

Founder: Cherkasy State Technological University

Media ID: R30-04616; №1916 from 30.05.2024

EDITORIAL BOARD

O. E. Pchelintseva	<i>Ukraine, Cherkasy State Technological University – Editor-in-Chief</i>
A. A. Barentsen	<i>Netherlands, University of Amsterdam, UVA</i>
D. Weiss	<i>Switzerland, University of Zurich</i>
B. Wiemer	<i>Germany, Mainz University</i>
D. P. Vojvodić	<i>Serbia, University of Novi Sad</i>
O. I. Vovk	<i>Ukraine, Bohdan Khmelnytsky Cherkasy National University</i>
S. Dickey	<i>USA, Kansas, University of Kansas</i>
V. Dubichynskyi	<i>Poland, University of Warsaw</i>
L. M. Koval	<i>Ukraine, Vasyl' Stus Donetsk National University</i>
M. Lazinski	<i>Poland, University of Warsaw</i>
O. Leszczak	<i>Poland, Jan Kochanowski University</i>
H. R. Mehlig	<i>Germany, Kiel University</i>
Y. M. Pletenetska	<i>Ukraine, National Aviation University, Kyiv</i>
V. M. Trub	<i>Ukraine, National Academy of Sciences of Ukraine, Institute of Ukrainian language</i>
V. I. Yarmak	<i>Ukraine, National Academy of Sciences of Ukraine, Potebnia Institute of Linguistics</i>
L. M. Usyk	<i>Ukraine, Cherkasy State Technological University (deputy Editor-in-Chief)</i>

Responsible for issue:

Olena Pchelintseva (Editor-in-Chief), Liudmyla Usyk (deputy Editor-in-Chief),
Dmytro Pidlasnyi (Technical Editor)

Recommended for publishing by a Scientific Council of Cherkasy State
Technological University (Transaction no. 12 from 16 June 2025, no. 7 from 15 December 2025)

Address of the editorial office: blvr. Shevchenka, 460, building 2, off. 210,
Cherkasy, 18006, Ukraine. E-mail: movacom@chdtu.edu.ua

Official website: movacom.chdtu.edu.ua

The journal is covered by the following services: Index Copernicus Journals Master List; Baidu Scholar; CNKI Scholar (China National Knowledge Infrastructure); CNPIEC – cnpLINKer; Dimensions; EBSCO; ERIHPlus; ExLibris; Google Scholar; J-Gate; KESLI-NDSL (Korean National Discovery for Science Leaders); MyScienceWork; Naver Academic; Naviga (Softweco); ProQuest; ReadCube; Semantic Scholar; TDOne (TDNet); WanFang Data; WorldCat (OCLC); X-MOL.

Creative Commons Attribution-NonCommercial 4.0 International License (CC BY NC)

© The authors, 2025

© Editorial, 2025

© Hrebenuk Alona, cover design, 2019

Міжнародний
науковий
журнал**LANGUAGE:****Codification • Competence • Communication****1-2(12-13)/2025**Заснований у серпні 2019 року
Виходить 1 раз на шість місяців**Засновник:** Черкаський державний технологічний університет
Ідентифікатор медіа: R30-04616; рішення №1916 від 30.05.2024**РЕДАКЦІЙНА КОЛЕГІЯ**

О. Е. Пчелінцева	<i>Україна, Черкаський державний технологічний ун-т (головний редактор)</i>
А. А. Барентсен	<i>Нідерланди, Амстердамський ун-т</i>
Д. Вайс	<i>Швейцарія, Цюрихський ун-т</i>
Б. Вімер	<i>Німеччина, Ун-т Майнца</i>
Д. П. Войводіч	<i>Сербія, Новосадський ун-т</i>
О. І. Вовк	<i>Україна, Черкаський національний ун-т ім. Б. Хмельницького</i>
С. Дікі	<i>США, Ун-т Канзасу</i>
В. Дубічинський	<i>Польща, Варшавський ун-т</i>
Л. М. Коваль	<i>Україна, Донецький національний ун-т ім. В. Стуса</i>
М. Лазинські	<i>Польща, Варшавський ун-т</i>
О. В. Лещак	<i>Польща, ун-т ім. Яна Кохановського</i>
Х. Р. Меліг	<i>Німеччина, Кільський ун-т ім. Крістіана Альбрехта</i>
Ю. М. Плетенецька	<i>Україна, Національний авіаційний ун-т</i>
В. М. Труб	<i>Україна, НАН України, Інститут української мови</i>
В. І. Ярмак	<i>Україна, НАН України, Інститут мовознавства ім. О. О. Потебні</i>
Л. М. Усик	<i>Україна, Черкаський державний технологічний ун-т (заступник головного редактора)</i>

Відповідальні за випуск:Олена Пчелінцева (головний редактор), Людмила Усик (заступник головного редактора),
Дмитро Підласий (технічний секретар редакції)Рекомендовано до друку вченою радою Черкаського державного
технологічного університету (протокол №12 від 16.06.2025, №7 від 15.12.2025)**Адреса редакції:** бульв. Шевченка, 460, корп. II, каб. 210, м. Черкаси, 18006, Україна.
E-mail: movacom@chdtu.edu.ua**Офіційний веб-сайт:** movacom.chdtu.edu.ua**Видання представлено:** Index Copernicus Journals Master List; Baidu Scholar; CNKI Scholar (China National Knowledge Infrastructure); CNPIEC – cnpLINKer; Dimensions; EBSCO; ERIHPlus; ExLibris; Google Scholar; J-Gate; KESLI-NDSL (Korean National Discovery for Science Leaders); MyScienceWork; Naver Academic; Naviga (Softweco); ProQuest; ReadCube; Semantic Scholar; TDOne (TDNet); WanFang Data; WorldCat (OCLC); X-MOL.**Creative Commons Attribution-NonCommercial 4.0 International License (CC BY NC)**

© Автори статей, 2025

© Редакція журналу, 2025

© Гребенюк А. Ю., дизайн обкладинки, 2019

THE CONTENTS

From the Editor	<i>Olena Pchelintseva</i>	5
LARGE LANGUAGE MODELS FOR LINGUISTS	<i>Yevheniia Dehtiarova</i> Using large language models for cross-lingual stylistic adaptation of children's fiction	7
	<i>Viktoriia Musiichuk</i> Pilot experiment on automatic detection of Vietnamese media narratives about the Russia-Ukraine war under limited resources	17
DISCURSIVE STUDIES	<i>Oleksandr Kapranov</i> Self-mentions in British prime minister's speeches on climate change	31
TERMINOLOGICAL STUDIES	<i>Liana Goletiani</i> Metaphor in Ukrainian terminology of EU secondary law: a comparative study	44
SCIENTIFIC LIFE	<i>Olena Pchelintseva</i> Intensive advanced training course «Large Language Model for Linguists»	60
ET CETERA (essays, reflections, impressions)	<i>Olena Denga, Sofiia Fiedosieieva</i> The national code of Ukrainian identity in the educational space: artistic practices of extracurricular activities	62

ЗМІСТ

Сторінка головного редактора	<i>Олена Пчелінцева</i>	5
ВЕЛИКІ МОВНІ МОДЕЛІ У ЛІНГВІСТИЧНИХ РОЗВІДКАХ	<i>Євгенія Дегтярьова</i> Використання великих мовних моделей для міжмовної стилістичної адаптації дитячого художнього тексту	7
	<i>Вікторія Мусійчук</i> Пілотний експеримент з автоматичного виявлення в'єтнамських медіанаративів про російсько-українську війну при обмежених ресурсах	17
ДИСКУРСИВНІ ДОСЛІДЖЕННЯ	<i>Олександр Капранов</i> Самозгадки у промовах прем'єр-міністра Великобританії про зміну клімату	31
ТЕРМІНОЛОГІЧНІ ДОСЛІДЖЕННЯ	<i>Ліана Голетіані</i> Метафора в українській термінології вторинного права ЄС: порівняльне дослідження	44
НАУКОВЕ ЖИТТЯ	<i>Олена Пчелінцева</i> Інтенсивний курс підвищення кваліфікації «Великі мовні моделі для лінгвістів»	60
ET CETERA (есеї, роздуми, враження)	<i>Олена Деньга, Софія Федосєєва</i> Національний код українства в освітньому просторі: мистецькі практики позааудиторної діяльності	62

FROM THE EDITOR

Another year has become part of history. A year that has witnessed transformations at every level of life; a year of rethinking both the world and ourselves; a year of unrealised plans and new, truly remarkable opportunities. And, as ever, our journal, despite its strictly linguistic focus, paradoxically reflects the realities unfolding around us: political manipulation, legal and environmental challenges in the EU and worldwide, and the rapid development of artificial intelligence technologies that are penetrating every sphere of academic, social, and personal life. Turning the pages of this issue, readers will immerse themselves in compelling research processes in which language emerges both as a tool for analysing reality and as a sensitive indicator of global change.



**O. E. Pchelintseva,
Editor-in-Chief**

Authors from Norway, Italy, and Ukraine seek to make sense of contemporary linguistic phenomena in a fast-moving world where technology, politics, law, culture, and identity are closely intertwined, forming a complex and multidimensional background for linguistic inquiry. The first thematic section of this issue addresses the interaction between language and cutting-edge technologies. The article by Yevheniia Dehtiarova demonstrates that large language models (LLMs) open up new opportunities for translators, while at the same time foregrounding the challenge of preserving individual authorial style, cultural specificity, and age-appropriate orientation of the text. This study carefully balances optimism about technological progress with a cautious recognition of the indispensable role of the human expert in the creative process. The theme of digital methods in the humanities is further elaborated in Viktoriia Musiichuk's article, which presents the results of a pilot experiment on the automatic detection of Vietnamese media narratives about the Russia-Ukraine war. The proposed hybrid methodology, combining classical narratology with NLP tools, demonstrates how linguistics can effectively respond to contemporary challenges by analysing global information flows and the mechanisms through which the image of war is constructed in foreign-language media spaces. Political discourse as a space for constructing collective meanings and positions is brought into focus in Oleksandr Kapranov's article. Based on speeches by the Prime Minister of the United Kingdom, Keir Starmer, on climate change, the author examines the use of self-mentions, showing how first-person pronouns contribute to shaping images of responsibility, solidarity, or distance between political leaders and their audiences. This study adds an important dimension to our understanding of rhetorical strategies in contemporary politics in the context of global environmental challenges. Legal discourse and issues of specialised translation are represented in Liana Goletiani's article devoted to metaphor in Ukrainian terminology of EU secondary law. This contribution convincingly demonstrates that metaphor in legal discourse is a crucial conceptual tool that requires careful scholarly reflection. The issue is complemented by a popular-science article by Olena Denga and Sofiia Fiedosieieva, which addresses the national cultural code of Ukrainianness in the educational space. Living language, music, theatrical performance, and personal engagement with cultural experience are shown to foster the development of linguistic and national identity among young people. This contribution resonates with particular warmth and symbolic meaning, reminding us of the humanistic mission of education even in the most challenging times.

My sincere gratitude goes to everyone who contributed to the creation of this issue under conditions of power outages, constant air-raid alerts, exhaustion, and the relentless flow of distressing news. The entire editorial team, layout designers, and technical editors worked in what can truly be described as extreme conditions of intense strain. Yet each of us – like all Ukrainians – continues to hold our own “front line,” making a personal contribution to a future Victory, a future Peace, and a future happy Ukraine.

СТОРІНКА ГОЛОВНОГО РЕДАКТОРА

Ще один рік став історією. Рік трансформацій на всіх рівнях життя, рік переосмислення світу та себе, рік нездійснених планів і нових неймовірних можливостей. І як завжди наш журнал – незважаючи на його суто лінгвістичний напрям – парадоксальним чином віддзеркалює те, що відбувається навколо: політичні маніпуляції, правові та екологічні проблеми у ЄС та світі, шалений розвиток технологій штучного інтелекту, що проникає в усі сфери наукового, соціального та особистого життя. Гортаючи сторінки цього випуску, ви зануритесь у неймовірно цікаві дослідницькі процеси, у яких мова постає і як інструмент аналізу реальності, і як чутливий індикатор глобальних змін. Автори з Норвегії, Італії та України намагаються осмислити сучасні мовні явища в умовах швидкоплинного світу, де технології, політика, право, культура й ідентичність тісно переплітаються, утворюючи складне, багатовимірне тло для лінгвістичних досліджень.

Перший тематичний блок випуску звертається до проблем взаємодії мови й новітніх технологій. Стаття Євгенії Дегтярьової переконує, що великі мовні моделі (LLM) відкривають нові можливості для перекладачів, водночас актуалізуючи питання збереження індивідуального авторського стилю, культурної специфіки та вікової адресованості тексту. Ця розвідка тонко балансує між оптимізмом щодо технологічного прогресу й обережним усвідомленням ролі людини-фахівця у творчому процесі. Тему цифрових методів у гуманітарних науках продовжує стаття Вікторії Мусійчук, у якій подано результати пілотного експерименту з автоматичного виявлення в'єтнамських медіанаративів про російсько-українську війну. Запропонована гібридна методологія, що поєднує класичну наратологію з інструментами NLP, демонструє, як лінгвістика може ефективно реагувати на виклики сучасності, аналізуючи глобальні інформаційні потоки та механізми формування образу війни в іншомовному медіапросторі. Політичний дискурс як простір конструювання колективних смислів і позицій опиняється в центрі уваги статті Олександра Капранова. На матеріалі промов прем'єр-міністра Великої Британії Кіра Стармера щодо зміни клімату автор досліджує функціонування самозгадок, показуючи, як за допомогою займенників першої особи вибудовується образ відповідальності, солідарності або дистанції між політичним лідером і аудиторією. Це дослідження додає важливий штрих до розуміння риторичних стратегій сучасної політики в контексті глобальних екологічних викликів. Юридичний дискурс і проблеми спеціального перекладу репрезентовані у статті Ліани Голетіані, що присвячена метафорі в українській термінології вторинного права ЄС. Ця робота переконливо демонструє, що метафора в юридичному дискурсі є важливим концептуальним інструментом, який потребує уважного наукового осмислення. Випуск доповнює науково-популярна публікація Олени Деньги та Софії Федосєєвої, що присвячена національному коду українства в освітньому просторі: живе слово, музика, театралізоване дійство та особистісне переживання культурного досвіду сприяють формуванню мовної й національної ідентичності молоді. Цей матеріал звучить особливо тепло й символічно, нагадуючи про гуманітарну місію освіти навіть у найскладніші часи.

Моя щира вдячність усім, хто доклав своїх зусиль до створення цього номеру в умовах відсутності світла, постійних повітряних тривог, перевантаження та гнітючих новин: весь колектив редакції, перекладачі, макетувальники та технічні редактори працювали фактично в екстремальних умовах граничного напруження. Але кожен за нас, так само як і всі українці, гідно тримає свій «фронт», робить свій внесок у майбутню Перемогу, майбутній Мир, майбутню щасливу Україну.

Олена Пчелінцева

УДК 81'25:82-93-053.2:004.8(44)

DOI: 10.24025/2707-0573.12.2025.343341

Євгенія Дегтярьова



**ВИКОРИСТАННЯ ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ
ДЛЯ МІЖМОВНОЇ СТИЛІСТИЧНОЇ АДАПТАЦІЇ
ДИТЯЧОГО ХУДОЖНЬОГО ТЕКСТУ**

У статті розглянуто потенціал великих мовних моделей (LLM) у міжмовній стилістичній адаптації дитячого художнього тексту. Аналіз зосереджено на перекладі з української на французьку з урахуванням стилістики, інтонації та вікової адресованості. Підкреслено можливості LLM як інструменту креативного перекладу в дитячій літературі.

