. Як наслідок, існування нейронної мережі безпосередньо пов'язане з контекстною інформацією.

- відмовостійкість (fault tolerance). При несприятливих умовах продуктивність нейронних мереж падає незначно. Враховуючи розподілений характер зберігання інформації в нейронній мережі, тільки серйозні пошкодження структури ШНМ можуть значно вплинути на її роботу.

- масштабованість (VLSI Implementability). Паралельна структура нейронних мереж потенційно прискорює вирішення деяких задач і забезпечує масштабованість нейронних мереж в рамках технології VLSI (very-large-scaleintegrated). Одною з переваг цієї технології є можливість представити достатньо складну поведінку з допомогою ієрархічної структури.

- одноманітність аналізу і проектування (Uniformity of analysis and design). Нейронні мережі являються універсальним механізмом обробки інформації. Це означає, що одне і те ж рішення нейронної мережі може використовуватися в багатьох предметних областях;

- аналогія з нейробіологією (Neurobiological analogy). Будова нейронних мереж визначається аналогією з людським мозком. В свою чергу, нейробіологи розглядають ШНМ як засіб моделювання фізичних явищ.

Також, ШНМ добре підходять для вирішення задач класифікації, оптимізації і прогнозування.

Розглянемо деякі проблеми, які вирішуються за допомогою використання нейронних мереж, що представляють інтерес для користувачів:

Класифікація образів. Завдання полягає у визначенні приналежності вхідного образу (наприклад, мовного сигналу чи рукописного символу), представленого вектором ознак, одному або декільком попередньо визначеним класам. Наприклад розпізнавання букв, розпізнавання мови, класифікація сигналу електрокардіограми, класифікація клітин крові.

Кластеризація (категоризація). При вирішенні задачі кластеризації, яка відома також як класифікація образів «без вчителя», відсутня навчальна вибірка з мітками класів. Алгоритм кластеризації заснований на подібності образів і розміщує близькі образи в один кластер. Відомі випадки застосування кластеризації для стиснення даних і дослідження властивостей даних.

Апроксимація функцій. Припустимо, що є навчальна вибірка (пари даних вхід-вихід), яка генерується невідомою функцією $F(x)$, спотвореною шумом. Завдання апроксимації полягає в знаходженні оцінки невідомої функції $F(x)$. Апроксимація функцій необхідна при рішенні численних інженерних і наукових задач моделювання.

Прогнозування. Завдання полягає в передбаченні деякого значення у деякий майбутній момент часу. Передбачення мають значний вплив на прийняття рішень у бізнесі, науці і техніці. Передбачення цін на фондовій біржі і прогноз погоди є типовими програмами прогнозування.

Оптимізація. Численні проблеми в математиці, статистиці, техніці, науці, медицині та економіці можуть розглядатися як проблеми оптимізації. Завданням алгоритму оптимізації є знаходження такого рішення, який задовольнить

обмеження системи і максимізує чи мінімізує цільову функцію. Задача комівояжера є класичним прикладом задачі оптимізації.

Таким чином перераховані переваги та властивості ШНМ доказують нагальну необхідність застосування їх підрозділами кібербезпеки для захисту інформаційного простору, організації превентивних заходів розвідки з метою попередження кіберзагроз.

Висновки до 2 розділу

Досліджено обчислювальну структуру нового типу – штучних нейронних мереж. Дослідження спричинене низкою успішних застосувань цієї нової технології, яка дозволила розробити ефективні підходи до вирішення проблем, що вважалися складними для реалізації на традиційних комп'ютерах. На назву “нейронні мережі” зараз претендують усі обчислювальні структури, які в тій чи іншій мірі моделюють роботу мозку. Штучний нейрон спочатку імітує властивості, що дані біологічному нейрону.

Підкреслено активність нейрона, яка повністю визначається його параметрами - вагами і його функцією активації. Існує безліч передавальних функцій, що застосовуються на практиці при розробці нейронних мереж, деякі з них служать для реалізації нелінійності системи. Вибір тієї чи іншої функції активації часто залежить від умов завдання і структури мережі. Так кожна нейронна мережа має не тільки свою архітектуру, а й тип, який краще підходить для вирішення конкретного завдання.

Наприклад, ШНМ згорткового типу набагато краще справляється з розпізнаванням образів, проблемами комп'ютерного зору та аналізом тональності тексту. За результатами аналізу принципів побудови нейронних мереж, завдань, які можуть виконувати ШНМ, обробляючи величезні масиви інформації, їх необхідно використовувати при розробленні моделей для застосування у кіберрозвідці при виявленні інформаційного впливу у кіберпросторі.

РОЗДІЛ 3

МЕТОДИКА ЗАСТОСУВАННЯ СЕНТИМЕНТНОГО АНАЛІЗУ ДЛЯ ВИЯВЛЕННЯ ІНФОРМАЦІЙНОГО ВПЛИВУ

3.1 Згорткові нейронні мережі як інструмент вирішення задач аналізу тональності тексту

Штучний інтелект і машинне навчання – це величезні галузі і сьогодні всіма передовими державами докладаються великі зусилля, щоб розширити їх застосування для вирішення багатьох проблем реального світу.

Одним із видів ШНМ є згорткова нейронна мережа (ЗНМ).

Згорткова нейронна мережа (ConvNet або CNN) – це штучна нейронна мережа (ANN), яка використовує алгоритми глибокого навчання для аналізу даних і їх класифікації.

CNN використовує принципи лінійної алгебри, такі як множення матриць, для виявлення шаблонів у зображенні. Оскільки ці процеси включають складні обчислення, вони потребують графічних процесорів (GPU - graphics processing unit) для навчання моделей.

Простими словами, CNN використовує алгоритми глибокого навчання, щоб приймати вхідні дані, і призначати важливість у формі упереджень і вагових коефіцієнтів різним аспектам цих даних. Початком для розроблення CNN слугували когнітрон і неокогнітрон.

Когніторон був розроблений на основі будови біологічної зорової кори, має ієрархічну принципово багат шарову архітектуру. Нейрони між шарами когнітрону пов'язані тільки локально і кожен шар реалізує різні рівні узагальнення. Вхідні шари сприймають прості образи, такі як лінії, великі однорідні ділянки, їх орієнтацію і локалізацію в просторі вхідних даних, в той час як глибокі шари

сприймають складніші абстрактні структури, незалежні від локалізації та інших простих ознак образу.

Когнітрон організовується з ієрархічно пов'язаних збудливих і гальмуючих шарів. Співвідношення збуджуючих і гальмуючих сигналів на вході нейрона визначає його стан збудження. Існують спрощені моделі когнітрону, що будуються з одновимірних шарів, але спочатку когнітрон конструювався як каскад двовимірних шарів. Пресинаптичний простір сигналів визначає виходи попереднього шару, постсинаптичний простір - входи наступного шару або площини [55].

Нейрон когнітрону сприймає не весь постсинаптичний простір сигналів, а тільки його частину, реалізується принцип локальної зв'язності. Область пресинаптичного простору сигналів, що утворюють постсинаптичний простір сигналів, що впливають на стан даного нейрона, називається його локальним рецептивним полем. Рецептивні поля близьких один до одного постсинаптичних нейронів, звані зонами конкуренції, перекриваються, тому активність даного пресинаптичного нейрона позначається на все більше поширення області постсинаптичних нейронів наступних шарів ієрархії. Розміри зон конкуренції обумовлюють кількість сприймання в просторі області ознак.

Неокогнітрон – згортова ієрархічна, багатошарова штучна нейронна мережа, заснована на принципах навчання без учителя. Розроблена професором Фукусімою у 1980 році і являє собою подальший розвиток когнітрону. Модель заснована на дослідях в галузі нейробіології і являє собою спрощену модель організації зорової кори ссавців.

Неокогнітрон є прямим розвитком ідеї, що лежать в основі когнітрону і точніше моделює структуру зорової кори головного мозку і є класифікатором, здатним до кращого розпізнавання об'єктів. Кожен шар неокогнітрона (рис.2.1) складається з площини простих S-нейронів і площини складних C-нейронів, що також організують локальну зв'язність. Локальне рецептивне поле на площині S-нейронів наступного шару формується пресинаптичними сигналами площини C-нейронів попереднього шару. Локальні ознаки способу сприймаються S-

нейронами, а спотворення локальних ознак компенсуються С-нейронами. В результаті цього процесу кожен шар після вхідного має своїм входом все більш узагальнену картину, утворену С-нейронами попередніх шарів. З кожним рівнем глибини первинні прості ознаки визначаються у більш складних з'єднаннях. Площину S-нейронів можна розглядати як один нейрон, ваги якого визначають ядро згортки, що застосовуються до попереднього шару у всіх можливих позиціях. Всі С-нейрони реагують на образ, відповідний ядру згортки, в їх рецептивному полі, тому він визначається інваріантно до його локалізації [55].

Сучасні глибокі ЗНМ засновані на ідеях, що лежать в основі неокогнітрона, і сьогодні застосовуються для вирішення широкого кола завдань: від промислових, корпоративних та дослідницьких до повсякденно побутових, включаючи завдання, які вирішуються мобільними пристроями.

Нейронні мережі згорткового типу (рис. 3.1) - використовують шари зі згортаючими фільтрами. Спочатку були винайдені для комп'ютерного зору. Проте, згодом стало відомо, що вони досягли чудових результатів під час аналізу тональності тексту.

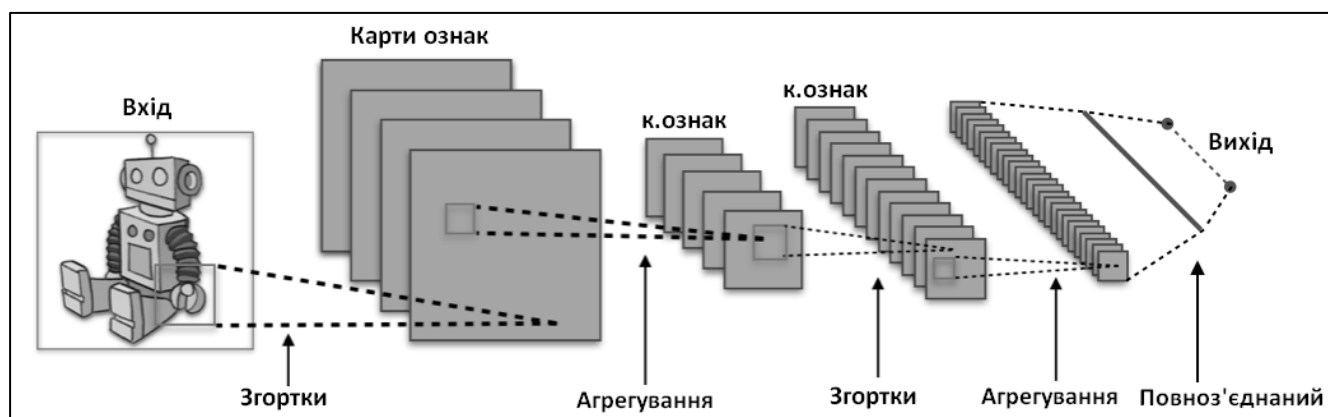


Рис. 3.1. Типова архітектура нейронної мережі згорткового типу

На відміну від мережі прямого поширення, де кожен вхідний нейрон з'єднується з вихідним нейроном в наступному шарі, в згорткових мережах для отримання вихідних значень застосовуються згортки над кожним вхідним шаром. В операції згортки використовується матриця ваг невеликого розміру, яка

проходить по всьому поточному шарі, формуючи після кожного зсуву сигнал активації для нейрона наступного шару з аналогічною позицією. Ця матриця називається ядром згортки, вона використовується для різних нейронів вихідного шару.

Головним шаром ЗНМ беззаперечно можна вважати шар згортки, робота якого є основою даної мережі. Параметри шару згортки складаються з набору фільтрів для навчання. Кожен фільтр має невеликі просторові габарити (ширину і висоту), але проходить по всій глибині вхідного об'єму. Наприклад, стандартний фільтр першого шару ЗНМ може бути розміром $[5 \times 5 \times 3]$. Під час проходження вперед, виконується ніби ковзання кожного фільтра по ширині і висоті вхідних даних і обчислюється скалярний добуток між записами фільтра і входом у будь-яке положення. Мережа навчає фільтри, що активуються. Потім мережа працює з цілим набором фільтрів в кожному шарі згортки, і кожен з цих фільтрів буде формувати окрему двовірну карту активації або карту ознак. Карти ознак складаються вздовж «глибини» вхідного об'єму і будуть формувати вихідний об'єм [56].

При роботі з вхідною інформацією високої розмірності, установка зв'язку між нейронами і всіма нейронами з попереднім об'ємом є недоцільною. Замість цього потрібно застосовувати кожен нейрон тільки до локальної області вхідного об'єму. Просторова протяжність його зв'язку є гіперпараметром і називається рецептивним полем (полем сприйнятливості). Ідея рецептивного поля полягає в тому, щоб об'єднати нейрони прихованого шару лише з такими нейронами попереднього шару, які входять в деяку маленьку область $n \times n$. Причому для кожного нейрона вибирається своя окрема область. Дана область і називається локальним рецептивним полем. Легко здогадатися, що площа рецептивного поля дорівнює площі фільтра. Просторова протяжність уздовж осі глибини завжди еквівалентна глибині вхідного об'єму. Кожен нейрон в згортковому шарі з'єднаний тільки з локальною вхідною областю.

Нейрони згорткових нейромереж залишаються незмінними: вони як і звичайні нейрони мереж обчислюють скалярний добуток їх ваги і вхідних даних з

подальшою нелінійністю, але їх підключення обмежується локальними просторовими вимірами [56].

Згорткові нейронні мережі схожі на мережі прямого поширення типу персептрон: нейрони даного типу також мають ті, яких навчають ваги і параметр зсуву. Однак ЗНМ беруть до уваги те, що вхідні дані становлять собою сукупність текстових даних, володіючи даною інформацією, можна закодувати певні властивості в архітектурі мережі, що дозволяє значно зменшити кількість параметрів нейронної мережі. Поруч із згортковими нейронними мережами часто згадують, рекурентні нейронні мережі тому варто розуміти їх відмінності.

Згорткова нейронна мережа:

- мережа має вхід фіксованого розміру і генерує виходи фіксованого розміру;
- відноситься до типу штучної нейронної мережі зі зворотним зв'язком - варіації багат шарових персептронів, які призначені для використання мінімальних обсягів попередньої обробки;
- використовують схему взаємодії між нейронами подібною до організації зорової кори тварин;
- підходять у застосуванні для обробки зображень, аналізу тексту й мови.

Рекурентна нейронна мережа:

- мережа може обробляти довільні розміри введення / виведення;
- на відміну від родинних нейронних мереж - може використовувати свою внутрішню пам'ять для обробки довільних послідовностей входів;
- рекурентні нейронні мережі використовують інформацію про часові ряди.
- підходять у застосуванні для аналізу тексту й мови.

Як правило будь-яка архітектура мережі має свої певні недоліки та переваги. Одним з методів вирішення проблеми недоліків конкретної архітектури – гібридизація мереж [57, 58]. Тут мається на увазі суміщення деяких архітектур в одну аби покращити ті чи інші якості.

Також можна зустріти і гібридизацію з не нейромережевими моделями, тобто зі зовсім іншими методами розпізнавання образів. Головним принципом при побудові гібридної архітектури являється використання сильної сторони тієї чи

іншої моделі та повного уникання слабких сторін. Іншими словами гібридні нейронні мережі являють собою об'єднання декількох нейронних мереж для вирішення певних завдань. Це дозволяє виконати розбиття задачі з високою складністю на більш прості підзадачі, а архітектура нейронної мережі може бути оптимізована під конкретну задачу.

Ще однією перевагою гібридних нейронних мереж є те, що при програмній реалізації швидкодія може варіюватися в широких межах залежно від обраної структури, що дозволяє застосовувати їх для вирішення завдань реального часу. Гібридні нейронні мережі об'єднують кілька різних за структурою і своїй природі «експертів», тому кожен з них може генерувати різні відповіді для одного і того ж вхідного набору даних. Для формування загальної відповіді гібридної нейронної системи найчастіше використовуються зважене підсумовування результатів або динамічний вибір однієї нейронної мережі, і використання отриманого нею результату в якості вихідного значення.

3.3 Алгоритми і правила навчання нейронної мережі

У ШНМ процес навчання може розглядатися як налаштування архітектури мережі і ваг зв'язків для ефективного виконання різних типів задач. Нейронна мережа повинна налаштувати ваги зв'язків по наявній навчальній вибірці. Функціонування мережі поліпшується у міру ітеративного налаштування вагових коефіцієнтів. Властивість мережі навчатися на прикладах робить їх більш привабливими в порівнянні з системами, які слідують певній системі правил функціонування, сформульованих експертами.

Для конструювання процесу навчання, перш за все, необхідно мати модель зовнішнього середовища, в якій функціонує нейронна мережа – знати доступну для мережі інформацію. Також необхідно визначити, як модифікувати вагові параметри мережі – які правила навчання управляють процесом налаштування. Алгоритм навчання означає процедуру, в якій використовуються правила навчання

для налаштування ваг. Найважливішою властивістю нейронних мереж є їх здатність навчатися на основі даних навколишнього середовища і в результаті навчання підвищувати свою продуктивність. Підвищення продуктивності відбувається з часом у відповідності з певними правилами.

Навчання нейронної мережі відбувається за допомогою інтерактивного процесу корегування синаптичних ваг і порогів. В ідеальному випадку нейронна мережа отримує знання про навколишнє середовище на кожній ітерації процесу навчання. З поняттям навчання асоціюється досить багато видів діяльності, тому складно дати цьому процесу однозначне визначення. Більше того, процес навчання залежить від точки зору на нього. Саме це робить практично неможливим появу будь-якого точного визначення цього поняття.

Наприклад, процес навчання з точки зору психолога докорінно відрізняється від навчання з точки зору шкільного вчителя. З позицій нейронної мережі, ймовірно, можна використовувати наступне визначення: навчання – це процес, у якому вільні параметри нейронної мережі настраюються за допомогою моделювання середовища, у яке ця мережа вбудована. Тип навчання визначається способом підстроювання цих параметрів. Це визначення процесу навчання нейронної мережі передбачає наступну послідовність подій:

- у нейронну мережу надходять стимули із зовнішнього середовища;
- у результаті реалізації першого пункту змінюються вільні параметри нейронної мережі;
- після зміни внутрішньої структури нейронна мережа відповідає на штрафи вже іншим чином.

Вищевказаний список чітких правил вирішення проблеми навчання нейронної мережі називається алгоритмом навчання. Не існує універсального алгоритму навчання, відповідного для всіх архітектур нейронних мереж. Існує лише набір засобів, представлений безліччю алгоритмів навчання, кожен з яких має свої переваги. Алгоритми навчання відрізняються один від одного способом налаштування синаптичних ваг нейронів.

Ще однією відмінною характеристикою є спосіб зв'язку навченою нейронної мережі із зовнішнім світом. У цьому контексті говорять про парадигму навчання, пов'язану з моделлю навколишнього середовища, в якій функціонує дана нейронна мережа. Розрізняють два способи навчання:

- з вчителем;
- без вчителя;

Процес навчання з учителем передбачає пред'явлення мережі вибірки навчальних прикладів. Кожен зразок подається на вхід нейронної мережі, потім проходить процедуру обробки всередині ШНМ. Після обчислення вихідного сигналу ШНМ порівнює отриманий результат з відповідним значенням цільового вектору, що представляє собою необхідний вихід мережі. Обчисливши помилку, відбувається зміна вагових коефіцієнтів зв'язків всередині мережі за обраним алгоритмом. Ваги підлаштовуються під кожен вектор до тих пір, поки помилка по всьому масиву вхідних даних не досягне заданого рівня.

Навчання без вчителя не вимагає знання правильної відповіді на кожен приклад навчальної вибірки. В цьому випадку розкривається внутрішня структура даних або кореляція в системі даних, що дозволяє розподілити образи за категоріями. Залежно від розв'язуваної задачі в навчальній вибірці використовуються ті чи інші типи даних і різні розмірності вхідних / вихідних сигналів.

Вхідні дані прикладів навчальної вибірки - зображення, таблиці чисел, розподілення. Типи вхідних даних - бінарні (0 і 1), біполярні (-1 і 1) числа, цілі або дійсні числа з деякого діапазону.

Вихідні сигнали мережі - вектори цілих або дійсних чисел.