Ключові слова: великі мовні моделі, дитяча література, інтонація, культурна трансляція, стилістична адаптація, українсько-французький переклад, GPT.

Вступ

Проблема міжмовної стилістичної адаптації художніх текстів для дітей набуває особливої актуальності в умовах активної глобалізації та зростання попиту на якісний переклад літератури. Збереження у перекладі стилю, інтонаційної палітри та культурних кодів оригіналу є надзвичайно важливим для дитячої аудиторії. У зв'язку з розвитком великих мовних моделей (LLM) виникає новий науковий і практичний виклик: чи можуть ці моделі стати ефективним інструментом не лише для автоматизованого перекладу, а й для стилістичної адаптації тексту, зокрема з урахуванням вікових, жанрових та культурних особливостей. Ця проблема безпосередньо пов'язана з важливими завданнями сучасної перекладознавчої науки й практики – зокрема, з удосконаленням інструментів для збереження художньої цінності текстів у процесі міжмовної трансформації.

Теоретичне підґрунтя

Проблема перекладу дитячої літератури осмислюється в контексті загальних теорій художнього перекладу (Lathey, 2015; Oittinen, 2000; Vardanian, 2020), дескриптивних підходів до перекладу (Toury, 1995), функціоналістських моделей (Nord, 1997) та концепцій культурної адаптації (Venuti, 1995). Дослідження міжмовної адаптації художніх текстів (Bassnett & Lefevere, 1998) підкреслює важливість динамічної еквівалентності й урахування культурних кодів. У той же час, у наукових дослідженнях бракує системних студій, що поєднують ці підходи з сучасними технологіями LLM.

© Дегтярьова Є., 2025

Останні публікації (Jiao et al., 2023; Ziems et al., 2023, 2024) демонструють, що великі мовні моделі здатні не лише виконувати буквальний переклад, а й генерувати варіанти тексту, зберігаючи стиль і тон. Однак питання їхнього використання саме для міжмовної стилістичної адаптації дитячих текстів залишається недостатньо розробленим, попри наявність окремих досліджень, присвячених специфіці перекладу дитячої літератури (Chen, Zhang, Li & Wang, 2025; Kong & Macken, 2025), які наголошують на важливості збереження стилю та культурних кодів для дитячої аудиторії.

Мета статті – дослідити потенціал великих мовних моделей для збереження та відтворення стилістичних особливостей у процесі перекладу дитячої художньої прози, а також виявити методичні підходи до оптимізації цього процесу.

Для досягнення мети роботи було сформульовано такі завдання:

- відібрати фрагменти українського дитячого тексту з яскравим стилем (лексика, синтаксис, образність, діалектизми);
- перекласти фрагменти українського дитячого тексту за допомогою LLM (ChatGPT);
- зіставити переклад із оригіналом та проаналізувати переклад з точки зору: збереження стилістичних особливостей (іронія, розмовність, інтонація, просторіччя), адаптації культурних елементів, відповідності дитячій читацькій аудиторії;
- виявити сильні й слабкі сторони LLM;
- обґрунтувати потенціал Instruction tuning / Fine-tuning / Custom prompting.

Методи та матеріал дослідження

У статті використано комплексний підхід до вивчення міжмовної стилістичної адаптації дитячого художнього тексту за допомогою великих мовних моделей, що включає такі основні методи:

1. Аналіз стилістики оригінального тексту – детальний аналіз стилістичних характеристик українського художнього твору, включно з лексикою, синтаксисом, образністю та іншими стилістичними маркерами, що визначають художній стиль та інтонаційну специфіку твору.
2. Метод експериментальної генерації тексту з використанням великих мовних моделей (застосування prompt engineering) – метод отримання перекладів шляхом інтерактивної роботи з моделлю, коригування та адаптації вхідних запитів (промптів).
3. Контрастивний аналіз – порівняння перекладених моделей текстів із оригіналом з метою виявлення відмінностей і відповідностей у стилістичних особливостях, тональності та смислового навантаженні.
4. Метод стилістичного контролю генерації тексту (стилометричний експеримент) – дослідження впливу різних налаштувань і параметрів промптів на стилістичні характеристики генерованих перекладів.
5. Метод оцінювання адаптації і навчання LLM – синтез отриманих результатів із погляду можливостей моделі до стилістичної адаптації, рекомендації щодо її подальшої оптимізації.

Матеріалом дослідження стали фрагменти з твору В. Нестайка «Тореодори з Васюківки» (Нестайко, 2004), які мають яскраву стилістичну своєрідність і культурні маркери.

Результати дослідження

Міжмовна стилістична адаптація – це процес перенесення художнього тексту з однієї мовної системи в іншу із збереженням не лише змісту, а й стилістичних, емоційно-експресивних та культурних особливостей оригіналу (Bassnett & Lefevere, 1998). У випадку дитячої літератури ця задача набуває особливої ваги, адже текст орієнтований на конкретну вікову аудиторію, має власний ритм, гумор, характерні мовні ігри та відображає певне культурне середовище.

Переклад уривків з повісті В. Нестайка «Тореодори з Васюківки» іншою мовою (зокрема, французькою) вимагає врахування таких аспектів:

- тон і стилістика дитячого мовлення (герої твору – школярі, тож текст насичений розмовними конструкціями, вигуками, жартами, характерними для дитячого середовища);

- культурно-мовні реалії (українські реалії (шкільні звичаї, побутові деталі, гумористичні алюзії) часто не мають прямих відповідників у цільовій культурі);

- ритм і звучання (дитячий текст має бути легким, динамічним, доступним для сприйняття й водночас образним).

Міжмовна стилістична адаптація передбачає не буквально відтворення, а гнучке використання мовних засобів у цільовій мові, щоб викликати подібний емоційний та стилістичний ефект у читача. Саме тут залучення великих мовних моделей (LLM) відкриває нові можливості: вони здатні не лише перекласти текст, а й запропонувати варіанти стилістичного перетворення, імітувати дитячу інтонацію, добирати близькі за звучанням чи стилем вирази та створювати культурно релевантні відповідники. Розглянемо це на прикладах.

Чату GPT було запропоновано для перекладу такий уривок: *Де Ява? Що з ним?.. Переекзаменовка...*

*Хвилюватись я почав ще, мабуть, уві сні. Бо коли прокинувся, то **ще не розтуливши очей**, уже відчув, **як холодно в мене під серцем**, як воно то **шалено б'ється-б'ється** то завмирає, неживе, а в цей час від потилиці через спину аж до п'ят біжать-біжать крижані бульбашки, – наче хто обливає мене холодною газованою водою, – і по всьому тілу **гусяча шкіра**.*

*О, як би я хотів не прокидатися сьогодні! Скільки б я дав, щоб уже було завтра! Але ні. Я відчував **неминучість і невблаганність долі**. Таке саме я відчував, коли сидів якимось, чекаючи черги, під дверима зубного лікаря і слухав як **противно дзижчить** отам у кабінеті бормашина а мати міцно тримала мене за руку, щоб я не **дремнув**.*

Чат GPT переклав цей фрагмент таким чином: *Où est Yava ? Qu'est-ce qu'il a ?... Le rattrapage...*

*J'ai commencé à m'inquiéter peut-être déjà dans mon sommeil. Car, en me réveillant, **avant même d'ouvrir les yeux**, j'ai **senti ce froid sous mon cœur**, ses battements **désordonnés**, violents, puis soudain figés, comme si mon cœur n'était plus vivant. En même temps, de ma nuque jusqu'aux talons couraient des bulles glacées – comme si quelqu'un m'arrosait avec de l'eau gazeuse bien froide – et **la chair de poule** me couvrait tout le corps.*

*Oh, comme j'aurais voulu ne pas me réveiller aujourd'hui ! Qu'est-ce que je n'aurais pas donné pour que ce soit déjà demain ! Mais non. Je sentais **l'inéluctabilité, l'implacabilité du destin**. La même chose que j'avais ressentie une fois en attendant mon tour devant la porte du dentiste, en entendant **le bourdonnement désagréable** de la fraise dans le cabinet, et ma mère me tenait fermement la main pour que je ne **m'enfuie** pas.*

Наведений уривок містить експресивні мовні одиниці, частину з яких чат GPT зміг вдало відтворити: *дремнути* – розмовне дієслово, що вживається в образному значенні «швидко втекти, рвонути») – *s'enfuir*; *гусяча шкіра* – усталений вираз для позначення фізіологічної реакції на холод або емоційний вплив (страх, захоплення) – *la chair de poule*; *неминучість і невідворотність долі* – абстрактне поняття, стилістично піднесене, метафора невідворотності життєвих обставин – *l'inéluctabilité, l'implacabilité du destin*. Також адекватно передано звуконаслідування *бормашина* *проти*во *дзижчить* – поєднання побутового слова (*бормашина*, тобто стоматологічний апарат) і звуконаслідування (*дзижчить*), що створює звуковий образ – *le bourdonnement désagréable de la fraise*.

В той же час було втрачено метафоричне значення виразів: *розтулити очі* – усталене словосполучення зі значенням «прокинутися, розплющити очі» – *ouvrir les yeux*; *шалено б'ється серце* – опис сильного емоційного збудження – *ses battements désordonnés*. Після коригування промпту чат GPT надав французькі відповідники зі схожим до оригіналу емоційним та стилістичним ефектом: *je sentais déjà ce froid bizarre au creux de l'estomac; mon cœur battait à toute vitesse*.

Цікавим є рішення чату GPT для стилістичної адаптації кличного відмінка у французькому тексті. Речення *Яво! Яво! Яво! Яво! Яво!..* чат переклав як *Yava ! Yava ! Yava ! Yava ! Yava !...* На зауваження, що не передано кличний відмінок, чат GPT відповів таке: «У французькій мові **кличний відмінок** (як в українському «Яво!») часто виражається через **ім'я в прямому відмінку** (тобто без змін, як у звичайному використанні), але часто для кличного відмінку додається виразна інтонація, яка передається іншими словами або **використанням звертання**.

Для кличного тону можна додати: «Yava ! Yava ! Yava ! Yava ! Yava !...». Простий варіант, що зберігає ім'я без змін. Тут важлива інтонація, яка передає кличний відтінок.

«Yava ! Yava ! Yava ! Yava ! Yava !... Hé, Yava !». «Hé» додає кличний відтінок – так можна звертатися до когось із кличним тоном, підкреслюючи, що це саме клич.

«Yava ! Yava ! Yava ! Yava ! Yava !... Eh, Yava !». «Eh» – це інший варіант, що також передає більш звертальну інтонацію, характерну для кличного відмінка».

Після запитання, як передати інтонацію на письмі, чат запропонував варіанти, які, на нашу думку, цілком підходять для перекладу: «Для того, щоб **передати інтонацію кличного тону** в письмовій формі французькою, можна використати кілька прийомів, які допоможуть передати **емоційний контекст** та **силу звертання**. Способи передати інтонацію на письмі:

1. **Повторення і окличні знаки**: використовувати **повторення і окличні знаки** для підкреслення кличу. Наприклад: «Yava ! Yava ! Yava ! Yava ! Yava !... Hé, Yava !» (тут повторення та **окличні знаки** підсилюють емоційність кличу.)

2. «Hé», «Eh» або «Où es-tu ?»: введення звертання «Hé» або «Eh» додає **кличну інтонацію**. Наприклад: «Yava ! Yava ! Yava ! Yava !... Eh ! Yava ! Où es-tu ?!» (тут «Eh!» додає емоційного звертання, а «Où es-tu?» підсилює почуття хвилювання.)

3. **Використання запитальних виразів або виразів на зразок «Eh bien?»**: щоб передати більш виразну інтонацію. «Yava ! Yava ! Yava ! Yava !

Yava !... Eh bien, Yava, où es-tu ?!» (це підсилює звернення і додає емоційного відтінку.)

4. Звукозміна за допомогою виділених букв (це може передати силу крику): наприклад, замість «Yava!» можна використовувати щось на кшталт «YAAAA-VA ! YAAAA-VA ! YAAAA-VA !...».

5. Множинні знаки оклику: використання множини окличних знаків (як «!») також дає відчуття напруги або кличу. «Yava ! Yava ! Yava !!! Yava !!!».

Розглянемо ще один показовий приклад стилістичної адаптації тексту чатом GPT: Я підбіг до школи тоді, коли **баба Маруся** вже ходила по подвір'ю, **калатала дзвоником**, і, мов не довіряючи йому, вигукнула:

– Дзвоніть! Дзвоніть! Екзамен почався! Почалися екзамені! Дзвоніть!
Я кинувся до неї:

– Ой **бабусю!** Ой не дзвоніть ще хвилиночку! Не дзвоніть! Яви десь нема! **Одну** хвилиночку тільки!

Баба Маруся **скрушно** похитала головою і якось винувато сказала.

– Не можу, серденько! Уже подзвонила. Пізно. **Було б раніш**, синочку. Де ж це він **заповітрився**? От біда! Ну, біжи в клас, а то й ти **не втрапиш**.

Переклад чату GPT: Je revins à l'école au moment où **grand-mère Maroussia** marchait dans la cour **en agitant la cloche**, et comme si elle ne faisait pas confiance au son, elle criait :

– **La cloche ! La cloche ! Les examens ont commencé ! Les examens ont commencé ! La cloche !**

Je me jetai sur elle :

– Oh, **grand-mère ! Attendez une minute, ne sonnez pas encore ! Yava a disparu ! Une seule minute !**

Grand-mère Maroussia secoua **tristement** la tête et dit d'un ton un peu coupable : – Je ne peux pas, mon petit ! J'ai déjà sonné. C'est trop tard. **Il aurait fallu** venir plus tôt, mon fils. Mais **où donc est-il passé ?** Quelle histoire ! Allez, file en classe, sinon tu **vas rater** toi aussi.

Перш за все ми звернули увагу чата на те, що **баба Маруся** це не *grand-mère Maroussia*. На що чат відповів: «Точно! «**Баба Маруся**» – це не просто «grand-mère Maroussia». Це вираз, який має більше побутовий і неформальний відтінок, часто використовується для позначення старшої жінки, не обов'язково бабусі, і може включати певну частину колориту або навіть поваги, характерного для традиційного українського села. Тому, перекладаючи «баба Маруся» французькою, важливо зберегти цей нюанс і, можливо, використати щось на зразок «**mamouchka Maroussia**» або «**baba Maroussia**», щоб передати образ старшої жінки в селі з характерним акцентом на її неформальність і теплоту».

На зауваження щодо того, що *mamouchka* це російське слово, чат відповів: «Так, ти правий! «**Mamouchka**» – це скоріше російський варіант (зниженого або люблячого звертання до бабусі), і він не зовсім підходить для українського контексту. Таким чином, «**baba Maroussia**» буде найбільш автентичним і вірним перекладом, оскільки він зберігає культурний контекст та звучання, близьке до української мови.

Далі ми спитали чат, чи зрозуміють французи слово *baba*. На що він відповів: «Так, слово «**baba**» може бути зрозумілим французам, але з певним нюансом. Воно не є типовим для французької мови, тому важливо враховувати контекст. У французькій мові «**baba**» має кілька значень (наприклад, відомий

десерт), але у контексті «**baba Maroussia**» це може бути сприйнято як дещо екзотичне або незвичне звертання. Однак, французи, ймовірно, зрозуміють це як народне, розмовне або зменшувальне звертання до старшої жінки, особливо якщо ви поясните в контексті тексту або історії. Вони також можуть звикнути до таких виразів через вплив інших мов, зокрема через російську або східноєвропейську культуру. Якщо хочеш бути більш певним, можна додати пояснення, коротке уточнення (анотація, зноска), наприклад: «**baba Maroussia (une femme âgée et chaleureuse, typique d'un village ukrainien)**» («баба Маруся (жінка похилого віку, тепла і характерна для українського села)»). Так французи краще зрозуміють значення цього слова в даному контексті. В результаті, «**baba Maroussia**» можна використовувати, і це буде зрозуміло, хоч і може додавати певну екзотичність або автентичний відтінок.