Для вирішення практичних завдань часто потрібні навчальні вибірки великого обсягу. Через жорстко обмежений обсяг оперативної пам'яті комп'ютера розмістити в ній великі навчальні вибірки неможливо. Тому вибірка ділиться на сторінки - групи прикладів. У кожен момент часу лише одна сторінка прикладів розташовується в пам'яті комп'ютера, інші - на жорсткому диску. Сторінки

послідовно завантажуються в пам'ять комп'ютера. Навчання мережі відбувається по всій сукупності сторінок прикладів, по всій навчальній вибірці.

В даний час відсутня універсальна методика побудови навчальних вибірок. Набір навчальних прикладів формується за бажанням користувача програми моделювання нейронних мереж індивідуально для кожної конкретної задачі, що вирішується.

Якщо в ненавчену нейронну мережу ввести вхідний сигнал одного з прикладів навчальної вибірки, то вихідний сигнал мережі буде істотно відрізнятись від бажаного вихідного сигналу, визначеного в навчальній вибірці. Функція помилки чисельно визначає подібність всіх поточних вихідних сигналів мережі і відповідних бажаних вихідних сигналів навчальної вибірки. Найбільш поширеною функцією помилки є середньоквадратичне відхилення.

Для навчання нейронних мереж можуть бути використані різні алгоритми. Можна виділити дві великі групи алгоритмів - градієнтні і стохастичні. Градієнтні алгоритми навчання мереж засновані на обчисленні приватних похідних функції помилки за параметрами мережі. Серед градієнтних розрізняють алгоритми першого і другого порядків. В стохастичних алгоритмах пошук мінімуму функції помилки ведеться випадковим чином.

При навчанні мереж, як правило, використовується один з двох наступних критеріїв зупинки: зупинка при досягненні деякого мінімального значення функції помилки, зупинка в разі успішного вирішення всіх прикладів навчальної вибірки.

Перед навчанням виконується ініціалізація нейронної мережі, тобто присвоювання параметрам мережі деяких початкових значень. Як правило, ці початкові значення - деякі малі випадкові числа. Для формування навчальних вибірок, ініціалізації і навчання в програмах моделювання нейронних мереж використовуються спеціальні процедури. Можливість використання багатосторінкового навчання є дуже важливою при вирішенні практичних завдань за допомогою нейронних мереж, що моделюються на звичайних комп'ютерах.

Навчання - це ітераційна процедура, яка при реалізації на звичайних комп'ютерах, вимагає значного часу. Алгоритми навчання істотно розрізняються

по швидкості збіжності. Одною з найважливіших характеристик програм для моделювання нейронних мереж є швидкість збіжності алгоритму (або алгоритмів) навчання, які реалізовані в програмі.

Теорія навчання розглядає три фундаментальні властивості, пов'язаних з навчанням за прикладами: ємність, складність зразків і обчислювальна складність.

Під ємністю розуміється, скільки зразків може запам'ятати мережа і які функції і межі прийняття рішень можуть бути на ній сформовані.

Складність зразків визначає число навчальних прикладів, необхідних для досягнення здатності мережі до узагальнення. Занадто мала кількість прикладів може викликати "перенавчання" мережі, коли вона добре дає собі раду на прикладах навчальної вибірки, але погано - на тестових прикладах, що підлягають тому ж статистичному розподілу.

Відомі 4 основні типи правил навчання: корекція за помилками, машина Больцмана, правило Хебба і метод зворотного поширення помилки.

Правило корекції за помилками [59]. При навчанні з учителем для кожного вхідного прикладу заданий бажаний вихід d . Реальний вихід мережі може не збігатися з бажаним. Принцип корекції за помилками при навчанні полягає у використанні сигналу $(d-y)$ для модифікації ваг, що забезпечує поступове зменшення помилки. Навчання має місце тільки в тому випадку, коли персептрон помиляється. Відомі різні модифікації цього алгоритму навчання.

Навчання Больцмана [60]. Являє собою стохастичне правило навчання, яке впливає з інформаційних теоретичних і термодинамічних принципів. Метою навчання Больцмана є таке налаштування вагових коефіцієнтів, при якому стан видимих нейронів задовольняє бажаний розподіл ймовірностей. Навчання Больцмана може розглядатися як спеціальний випадок корекції за помилками, в якому під помилкою розуміється розбіжність кореляцій станів в двох режимах.

Правило Хебба [61]. Найстарішим навчальним правилом є постулат навчання Хебба. Хебб спирався на наступні нейрофізіологічні спостереження: якщо нейрони з обох сторін синапсу активізуються одночасно і регулярно, то сила синаптичного зв'язку зростає. Важливою особливістю цього правила є те, що зміна синаптичної

ваги залежить тільки від активності нейронів, які пов'язані з даним синапсом. Це істотно спрощує ланцюги навчання.

Для того щоб перевірити навички, набуті нейронною мережею в процесі навчання, використовується імітація функціонування мережі. У мережу вводиться деякий сигнал, який, як правило, не збігається ні з одним з вхідних сигналів прикладів навчальної вибірки. Далі аналізується отриманий вихідний сигнал мережі. Тестування навченої мережі може проводитися на одиночних вхідних сигналах, або на контрольній вибірці, яка має структуру, аналогічну навчальній вибірці.

Метод зворотного поширення помилки (Backprop) [62]. Метод є класичним методом навчання з учителем, заснований на правилі корекції помилок і був розроблений як метод навчання багатoshарового персептрона. Основна ідея полягає в поширенні сигналів помилки після її обчислення на виході мережі в напрямку, зворотньому прямому поширенню сигналів під час звичайного обчислювального процесу, від виходів мережі до її входів. При зворотньому проході синаптичні ваги налаштовуються з метою мінімізації помилки. Фактично, реалізується стохастичний градієнтний спуск, відбувається рух в багатовимірному просторі ваг до мінімуму помилки в сторону, протилежну градієнту. В процесі навчання циклічно знаходяться рішення на одно-критеріальні завдання оптимізації. Для можливості реалізації методу передавальна функція нейронів повинна бути диференційована. Оскільки метод є модифікацією класичного методу градієнтного спуску, то може розглядатися як градієнтний спуск по поверхні помилки [63].

3.3 Методика застосування сентиментного аналізу для виявлення інформаційного впливу соціальних мереж

Використовуючи метод зворотного поширення помилки (Backprop) та враховуючи результати дослідження в попередніх розділах, створений алгоритм програми, суть якого наступна:

- основа алгоритму – це список емоційних слів. Кожному слову в списку присвоюється значення інтенсивності емоції, що виражається ним від 2 до 5;
- в аналіз вбудований навчальний алгоритм для того, щоб оптимізувати значення емоцій у словнику, розміченому людиною;
- є також алгоритм виправлення правопису. Програма автоматично видаляє літери, що повторюються більше двох разів, видаляє подвоєну літеру, якщо таке поєднання у мові зустрічається рідко, видаляє подвоєну літеру в слові з помилкою, якщо після виправлення вийде нормативне слово без помилки (приклад - ппривітт);
- програма використовує список підсилювальних слів, які можуть збільшувати або зменшити тональність. Значення кожного слова підвищується на 1 або 2 одиниці (наприклад, із прислівниками) або зменшується на 1 одиницю (зі словом "якось");
- використовується список слів, які висловлюють заперечення. Вони змінюють значення тональності на протилежне. Розпізнавання випадків, коли присутність негативних слів не змінює тональність, було вбудовано, оскільки такі випадки траплялися рідко в пробному наборі даних;
- використовується список із емотиконами. Оскільки емотикони більш незалежні від предметної області, ніж слова, їх можна використовувати як відносно надійний показник тональності в реченнях. Англійські слова «emotion» і «icon» утворюють термін emoticon. Короткі рядки символів, літер або цифр призначені для відображення міміки та пози. Основні з них наведені в таблиці 4.1

– Таблиця 4.1

Перелік емотиконів, які використовуються в моделі

Ікона										Емодзі	Значення
:~)	:~:]	:~>	8-)8)	:~:}	:o)	:v)	:^)	=]	=)	😊☐😊😊😊😊😊	Усміхнене, щасливе обличчя

Ікона							Емодзі		Значення			
:D :D	8-D 8Д	=D	=3	Б^Д	с:	С:			Сміючись, великий оскал, посміхаючись в окулярах ¹			
x-D xD	X-D XD								Сміється			
:-))									Дуже щасливе або друге підборіддя			
:(:(	:- с:с	:-< :<	:-[:[	:-	:{	: @	:(	;(			Хмурий, сумний,	
:'-(:'(	:=(								Плач			
:'-):' :'-):'	:"Д								Сльози щастя			
>:(>:(	>:[								Сердитий			
D-' D-'	Д:<	D:	Д8	D;	D=	DX				Жах, відраза, смуток, велике збентеження (справа наліво)		
:O :O	:- o:o	:- 0:0	8-0	>:O	=O =o =0					Здивування, шок, позіхання		
:~* :*	:x								Поцілунок			
;-) ;)	*-) *)	;-] ;]	;-^) ;>	:-, ;-	;D	;3				Підморгування, усмішка		
:-П :P	X-P XP	x-p xp	:-p :p	:- P:P	:- p:p	:- 6:6	d:	=p	>:P			нахабний/грайливий,

Ікона										Емодзі	Значення
:~:/	:~.	>:\	>:/	:\	=/	=\	:L	=L	:S	☐☐😞😞	Скептичний, роздратований, не визначився, неспокійний, нерішучий
:- :										😐😐	Пряме обличчя без виразу, нерішучість
:\$	://) ://3									😐😐😐	Збентежений
O:-)O:)	O:-3 O:3	O:-)O:)	O;^)							😐👼	Ангел, ореол, святий невинний
>:-) >:)	}:~) }:~)	3:-)3:)	>:-) >:)	>:3 ;3						😼	Зло, диявольське
; ~)	~O	B~)								😎😐	Круто, нудно, позіхаючи
%~) %)										😐😐👉😐😐	П'яний, розгублений
:-###.. :###..										☐😐😐	Бути хворим [†]
';:-	';:-									☐	Скептицизм, зневіра, несхвалення
:E										😐	Гримасуючий, нервовий, незграбний

- у запитаннях ігноруються негативні слова;
- якщо в реченні зустрічаються емоційні слова, які вимагають виправлення (з повторюваною більше двох разів буквою), у таких слів значення тональності підвищується на 1 одиницю. У повідомленнях це вважається звичайним засобом вираження емоцій та сили емоцій. Однак одна зайва буква вважається друкарською помилкою;
- якщо пропозиція закінчується на знак оклику, значення тональності підвищується на 2 одиниці.