Щоб перевірити адекватність адаптації тексту, ми попросили чат GPT зробити зворотний переклад, який вийшов таким: *Я підбіг до школи саме в той момент, коли баба Маруся ходила по шкільному подвір'ю з дзвіночком і, наче не довіряючи самому звуку, вигукувала:*

– *Дзвоник! Дзвоник! Почались іспити! Іспити почались! Дзвоник!*

Я кинувся до неї:

– *Ой, бабусю! Не дзвонить ще хвилиночку! Не дзвонить! Яви десь нема!*

Тільки хвилиночку!

Баба Маруся сумно похитала головою і якимось винувано сказала:

– *Не можу, серденько! Уже подзвонила. Пізно. Тобі б раніше, синочку. Де ж він запропастився? От біда! Ну, біжи в клас, а то й ти не встигнеш.*

Як бачимо, випущеними виявилися емоційно марковане дієслово *калатати* і вигук *ой*, натомість відбулося додавання прикметника, який відсутній в оригіналі, *шкільне (подвір'я)*. Також втрачено експресивність деяких лексичних одиниць: *сумно* замість *скрушно*, *тобі б раніше* (було б *раніше*), *запропастився* (заповітрився), *не встигнеш* (не *втрапиш*). До того ж, *дзвоник* став *дзвіночком*, а в мовленні героїні не збережено автентичність її мовлення (просторіччя *Дзвоник! Дзвоник!* передано нормативним *Дзвоник! Дзвоник!*).

З наведених прикладів видно, що переклад GPT загалом зберігає зміст, логіку та настрій оригіналу, але стилістично звучить трохи стриманіше, більш літературно і менш розмовно.

Оригінал має яскраву розмовну інтонацію, дитяче, емоційно безпосереднє мовлення, гумор (наприклад: *гецнешся об щось ліктем, мотонув назад, заглянувши у найглухіші закутки* тощо), тому потреба в стилістичній адаптації твору не підлягає сумніву.

Хоча LLM вдалося точно передати зміст і логічну структуру твору, але помічено достатньо труднощів перекладу:

1. Мовна стилістика: поєднання дитячої наївності та іронії. Оригінал написано мовою дитини – щирою, образною, зі змінами темпу, емоційними сплесками. Передати це французькою – неабияке завдання, оскільки французька літературна мова часто виглядає стриманішою. Французький дитячий наратив менш схильний до фраз типу *А я як...*, *Та він же...*, *Граптом!...*, які притаманні українській розмовній стилістиці.

2. Колоритні діалектизми та просторіччя. Такі елементи, як: *ма'ть* (просторічне *мабуть*), *ех ти*, *мама – депутат!*, *Сопуть учні*, *носами в диктант упхавшись...* викликають труднощі, оскільки потребно передати соціолект або регіональний колорит, не втрачаючи сенсу й гумору.

3. Гумор і сарказм. Гумор у тексті часто базується на абсурді, перебільшенні, дитячому сприйнятті ситуацій. Наприклад: *як булька на воді, щоб ти так усе життя злазила!*.

4. Культурні реалії. *Мама – депутат* – у радянському контексті це не просто посада, а символ соціального статусу. У французькій культурі *député* не викликає такого враження; *екзамен, переекзаменовка* – радянська шкільна система не має прямих відповідників у французькій системі, перездачі у початковій школі не існують.

5. Передача внутрішніх монологів. Репліки типу: *Може, він уже й неживий?; Хоч ріжете – не повірю!* потребують інтонаційної точності, оскільки це емоційні монологи дитини, часто з перебільшеннями, страхами, наївністю.

Тому вважаємо за доцільне застосування таких методів покращення стилістичної точності перекладу за допомогою підказок (prompts) і додаткового навчання (fine-tuning або instruction tuning):

1. Контекстні стилістичні підказки (custom prompts). У prompt потрібно додати інструкцію щодо стилістики, жанру, цільової аудиторії або конкретного стилістичного прийому. Приклад prompt-a: Traduis le texte suivant en français en gardant le ton émotionnel d'un roman jeunesse ukrainien. Respecte le style oral, les expressions enfantines, et la dynamique entre les personnages. Додаткові уточнення: Utilise un langage familier mais pas vulgaire. Garde le rythme du récit et les répétitions stylistiques.

2. Створення стилістичних пар для instruction tuning. Необхідно підготувати 20–100 пар на зразок:

Input: Уривок українською + інструкція до стилю.

Output: Французький переклад у відповідному стилі.

Це дозволяє налаштувати модель або створити власного асистента стилістичного перекладу. Наприклад:

{ «input»: «Переклади як уривок дитячої пригодницької повісті з живою розмовною мовою. Текст: *Ех ти, мама-депутат!*,

«output»: *Ah toi, la maman députée !* }

3. Style conditioning через few-shot prompting. У prompt потрібно додати кілька прикладів перекладу з поясненням стилістичних рішень, і лише тоді – новий уривок для перекладу.

Exemple 1 :

Texte : *Ех ти, мама-депутат!*

Traduction : *Ah toi, la maman députée !* (style familier, affectif).

Exemple 2 :

Texte : *Ява крутанувся на місці і прожогом вибіг з класу.*

Traduction : *Yava a fait volte-face et a filé hors de la classe* (rythme rapide, émotion forte).

Nouveau texte : ...

Traduction :

4. Fine-tuning на корпусі перекладів. Необхідно сформувати корпус із 500 або більше узгоджених пар (UA→FR) у заданому стилі, й модель донавчається (fine-tuned) для кращого збереження стилістики саме такого тексту. Таким чином модель «вчиться» не лише перекладати, а й імітувати заданий стиль (дитячий, пригодницький, емоційний). Можна використовувати open-source моделі (Mistral, LLaMA) для локального fine-tuning.

5. Стилiстичне постредагування за допомогою prompt-рефiнiнгу. Перекладений текст потрібно подати на повторну обробку з фокусом

виключно на стилі. Приклад: Voici un texte traduit. Réécris-le en renforçant le style oral et enfantin, sans changer le sens.

Висновки

Аналіз отриманих текстів показав, що LLM здатні відтворювати загальний стиль оригіналу, зокрема легкість викладу, гумор та ритмічність діалогів, проте без спеціальних підказок (prompt engineering) вони часто втрачають дрібні стилістичні нюанси та культурно марковані елементи.

Експеримент з модифікованими інструкціями до моделі (few-shot prompting із зазначенням жанру, цільової вікової групи та бажаного стилю) продемонстрував помітне поліпшення якості перекладу: зросла відповідність тональності оригіналу, збережено гумористичні елементи та ігровий характер мовлення персонажів.

Таким чином, LLM не можна розглядати як «чорний ящик» для перекладу; ефективність їх використання значною мірою залежить від точного формулювання завдань та контексту.

Дослідження довело, що великі мовні моделі мають значний потенціал у сфері міжмовної стилістичної адаптації дитячих художніх текстів. Водночас для досягнення високої якості перекладу необхідно застосовувати методи інженерії підказок та попереднього аналізу стилістичних особливостей оригіналу. LLM можуть стати ефективним інструментом творчого перекладу, коли їх застосування супроводжується людським контролем і критичною оцінкою результатів. Таке поєднання автоматизованого перекладу та експертного редагування дозволяє зберегти у перекладі легкість, гумор та інтонаційне багатство оригіналу, що є ключовим для дитячої аудиторії.

Перспективи подальших досліджень полягають у розробленні спеціалізованих протоколів для навчання й тестування LLM на корпусах дитячої літератури, а також у вивченні ефективності поєднання автоматизованих і людських підходів до редагування та адаптації перекладів.

Бібліографія

- Нестайко, В. З. (2004). *Топеадори з Васюківки*. А-БА-БА-ГА-ЛІА-МА-ГА.
- Bassnett, S., & Lefevere, A. (1998). *Constructing cultures: Essays on literary translation*. Multilingual Matters. <https://doi.org/10.21832/9781853593529>
- Chen, L., Zhang, Y., Li, X., & Wang, J. (2025). *Classic4Children: Adapting Chinese literary classics for children with large language model*. arXiv. <https://doi.org/10.48550/arXiv.2502.01090>
- Jiao, W., Wang, J., Huang, S., & Tu, Z. (2023). *Is ChatGPT a good translator? A preliminary study*. arXiv. <https://arxiv.org/abs/2301.08745>
- Kong, L., & Macken, L. (2025). Can Peter Pan survive MT? A stylometric study on translating children's literature with large language models. *Journal of Translation and Language Technologies*, 18(1), 45–68. <https://doi.org/10.1075/jtlt.18.1.03kon>
- Lathey, G. (2015). *Translating children's literature*. Routledge.
- Oittinen, R. (2000). *Translating for children*. Routledge. <https://doi.org/10.4324/9780203902004>
- Toury, G. (1995). *Descriptive translation studies and beyond*. John Benjamins Publishing Company.
- Vardanian, M. (2020). Translating children's literature of Ukrainian diaspora as an implementation of the educational ideal of Ukrainian abroad. In V. Hamaniuk, S. Semerikov, & Y. Shramko (Eds.), *SHS Web of Conferences: The International Conference on History, Theory and Methodology of Learning (ICHTML 2020)* (Vol. 75). Kryvyi Rih, Ukraine, May 13–15, 2020. https://www.shs-conferences.org/articles/shsconf/pdf/2020/03/shsconf_ichtml_2020_01002.pdf
- Venuti, L. (Ed.). (2000). *The translation studies reader*. Routledge.

- Ziems, C., Mirzadeh, S. I., Ren, X., & Yang, D. (2024). Can large language models capture authorial style? In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (ACL 2024)*.
- Ziems, C., Mirzadeh, S. I., Zhang, X., Ren, X., & Yang, D. (2023). *StyLO: Stylometric analysis of large language models*. arXiv. <https://arxiv.org/abs/2303.11366>

References

Nestayko, V. Z. (2004). Toreadory z Vasyukivky. A-BA-BA-HA-LA-MA-HA

Резюме

Дегтярьова Євгенія

ВИКОРИСТАННЯ ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ ДЛЯ МІЖМОВНОЇ СТИЛІСТИЧНОЇ АДАПТАЦІЇ ДИТЯЧОГО ХУДОЖНЬОГО ТЕКСТУ

Постановка проблеми. У перекладі дитячої літератури важливу роль відіграє не лише точність передачі змісту, а й стилістична відповідність оригіналу: інтонаційна виразність, лексична адаптованість, вікова адресованість. У межах міжмовної адаптації ці чинники вимагають тонкого балансування між культурною специфікою та читабельністю. Сучасні великі мовні моделі (LLM) – такі як GPT – відкривають нові можливості для автоматизованого й напівавтоматизованого перекладу, але постає питання: наскільки ефективними вони є саме в контексті стилістичної адаптації дитячого художнього тексту?

Мета статті. Метою статті є дослідження потенціалу великих мовних моделей у здійсненні міжмовної стилістичної адаптації українських дитячих художніх текстів французькою мовою з урахуванням інтонаційної, стилістичної та вікової специфіки.

Методи дослідження. У роботі використано методи контрастивного аналізу, стилістичного аналізу оригіналу та перекладів, а також експериментального перекладу з використанням LLM (ChatGPT-4). Результати оцінено з точки зору відповідності мовленнєвих реєстрів, передачі інтонації, збереження наративної структури й емоційного тону.

Основні результати дослідження. LLM демонструють високий рівень граматичної і лексичної точності, здатність до стилістичної варіативності та адаптації дитячої лексики. Однак модель має тенденцію до нейтралізації локальних культурних маркерів, може згладжувати індивідуальний стиль автора й потребує чітких промптів для збереження жанрової специфіки. Водночас вона здатна генерувати декілька варіантів перекладу з різними стилістичними акцентами, що розширює можливості перекладача.

Висновки і перспективи. LLM можуть стати ефективним інструментом допомоги в перекладі дитячої літератури, особливо на етапі первинної адаптації або створення варіантів. Однак для повноцінного стилістично точного перекладу потрібне залучення перекладача-фахівця. Подальші дослідження доцільно спрямувати на розробку промптів, орієнтованих на вікову специфіку та культурну адаптацію, а також на удосконалення систем оцінювання стилістичної адекватності LLM-перекладу.

Ключові слова: великі мовні моделі, дитяча література, інтонація, культурна трансляція, стилістична адаптація, українсько-французький переклад, GPT.

Abstract

Dehtiarova Yevheniia

**USING LARGE LANGUAGE MODELS FOR CROSS-LINGUAL
STYLISTIC ADAPTATION OF CHILDREN'S FICTION**

Background. In translating children's literature, not only semantic accuracy but also stylistic fidelity to the original, particularly, intonational expressiveness, lexical accessibility, and age appropriateness, play a crucial role. Within cross-lingual adaptation, these factors require a delicate balance between cultural specificity and readability. Contemporary large language models (LLMs), such as GPT, offer new possibilities for automated and semi-automated translation, raising the question: how effective are they in preserving stylistic features when translating children's fiction?

Purpose. This article explores the potential of large language models for cross-lingual stylistic adaptation of Ukrainian children's literary texts into French, taking into account intonational, stylistic, and age-related specificities.

Methods. The study employs methods of contrastive analysis, stylistic examination of original texts and their translations, and experimental translation using an LLM (ChatGPT-4). The outcomes are evaluated in terms of register appropriateness, intonation transfer, narrative structure retention, and emotional tone.

Results. LLMs demonstrate a high level of grammatical and lexical precision, stylistic flexibility, and adaptability to child-oriented vocabulary. However, they tend to neutralize local cultural markers, may flatten the author's unique style, and require well-designed prompts to preserve genre-specific features. At the same time, LLMs can generate multiple stylistically diverse translation options, enhancing the translator's creative scope.

Discussion. LLMs can serve as effective tools for assisting in the translation of children's literature, particularly in the initial stages of adaptation or in creating alternative versions. However, achieving stylistically precise and nuanced translation still requires the expertise of a professional human translator. Future research should focus on developing age- and culture-sensitive prompting strategies and improving evaluation frameworks for assessing stylistic adequacy in LLM-generated translations.

Keywords: large language models, children's literature, intonation, cultural transfer, stylistic adaptation, Ukrainian-French translation, GPT.

Відомості про автора

Дегтярьова Євгенія Олександрівна, канд.філол.н., доцент, кафедра теорії, практики та перекладу французької мови КПІ ім. Ігоря Сікорського (Україна), e-mail: degtereva1981@gmail.com

Dehtiarova Yevheniia Oleksandrivna, PhD in Philology, Associate Professor, Department of Theory, Practice and Translation of the French Language, Igor Sikorsky Kyiv Polytechnic Institute (Ukraine), e-mail: degtereva1981@gmail.com

ORCID 0000-0002-6381-4875

Надійшла до редакції 26 вересня 2025 року
Прийнято до друку 23 листопада 2025 року

УДК 81`322: 81`42: 811.612.91

DOI: 10.24025/2707-0573.12.2025.345012

Вікторія Мусійчук

**ПЛОТНИЙ ЕКСПЕРИМЕНТ З АВТОМАТИЧНОГО
ВИЯВЛЕННЯ В'ЄТНАМСЬКИХ МЕДІАНАРАТИВІВ
ПРО РОСІЙСЬКО-УКРАЇНСЬКУ ВІЙНУ
ПРИ ОБМЕЖЕНИХ РЕСУРСАХ**



У статті описано проведення експерименту з автоматизованого виділення наративів про російсько-українську війну з в'єтнамських медіатекстів. Традиційний наративний аналіз поєднано з цифровими методами гуманітарних наук. Дослідження базується на абдуктивному підході, що поєднує двосторонню обробку даних. Послідовні етапи експерименту передбачають використання сучасних інструментів обробки природної мови, адаптованих спеціально для в'єтнамської мови та в умовах обмежених ресурсів.

Ключові слова: в'єтнамська мова, медіа, наративи, PhoBERT, кластеризація, російсько-українська війна.

1. Вступ

Наративи у медіатекстах відіграють ключову роль у формуванні громадської думки щодо міжнародних конфліктів, зокрема російсько-української війни. У в'єтнамському медіапросторі ці наративи характеризуються поляризацією та неоднозначністю. Державні ЗМІ, декларуючи нейтральну позицію В'єтнаму, водночас послуговуються проросійськими інтерпретаціями подій та термінології. Соціальні мережі наповнені як проросійськими, так і проукраїнськими настроями. Тому системне вивчення природи формування в'єтнамських медіанаративів є надзвичайно актуальним як з лінгвістичної точки зору, так і для подальшого застосування результатів цих досліджень для вироблення ефективних стратегій комунікації на міжнародному дипломатичному рівні.

Дослідження в'єтнамських медіанаративів російсько-української війни передбачає створення корпусу новин, повідомлень, статей, дописів з сайтів ЗМІ, соцмереж, блогерських платформ. З 2022 року донині (а для якісного дослідження варто брати період з 2014 року) – це має бути величезний обсяг даних. Сучасні наративні дослідження для опрацювання великих масивів даних передбачають використання методів цифрової гуманітаристики та штучного інтелекту, а також поєднання якісних підходів з кількісними. Такі

© Мусійчук В., 2025

дослідження мають загальні обмеження, зокрема, це застосування ресурсомістких, як технічно, так і фінансово, технологій. Основою для дослідження наративів є комплексний підхід до збору і опрацювання корпусу текстів, зокрема відбір релевантних джерел із очищенням та структуризацією даних. Необхідно зазначити, що для в'єтнамської мови уже на цьому етапі постає чимало викликів: відсутність анотованих датасетів; обмежена кількість як загальномовних, так і тематичних медійних корпусів; обмежений доступ до багаторівневих даних на в'єтнамських медійних сайтах; окрім звичайного очищення даних, в'єтнамські тексти вимагають додаткової токенизації для виділення слів, адже пробіли в тексті стоять не між словами, а між морфемами; більшість ШІ-моделей, NLP-інструментів навчені на англійськомовних корпусах, і їхня точність для в'єтнамської мови є істотно нижчою; розробка NLP-інструментів для в'єтнамської мови ведеться не так активно, а попередні розробки дуже швидко застарівають та не можуть бути застосовані через несумісність з сучасними системами. Така обмеженість ресурсів ускладнює масштабні дослідження. Наступний крок наративного аналізу – власне виокремлення та інтерпретація наративів, що базується на аналізі текстового корпусу – також потребує особливого підходу з урахуванням особливостей для в'єтнамського мовного матеріалу. Отже, завданням цієї статті є пристосування наявних методів та інструментів автоматичного аналізу в'єтнамськомовних текстів з метою виробити алгоритм для виявлення наративів в умовах обмежених ресурсів.

2. Теоретичне підґрунтя

Наративний аналіз об'єднує структуралістські, критично-дискурсивні підходи, фрейм-аналіз та комп'ютерний аналіз. Спадщина класичної наратології бере свій початок у структуралізмі, зокрема у роботах В. Проппа, який виокремив 31 функцію казки, що стали основою для виявлення універсальних наративних патернів (Propp, 1968). Р. Барт розвинув цю модель, запропонувавши трирівневу систему аналізу: рівень функцій, рівень дій та рівень наративу, де кожен рівень інтегрує попередній у складнішу семантичну структуру (Barthes, 1975). А. Греймас переосмислив проппівські персонажі в актантну модель (суб'єкт-об'єкт, відправник-отримувач, помічник-опонент), яка дозволяє аналізувати глибинні семантичні структури незалежно від поверхневих лінгвістичних форм (Greimas, 1983).

Критичний дискурс-аналіз, представлений у роботах Н. Ферклафа та Т. ван Дейка, трактує наратив як соціальну практику, що відтворює та легітимізує владні відносини через мовні вибори (Fairclough, 1995; van Dijk, 1998). Ферклаф запропонував трикомпонентну модель аналізу: текст (опис лінгвістичних властивостей), дискурсивна практика (інтерпретація процесів виробництва/споживання) та соціальна практика (пояснення ідеологічних наслідків) (Fairclough, 1989, 1995). Ван Дейк додав соціо-когнітивний вимір, акцентуючи увагу на ментальних моделях та ідеологічних схемах (van Dijk, 2008).

Теорія фреймінгу (Entman, 1993; Goffman, 1974) пояснює, як медіа подають події в обмежених рамках (фреймах), вибірково акцентуючи певні аспекти реальності: визначають проблему, діагностують причини, виносять моральні судження та пропонують рішення (Entman, 1993). Комп'ютерний наративний аналіз інтегрує зазначені вище теорії з NLP-методами (Piper, 2021). Методи на основі графів, як-от наративні мапи (Narrative Maps) та наративні стежки (Narrative Trails), використовують семантичні векторні представлення

текстів, щоб вибудовувати послідовні й цілісні сюжетні лінії, оптимізуючи їхню змістову узгодженість (Keith, 2020; German, 2025). Водночас подієцентричні підходи (event-centric frameworks) зосереджуються на виокремленні подій, дійових осіб та їхніх взаємозв'язків, що дає змогу виявляти характер фреймування цих подій у медійному дискурсі (Das, 2024; Levi, 2020).

Для нашого завдання автоматизованого виявлення наративів доцільним видається подієцентричний підхід, заснований на виділенні ключових слів, з елементами фреймування та кластеризації на основі ембедінгу (семантичних векторів).

3. Огляд літератури

Дослідження сприйняття російсько-української війни у в'єтнамському медіапросторі є нечисленними. Аналіз понад двохсот дописів в'єтнамського сегменту мережі Facebook протягом першого року конфлікту, проведений Май Тхі Данг Тху, виявив вісім ключових наративних фреймів, серед яких домінували антизахідні (США та Європа), антиукраїнські та проросійські настрої, що відображають національні настрої та геополітичні чинники. Підкреслено критичну роль соціальних медіа як впливового інструменту для просування геополітичних наративів і формування громадської думки у В'єтнамі. Дослідниця використала ручний контент-аналіз для виявлення прихованих конотацій та організації їх у фрейми, а також кількісні методи, як-от частотний аналіз, дослідження динаміки та кореляції між фреймами (Mai, 2025).

У двох працях в'єтнамських дослідників використано метод дискурс-аналізу для виявлення особливостей проросійських наративів у в'єтнамському кіберпросторі після початку повномасштабного вторгнення Росії в Україну. Аналіз 28 активних Facebook-груп засвідчив їхню функцію «ехокамер», що систематично посилюють дискурс, який легітимізує російську агресію, позитивно конотуючи Москву та негативно репрезентуючи Україну й Захід. Дослідники пов'язують цю тенденцію з комплексом факторів, зокрема історичною ностальгією за радянським періодом та стійкими антизахідними сентиментами. Такі онлайн-наративи сприяють формуванню пропагандистського середовища, яке дозволяє в'єтнамському уряду, зокрема його консервативному крилу, утримувати нейтральну позицію у публічному дискурсі, унеможливаючи домінування проукраїнських інтерпретацій (Hoang, 2022, 2025).

То Мінь Шон дослідив полярність громадської думки щодо конфлікту між Росією та Україною, зважаючи на в'єтнамську офіційну політику «принципового нейтралітету». Ця позиція є дипломатично складною через міцне партнерство з Росією та історичну ностальгію за радянською допомогою В'єтнаму, що й призвело, на думку дослідника, до внутрішнього розколу громадської думки. Хоча частина населення висловлює пропутінські настрої, інші занепокоєні питанням національного суверенітету та міжнародного права (То, 2022).

Загалом, наявні дослідження переважно сфокусовані на аналізі соцмереж з використанням якісних методів дослідження, як-от контент-аналіз та дискурс-аналіз, що є виправданим на невеликих вибірках даних. Однак для великих масивів даних не обійтися без методів комп'ютерної обробки та аналізу. Отже, метою цієї статті є тестування методології автоматичного виявлення в'єтнамських медіанаративів про російсько-українську війну за

умов обмежених ресурсів. Для пілотного експерименту зібрано невеликий корпус з новинного сайту «В'єтнамської інформаційної агенції» (Báo tin tức): 80 новин за період січень-квітень 2022 року та 80 новин за період січень-квітень 2025 року, відібраних за пошуковим запитом «Україна».

4. Методи та матеріал дослідження

Дослідження спроектовано як пілотний експеримент для тестування методології автоматичного видобування медіанаративів при обмежених ресурсах. Корпус складається зі 160 в'єтнамських новинних текстів з сайту «В'єтнамської інформаційної агенції», вибраних за пошуковим словом «Україна» у кількох варіаціях його написання в'єтнамською мовою, а саме *Ukraine, Ukraina, Ucraina*. Усі новини про російсько-українську війну та інші новини про Україну було розмежовано, про що йтиметься нижче. Вибір джерела для корпусу зумовлений тим, що «В'єтнамська інформагенція» є офіційним державним інформаційним органом, який є джерелом новин для багатьох інших в'єтнамських ЗМІ. Вибір періоду чотири місяці 2022 року та чотири місяці 2025 року дозволив забезпечити на невеликому обсязі даних тематичне та подієве розмаїття (настрої перед війною, масова евакуація та міграція, окупація, воєнні дії, обстріли тилкових населених пунктів, дипломатичні перемовини тощо). Вибір розміру корпусу визначений метою швидкого тестування, можливості ручної перевірки та практичними обмеженнями (приблизно 12 гігабайтів оперативної пам'яті, максимальна тривалість безперервної роботи 12 годин, та поточні обмеження на пропускну спроможність при завантаженні/вивантаженні даних з Google Drive) безкоштовної версії хмарної платформи для розробки на мові Python – Google Colab (Google Colaboratory), що використовувалася для автоматичної обробки.

Корпус було зібрано методом вебскрейпінгу з використанням платформи ParseHub (ParseHub). Ця платформа дозволяє збирати тексти, одночасно очищуючи від зайвих артефактів, що полегшує подальшу роботу з корпусом.

Базовий препроцесинг було здійснено у середовищі Google Colab за допомогою бібліотеки Underthesea-Vietnamese NLP Toolkit (Underthesea, 2025). Інструмент виконує токенизацію у вигляді об'єднання морфем у слова. Цей крок є критично важливим для подальшого аналізу, адже в'єтнамська мова не має пробілів між словами на письмі, натомість має пробіли між морфемами. Таким чином, проведення такої токенизації є обов'язковим для подальшої можливості виділення ключових слів, тематичної рубрикації та інших необхідних для наративного дослідження функцій. З іншого боку, в'єтнамський текст не потребує лематизації, адже в'єтнамська мова не має словозміни.

Виділення наративів відбувалося на основі подієцентричного підходу, який спирається на ключові слова для подій, персонажів та тем. Також було здійснено групування виділених фрагментів та відповідне фреймування й кластеризацію для подальшої можливості інтерпретації наративів.

Загальний хід дослідження ґрунтується на абдуктивному підході, запропонованому Ч. Пірсом, який поєднує переваги індукції та дедукції (Burch, 2024). У цьому дослідженні це реалізовано двома взаємодоповнювальними напрямками:

- 1) індуктивний: ключові слова → теми текстів → великі тематичні групи;
- 2) дедуктивний: початкова кластеризація → конкретні прояви тем.

Для автоматичного видобування ключових слів і фраз із текстів використано KeyBERT, який дозволяє це робити без попереднього навчання (Grootendorst, 2020). На відміну від традиційних статистичних методів, які спираються на частотність слів, KeyBERT використовує семантичні ембедінги (числові представлення слів у багатовимірному просторі), які захоплюють змістовне значення слів. Алгоритм KeyBERT працює в кілька етапів. По-перше, весь документ кодується в один вектор (багатовимірну точку) за допомогою моделі Sentence-BERT, що представляє загальний зміст документа. Далі вилучаються всі окремі слова та короткі фрази, й кожна із них також кодується у вектор. Алгоритм порівнює кожного кандидата із вектором цілого документа, обчислюючи їхню семантичну подібність. Найбільш подібні до документа слова та фрази обираються як ключові. Основною перевагою KeyBERT є те, що він не потребує попереднього навчання на спеціалізованих корпусах та підтримує будь-яку мову, для якої є модель BERT, що робить його особливо корисним для низькоресурсних мов, як-от в'єтнамська.

У нашому дослідженні для цієї мети використано модель VoVanPhuc/sup-SimCSE-VietNameese-phobert-base – це передова модель штучного інтелекту, спеціалізована на створенні векторних представлень речень для в'єтнамської мови (VoVanPhuc, 2024). Вона базується на моделі PhoBERT, яку було натреновано на в'єтнамських текстах, що забезпечує їй розуміння як повсякденної, так і професійної лексики (Nguyen, 2020). Крім того, сентимент-аналіз з PhoBERT показав точність 93.12% у класифікації новин на позитивні / негативні / нейтральні (Nguyen, 2021), демонструючи потенціал моделі для аналізу наративів. PhoBERT перетворює кожен текст на числовий вектор розмірністю 768, що дозволяє порівнювати тексти математичним шляхом: тексти зі схожим змістом матимуть числові представлення, близькі одне до одного у цьому просторі. SimCSE-Vietnamese – це окрема модель, оптимізована для порівняння семантичної подібності текстів у в'єтнамській мові. Вона натренована на множині пар текстів, де парафрази позначені як подібні, а семантично не пов'язані тексти – як неподібні. Це навчання робить модель особливо чутливою до справжніх змістовних схожостей. Використання PhoBERT+SimCSE забезпечило дві ключові переваги для цього дослідження: семантичні представлення, адаптовані до в'єтнамської мови, і можливість швидкої обробки в Google Colab.

Для кластеризації текстів у цьому експерименті використано два комплементарні методи: K-means та HDBSCAN (Hierarchical Density-Based Spatial Clustering of Applications with Noise). K-means – це алгоритм, що поділяє тексти на певну кількість груп (кластерів). Алгоритм працює ітеративно: спочатку він випадково вибирає кілька «центрів» (чисельні точки в семантичному просторі, які представляють центр кожної групи), до яких зараховуються тексти. Після цього алгоритм перераховує позиції центрів на основі всіх текстів, зарахованих до кожної групи, й перерозподіляє групи. Процес продовжується, поки центри не перестануть суттєво змінюватися, що свідчить про стабільність групування (Abiodun, 2023).