Враховуючи результати дослідження, методика виявлення інформаційного впливу [64] за допомогою застосування сентиментного аналізу контенту в соціальних мережах є наступною.

1. Визначення мети (теми) сентиментного аналізу (тематичного напрямку).
2. Підготовка (збір) контенту соціальних мереж (створення фалів для моделі).
3. Формування набору вхідних даних (бази даних):
 - тональних словників
 - створення списку емоційних слів;
 - створення списку слів, які підвищують або понижають тональність,
 - створення заперечувальних слів;
 - створення списку емотиконів.
 - контрольних словників для навчання моделі (навчальної вибірки);
4. Застосування алгоритму виправлення правопису.
5. Векторизація признаков навчальної вибірки і присвоєння словам ваги
6. Вибір алгоритму класифікації і навчання класифікатора.
7. Перевірка контрольних даних експертами і оцінювання узгодженості висновків експертів .
8. Побудова моделі сентиментного аналізу за тематичним напрямком.
9. Вибір алгоритму навчання моделі.
10. Навчання моделі.
11. Збереження навченої моделі в базі для використання в майбутньому за тематичним напрямком

Висновки до 3 розділу

Таким чином, згорткові нейронні мережі знайшли своє застосування не тільки під час вирішення задач розпізнавання образів, але й для обробки текстів та мови з метою визначення їх емоційного забарвлення.

Для визначення тональності тексту та його класифікації згорткові нейронні мережі краще впораються з такими завданнями, ніж рекурентні, за рахунок того, що операція згортки чергується з операцією підвибірки. Підвибірка застосовується для зменшення загального розміру вхідних даних і збільшення ступеня інваріантності застосовуваних до них фільтрів. Суттю підвибірки є вибір максимального нейрона з декількох сусідніх. Потім, цей нейрон приймаємо за елемент наступної, але вже зменшеної карти ознак. Дана операція допомагає збільшити інваріантність до масштабу вхідних даних. Після проходження декількох шарів карта ознак вироджується в вектор або скаляр, і таких карт ознак стає сотні.

Визначена методика виявлення інформаційного впливу і розглянуті технічні аспекти побудови моделі сентиментного аналізу тексту із застосування мови програмування Python на основі ЗНМ. Сформульована гіпотеза ефективності застосування моделі сентиментного аналізу підрозділами кібербезпеки при визначенні потенційних загроз інформації в соціальних мережах підтверджена аналізом висновків фахівців у літературних джерелах

Таким чином, для визначення емоційного забарвлення контенту соціальних мереж з метою виявлення та попередження інформаційного впливу і підвищення ефективності управління кібербезпекою потрібно розробляти моделі сентиментного аналізу на основі ЗНМ та проводити їх навчання за кожним тематичним напрямком, основою яких є тематичні словники.

РОЗДІЛ 4

РОЗРОБКА МОДЕЛІ НА ОСНОВІ НЕЙРОННОЇ МЕРЕЖІ ТА ЇЇ НАВЧАННЯ ДЛЯ ЗДІЙСНЕННЯ АНАЛІЗУ ТОНАЛЬНОСТІ ТЕКСТУ В СОЦІАЛЬНИХ МЕРЕЖАХ

4.1 Обґрунтування вибору мови програмування

Для розробки програмного продукту була обрана мова програмування Python. Python - інтерпретована об'єктно-орієнтована мова програмування високого рівня зі строгою динамічною типізацією. Розроблена в 1990 році Гвідо ван Россумом. Структури даних високого рівня разом із динамічною семантикою та динамічним зв'язуванням роблять її привабливою для швидкої розробки моделей, а також як засіб поєднування наявних компонентів. Python підтримує модулі та пакети модулів, що сприяє модульності та повторному використанню коду. Інтерпретатор Python та стандартні бібліотеки доступні як у скомпільованій, так і у вихідній формі на всіх основних платформах.

В мові програмування Python підтримується кілька парадигм програмування, зокрема: об'єктно-орієнтована, процедурна, функціональна та аспектно-орієнтована.

Серед основних її переваг можна назвати переваги, які, крім того, відповідають визначеним у першому розділі критеріям [65]:

- чистий синтаксис (для виділення блоків слід використовувати відступи);
- переносність програм (що властиве більшості інтерпретованих мов);
- масштабованість;
- стандартний дистрибутив має велику кількість корисних модулів (включно з модулем для розробки графічного інтерфейсу);
- можливість використання Python в діалоговому режимі (дуже корисне для експериментування та розв'язання простих задач);

- стандартний дистрибутив має просте, але разом із тим досить потужне середовище розробки, яке написане на мові Python;
- зручний для розв'язання математичних проблем (має засоби роботи з комплексними числами, може оперувати з цілими числами довільної величини, у діалоговому режимі може використовуватися як потужний калькулятор);
- відкритий код (можливість редагувати його іншими користувачами).

Python має ефективні структури даних високого рівня та простий і ефективний підхід до об'єктно-орієнтованого програмування. Елегантний синтаксис Python, динамічна обробка типів, а також те, що це інтерпретована мова, роблять її ідеальною для написання скриптів та швидкої розробки прикладних програм у багатьох галузях на більшості платформ.

Інтерпретатор мови Python і багата стандартна бібліотека (як вихідні тексти, так і бінарні дистрибутиви для всіх основних операційних систем) можуть бути отримані з сайту Python www.python.org, і можуть вільно розповсюджуватися. Цей самий сайт має дистрибутиви та посилання на численні модулі, програми, утиліти та додаткову документацію. Інтерпретатор мови Python може бути розширений функціями та типами даних, розробленими на C чи C++ (або на іншій мові, яку можна викликати із C). Python також зручна як мова розширення для прикладних програм, що потребують подальшого налагодження.

4.2 Результат навчання нейронної мережі

Для навчання нейронної мережі був обраний алгоритм RMSprop. RMSprop слід розглядати як адаптацію алгоритму Rprop для міні-пакетного навчання.

Rprop - алгоритм, який використовується для повної пакетної оптимізації в нейронних мережах. Rprop намагається вирішити проблему того, що градієнти можуть сильно відрізнятися за величиною. Деякі градієнти можуть бути крихітними, а інші - величезними, що призводить до дуже складної проблеми - намагаючись знайти єдиний глобальний рівень навчання для алгоритму. Якщо використовувати повноцінне навчання, тоді можна впоратися з цією проблемою,

лише використовуючи знак градієнта. Такми чином, можна гарантувати, що всі оновлення ваг однакового розміру. Rprop поєднує ідею лише використання знака градієнта з ідеєю адаптації розміру кроку індивідуально для кожної ваги.

Отже, замість того, щоб дивитися на величину градієнта, потрібно розглянути розмір кроку, визначений для конкретної ваги. І цей розмір кроків адаптується індивідуально з часом, так навчання буде прискорюватись у потрібному напрямку дослідження.

Для регулювання розміру кроку деякої ваги використовується наступний алгоритм:

1. Спочатку розглядають ознаки двох останніх градієнтів для ваги.
2. Якщо вони мають однаковий знак, це означає, дослідження іде в правильному напрямку, і слід прискорити його невеликою часткою, тобто необхідно збільшити розмір кроку мультиплікативно (наприклад, у 1,2 рази). Якщо вони різні, це означає, що зроблено занадто великий крок і перестрибнули локальні мінімуми, тому слід зменшити розмір кроку мультиплікативно (наприклад, на 0,5).

3. Потім потрібно обмежити розмір кроку між деякими двома значеннями. Ці значення дійсно залежать від програми та набору даних, хороші значення, які можуть бути за замовчуванням, становлять 50-й та мільйонний, що є гарним початком.

4. Тепер можна застосувати оновлені ваги.

Rprop насправді не працює з дуже великими наборами даних, тому потрібно проводити оновлення міні-пакетних ваг. Причина цього в тому, що він порушує центральну ідею стохастичного градієнтного спуску, і коли існує достатньо мала швидкість навчання, Rprop усереднює градієнти на наступних міні-партіях. За допомогою Rprop збільшується вага в 9 разів і зменшується лише один раз, тому вага збільшується набагато швидше.

Щоб поєднати надійність Rprop (просто використовуючи знак градієнта), потрібно дивитися на Rprop з іншої точки зору. Rprop еквівалентний використанню градієнта, але також ділиться на величину градієнта, тому отримується однакова

величина, незалежно від того, наскільки великий цей градієнт. Проблема міні-партій полягає в тому, що потрібно кожен раз ділити на різні градієнти, то чому б не зробити число, на яке ділиться, подібним для сусідніх міні-партій? RMSprop є збереженням ковзаючого середнього квадратичного градієнта для кожної ваги. В подальшому градієнт ділиться на квадратний корінь середнього квадрата. Тому його називають RMSprop (середній квадрат кореня). З математичної точки зору це виглядає так:

$$E[g^2]_t = \beta E[g^2]_{t-1} + (1 - \beta) \left(\frac{\delta C}{\delta w} \right)^2 \quad (4.1)$$

$$w_t = w_{t-1} - \frac{\eta}{\sqrt{E[g^2]_t}} \frac{\delta C}{\delta w} \quad (4.2)$$

де $E[g]$ - ковзна середня площа градієнтів.

dC / dw - градієнт функції витрат відносно ваги.

η - рівень навчання.

β - параметр ковзного середнього (хороше значення за замовчуванням - 0,9).

Як видно з наведеного рівняння, відбувається адаптація швидкості навчання, шляхом ділення на корінь градієнта квадрата, але оскільки у нас є лише оцінка градієнта на поточну міні-партію, замість цього потрібно використовувати ковзну середню. Значення за замовчуванням для ковзного середнього параметра становить 0,9. Даний алгоритм працює дуже добре для більшості програм.

Програмний продукт (ДОДАТОК А, Б) призначений для виконання аналізу тональності тексту. Являє собою нейронну мережу згорткового типу. Працює на базі безкоштовного віртуального середовища Google Colaboratory (рис. 4.1.)