У цьому дослідженні метод кластеризації HDBSCAN застосовано як експериментальний інструмент для групування в'єтнамських текстів про російсько-українську війну за їхньою семантичною подібністю, використовуючи ембедінги PhoBERT/SimCSE як вхідні дані. HDBSCAN реалізує кластеризацію на основі щільності з автоматичним виявленням

викидів, а PhoBERT-ембедінги забезпечують семантичну якість (Nguyen et al., 2023). HDBSCAN належить до щільнісних методів кластеризації: він шукає в багатовимірному просторі ділянки, де тексти розташовані «густо» (мають подібні значення ембедінгів), і відокремлює їх від зон низької щільності, які інтерпретуються як шум або маргінальні випадки (McInnes, 2017). На відміну від K-means, HDBSCAN не вимагає наперед задавати кількість кластерів, а натомість має параметр мінімального розміру кластера, що краще відповідає ситуації, коли невідомо, скільки саме різних наративних груп є в корпусі.

GPT-4 (ChatGPT) було використано як допоміжний інструмент для розподілу ключових фраз на події, персонажі та теми з подальшою ручною верифікацією, а також для якісної анотації попередньо виділених кластерів.

Дослідження організовано як паралельна послідовність двох напрямів обробки. Для обох напрямів спочатку створено корпус в'єтнамських медіатекстів та здійснено препроцесинг вихідних даних та ембедінг через PhoBERT+SimCSE. Далі у пешому напрямі йде видобування ключових фраз за допомогою KeyBERT та групування з GPT-4. У другому напрямі – кластеризація через K-means і HDBSCAN та інтерпретація за допомогою GPT-4.



Схема 1. Алгоритм автоматичного виявлення наративів у в'єтнамських медіатекстах з абдуктивним підходом

5. Результати дослідження

5.1. Загальні характеристики корпусу та розподіл даних

Для дослідження було створено експериментальний мінікорпус із 160 в'єтнамських новинних текстів, видобутих за допомогою платформи ParseHub. Корпус було збережено у форматі .csv з усіма вихідними даними (як-от дата, лінк, автор, заголовок), де кожна новина знаходиться в окремому рядку. Такий формат дозволяє в одну операцію обробляти багато текстів, водночас не зливаючи усі тексти в один. Попередню обробку тексту було проведено за допомогою інструменту Underthesea, а саме його частини word_tokenizer, для розбиття тексту на слова. Проблема виділення меж слова у в'єтнамській мові така складна й неоднозначна, що жоден автоматичний інструмент досі не може вирішити це завдання без похибок. Зокрема, на

нашому матеріалі серед іншого було помічено, що після обробки в одне слово об'єднувалися іноземні імена та прізвища або власні назви з кількох слів (наприклад, *Donald_Trump*, *Eurovision_News*), що для нашого дослідження було не вадою, а перевагою.

5.2. Від ключових слів до тематичних груп

У першій частині дослідження кожен текст було піддано структурованому аналізу за допомогою методу KeyBERT, що функціонує на основі моделі PhoBERT SimCSE-Vietnamese. Цей відбір інструментарію було обґрунтовано тим, що SimCSE-Vietnamese є передовою моделлю, спеціалізованою на створенні контекстуально чутливих векторних представлень речень для в'єтнамської мови з використанням методу Simple Contrastive Learning of Sentence Embeddings на основі архітектури PhoBERT. На цьому етапі з корпусу текстів було автоматично виділено всі можливі n-грами довжиною від 1 до 3 слів, після чого відібрано найчастотніші з них разом із відповідними векторними представленнями. Це дозволило встановити базові семантичні одиниці наративного дискурсу. Було проекспериментовано видобувати 50, 100 та 200 ключових фраз. Необхідно зазначити, що збільшення кількості ключових слів мало тенденції до повторів. Наприклад, один персонаж повторювався кілька разів у сполученні з різними діями або одна й та сама подія подавалася у різних граматичних вираженнях:

<i>tham_gia của mỹ</i>	0.4738	‘участь Америки’,
<i>mỹ đã tham_gia</i>	0.4742	‘Америка брала участь’,
<i>kirill_dmitriev đưa ra</i>	0.4584	‘Кіріл Дмитрієв запропонував’,
<i>dmitriev cho biết</i>	0.4677	‘Дмитрієв сказав’,
<i>kirill_dmitriev báo_cáo</i>	0.4756	‘Кіріл Дмитрієв повідомив’.

Крім того, на початковому етапі виділення ключових фраз відбувалось без використання стоп-слів, тому до списків потрапили такі нерелевантні сполуки:

<i>viện trợ quân sự đã</i>	0.5660	‘військова допомога вже’ – <i>đã</i> є граматичним показником минулого часу та не має значення без дієслова,
<i>hỗ trợ quân sự hơn</i>	0.5643	‘військова підтримка більше ніж’ – <i>hơn</i> є граматичним показником ступеню порівняння,
<i>chuyên gia quân sự và</i>	0.5087	‘військові експерти та’ – <i>và</i> є сполучником.

Отже, використання списку стоп-слів для цього етапу є вкрай важливим. При чому, окрім звичайних для таких списків службових слів, імовірно варто додати й повнозначні слова, які часто трапляються у новинах, але не завжди мають вагоме значення для виділення ключових фраз, як-от: *sказав*, *повідомив* тощо.

Ще одним недоліком цієї фази експерименту є те, що необхідно заздалегідь визначити кількість ключових слів і фраз. Для великих масивів даних це може бути незручно.

Наступним завданням було перетворити ключові фрази з KeyBERT на наративні мітки, зокрема події й персонажі. Цей підхід дав можливість згрупувати схожі ключові слова і фрази, спростити їх та систематизувати наративи на двох рівнях абстракції. Для виконання завдання було використано ChatGPT. У результаті було створено списки слів на позначення подій та персонажів. Наприклад, ключові слова на позначення подій: *tấn công* – *атака*, *pháo kích* – *обстріл*, *xâm lược* – *вторгнення*, *di tản* – *евакуація*, *viện trợ nhân*

đạo – гуманітарна допомога, *đàm phán* – переговори, *chiến sự* – військові дії. Серед персонажів, або акторів, зокрема було виділено такі: *Ukraine* – Україна, *Nga* – Росія, *Zelensky* – Зеленський, *Putin* – путін, *LHQ* – ООН, *quân đội* – армія, *nạn nhân* – жертва, *người dân* – цивільне населення, *người tị nạn* – біженець.

Далі було проведено повторну вибірку ключових слів через KeyBERT для кожного конкретного тексту (5 ключових слів на текст), а також перевірку на наявність заданих подій та персонажів у кожному тексті. На основі цих результатів за допомогою GPT-4 було виділено тематичні фрейми для кожної статті, як-от «Україна – жертва, яку підтримує світ», «Росія – агресор», «Є надія на розв’язання конфлікту», «Війна впливає на світову економіку» тощо. Ці фрейми було згруповано в тематичні групи, як-от: «Агресія та захист», «Гуманітарна криза», «Дипломатія та переговори», «Економічні наслідки».

Усі автоматично отримані результати ми також перевіряли вручну для верифікації, результат був задовільним. Більшість похибок припадала на визначення ключових слів, що зумовлено технічною недосконалістю. Групування й систематизація на абстрактному рівні за допомогою автоматизованих засобів та ручної верифікації збіглися повністю. Таким чином, можемо засвідчити ефективність індуктивного підходу з використанням доступних ресурсів.

5.3. Від кластерів до наративів

У другій частині дослідження було проведено кластеризацію усіх текстів за змістом. За допомогою K-means алгоритму тексти було розбито на кластери відповідно до семантичної подібності їхніх PhoBERT-ембедінгів. Обмеженням K-means є те, що від дослідника одразу вимагається визначитися з кількістю кластерів. Таким чином, занадто велика кількість може призвести до розпорошення результатів, а занадто мала – до можливої втрати важливих тематичних груп. Алгоритму було задано розбити корпус на 5 кластерів. За допомогою GPT проведено інтерпретацію кожного кластера, а саме – найменування кластера, визначення провідної теми, повторюваних ключових лексем та ймовірного ключового наративу. У результаті виділено такі кластери: «Військові дії», «Гуманітарна криза та евакуація», «Дипломатичні зусилля та міжнародна реакція», «Економічні наслідки та енергетична криза», «Інциденти та межі конфлікту». Наприклад, кластер «Військові дії» об’єднав тексти, які описують активні військові операції, обстріли населених пунктів та окупацію територій. Тема кластера: «Агресія та захист». Ключові слова: *tấn công* – атака, *chiến sự* – воєнні дії, *quân sự* – військовий, *tên lửa* – ракета. Провідний наратив: «Росія наступає – Україна захищається». Найбільш неоднозначним виявився кластер «Інциденти на межі конфлікту», який містив тексти про атаки на об’єкти цивільної інфраструктури, тероризм, порушення міжнародного права, а також повідомлення, що не стосувалися війни безпосередньо.

На практиці K-means виявився цінним інструментом для цього дослідження, однак метод вимагає попередньої вказівки кількості кластерів, що є викликом, коли заздалегідь невідомо, скільки різних типів наративів присутні в даних. З метою усунути ці прогалини було проведено додаткову кластеризацію за допомогою алгоритму HDBSCAN з параметром мінімального розміру кластера (*min_cluster_size*), який було варійовано від 10 до 3 залежно від щільності даних у векторному просторі. Алгоритм HDBSCAN виявився

оптимальним для цього завдання, оскільки він дозволяє автоматично визначати кількість кластерів на основі щільності точок (текстів) та природним чином ідентифікувати «шумові» спостереження (позначені як кластер -1), які не належать до жодної стійкої групи. Перший запуск алгоритму з параметром `min_cluster_size=10` класифікував усі новини як шум (-1), що вказувало на недостатню щільність кластерів при такому гіперпараметрі. Після зниження порогу до `min_cluster_size=3`, алгоритм виявив 6 основних кластерів (позначених від 0 до 5) плюс категорія шуму (-1). Ці «шумові» тексти виявилися корисними для наративного аналізу, оскільки часто містили або дуже специфічні сюжети, або змішані фрейми, які не вписувалися в жоден із домінантних кластерів. Таким чином, HDBSCAN виконував роль інструмента для виявлення як ядрових, так і периферійних наративів, доповнюючи більш жорсткий K-means у змішаному кластеризаційному підході.

У результаті інтерпретації моделлю GPT кластерам присвоєно такі назви: «Глобальні економіко-гуманітарні наслідки війни», «Біженці та гуманітарна допомога», «Геополітичні наслідки війни», «Дипломатичне посередництво», «Нейтральні гравці у контексті війни», «Розвиток подій на фронті». До прикладу, кластер «Глобальні економіко-гуманітарні наслідки війни» об'єднав тексти, що розглядають не військові дії напряму, а вторинні наслідки війни, російсько-українська війна згадується як контекст, що впливає на ринки, торгівлю, виробництво та постачання продовольства, зокрема на в'єтнамський експорт, інфляцію, ціни на енергоносії. Провідні наративи: «Війна впливає на життя людей в усьому світі», «Війна призводить до негативних наслідків у світовій економіці».

До кластеру «шуму» в тому числі потрапили тексти, що не стосувалися війни, зокрема частина новин січня 2022 року. Тому було проведено додаткову класифікацію усіх текстів на ті, що пов'язані з війною, та непов'язані. Валідацію кластеризації здійснено вручну, результати задовільні.

Результати кластеризації за допомогою K-means та HDBSCAN показали суттєвий збіг у виявлених змістовних осях конфлікту, але з різними рівнями деталізації. Зокрема, в обох класифікаціях видно перевагу у в'єтнамських новинах висвітлення не безпосереднього аналізу військових дій, а впливу війни на місцеву та глобальну економіку, політику та гуманітарну сферу. Таким чином, застосування такої перехресної кластеризації слугує додатковою верифікацією отриманих результатів.

6. Висновки

Проведене пілотне дослідження продемонструвало ефективність обраної методології для автоматизованого виявлення в'єтнамських медіанаративів про російсько-українську війну за умов обмежених ресурсів. Обробка експериментального корпусу текстів за допомогою Underthesea, KeyBERT, PhoBERT / SimCSE, K-means, HDBSCAN та GPT дозволила виявити основні наративні вісі: військові дії, гуманітарна криза, дипломатія та глобальні наслідки. Ключовою знахідкою стало домінування гуманітарного фрейму у в'єтнамських медіа. В'єтнамські офіційні ЗМІ демонструють нейтральну позицію, акцентуючи на людських стражданнях та практичних наслідках для Азії та світу, що узгоджується з державною політикою «бамбукової дипломатії» – гнучкого балансування між великими державами.

Застосування абдуктивного підходу показало валідність використаних методів, адже в обох напрямках дослідження отримано зіставні результати. Крім того, ручна перевірка також підтвердила отримані дані та інтерпретації. Методологія також виявилася стійкою до ресурсно обмежених умов, зокрема в середовищі Google Colab, та цілком придатною для розширення корпусу на локальних GPU, Colab Pro або з використанням API від LLM. Отже, цей алгоритм може бути випробуваний для автоматичного виявлення наративів у великих обсягах даних.

Подальші дослідження можуть стосуватися зіставлення корпусів офіційних ЗМІ, соціальних мереж та експертних повідомлень, з визначенням NER для персонажів (PhoBERT+ViSOBERT) та динамічним трекінгом наративів через LSTM-моделі, із застосуванням цифрових гуманітарних наук для аналізу пропагандистських дискурсів. Цікавою може бути міжмовна компаративістика (в'єтнамськомовні vs. англомовні, україномовні медіа), однак для інших мов інструменти мають бути адаптовані, адже в цьому дослідженні завданням було розробити дієздатний алгоритм саме для в'єтнамської мови.

Дослідження підтверджує, що цифрові гуманітарні науки роблять наративний аналіз доступним для великих масивів даних. Методологія відкриває шлях до масштабних досліджень в'єтнамських дискурсів, де автоматичні інструменти та ШІ слугують прискорювачами та підсилювачами людського аналізу.

Подяка

Це дослідження й експеримент було проведено під наставництвом Наталії Чейлитко в рамках курсу «Великі мовні моделі для лінгвістів», організованого Єнським університетом імені Фрідріха Шіллера 11.02.-24.06.2025.

Primary Sources

Báo tin tức. <https://baotintuc.vn>
Chat GPT. <https://chatgpt.com/>
Google Colaboratory. *Colab FAQ: Resource limits.* <https://research.google.com/colaboratory/faq.html>
ParseHub. <https://www.parsehub.com/>
Underthesea – Vietnamese NLP Toolkit (2025). <https://github.com/undertheseanlp/underthesea/tree/main>
VoVanPhuc. (2024). *sup-SimCSE-Vietnamese-phobert-base.* <https://huggingface.co/VoVanPhuc/sup-SimCSE-VietNameese-phobert-base>

References

- Abiodun, M. I., Absalom, E. E., Laith, A., Belal, A. & Jia, H. (2023). K-means clustering algorithms: A comprehensive review, variants analysis, and advances in the era of big data. *Information Sciences*, 622, 178–210. <https://doi.org/10.1016/j.ins.2022.11.139>
- Barthes, R. (1975). An introduction to the structural analysis of narrative. *New Literary History*, 6(2), 237–272. <https://doi.org/10.2307/468419>
- Burch, R. (2024). Charles Sanders Peirce, *The Stanford Encyclopedia of Philosophy* (Summer 2024 Edition), Edward N. Zalta & Uri Nodelman (eds.). <https://plato.stanford.edu/archives/sum2024/entries/peirce/>
- Das, R., Chandra, A., Lee, I-Ta & Pacheco, M. L. (2024). Media framing through the lens of event-centric narratives. In *Proceedings of the 6th Workshop on Narrative Understanding*, 85–98. <https://doi.org/10.18653/v1/2024.wnu-1.15>

- Entman, R. M. (1993). Framing: Toward clarification of a fractured paradigm. *Journal of Communication*, 43(4), 51–58. <https://doi.org/10.1111/j.1460-2466.1993.tb01304.x>
- Fairclough, N. (1989). *Language and power*. London: Longman.
- Fairclough, N. (1995). *Critical discourse analysis: The critical study of language*. London: Longman.
- German, F., Keith, B. & North, C. (2025). Narrative trails: A method for coherent storyline extraction via maximum capacity path optimization. In *Proceedings of the Text2Story'25 Workshop, Lucca (Italy)*, 10.04.2025. <https://doi.org/10.48550/arXiv.2503.15681>
- Goffman, E. (1974). *Frame analysis: An essay on the organization of experience*. Harvard University Press.
- Greimas, A. J. (1983). *Structural semantics: An attempt at a method*. University of Nebraska Press.
- Grootendorst, M. (2020). *KeyBERT: Minimal keyword extraction with BERT*. Python Package. <https://github.com/MaartenGr/KeyBERT>
- Hoang, T. H. & Dien Nguyen, A. L. (2022). The Russia-Ukraine war: Unpacking online pro-Russia narratives in Vietnam. *ISEAS Perspective*, 44, 1–14. https://www.iseas.edu.sg/wp-content/uploads/2022/03/ISEAS_Perspective_2022_44.pdf
- Hoang, T. H. (2025). Vietnam's mediascape amid the war in Ukraine: Between method and mayhem. In *Fulcrum*. <https://fulcrum.sg/vietnams-mediascape-amid-the-war-in-ukraine-between-method-and-mayhem/>
- Keith, B. & Mitra, T. (2020). Narrative maps: An algorithmic approach to represent and extract information narratives. In *Proceedings of Text2Story 2020*. <https://doi.org/10.48550/arXiv.2009.04508>
- Levi, E., Mor, G., Shenhav, S.R., & Sheaffer, T. (2020). CompRes: A dataset for narrative structure in news. In *Proceedings of LREC 2020*. ArXiv, abs/2007.04874.
- Mai, T. D. T. (2025). News framing on social media: A case study of Russia-Ukraine war narration on Facebook in Vietnam. *Keio Communication Review*, 47(3), 33–61. https://koara.lib.keio.ac.jp/xoonips/modules/xoonips/download.php/AA00266091-20250300-0033.pdf?file_id=187546
- McInnes, L., & Healy, J. (2017). Accelerated hierarchical density clustering. In *IEEE International Conference on Data Mining Workshops (ICDMW)*, 33–42. <https://doi.org/10.1109/ICDMW.2017.12>
- Nguyen, D. Q., & Nguyen, A. T. (2020). PhoBERT: Pre-trained language models for Vietnamese. In *Findings of the Association for Computational Linguistics: EMNLP 2020*, 1037–1042. <https://doi.org/10.18653/v1/2020.findings-emnlp.92>
- Nguyen, N. T., Nguyen, D. V., Phan, K. T. K., & Nguyen, N. L.T. (2023). Abusive span detection for Vietnamese narrative texts. In *Proceedings of the 12th International Symposium on Information and Communication Technology (SOICT '23)*, 471–478. <https://doi.org/10.1145/3628797.3628921>
- Nguyen, S. T., Nguyen, N. L., Trang, T., Nguyen, T. H., Duong, T. T. P. & Tuan, N. (2021). Stock article title sentiment-based classification using PhoBERT. In *Proceedings of the 2nd International Conference on Human-centered Artificial Intelligence (Computing4Human 2021)*. <https://ceur-ws.org/Vol-3026/paper25.pdf>
- Piper, A., So, R. J. & Bamman, D. (2021). Narrative theory for computational narrative understanding. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, 298–311. 10.18653/v1/2021.emnlp-main.26
- Propp, V. (1968). *Morphology of the folktale* (2nd ed., L. Scott, Trans.). University of Texas Press.
- To, M.S. (2022). Explaining the Vietnamese public's mixed responses to the Russia-Ukraine crisis. In *The Diplomat*. <https://thediplomat.com/2022/03/explaining-the-vietnamese-publics-mixed-responses-to-the-russia-ukraine-crisis/>
- van Dijk, T. A. (1998). *Ideology: A multidisciplinary approach*. London: Sage.
- van Dijk, T. A. (2008). *Discourse and context: A sociocognitive approach*. Cambridge, New York: Cambridge University Press.

Резюме

Мусійчук Вікторія

ПІЛОТНИЙ ЕКСПЕРИМЕНТ З АВТОМАТИЧНОГО ВИЯВЛЕННЯ В'ЄТНАМСЬКИХ МЕДІАНАРАТИВІВ ПРО РОСІЙСЬКО-УКРАЇНСЬКУ ВІЙНУ ПРИ ОБМЕЖЕНИХ РЕСУРСАХ

Постановка проблеми. Наративи у медіатекстах відіграють ключову роль у формуванні громадської думки щодо міжнародних конфліктів, зокрема російсько-української війни. Системне вивчення природи формування в'єтнамських медіанаративів є актуальним як з лінгвістичного погляду, так і для застосування результатів для вироблення ефективних стратегій міжнародної комунікації. Традиційний наративний аналіз неефективний для великих корпусів через ресурсомісткість, особливо для низькоресурсних мов як в'єтнамська (відсутність анотованих датасетів, складна токенизація морфем, обмежений доступ до багаторівневих даних). Попередні дослідження обмежені якісним дискурс-аналізом соцмереж, без автоматизованих методів NLP для масштабування.

Мета статті. Розробити та протестувати гібридну методологію автоматизованого видобування наративів з в'єтнамських медіатекстів за умов обмежених ресурсів, поєднуючи класичну наратологію з цифровими гуманітарними методами (NLP, кластеризація), для ідентифікації подієцентричних осей (події, персонажі, фрейми) та їхньої інтерпретації.

Методи дослідження. Пілотний експеримент на корпусі 160 новин з Báo tin tức, зібраних ParseHub, токенизований Underthesea. Абдуктивний підхід (Burch, 2024): 1) індуктивний – KeyBERT+PhoBERT/SimCSE-Vietnamese (ембедінг, ключові фрази), GPT-4 (групування на події/персонажів/теми); 2) дедуктивний – K-means/ HDBSCAN+PhoBERT/SimCSE-Vietnamese (ембедінг, кластеризація), GPT-4 (інтерпретація). Ручна верифікація результатів.

Основні результати дослідження. У першому напрямі експерименту за допомогою KeyBERT виявлено ключові фрази, які далі було розподілено за наративними мітками (події, персонажі), згруповано у тематичні групи та наративні фрейми за допомогою GTP. У другому напрямі проведено паралельну кластеризацію двома інструментами: K-means та HDBSCAN. Виявлені кластери інтерпретовано та виділено наративні фрейми й ключові терміни. У результаті зіставлення двовекторного підходу виявлено збіг у видобутих змістовних осях та наративних фреймах. Зокрема, в обох напрямках очевидна перевага висвітлення у в'єтнамських новинах не безпосереднього аналізу військових дій, а впливу війни на місцеву та глобальну економіку, політику та гуманітарну сферу.

Висновки і перспективи. Дослідження продемонструвало ефективність обраної методології для автоматизованого виявлення в'єтнамських медіанаративів про російсько-українську війну за умов обмежених ресурсів. Застосування абдуктивного підходу показало валідність методів, адже в обох напрямках дослідження отримано корелюючі результати. Методологія виявилася стійкою до ресурсно обмежених умов, зокрема в середовищі Google

Colab, та придатною для розширення корпусу. Подальші дослідження можуть стосуватися зіставлення різних корпусів, застосування NER для персонажів та динамічного трекінгу наративів через LSTM-моделі, сприяючи аналізу пропагандистських дискурсів.

Ключові слова: в'єтнамська мова, медіа, наративи, PhoBERT, кластеризація, російсько-українська війна.

Abstract

Musiichuk Viktoriia

PILOT EXPERIMENT ON AUTOMATIC DETECTION OF VIETNAMESE MEDIA NARRATIVES ABOUT THE RUSSIA-UKRAINE WAR UNDER LIMITED RESOURCES

Background. Narratives in media texts play a crucial role in shaping public opinion on international conflicts, including the Russia–Ukraine war. Systematic investigation of how Vietnamese media narratives are constructed is relevant both from a linguistic perspective and for developing effective strategies of international communication. Traditional narrative analysis is inefficient for large corpora due to its resource intensity, especially for low-resource languages such as Vietnamese (lack of annotated datasets, complex morpheme tokenization, limited access to multi-layered data). Existing studies are largely confined to qualitative discourse analysis of social media and do not employ scalable NLP-based automation.

Purpose. The aim of the article is to develop and test a hybrid methodology for the automated extraction of narratives from Vietnamese media texts under limited computational resources, combining classical narratology with digital humanities methods (NLP, clustering) in order to identify event-centric narrative axes (events, characters, frames) and provide their interpretation.

Methods. The study presents a pilot experiment on a corpus of 160 news items from Báo tin tức, collected via ParseHub and tokenized with Underthesea. An abductive approach (Burch, 2024) was implemented along two complementary strands: (1) an inductive strand using KeyBERT + PhoBERT / SimCSE-Vietnamese (text embeddings, keyphrase extraction) and GPT-4 (grouping into events, characters, and themes); (2) a deductive strand using K-means / HDBSCAN + PhoBERT / SimCSE-Vietnamese (embeddings, clustering) and GPT-4 (cluster interpretation). All automatic outputs were subjected to manual verification.

Results. In the first strand, KeyBERT was used to extract keyphrases, which were subsequently mapped onto narrative labels (events, characters), aggregated into thematic groups and narrative frames with the assistance of GPT-4. In the second strand, parallel clustering was performed with K-means and HDBSCAN. The resulting clusters were interpreted and associated with narrative frames and core lexical items. Comparison of the two-vector approach revealed convergence in the extracted semantic axes and narrative frames. In both strands, Vietnamese news were found to prioritise coverage of the war's impact on local and global economies, politics, and the humanitarian sphere over detailed analysis of military operations.

Discussion. The study demonstrates the effectiveness of the proposed methodology for automated detection of Vietnamese media narratives about the Russia–Ukraine war under limited resources. The abductive design proved methodologically valid,

as both strands produced consistent and mutually reinforcing results. The workflow is robust in resource-limited environments such as Google Colab and is scalable to larger corpora. Future research may include systematic comparison of different corpora, integration of NER for character extraction, and dynamic narrative tracking using LSTM-based models, thereby contributing to the analysis of propagandistic discourses.

Keywords: Vietnamese language, media, narratives, PhoBERT, clustering, Russia–Ukraine war.

Відомості про автора

Мусійчук Вікторія Анатоліївна, кандидат філологічних наук, старший науковий співробітник, завідувач відділу Азіатсько-Тихоокеанського регіону Інституту сходознавства ім. А. Ю. Кримського НАН України, e-mail: v.musiychuk@nas.gov.ua

Musiichuk Viktoriia, Ph.D. Philology, Senior researcher, A. Krymskyi Institute of Oriental Studies National Academy of Sciences of Ukraine, Head of Asia-Pacific Department, e-mail: v.musiychuk@nas.gov.ua

ORCID 0000-0002-4456-4487

Надійшла до редакції 30 листопада 2025 року

Прийнято до друку 25 грудня 2025 року

UDC 81'373 [811.111]

DOI: 10.24025/2707-0573.12.2025.343342

Oleksandr Kapranov



SELF-MENTIONS IN BRITISH PRIME MINISTER'S SPEECHES ON CLIMATE CHANGE

This contribution involves a quantitative study whose research aim is to learn about the occurrence of self-mentions, i.e. words like “I”, “my”, “me”, “mine”, “myself”, “we”, “our”, “us”, “ours”, and “ourselves”, in a corpus of speeches on the topic of climate change delivered by the current British prime minister Keir Starmer. It is argued in the study that a close examination of self-mentions could contribute to our understanding of Starmer’s stance on climate change, which is a topical issue in the British corporate and political discourses. The results of the corpus analysis reveal that Starmer’s speeches on the issue of climate change are marked by the frequently occurring self-mentions “we” and “our”. These findings and their interpretation through the lens of climate change communication are further presented in the article.

Keywords: climate change discourse, Keir Starmer, political speeches, prime minister, self-mentions, the UK.

1. Introduction

Since the early 2000s, the issue of climate change has been one of the highly debatable topics in political discourses (Boykoff, 2008; Fløttum, 2010, 2014, 2019; Fløttum & Gjerstad, 2017; Kapranov, 2015a, 2018; Willis, 2017; Wright & Irwin, 2025) as well as in corporate communication (Burns & Cowlshaw, 2014; Dahl & Fløttum, 2019; Ferguson, de Aguiar, & Fearfull, 2016; Kapranov, 2015b, 2017; Moorcroft, Hampton, & Whitmarsh, 2025). In the United Kingdom (the UK), particularly, the issue of climate change has become a contentious problem that is regularly elucidated by the British media, the public at large, and a number of climate change protest movements (Gavin & Marshall, 2011; Hayes & O’Neill, 2021; Kapranov, 2024a, 2024b; Swain, 2025).

Whilst there are multiple and, often, competing discursive voices on the issue of climate change in the UK, it seems quite sensible to heed to the discourse on climate change by the ruling party, which is, presently, the Labour Party (Nisbett et al., 2025). The current Labour government is led by Keir Starmer, the UK’s prime minister, who was sworn into office on 5 July 2024. Whereas his premiership is, so far, rather short, it seems, nevertheless, pertinent to look at Starmer’s discourse on climate change in more detail, given that Starmer determines, to a substantial degree, the UK’s government policies associated with the issue of climate change (Diamond, Richards, & Warner, 2025; Kapranov, 2025a). In this regard, it should be

mentioned that there are several recent studies on the lexico-syntactic peculiarities of Starmer's climate change discourse (Kapranov, 2025b, 2025c). At the same time, however, such aspect of Starmer's discourse on climate change as self-mentions has not been analysed yet.

In order to bridge the gap in our knowledge about Starmer's discourse on climate change, this contribution involves a quantitative study that examines the frequency of self-mentions, which are manifested by the first-person pronouns singular ("I", "my", "me", "myself", and "mine") and plural ("we", "our", "us", "ourselves", and "ours"), in his speeches on climate change. It is argued in the study that self-mentions may help to uncover the way Starmer frames the issue climate change, for instance, whether or not he presents himself as a team player, who looks at climate change as a challenge posed to the entire government and/or the entire British nation, or, alternatively, his use of self-mentions may reveal that he portrays himself as the sole fighter against the negative consequences of climate change. In this light, the study aims at answering the following **research question (RQ)**:

RQ: What is the frequency of self-mentions (i.e., "I", "my", "me", "mine", "myself", "we", "our", "us", "ours", and "ourselves") in Keir Starmer's political speeches on climate change?

Further, this contribution is organised as follows. First, the literature on self-mentions in political discourse is outlined. Second, the present quantitative study is given in conjunction with its methodology, results and discussion. Finally, the major findings are summarised and conclusions are presented.