Набір даних, що використовувався для навчання та перевірки роботи мережі, містить 31 185 російськомовних публічних постів з соціальної мережі Facebook (data.csv). Даний набір даних представлений у вільному доступі на сайті <http://text-machine.cs.uml.edu>, звідки його успішно завантажили на персональний Google диск для подальшого використання.

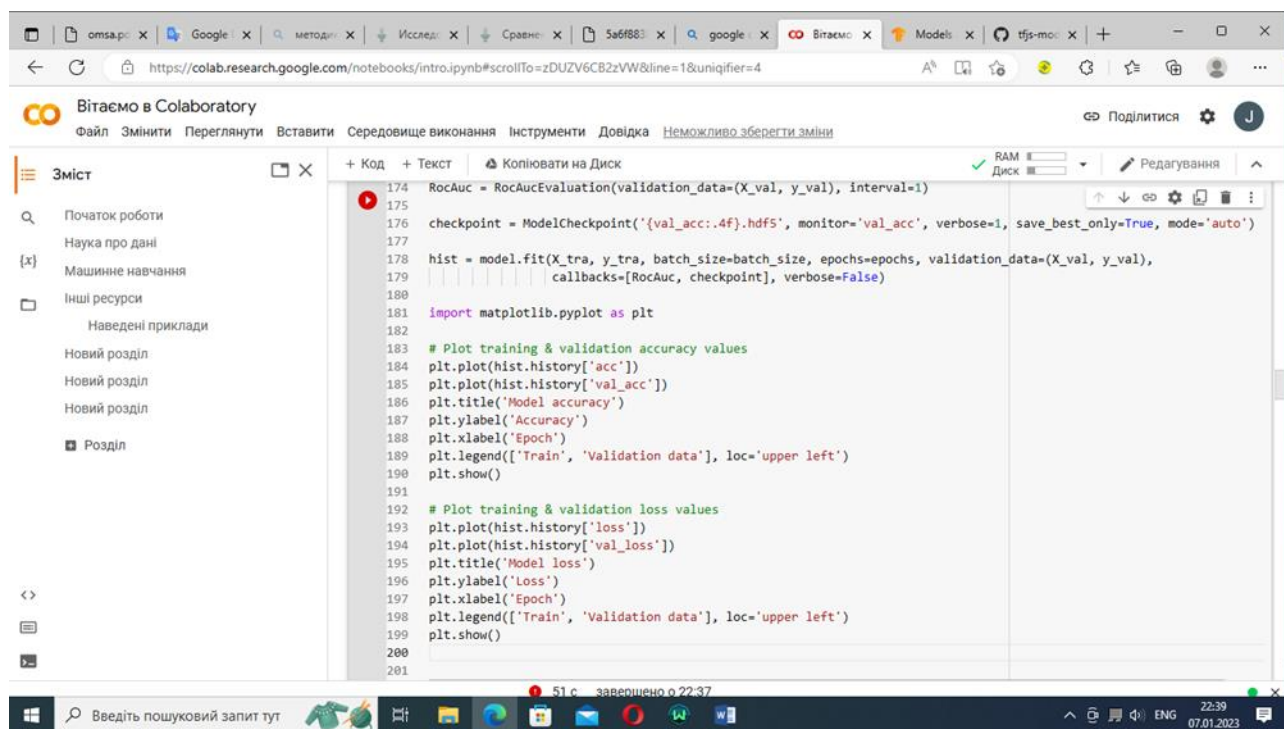


Рис. 4.1. Вікно завантаженої програми в Google Colaboratory

Для навчання мережі дані розділені на тренувальні (23370) та тестові (7790). Навчання нейронної мережі – найперше, що потрібно зробити для того аби вона якомога точніше вирішувала поставлену задачу.

Отже, ми подаємо ці тренувальні дані на нашу модель, так як, не можна подавати сирий текст безпосередньо в моделі глибокого навчання, текстові дані повинні бути попередньо приведені до векторної форми. Це приведення виконується технологією Tokenizer.

Для нормалізації даних ми здійснили embedding за допомогою технології FastText. Таким чином, зменшили кількість даних для навчання мережі.

Далі текстові дані вже у векторній формі перемножуються на ваги нейронів, сумуються, проганяються через функцію активації для нормалізації вхідних даних, наприклад якщо на вході велике число, то попустивши його через функцію активації, на виході отримаємо число у потрібному нам діапазоні. Для своєї мережі ми обрали функцію активації ELU.

ELU - це функція, яка, як правило, швидше переходить до нуля і дає більш точні результати. На відміну від інших функцій активації, ELU має додаткову константу альфа, яка повинна бути додатною цифрою [63].

Слід зауважити, що нейрони вхідного шару лише передають сигнал, що на них подається, нейронам прихованого шару. Таких прихованих шарів може бути доволі багато, залежно від того яку задачу вирішує мережа, у нашому випадку таких шарів 4. У цих шарах і відбувається основний об'єм роботи. Далі нейрони прихованих шарів передають оброблений сигнал нейронам вихідного шару. Вихідний шар ще раз оброблює сигнал та видає результат. Таку процедуру ми виконували упродовж 20 епох (разів). На 7-й епосі ми отримали найкращу точність даних за допомогою метрики якості роботи (AUC). Якщо перевищити кількість епох, мережа перенавчиться і буде видавати помилковий результат. Для того, щоб уникнути цього, використана технологія Dropout.

Dropout - це підхід до регуляризації в нейронних мережах, який допомагає зменшити взаємозалежне навчання серед нейронів (рис. 4.2).

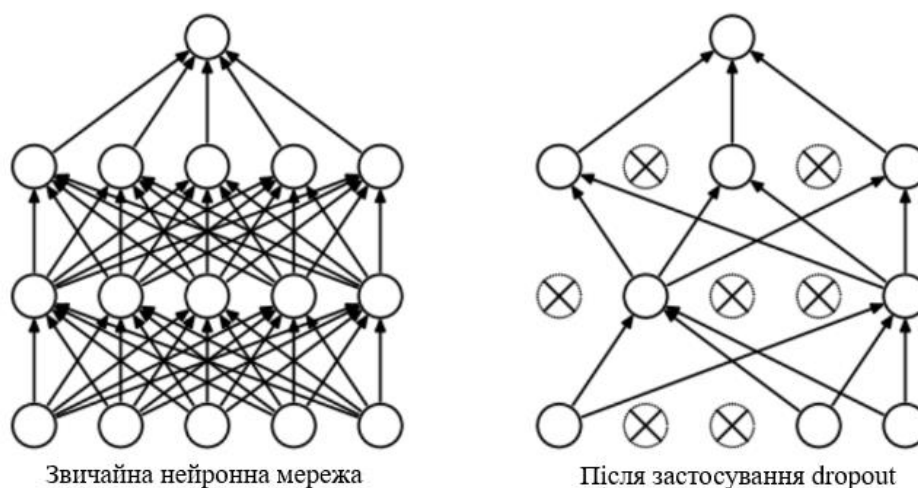


Рис.4.2. Застосування технології Dropout

Після навчання мережі ми здійснювали її тестування на іншій частині даних.

Процес тестування полягав у порівнянні результатів, котрі видає мережа з тими, що відомі нам. Якщо похибка допустима - ми зберігаємо коефіцієнти ваг, у протилежному випадку продовжуємо навчання.

Основні бібліотеки, що використовувались:

- numpy - це основний пакет наукових обчислень;
- matplotlib - відома бібліотека для побудови графіків;
- scikit-learn - це безкоштовна бібліотека машинного навчання програмного забезпечення для мови програмування Python;
- бібліотека глибокого навчання Keras пропонує деякі основні інструменти, які допоможуть нам підготувати дані;
- FastText - це бібліотека для здійснення embedding даних. Бібліотека дозволяє створити алгоритми машинного навчання з вчителем або без вчителя для отримання векторних представлень слів. FastText дозволяє працювати з даними, що представлені на 294 мовах.

Для врахування висловлювань в соціальних мережах українською мовою необхідно створити додаткові словники, без яких похибка виявлення інформаційного впливу за допомогою створеної моделі була б великою.

Такі словники необхідні були для врахування підвищеної або понижувальної тональності висловів, заперечувальних слів і часток, запитальних речень.

Приклади таких словників слів з їхніми коефіцієнтами:

а) підвищеної або понижувальної тональності (BoosterWordList):

- абсолютно 2;
- дуже 2;
- злегка 1;
- настільки 1;
- несказанно 3;
- особливо 1;
- просто 1;
- трохи -1;

б) емоційної лексики (EmotionLookupTable):

- агрес* -4;
- агресивн* -3;

- азарт* 3;
- акуратн* 2;
- ангельськ* 3;
- апетитн* 3;
- афіген* 4;
- беззакон* -2;
- безнаді* -3;
- безпереч* 1;
- безперіч* 1;
- мімішн* 3;
- мінорн* -2;
- міф* -2;

в) запитальних слів (QuestionWords):

- де;
- коли;
- навіщо;
- хто;
- чому;
- що;

г) слів, за допомогою яких зміст міняється на протилежний (NegatingWordList):

- без;
- не;
- нема;
- немає;
- ні;
- ніколи.

Висновки до розділу

Розроблена на базі безкоштовного віртуального середовища Google Colaboratory і створена модель сентиментного аналізу являє собою програмний продукт (нейронну мережу згорткового типу), що призначений для здійснення сентиментного аналізу текстів. Для навчання та тестування моделі нейронної мережі було завантажено набір даних, а саме коментарів соціальної мережі Facebook. Навчання здійснювалося протягом 20 епох, у результаті ми досягли точності близько 86%, що є дуже гарним результатом для нейронних мереж.

Таким чином, удосконалена новими словниками модель дозволила врахувати при виявленні інформаційного впливу україномовний контент соціальної мережі.

Досліджено застосування мови програмування Python, яка має ефективні структури даних високого рівня та простий, але ефективний підхід до об'єктно-орієнтованого програмування. Досконалий синтаксис Python, динамічна обробка типів, а також те, що це інтерпретована мова, роблять її ідеальною для написання скриптів та швидкої розробки прикладних програм у багатьох галузях на більшості платформ. Застосований алгоритм для оптимізації ваг RMSprop, як вдосконалена версія алгоритму Backpropagation підтвердив свою ефективність при навчанні моделі сентиментного аналізу для виявлення емоційного забарвлення з метою визначення інформаційного впливу в соціальних мережах.