2. Self-Mentions in Political Discourse: An Outline of the Literature

One of the definitions of self-mentions that gained currency in discourse studies belongs to Hyland (2001, 2004), who regards it as the use of the first-person pronouns in singular and plural, respectively, such as "I", "my", "me", "mine", "myself", "we", "our", "us", "ours", and "ourselves", in order to manifest the authorial voice, i.e. to mark explicitly the presence of the person who has authored the text and/or speech. Hyland's approach to self-mentions has been embraced by researchers in the fields of academic writing (Carrió-Pastor, 2020; Kapranov, 2021), corporate communication (Nervino, Wang, & Ma, 2025), and political discourse (Balog, 2022; Coesemans & De Cock, 2017; Kaewrungruang & Yaoharee, 2018; Liu & Liu, 2020; Vuković, 2012). Given that there are previous studies that provide detailed analyses of research on self-mentions in political discourses (see, for example, Albalat-Mascarell and Carrió-Pastor (2019)), the present literature outline does not pretend to be exhaustive. However, it summarises research on self-mentions in a variety of political discourses published in 2024-2025 (Abdulla & Ahmed, 2024; Junianto, 2025; Kapranov, 2024c; Korostenskienė & Žebelytė, 2024; Liu, 2024, 2025; Romadlani, 2024; Williams & Wright, 2024).

In this regard, a study of self-mentions in political discourse by Liu (2024) demonstrates that the use of self-mentions constitutes a critical aspect of institutional identity. Specifically, Liu (2024) asserts that self-mentions in political discourse signal how speakers perform actions in their particular institutional roles. Furthermore, Liu (2024) contends that self-mentions is a salient discursive means of manifesting and construing institutional identity within the context of political press conferences. Liu (2024) notes that political journalists in China use the Chinese analogue of the self-mention "we" in order to show their alignment with the

discursive tonality of the Chinese authorities. The use of the self-mention “we”, however, is less frequent in the context of independent press conferences.

Notably, Liu (2025) makes similar observations on the use of self-mentions in her most recent article published in 2025. In the article, Liu (2025) mentions that Chinese interpreters have to navigate a complex discursive and, importantly, ideological space afforded by political press conferences. Liu (2025) suggests that the interpreters have to convey journalists’ questions by transforming their word choices and question forms, inclusive of self-mentions. Liu’s (2025) study is reminiscent of that conducted by Phanthaphoommee and Munday (2024). Particularly, Phanthaphoommee and Munday (2024) investigate the way self-mentions are translated in political discourse by the Thai Prime Minister Prayut Chan-o-cha. The findings of their study reveal that the translations convey the explicitness of agency and exhibit the variation of stylistic choices involved in the rendering of political texts. Phanthaphoommee and Munday (2024) emphasise that self-mentions in the translation of political texts contribute to demonstrating the interpersonal overtones and presenting the image of the Thai government in a positive manner, which showcases its legitimacy and appeals to the international community.

Identically to Liu (2024, 2025), Korostenskienė and Žebelytė (2024) explore the use of self-mentions, as well as hedges in the context of political press conferences. These authors have found that self-mentions are actively employed at press conferences in the USA. Specifically, Korostenskienė and Žebelytė (2024) have established that self-mentions frequently occur in introductory phrases and comment clauses as the combinations “pronoun + think”, “pronoun + know”, and “pronoun + believe” (Korostenskienė & Žebelytė, 2024, p. 65). However, such combinations as “pronoun + feel”, “pronoun + view”, and “pronoun + imagine” are infrequent.

Also set in the context of political discourse in the USA, an article by Romadlani (2024) looks at self-mentions in Biden’s inaugural speech. Romadlani (2024) reports that Biden’s inaugural speech is marked by the frequent occurrence of the self-mention “we” (72% of all self-mentions in the speech), whilst the use of the self-mention “I” is less frequent (28%). Romadlani (2024) observes that Biden uses “I” in order to express his gratitude to the voters, whereas the self-mention “we” is employed to construe the sense of unity and togetherness. It should be noted that Biden’s use of self-mentions is analysed in a similar study conducted by Abdulla and Ahmed (2024). They have discovered that the use of self-mentions is a crucial part of Biden’s inaugural address. Identically to Abdulla and Ahmed (2024), as well as to Romadlani (2024), a study by Junianto (2025) testifies that self-mentions in Biden’s political speeches contribute to his effort to build rapport with the audience and assert his credibility. In contrast to Biden, who seems to use the self-mention “I” sparingly, a recent study by Kapranov (2024c) illustrates that King Charles III employs “I” rather frequently, at least as far as his speeches on the issue of climate change are concerned. It has been established in the study that King Charles III utilises “I” in order to impart a personalised dimension to his discourse on climate change (Kapranov, 2024c).

As far as the use of self-mentions in British political discourses is concerned, Williams and Wright (2024) argue that British officials employ self-mentions in order to allocate or rather re-allocate responsibility. Williams and Wright (2024) have discovered that the most frequently occurring self-mention in the British political discourse on the topic of COVID-19 is represented by “we”. According to

Williams and Wright (2024), the self-mention “we” is strategically deployed by British politicians in such a way that its use creates inherent ambiguity. This rhetorical move is done in order to mitigate the politicians’ own responsibility for controlling the spread of the virus.

As shown by the literature outline, there is a substantial bulk of fairly recent studies, which have been published in 2024-2025, whose research aims involve analyses of self-mentions in various types of political discourses. Presently, however, there are no published studies on self-mentions in climate change discourse by the current British prime minister Keir Starmer. Further, in section 3 of the article, a quantitative study on self-mentions in Starmer’s climate change discourse is presented.

3. The Present Study: Its Theoretical Framework, Corpus, and Methodology

First of all, the theoretical framework of the present study should be outlined. The study is grounded in the definition of self-mentions provided by Hyland (2001, 2002, 2004, 2016). As mentioned in the introduction, Hyland (2001, 2002, 2004, 2016) defines self-mentions as the first-person pronouns that manifest the authorial presence in a text. In addition, the theoretical framework of the study is informed by the publication by Albalat-Mascarell and Carrió-Pastor (2019), who posit that self-mentions are a pragmatic means of the authorial self-projection, which is associated, typically, with a positive image of self that renders credibility. The present study shares Albalat-Mascarell and Carrió-Pastor’s (2019) view of self-mentions as a means of establishing the politician’s credibility as a competent and authoritative political actor.

In this light, the study aims at (i) collecting the corpus of Starmer’s political speeches on the issue of climate change (ii) analysing the corpus quantitatively in order to answer the RQ in the study (see introduction). The procedure of corpus collection follows the prior publications by Kapranov (2025a, 2025b, 2025c), which employ a computer-mediated search for Starmer’s speeches at <https://www.gov.uk/government/> by means of applying the following keywords: *anthropogenic climate change, climate change, climate change adaptation, climate change demonstration, climate change event, climate change mass media coverage, climate change mitigation, climate change policy, climate change protest, climate risk/risks, CO2 absorption, CO2 capture and storage, CO2 emission/emissions, CO2 emission reduction/reductions, extreme weather event/events, extreme drought, extreme rain/rainfall, global warming, green energy, greenhouse gasses/GHG, green technology, Keir Starmer, net zero, rise in sea level/levels, speech, wind energy, wind farm, the consequences of climate change, and health effects of climate change*. The search has resulted in 11 speeches (the total number of words = 14 330, mean = 1 302.7, standard deviation = 699.9) delivered by Starmer on the topic of climate change.

Methodologically, the study involves a quantitative procedure of computing the frequency of self-mentions in the corpus. To that end, the computer program AntConc (Anthony, 2022) is used in order to calculate the total number of occurrences of self-mentions in the corpus. To do so, 11 speeches in the corpus are processed separately in AntConc and, thereafter, the means and standard deviations of each type of self-mentions are computed (see Kapranov and Voloshyna (2023) as well as Kapranov (2025a, 2025b, 2025c) for more details concerning the corpus

analysis). The results of the quantitative analysis and their discussion are presented in subsection 3.1 of the article.

3.1. Results and Discussion

The results of the corpus analysis indicate that there are 548 self-mentions in total (mean (M) = 68.5, standard deviation (SD) = 83.7). The descriptive statistics of the types of self-mentions, inclusive of their total number (in absolute values) per type, means, standard deviations, as well as their maximum (Max) and minimum (Min) occurrences are given in Table 1 below.

Table 1. Self-Mentions in the Corpus

#	Self-Mentions	Total N	M	SD	Max	Min
1	I	95	9.5	9.7	28	1
2	Me	5	2.5	0.5	3	2
3	Mine	-	-	-	-	-
4	My	27	3.9	2.4	7	1
5	Myself	1	0	0	1	1
6	We	244	22.2	20.9	65	1
7	Us	18	3.0	2.2	7	1
8	Our	155	15.5	14.9	54	3
9	Ours	-	-	-	-	-
10	Ourselves	3	1.5	0.5	2	1

The results of the corpus analysis reveal that Starmer's speeches on the issue of climate change involve the following types of self-mentions: "I", "me", "my", "myself", "we", "us", "our", and "ourselves". Judging from these findings, Starmer and his speechwriters do not employ the self-mentions "mine" and "ours", respectively. Instead, they seem to capitalise on the frequent use of the self-mentions "we", "our", and "I", as illustrated by Figure 1, which plots the normalised frequencies (per 1 000 words) of the self-mentions in the present corpus.

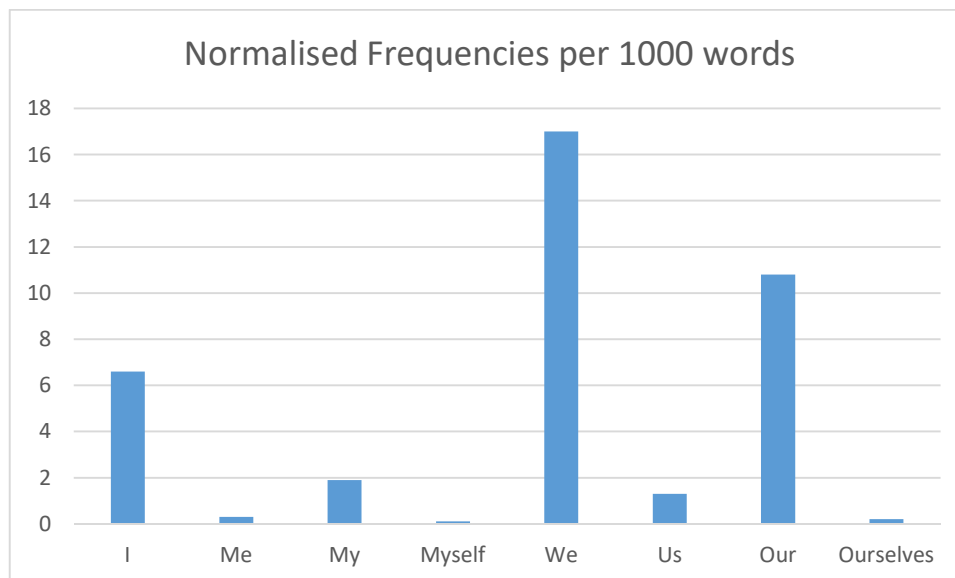


Figure 1. The Normalised Frequencies of Self-Mentions

As exemplified by Figure 1, the self-mentions “we” and “our” seem to dominate the corpus of speeches on climate change by Starmer. Let us discuss these self-mentions and illustrate their use. The self-mention “we” both in absolute values (see Table 1) and in normalised values per 1 000 words (see Figure 1) is the most frequently occurring self-mention in the corpus. This finding lends indirect support to the prior literature (Abdulla & Ahmed, 2024; Junianto, 2025; Liu, 2024; Phanthaphoommee & Munday, 2024; Romadlani, 2024; Williams & Wright, 2024), which posits that “we” is a highly frequent self-mention in political discourses, especially in the Anglophone ones. As shown by Table 1, Starmer’s speeches on the issue of climate change are clearly marked by his preference to use the self-mention “we”, as illustrated by excerpt (1) below.

(1) To deliver on the Paris Agreement and keep 1.5 degrees within reach, in the first 100 days of my government **we** launched Great British Energy – to create clean British power. **We** created a National Wealth Fund – to invest in the green industries and jobs of the future. **We** scrapped the ban on onshore wind, committed to no new North Sea oil and gas licences and closed the UK’s final coal power plant at the end of September – becoming the first G7 economy to phase out coal power. (Starmer, 2024a)

In (1), the self-mention “we” is employed in its exclusive form. In this regard, it should be observed that the literature (Hyland, 2001, 2002, 2004; Kapranov, 2024c) distinguishes between two forms of the self-mention “we”, the so-called inclusive “we” and exclusive “we”, respectively. The inclusive “we” presupposes such communicative situations, in which the speaker involves the addressee into the same domain, i.e. creates discursive togetherness with the audience, in which the speaker and the audience are treated as one collective body of people. In contrast to the inclusive “we”, the exclusive “we” pertains to such communicative situations, in which the speaker does not involve the addressee into the “communal we”, but emphasises the fact that the speaker and the group of people that the speaker represents are different from the audience. Starmer’s use of the exclusive “we” is evident from (1), in which all of the instances of “we” are exclusive.

Now, let us explain why “we” in excerpt (1) can be considered exclusive. Excerpt (1) is taken from Starmer’s speech on climate change, which he delivered at the UN Climate Change Conference in Baku (Azerbaijan) on 12 November 2024. Obviously, Starmer addresses the international audiences there, which are comprised of the attendees from all over the world. Accordingly, we may contend that his use of “we” refers to his government and his party, thus excluding the participants who attend the UN Climate Change Conference in Baku. In other words, he explains to the audience what he and his government are doing and refers to that as a communal effort, or a team effort for that matter, which excludes the audience. Hence, the use of “we” in (1) is exclusive. In this regard, the analysis of the corpus shows that nearly all the instances of his use of “we” in relation to the issue of climate change are exclusive. Only on rare occasions, he employs the inclusive “we” in a rhetorical attempt to address the whole of humankind or to emphasise that the UK and one more or several countries are involved in climate change-related measures to mitigate the negative consequences of climate change.

In conjunction with the frequency of “we”, it could be argued that a relatively high occurrence of the self-mention “our” seems to be quite logical in the present corpus. Its use is illustrated by excerpt (2) below.

(2) That is in addition to the recent investment in carbon capture in Teesside and Merseyside, which will create 4,000 jobs, and the investment announced by My Rt Hon Friend the Chancellor for 11 new green hydrogen projects across Britain. Tackling climate change is, of course, a global effort so at the G20, together with Brazil and 10 other countries, I launched **our** global clean power alliance to speed up the international roll-out of clean power, accelerate investment and cut emissions around the world. (Starmer, 2024b)

In (2), we can argue that the self-mention “our” is used in its inclusive form, given that Starmer talks about a global alliance with Brazil and other G20 countries whose aim is to tackle the climate crisis. As mentioned before, however, the inclusive use of “our” and “we” is rather infrequent in the corpus.

Finally, let us discuss another major finding that consists in Starmer’s relatively frequent use of the self-mention “I”. Whilst “I” is the third frequently occurring self-mention in the corpus after “we” and “our”, Starmer’s political speeches on the issue of climate change can hardly be defined in terms of the “I” stance. This finding seems to be in contrast to the prior study conducted by Kapranov (2024c), who has discovered that the currently reigning British monarch King Charles III uses or, perhaps, overuses the self-mention “I” in his discourse on climate change. Unlike King Charles III, Starmer employs the self-mention “I” more frugally. Starmer’s use of “I” is emblematised by excerpt (3) below.