ВИСНОВКИ

Розкриті сутність, зміст та основні поняття сентиментентного аналізу текстів, як область комп'ютерної лінгвістики, що присвячена автоматичному виявленню оцінок, емоцій людини щодо згадуваних в тексті сутностей, таких як, наприклад, продукт, послуга, подія, організація, особистість і т.д. або виявлення загальної оцінки, коли просто аналізується емоційність висловлювання.

Зроблений аналіз підходів і задач сентиментентного аналізу в кібербезпеці, проаналізовані методи визначення емоційного забарвлення інформації, розміщеної в соціальних мережах, визначені недоліки та переваги застосування сентиментентного аналізу в управлінні кібербезпекою з метою організації превентивних заходів попередження загроз інформаційній безпеці.

Зроблені висновки про необхідність розроблення програмного забезпечення для ефективного застосування підрозділами кібербезпеки сентиментентного аналізу у превентивній роботі виявлення загроз інформаційній безпеці суспільства. Наявність інформації, отриманої з соціальних мереж, дозволяє за допомогою розроблених програмних продуктів на основі штучного інтелекту виявляти неявний негативний вплив для його подальшої нейтралізації.

Визначені принципи побудови моделі сентиментентного аналізу із застосуванням системного аналізу, проаналізована архітектура штучних нейронних мереж, визначені переваги їх застосування у кібербезпеці при організації кіберрозвідки.

Визначена методика виявлення інформаційного впливу і розглянуті технічні аспекти побудови моделі сентиментентного аналізу тексту із застосування мови програмування Python для визначення емоційного забарвлення змісту інформації в соціальних мережах, сформульована гіпотеза ефективності застосування моделі сентиментентного аналізу підрозділами кібербезпеки при визначенні потенційних загроз інформації в соціальних мережах. З метою підвищення ефективності

управління кібербезпекою розглянуто переваги моделей, розроблених на основі штучних нейронних мереж та запропоновано застосування сформованої методики.

Детально було розглянуто метод, заснований на навчанні з учителем, так як він є одним з найпопулярніших підходів аналізу тональності та використовувався мною у практичній частині даної роботи.

Для обґрунтування висновків і підтвердження гіпотез, визначених у третьому розділі, було розроблено програмне забезпечення на мові Python для здійснення сентиментного аналізу текстів. На базі безкоштовного віртуального середовища Google Colaboratory створено програмний продукт (нейронну мережу згорткового типу), що призначений для здійснення сентиментного аналізу текстів. Для навчання та тестування нейронної мережі було завантажено набори даних, які були в загальному доступі на сайті <http://text-machine.cs.uml.edu>. Завдання складалося у визначенні позитивного, негативного, нейтрального забарвлення постів, у яких виявлено відношення до подій в Україні у 2014 році з виключенням виразів подяки, привітань та поздоровлень, які будуть засмічувати результат.

Із 31 185 постів із статистикою Каппа 0,58, що за версією Флейса є задовільним результатом, для навчання мережі дані розділили на тренувальні (23370) та тестові (7790). З метою оптимізації ваг був обраний алгоритм RMSprop (Root Mean Square propagation – середньоквадратичне розповсюдження).

Навчання здійснювалося протягом 20 епох і в результаті було досягнуто точності близько 86%, що є дуже гарним результатом для нейронних мереж та підтверджує валідність отриманої моделі. Її застосування в управлінні кібербезпекою дозволить своєчасно в автоматизованому режимі визначати потенційні загрози на основі моніторингу інформації, розміщеної в соціальних мережах в кіберпросторі.

Основні результати роботи

1. В магістерській роботі вперше запропоновано застосування сентиментного аналізу в кібербезпеці для визначення емоційного забарвлення текстів в соціальних мережах з метою організації превентивних заходів в управлінні кібербезпекою. Робота охоплює методи та підходи, які обіцяють безпосередньо активувати системи пошуку інформації, орієнтовані на емоційне

забарвлення змісту. Увага зосереджена на методах, які спрямовані на вирішення проблем кібербезпеки, пов'язаних із додатками в соціальних мережах, методах що оцінюють настрої.

2. Удосконалена модель сентиментного аналізу, яка розроблена на мові програмування Python.

3. Дістали подальшого розвитку застосування методів контент-аналізу при організації управління кібербезпекою з метою проведення превентивних заходів попередження негативного інформаційного впливу на особистість, суспільство, державу.

4. Проаналізований стан інформаційного захисту в Україні та визначена необхідність застосування штучного інтелекту для побудови моделей сентиментного аналізу в кібербезпеці при аналізі кіберзагроз.

Практичне значення одержаних результатів.

Запропонована модель сентиментного аналізу на основі штучних нейронних мереж дозволить своєчасно виявити негативний вплив контенту соціальної мережі та попередити загрози кіберпростору держави.

Проаналізувавши проведену роботу та дослідження, можна зробити висновок, що поставлені цілі були успішно досягнуті, створена нейронна мережа успішно впоралася з поставленим їй завданням. Методику її створення, навчання слід застосувати при організації заходів кібербезпеки.

Подальші дослідження застосування сентиментного аналізу, як одного із превентивних способів кіберзахисту дозволять своєчасно виявити кіберзагрози, відстежувати настрої та моніторити підготовку терористичної діяльності в кіберпросторі.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Про рішення Ради національної безпеки і оборони України від 14 травня 2021 року «Про Стратегію кібербезпеки України»: Указ Президента України від 26.08.2021р. № 447. [Електронний ресурс] – Режим доступу: <https://zakon.rada.gov.ua/go/n0055525-2>
2. Bing Liu. Sentiment Analysis and Opinion Mining / Liu Bing. – New Jersey - Morgan & Claypool Publishers, 2012. – 167 p.
3. [Liu Yang, Huang Xiangji, An Aijun, Yu Xiaohui: ARSA: a sentiment-aware model for predicting sales performance using blogs. SIGIR 2007: pp. 607-614.
4. O'Connor Brendan, Ramnath Balasubramanyan, Bryan R. Routledge, and Noah A. Smith. From Tweets to Polls: Linking Text Sentiment to Public Opinion Time Series. In Proceedings of the International AAAI Conference on Weblogs and Social Media (ICWSM 2010), 2010.
5. Tumasjan, Andranik, Timm O. Sprenger, Philipp G. Sandner, and Isabell M. Welp. Predicting elections with twitter: What 140 characters reveal about political sentiment. In proceedings of the International Conference on Weblogs and Social Media (ICWSM-2010), 2010.
6. Joshi Mahesh, Dipanjan Das, Kevin Gimpel, and Noah A. Smith. Movie reviews and revenues: An experiment in text regression. In Proceedings of the North American Chapter of the Association for Computational Linguistics Human Language Technologies Conference (NAACL 2010), 2010.
7. Sadikov Eldar, Aditya Parameswaran, and Petros Venetis. Blogs as predictors of movie success. In Proceedings of the Third International Conference on Weblogs and Social Media (ICWSM-2009), 2009.
8. Аркуша Д. О., Щавінський Ю. В., Лисенко А. В. Використання сентиментного аналізу тексту в кібербезпеці. /Актуальні проблеми кібербезпеки: Матеріали Всеукраїнської науково-практичної конференції (м. Київ, 27 жовтня 2022 року). Навчально-науковий інститут захисту інформації, Державний університет телекомунікацій. Київ, – 2022. – С. 200-202. URL: https://dut.edu.ua/uploads/p_2121_20358827.pdf#page=171

9. Bo Pang, Lillian Lee, Shivakumar Vaithyanathan [Thumbs up? Sentiment Classification using Machine Learning Techniques](#) // – 2002. – С. 79–86.
10. Peter Turney Thumbs Up or Thumbs Down? Semantic Orientation Applied to Unsupervised Classification of Reviews // Proceedings of the Association for Computational Linguistics. – 2002. – С. 417–424. – arXiv: LG/0212032
11. Khurshid Ahmad Affective Computing and Sentiment Analysis: Metaphor, Ontology, Affect and Terminology / Ahmad Khurshid – Berlin – Springer Science & Business Media, 2011 – 164 p.
12. Erik Cambria. SenticNet 2: A semantic and affective resource for opinion mining and sentiment analysis // Proceedings of AAAI FLAIRS : конференція. – 2012. – P. 202–207.
13. Про основні засади забезпечення кібербезпеки України: Закон України від 07.11.2017 р. № **2163-VIII**. Дата оновлення: 17.08.2022. URL: <https://zakon.rada.gov.ua/laws/show/2163-19#Text> (дата звернення: 07.01.2023)
14. Сарбасова, А. Н. Исследование методов сентимент-анализа русскоязычных текстов / А. Н. Сарбасова. – Текст : непосредственный // Молодой ученый. – 2015. – № 8 (88). – С. 143-146. – URL: <https://moluch.ru/archive/88/17413/> (дата звернення: 07.01.2023).
15. P. Willett. The Porter stemming algorithm: then and now // Program: Electronic Library and Information Systems. – 2006. – Vol. 40, iss. 3. – P. 219–223.
16. Thelwall, M., Buckley, K., Paltoglou, G. Cai, D., & Kappas, A. Sentiment strength detection in short informal text. Journal of the American Society for Information Science and Technology, 61(12), 2010. pp. 2544-2558.
17. Balahur, A., Steinberger, R., Kabadjov, M., Zavarella, V., Goot, E. v. d., Halkia, M., Pouliquen, B., & Belyaeva, J. Sentiment analysis in the news. In Proceedings 57 of the international conference on language, resources and evaluation, 2010. pp. 2216-2220. Valletta, Malta.
18. Pang B., & Lee L. Sentimental education: Sentiment analysis using subjectivity summarization based on minimum cuts. In Proceedings of ACL 2004, pp. 271-278. New York: ACL Press.

19. Gamon, M., Aue, A., Corston-Oliver, S., & Ringger, E. Pulse: Mining customer opinions from free text (IDA 2005). Lecture Notes in Computer Science, 3646, 2005. pp. 121-132.
20. Pang B., Lee L. Opinion Mining and Sentiment Analysis. Foundations and Trends in Information Retrieval, v.2 n.1-2, January, 2008, pp. 1-135.
21. Riloff, E., Patwardhan, S., & Wiebe, J. Feature subsumption for opinion analysis. Proceedings of the Conference on Empirical Methods in Natural Language Processing, 2006. pp. 440-448.
22. Ю. В. Рубцова. Построение корпуса текстов для настройки тонового классификатора // Программные продукты и системы, 2015, № 1(109), –С.72–78.
23. Erik Cambria, Amir Hussain, Catherine Havasi, and Chris Eckl. Common Sense Computing: from the Society of Mind to Digital Intuition and Beyond // Biometric ID Management and Multimodal Communication Lecture Notes in Computer Science : журнал. – 2009. – P. 252-259.
24. Павлов, Ю. Н. Сравнение методов оценки тональности текста / Ю. Н. Павлов, К. А. Майструк. – Текст : непосредственный // Молодой ученый. – 2016. – № 12 (116). – С. 59-64. – URL: <https://moluch.ru/archive/116/31521/> (дата обращения: 07.01.2023).
25. Bo Pang, Lillian Lee. Seeing stars: exploiting class relationships for sentiment categorization with respect to rating scales // In Proceedings of the 43rd annual meeting of the Association for Computational Linguistics (ACL) : журнал. – 2005. – No. June 25–30. – P. 115–124.
26. Bo Pang, Lillian Lee. A Sentimental Education: Sentiment Analysis Using Subjectivity Summarization Based on Minimum Cuts // Proceedings of the Association for Computational Linguistics (ACL) : журнал. – 2004. – P. 271–278.
27. Stefano Baccianella. Sentiwordnet 3.0: An enhanced lexical resource for sentiment analysis and opinion mining // Proceedings of LREC : конференция. – 2010. – P. 2200–2204.

28. Qiu, G. Opinion Word Expansion and Target Extraction through Double Propagation/ Guang Qiu, Bing Liu, Jiajun Bu, Chun Chen // Computational linguistics. – 2011. – Vol. 37. – No. 1. – Pages 9–27.
29. Wilson, T. Recognizing Contextual Polarity: An Exploration of Features for Phrase-Level Sentiment Analysis / Theresa Wilson, Janyce Wiebe, Paul Hoffmann// Computational linguistics. – 2009. – Vol. 35. – No. 3. – Pages 399–433.
30. Dras, M. Evaluating Human Pairwise Preference Judgments [Оцінювання парних оціночних суджень експертів] / Mark Dras // Computational linguistics. – 2015. – Vol. 41. – No. 2. – Pages 337–345.
31. Табоада, Мейт, Джуліан Брук, Мілан Тофілоскі, Кімберлі Волл та Манфред Стеде (2011) [Методи аналізу настроїв, засновані на лексиконах](#). Обчислювальна лінгвістика 37 (2): 267-307.
32. Hayes, Andrew F. & Krippendorff, Klaus (2007). Answering the call for a standard reliability measure for coding data. Communication Methods and Measures, 1, 2007. pp. 77–89.
33. Бабаскіна І. О. Дослідження методів аналізу тональностей тексту : пояснювальна записка до атестаційної роботи здобувача вищої освіти на другому (магістерському) рівні, спеціальність 121- Інженерія програмного забезпечення / І. О. Бабаскіна ; М-во освіти і науки України, Харків. нац. ун-т радіоелектроніки. – Харків, 2019. – 62 с.
34. Рябишев О. В. Дослідження методів аналізу емоційного забарвлення тексту українською мовою : пояснювальна записка до атестаційної роботи здобувача вищої освіти на другому (магістерському) рівні, спеціальність 121 – Інженерія програмного забезпечення / О. В. Рябишев ; М-во освіти і науки України, Харків. нац. ун-т радіоелектроніки. – Харків, 2021. – 86 с.
35. Зубрицький, А. Ю. Інтелектуальна система дослідження та аналізу текстів: магістерська дис.: 121 Інженерія програмного забезпечення / Зубрицький Андрій Юрійович. - Київ, 2019. - 84 с.

36. Про рішення Ради національної безпеки і оборони України від 14 вересня 2020 року “Про Стратегію національної безпеки України”: Указ Президента України від 14 вересня 2020 року № 392/2020.
37. Kyva, Vladyslav & Sudnikov, Yevhen & Voitko, Oleksandr. (2018). Методи розвідки кіберпростору. Сучасні інформаційні технології у сфері безпеки та оборони. 33. 2018. С. 45-52. <https://doi/10.33099/2311-7249/2018-33-3-45-52> .
38. Козицька, О. Кіберрозвідка як новий напрям оперативно-розшукової діяльності. Матеріали конференцій МЦНД, Жовтень 2020, с. 107-8, url: <https://doi:10.36074/23.10.2020.v1.13> .
39. Сідченко, С. О., Белімов, В. В. & Хударковський К. І. Можлива організаційна структура підрозділу розвідки у кібернетичному просторі Системи озброєння і військова техніка, 3 (7), 2006. с.47-50.
40. Сидченко С. А., Петров, В. Л., Белимов В.В. & Залкин С. В. Основные характеристики разведки информации в кибернетическом пространстве. Радиоэлектронні і комп’ютерні системи. (3), 2006. с. 90-95.
41. Taboada, M. and J. Grieve. Analyzing Appraisal Automatically. American Association for Artificial Intelligence. of AAAI Spring Symposium on Exploring Attitude and Affect in Text. Stanford. March 2004. pp.158-161.
42. Kruse, Rudolf,; Borgelt, Christian; Klawonn, F.; Moewes, Christian; Steinbrecher, Matthias; Held, Pascal. Computational intelligence : a methodological introduction. Springer. 2013. 234-255 pp
43. Kushal Dave, Steve Lawrence, and David M. Pennock. Mining the peanut gallery: Opinion extraction and semantic classification of product reviews. In Proceedings of WWW, pages 519–528, 2003.
44. Qiaozhu Mei, Xu Ling, Matthew Wondra, Hang Su, and ChengXiang Zhai. Topic sentiment mixture: Modeling facets and opinions in weblogs. In Proceedings of WWW, ACM Press New York, NY, USA, 2007. pages 171–180,
45. Martineau Justin, and Finin Tim. Delta TFIDF: An Improved Feature Space for Sentiment Analysis. Third AAAI International Conference on Weblogs and Social Media, May 2009, San Jose CA.

46. Domingos, Pedro & Michael Pazzani. On the optimality of the simple Bayesian classifier under zero-one loss. *Machine Learning*, 29. 1997. :103-137
47. Bruss, F. Thomas. "250 years of 'An Essay towards solving a Problem in the Doctrine of Chance. By the late Rev. Mr. Bayes, communicated by Mr. Price, in a letter to John Canton, A. M. F. R. S.' *Jahresbericht der Deutschen Mathematiker-Vereinigung*, Springer Verlag, Vol. 115, Issue 3-4 (2013), 129-133. url: <https://DOI.10.1365/s13291-013-0077-z>
48. I.Sutskever, J.Martens, G.Dahl, G.Hinton. On the importance of initialization and momentum in deep learning. *J. of Machine Learning Research*, 2013, V. 28, No. 3, pp. 1139-1147.
49. Субботін С.О. Подання й обробка знань у системах штучного інтелекту та підтримки прийняття рішень. Запоріжжя: ЗНТУ, 2008. - 341 с.
50. Руденко О. Г., Бодянський Є. В. Штучні нейронні мережі: Навчальний посібник. - Харків: ТОВ "Компанія СМІТ", 2006. - 404 с
51. Каллан Р. Основные концепции нейронных сетей = The Essence of Neural Networks First Edition. – «Вильямс», 2001. – 288 с.
52. N.Srivastava, G.Hinton, A.Krizhevsky, I.Sutskever, R.Salakhutdinov. Dropout: A Simple Way to Prevent Neural Networks from Overfitting. *J. of Machine Learning Research*, 2014, V.15, pp.1929-1958.
53. Круглов, В.В. Искусственные нейронные сети. Теория и практика: моногр. / В.В. Круглов, В.В. Борисов. – М.: Горячая линия - Телеком; Издание 2-е, стер., 2002. - 382 с.
54. Шмідхубер, Юрген. Глибоке навчання в нейронних мережах: огляд". Нейронні мережі – 61. 2015. с. 85–117. arXiv:1404.7828. url: <https://doi:10.1016/j.neunet.2014.09.003>
55. Бунаков, В. Е. Нейронная физика. Учебное пособие: моногр. / В.Е. Бунаков, Л.В. Краснов. – М.: Издательство Санкт-Петербургского университета, 2015. - 200 с.
56. Ciresan, Dan; Meier, Ueli; Schmidhuber, Jürgen. Multi-column deep neural networks for image classification. 2012 IEEE Conference on Computer Vision and

- Pattern Recognition[en] (New York, NY: Institute of Electrical and Electronics Engineers (IEEE)). 2012. arXiv:1202.2745v1. doi:10.1109/CVPR.2012.6248110.
57. Гніденко В. А. Гібридні нейронні мережі з імунним навчанням для розв’язання задачі прогнозування часових рядів / В. А. Гніденко, В. О. Звонкова // Радіoeлектроніка та молодь в ХХІ столітті : матеріали 24 Міжнар. молодіж. форуму, 7-9 квіт. 2020 р. – Харків : ХНУРЕ, 2020. – Т. 5. – С. 169-170.
 58. Чумаченко, О. І. Структурно-параметричний синтез гібридних нейронних мереж : автореф. дис. ... д-ра техн. наук. : 05.13.23 – системи та засоби штучного інтелекту / Чумаченко Олена Іллівна. – Київ, 2019. – 42 с.
 59. Бодянский Е.В. Искусственные нейронные сети: архитектуры, обучение, применения / Е.В. Бодянский, О.Г. Руденко - Харьков: ТЕЛЕТЕХ, 2004. - 369с.
 60. Golden R.M. Mathematical Methods for Neural Network Analysis and Design. – Cambridge, Massachusetts: The MIT Press, 1996. – 420 p.
 61. Caporale N. “Spike timing-dependent plasticity: a Hebbian learning rule”. Annual Review of Neuroscience. 31: 2008. p.p. 25-46. url: <https://DOI:10.1146/annurev.neuro.31.060407.125639> .
 62. Rumelhart, David E.; Williams, Ronald J. Learning Internal Representations by Error Propagation. In: Parallel Distributed Processing. Cambridge, 1986, MA: MIT Press 1: 1986. p.p. 318-362.
 63. Хайкин С. Нейронные сети. Полный курс / С. Хайкин – М.: Вильямс, 2006 – 1104 с.
 64. Аркуша Д.О. Застосування нейронних мереж для сентиментного аналізу в кібербезпеці. / Проблеми комп’ютерної інженерії: Матеріали III-ї науково-практичної конференції (м. Київ, 01 грудня 2022 року). Навчально-науковий інститут інформаційних технологій, Державний університет телекомунікацій. Київ, – 2022. – 137 с.
 65. Larry Constantine. Software development. StP, Peter, 2004. - 592 p

ДОДАТОК А

```
pip install PyDrive
```

```
import os
```

```
from pydrive.auth import GoogleAuth
```

```
from pydrive.drive import GoogleDrive
```

```
from google.colab import auth
```

```
from oauth2client.client import GoogleCredentials
```

```
import io
```

```
import zipfile
```

```
from google.colab import files
```

```
auth.authenticate_user()
```

```
gauth = GoogleAuth()
```

```
gauth.credentials = GoogleCredentials.get_application_default()
```

```
drive = GoogleDrive(gauth)
```

```
download_RuSentiment = drive.CreateFile({'id': '1ahxKR6n6KpsIkj7YLimcZvR  
OK7n5bYzV'})
```

```
download_RuSentiment.GetContentFile('data.csv')
```

```
download_fasttext = drive.CreateFile({'id': '1EU0d1IO-  
kCK6LVdHcl5yAlXoa_OwG0jF'})
```

```
download_fasttext.GetContentFile('fasttext.zip')
```

```
!unzip 'fasttext.zip'
```

```
Archive: fasttext.zip
```

```
inflating: fasttext.min_count_100.vk_posts_all_443550246.300d.vec
```

```
!ls
```

```
adc.json fasttext.min_count_100.vk_posts_all_443550246.300d.vec sample_data
```

```
data.csv fasttext.zip
```



```

import pandas as pd
import numpy as np

from sklearn.model_selection import cross_val_score
from sklearn.model_selection import train_test_split
from sklearn.metrics import roc_auc_score, accuracy_score

from scipy.sparse import hstack
import matplotlib.pyplot as plt

np.random.seed(42)

from keras.models import Model
from keras.layers import Input, Embedding, Dense, Conv2D, MaxPool2D
from keras.layers import Reshape, Flatten, Concatenate, Dropout, SpatialDropout
1D

from keras.preprocessing import text, sequence
from keras.callbacks import Callback
from keras.callbacks import ModelCheckpoint

from keras.models import load_model
import pickle

import warnings
warnings.filterwarnings('ignore')

import os
%matplotlib inline

data = pd.read_csv('./data.csv')

```

```
train, test = train_test_split(data)
data.head(10)
```

```
print('train shape is', train.shape)
print('test shape is', test.shape)
train shape is (23370, 2)
test shape is (7790, 2)
```

```
def convert_data(data):
    columns=['text', 'neutral', 'negative', 'positive', 'skip', 'speech']
    grouped_df = data.groupby('label')
    df_new = pd.DataFrame()
    df_new['text'] = data['text']

    for c in grouped_df.groups.keys():
        df_new.loc[grouped_df.groups[c], c] = 1

    df_new = df_new.fillna(0)
    return df_new
```

```
train_new = convert_data(train)
test_new = convert_data(test)
train_new = train_new[train_new.skip != 1][train_new.speech != 1][['text','neutral',
'negative', 'positive']]
test_new = test_new[test_new.skip != 1][test_new.speech != 1][['text','neutral', 'negative', 'positive']]
train_new.head(50)
```

```
X_train = train_new["text"].fillna("fillna").values
y_train = train_new[['neutral', 'negative', 'positive']].values
```

```
X_test = test_new["text"].fillna("fillna").values
y_test = test_new[['neutral', 'negative', 'positive']].values
```

```
max_features = 59594
maxlen = 200
embed_size = 300
```

```
tokenizer = text.Tokenizer(num_words=max_features)
tokenizer.fit_on_texts(list(X_train) + list(X_test))
X_train = tokenizer.texts_to_sequences(X_train)
X_test = tokenizer.texts_to_sequences(X_test)
x_train = sequence.pad_sequences(X_train, maxlen=maxlen)
x_test = sequence.pad_sequences(X_test, maxlen=maxlen)
```

```
x_train.shape, x_test.shape
((17393, 200), (5872, 200))
```

```
EMBEDDING_FILE = './fasttext.min_count_100.vk_posts_all_443550246.300d.
vec'
```

```
def get_coefs(word, *arr): return word, np.asarray(arr, dtype='float32')
```

```
embeddings_index = dict(get_coefs(*o.rstrip().rsplit(' ')) for o in open(EMBEDDI
NG_FILE, encoding="utf8"))
```

```
word_index = tokenizer.word_index
nb_words = min(max_features, len(word_index))
embedding_matrix = np.zeros((nb_words, embed_size))
for word, i in word_index.items():
    if i >= max_features: continue
    embedding_vector = embeddings_index.get(word)
```

```

    if embedding_vector is not None: embedding_matrix[i] = embedding_vector
embedding_matrix.shape
(59594, 300)

```

```

class RocAucEvaluation(Callback):
    def __init__(self, validation_data=(), interval=1):
        super(Callback, self).__init__()

        self.interval = interval
        self.X_val, self.y_val = validation_data

    def on_epoch_end(self, epoch, logs={}):
        if epoch % self.interval == 0:
            y_pred = self.model.predict(self.X_val, verbose=0)
            score = roc_auc_score(self.y_val, y_pred)
            print("\n ROC-AUC - epoch: %d - score: %.6f \n" % (epoch+1, score))

filter_sizes = [1,2,3,5]
num_filters = 32

def get_model():
    inp = Input(shape=(maxlen, ))
    x = Embedding(max_features, embed_size, weights=[embedding_matrix])(inp)
    x = SpatialDropout1D(0.4)(x)
    x = Reshape((maxlen, embed_size, 1))(x)

    conv_0 = Conv2D(num_filters, kernel_size=(filter_sizes[0], embed_size), kern
el_initializer='normal',
                                activation='elu')(x)

```

```

conv_1 = Conv2D(num_filters, kernel_size=(filter_sizes[1], embed_size), kernel_initializer='normal',
                activation='elu')(x)
conv_1 = Dropout(rate=0.2)(conv_1)

conv_2 = Conv2D(num_filters, kernel_size=(filter_sizes[2], embed_size), kernel_initializer='normal',
                activation='elu')(x)
conv_2 = Dropout(rate=0.2)(conv_2)

conv_3 = Conv2D(num_filters, kernel_size=(filter_sizes[3], embed_size), kernel_initializer='normal',
                activation='elu')(x)
conv_3 = Dropout(rate=0.2)(conv_3)

maxpool_0 = MaxPool2D(pool_size=(maxlen - filter_sizes[0] + 1, 1))(conv_0)
maxpool_1 = MaxPool2D(pool_size=(maxlen - filter_sizes[1] + 1, 1))(conv_1)
maxpool_2 = MaxPool2D(pool_size=(maxlen - filter_sizes[2] + 1, 1))(conv_2)
maxpool_3 = MaxPool2D(pool_size=(maxlen - filter_sizes[3] + 1, 1))(conv_3)

z = Concatenate(axis=1)([maxpool_0, maxpool_1, maxpool_2, maxpool_3])
z = Flatten()(z)
z = Dropout(rate=0.2)(z)

outp = Dense(3, activation="sigmoid")(z)

model = Model(inputs=inp, outputs=outp)
model.compile(loss='binary_crossentropy',
              optimizer='RMSprop',
              metrics=['accuracy'])

return model

```

```

model = get_model()

batch_size = 256
epochs = 20

X_tra, X_val, y_tra, y_val = train_test_split(x_train, y_train, train_size=0.9, random_state=42)

RocAuc = RocAucEvaluation(validation_data=(X_val, y_val), interval=1)

checkpoint = ModelCheckpoint('{val_acc:.4f}.hdf5', monitor='val_acc', verbose=1, save_best_only=True, mode='auto')

hist = model.fit(X_tra, y_tra, batch_size=batch_size, epochs=epochs, validation_data=(X_val, y_val),
                callbacks=[RocAuc, checkpoint], verbose=False)

import matplotlib.pyplot as plt

# Plot training & validation accuracy values
plt.plot(hist.history['acc'])
plt.plot(hist.history['val_acc'])
plt.title('Model accuracy')
plt.ylabel('Accuracy')
plt.xlabel('Epoch')
plt.legend(['Train', 'Validation data'], loc='upper left')
plt.show()

# Plot training & validation loss values
plt.plot(hist.history['loss'])
plt.plot(hist.history['val_loss'])

```

```
plt.title('Model loss')
plt.ylabel('Loss')
plt.xlabel('Epoch')
plt.legend(['Train', 'Validation data'], loc='upper left')
plt.show()
```

```
y_pred = model.predict(x_test, batch_size=1024)
```

```
score = roc_auc_score(y_test, y_pred)
print("\n ROC-AUC - test - score: %.6f \n" % (score))
ROC-AUC - test - score: 0.852162
```

```
!ls
0.7755.hdf5  adc.json
0.8197.hdf5  data.csv
0.8293.hdf5  fasttext.min_count_100.vk_posts_all_443550246.300d.vec
0.8343.hdf5  fasttext.zip
0.8377.hdf5  sample_data
0.8404.hdf5
best_model = get_model()
best_model.load_weights('0.8404.hdf5')
```

```
y_pred = best_model.predict(x_test, batch_size=1024)
score = roc_auc_score(y_test, y_pred)
print("\n ROC-AUC - test - score: %.6f \n" % (score))
ROC-AUC - test - score: 0.876336
```

ДОДАТОК Б

```

from __future__ import division
from collections import defaultdict
from math import log

def train(samples):
    classes, freq = defaultdict(lambda:0), defaultdict(lambda:0)
    for feats, label in samples:
        classes[label] += 1          # count classes frequencies
        for feat in feats:
            freq[label, feat] += 1   # count features frequencies

    for label, feat in freq:         # normalize features
        freq[label, feat] /= classes[label]
    for c in classes:               # normalize classes frequencies
        classes[c] /= len(samples)

    return classes, freq            # return P(C) and P(O|C)

def classify(classifier, feats):
    classes, prob = classifier
    return min(classes.keys(),      # calculate argmin(-log(C|O))
               key = lambda cl: -log(classes[cl]) + \
                               sum(-log(prob.get((cl,feat), 10**(-7)))) for feat in
               feats))

```