(3) Meanwhile climate change hits economic growth, leaves us exposed to catastrophic flooding and both of these forces drive unsustainable levels of migration. It all manifests in a feeling amongst very many people that the system isn’t working for them. That it’s time to take back control of our lives, our borders, our livelihoods. And **I** hear that. People want action – and they want change. So, we will not dismiss those concerns. We will answer them. And we will do it by acting at home, absolutely, and also by using our strength abroad. Because in this new era we need to do things differently. (Starmer, 2024c)

In (3), Starmer resorts to the use of “I” in an attempt to amplify the fact that he actually listens to people’s concerns in relation to the current climate crisis. Arguably, he tries to emphasise that not only his government and the Labour Party attend to the vox populi, but he himself listens to the voters as the British prime minister as far as the issue of climate change is concerned. In other words, whilst normally Starmer showcases the team effort of tackling the climate crisis and, accordingly, employs “we”, he also avails himself of the use of “I” when it is needed to depict a picture of him as the party leader and, concurrently, as the incumbent British prime minister who is aware of the magnitude of the climate crisis.

4. Conclusions

This contribution involves a quantitative study that aims to learn about Keir Starmer’s use of self-mentions in the context of his political speeches on climate change. Having analysed the corpus of his speeches, it seems possible to conclude that Starmer employs the self-mention “we” and “our” the most. This finding is supported by a number of prior studies, which point to the frequent occurrence of “we” in political discourses. Evidently, Starmer’s frequent use of “we” in his speeches on climate change aligns his climate change discourse with the analogous

discourses by Anglophone and international political actors. Furthermore, the frequent use of “we” by Starmer is indicative of his rhetorical strategy to portray himself as a team player, who is engaged in combatting the climate crisis together with his government and the Labour Party. In other words, Starmer’s strategic choice of using the self-mention “we” provides a deeper insight into the way he sees and talks about the issue of climate change, which he, apparently, regards as a challenge that requires a substantial team effort. We may conclude this study by contending that Starmer deploys self-mentions, especially, the self-mention “we” and its forms in order to portray his public image as a team player. Presumably, his image as a team player is built by his communication team in order to present a relatable and positive image that could resonate with the public at large, at least as far as the issue of climate change is concerned.

Acknowledgements

The author is appreciative of the editor and two anonymous reviewers.

Primary Sources

<https://www.gov.uk/government/>

References

- Abdulla, F., & Ahmed, S. (2024). A stylistic analysis of Joe Biden's inaugural political speech. *International Journal of Linguistics, Literature & Translation*, 7(7), 1–7. <https://doi.org/10.32996/ijlnt.2024.7.7.1>.
- Albalat-Mascarell, A., & Carrió-Pastor, M. L. (2019). Self-representation in political campaign talk: A functional metadiscourse approach to self-mentions in televised presidential debates. *Journal of Pragmatics*, 147, 86–99. <https://doi.org/10.1016/j.pragma.2019.05.011>.
- Anthony, L. (2022). *AntConc Version 4.0.11*. Tokyo: Waseda University.
- Balog, P. (2022). An analysis of Queen Elizabeth II’s coronavirus speech using Hyland’s metadiscourse theory. *International Journal of Arts, Sciences and Education*, 3(1), 107–120.
- Boykoff, M. T. (2008). The cultural politics of climate change discourse in UK tabloids. *Political Geography*, 27(5), 549–569. <https://doi.org/10.1016/j.polgeo.2008.05.002>.
- Burns, P. M., & Cowlshaw, C. (2014). Climate change discourses: How UK airlines communicate their case to the public. *Journal of Sustainable Tourism*, 22(5), 750–767. <https://doi.org/10.1080/09669582.2014.884101>.
- Carrió-Pastor, M. L. (2020). Variation in the use of self-mentions in different specific fields of knowledge in academic English. In M. L. Carrió-Pastor (ed.) *Corpus Analysis in Different Genres* (pp. 13–32). New York: Routledge. <https://doi.org/10.4324/9780367815905>.
- Coesemans, R., & De Cock, B. (2017). Self-reference by politicians on Twitter: Strategies to adapt to 140 characters. *Journal of Pragmatics*, 116, 37–50. <https://doi.org/10.1016/j.pragma.2016.12.005>.
- Dahl, T., & Fløttum, K. (2019). Climate change as a corporate strategy issue: A discourse analysis of three climate reports from the energy sector. *Corporate Communications: An International Journal*, 24(3), 499–514. <https://doi.org/10.1108/CCIJ-08-2018-0088>.
- Diamond, P., Richards, D., & Warner, S. (2025). Change and continuity in British politics: Can the Starmer government’s approach to governance resolve the crisis in the British state without radical reform?. *The Political Quarterly*, 96(1), 140–148. <https://doi.org/10.1111/1467-923X.13474>.
- Ferguson, J., Sales de Aguiar, T. R., & Fearfull, A. (2016). Corporate response to climate change: language, power and symbolic construction. *Accounting, Auditing & Accountability Journal*, 29(2), 278–304. <https://doi.org/10.1108/AAAJ-09-2013-1465>.

- Fløttum, K. (2010). A linguistic and discursive view on climate change discourse. *ASp. la revue du GERAS*, 58, 19–37. <https://doi.org/10.4000/asp.1793>.
- Fløttum, K. (2014). Linguistic mediation of climate change discourse. *ASp. la revue du GERAS*, 65, 7–20. <https://doi.org/10.4000/asp.4182>.
- Fløttum, K. (2019). A cross-disciplinary perspective on climate change discourse. In I. Simonnæs, Ø. Andersen & K. Schubert (eds.), *New Challenges for Research on Language for Special Purposes* (pp. 21–37). Berlin: Frand & Timme.
- Fløttum, K., & Gjerstad, Ø. (2017). Narratives in climate change discourse. *Wiley Interdisciplinary Reviews: Climate Change*, 8(1), 1–15. <https://doi.org/10.1002/wcc.429>.
- Gavin, N. T., & Marshall, T. (2011). Mediated climate change in Britain: Scepticism on the web and on television around Copenhagen. *Global Environmental Change*, 21(3), 1035–1044. <https://doi.org/10.1016/j.gloenvcha.2011.03.007>.
- Hayes, S., & O'Neill, S. (2021). The Greta effect: Visualising climate protest in UK media and the Getty images collections. *Global Environmental Change*, 71, 1–26. <https://doi.org/10.1016/j.gloenvcha.2021.102392>.
- Hyland, K. (2001). Humble servants of the discipline? Self-mention in research articles. *English for Specific Purposes*, 20(3), 207–226. [https://doi.org/10.1016/S0889-4906\(00\)00012-0](https://doi.org/10.1016/S0889-4906(00)00012-0).
- Hyland, K. (2002). Options of identity in academic writing. *ELT Journal*, 56(4), 351–358. <https://doi.org/10.1093/elt/56.4.351>.
- Hyland, K. (2004). Disciplinary interactions: Metadiscourse in L2 postgraduate writing. *Journal of Second Language Writing*, 13(2), 133–151. <https://doi.org/10.1016/j.jslw.2004.02.001>.
- Hyland, K. (2016). Writing with attitude: Conveying a stance in academic texts. In E. Hinkel (ed.), *Teaching English Grammar to Speakers of Other Languages* (pp. 246–265). New York: Routledge. <https://doi.org/10.4324/9781315695273>.
- Junianto, M. I. F. (2025). Metadiscourse markers as persuasive strategies in Joe Biden's 2024 State of the Union Address. *Caruban Proceeding*, 1(1), 110–126.
- Kaewrungruang, K., & Yaoharee, O. (2018). The use of personal pronoun in political discourse: A case study of the final 2016 united states presidential election debate. *Reflections*, 25(1), 85–96.
- Kapranov, O. (2015a). Conceptual metaphors in Ukrainian prime ministers' discourse involving renewables. *Topics in Linguistics*, 16(1), 4–16. <https://doi.org/10.2478/topling-2015-0007>.
- Kapranov, O. (2015b). Do international corporations speak in one voice on the issue of global climate change: The case of British Petroleum and The Royal Dutch Shell Group. In C. Can, A. Kilimci, & K. Papaja (eds.) *Social Sciences and Humanities: A Global Perspective* (pp. 306–322). Ankara: Detay Yayıncılık.
- Kapranov, O. (2017). Conceptual metaphors associated with climate change in corporate reports in the fossil fuels market: Two perspectives from the United States and Australia. In K. Fløttum (ed.), *The Role of Language in the Climate Change Debate* (pp. 90–109). London: Routledge.
- Kapranov, O. (2018). The framing of climate change discourse by Statoil. *Topics in Linguistics*, 19(1), 54–68. <https://doi.org/10.2478/topling-2018-0004>.
- Kapranov, O. (2021). Self-mention in academic writing by in-service teachers of English: Exploring authorial voices. In M. Burada, O. Tatu, & R. Sinu (eds.), *Exploring Language Variation, Diversity and Change* (pp. 76–100). Newcastle upon Tyne: Cambridge Scholar Publishing.
- Kapranov, O. (2024a). Between a burden and green technology: Rishi Sunak's framing of climate change discourse on Facebook and X (Twitter). *Information & Media*, 99, 85–105. <https://doi.org/10.15388/Im.2024.99.5>.
- Kapranov, O. (2024b). Spraying paint on Stonehenge: The framing of climate change protest by the leading Anglophone media. *Culture. Society. Economy. Politics*, 4(2), 10–26. <https://doi.org/10.2478/csep-2024-0008>.
- Kapranov, O. (2024c). Self-mentions in climate change discourse by King Charles III. *Journal of Contemporary Philology*, 7(1), 29–46. <https://doi.org/10.37834/JCP2471029k>.

- Kapranov, O. (2025a). Keir Starmer's framing of climate change: A qualitative investigation of political speeches. *Acta Marisiensis. Philologia*, 7, 36–48. <https://doi.org/10.62838/amph-2025-0131>.
- Kapranov, O. (2025b). Syntactic means in Keir Starmer's speeches on climate change: An Investigation of dependent clauses. *Annals of the Faculty of Philology (Анали Филолошког факултета)*, 37(1), 13–28. <https://doi.org/10.18485/analiff.2025.37.1.1>.
- Kapranov, O. (2025c). Discourse markers in Keir Starmer's speeches on climate change. *Studies in Linguistics, Culture, and FLT*, 13(1), 8–26. <https://doi.org/10.46687/CUPP2991>.
- Kapranov, O., & Voloshyna, O. (2023). Learning English under the sounds of air raid sirens: Analysing undergraduate EFL students' sustainable learning practices. *Sustainable Multilingualism/Darnioji daugiakalbystė*, 23, 1–24. <https://doi.org/10.2478/sm-2023-0011>
- Korostenskienė, J., & Žebelytė, A. (2024). Quod licet Iovi: Hedges in US politicians' press conference speeches. *Brno studies in English*, 50(2), 51–84. <https://doi.org/10.5817/BSE2024-2-3>.
- Moorcroft, H., Hampton, S., & Whitmarsh, L. (2025). Climate change and wealth: understanding and improving the carbon capability of the wealthiest people in the UK. *PLOS Climate*, 4(3), 1–25. <https://doi.org/10.1371/journal.pclm.0000573>.
- Liu, R. Y. (2024). Constituting institutional identity in political discourse: The use of the first-person plural pronoun in China's press conferences. *Language in Society*, 53(3), 421–444. <https://doi.org/10.1017/S0047404523000386>.
- Liu, R. Y. (2025). Interpreters as spin doctors: The interactional role of interpreters in China's political press conferences. *The International Journal of Press/Politics*, 30(1), 423–444. <https://doi.org/10.1177/19401612231204514>.
- Liu, S. L., & Liu, Y. L. (2020). A comparative study of interactional metadiscourse in English speeches of Chinese and American stateswomen. *DEStech Transactions on Social Science Education and Human Science*, 1, 364–368. <https://doi.org/10.12783/dtssehs/icesd2020/34442>.
- Nervino, E., Wang, J., & Ma, C. (2025). Self-mentions and stakeholders in climate change discourse: The case of banks. *Language and Dialogue*, 15(1), 56–80. <https://doi.org/10.1075/ld.00187.ner>.
- Nisbett, N., Spaiser, V., Leston-Bandeira, C., & Valdenegro, D. (2025). Climate action or delay: the dynamics of competing narratives in the UK political sphere and the influence of climate protest. *Climate Policy*, 25(4), 513–526. <https://doi.org/10.1080/14693062.2024.2398169>.
- Phanthaphoommee, N., & Munday, J. (2024). Pronoun shifts in political discourse: The English translations of the Thai Prime Minister Prayut Chan-o-cha's statements on the international stage. *Babel*, 70(6), 825–851. <https://doi.org/10.1075/babel.00388.pha>.
- Romadlani, M. M. I. (2024). Personal pronouns in Biden's inaugural speech: A critical discourse perspective. *Journal of Language and Literature*, 24(1), 252–262. <https://doi.org/10.24071/joll.v24i1.6330>.
- Starmer, K. (2024a). *Prime Minister's National Statement at COP29: 12 November 2024*. Accessed on 1.06.2025 at <https://www.gov.uk/government/speeches/prime-ministers-national-statement-at-cop29-12-november-2024>.
- Starmer, K. (2024b). *PM Statement to the House of Commons on the G20 and COP29 Summits: 21 November 2024*. Accessed on 1.06.2025 at <https://www.gov.uk/government/speeches/pm-statement-to-the-house-of-commons-on-the-g20-and-cop29-summits-21-november-2024>.
- Starmer, K. (2024c). *PM Speech at Lord Mayor's Banquet: 2 December 2024*. Accessed on 1.06.2025 at <https://www.gov.uk/government/speeches/pm-speech-at-lord-mayors-banquet-2-december-2024>.
- Swain, K. A. (2025). Media framing of climate change mitigation and adaptation. In M. Lackner, B. Sajjadi, & W.-Y. Chen (eds.) *Handbook of Climate Change Mitigation and Adaptation* (pp. 4165–4233). Cham: Springer Nature. https://doi.org/10.1007/978-3-031-84483-6_6.

- Vuković, M. (2012). Positioning in pre-prepared and spontaneous parliamentary discourse: Choice of person in the Parliament of Montenegro. *Discourse & Society*, 23(2), 184–202. <https://doi.org/10.1177/0957926511431507>.
- Williams, J., & Wright, D. (2024). Ambiguity, responsibility and political action in the UK daily COVID-19 briefings. *Critical Discourse Studies*, 21(1), 76–91. <https://doi.org/10.1080/17405904.2022.2110132>.
- Willis, R. (2017). Taming the climate? Corpus analysis of politicians' speech on climate change. *Environmental Politics*, 26(2), 212–231. <https://doi.org/10.1080/09644016.2016.1274504>.
- Wright, K., & Irwin, S. (2025). Not talking about climate change: everyday interactions, relational work and climate silences. *Sociology*, 59(3), 406–423. <https://doi.org/10.1177/00380385241289743>.

Резюме

Капранов Олександр

САМОЗГАДКИ У ПРОМОВАХ ПРЕМ'ЄР-МІНІСТРА ВЕЛИКОБРИТАНІЇ ПРО ЗМІНУ КЛІМАТУ

Постановка проблеми. У Великій Британії питання зміни клімату стало спірною проблемою, яку регулярно висвітлюють британські ЗМІ, широка громадськість та низка протестних рухів. Незважаючи на те, що у Великій Британії звучать конкуруючі дискурсивні голоси щодо зміни клімату, видається цілком розумним прислухатися до дискурсу щодо зміни клімату, який проводить нинішній лейбористський уряд на чолі з Кіром Стармером. Хоча його прем'єрство недовге, тим не менш, видається доречним розглянути дискурс Стармера щодо зміни клімату детальніше, враховуючи, що саме Стармер визначає політику уряду Великої Британії, пов'язану з цим питанням. Варто зазначити, що існує кілька нещодавніх досліджень лексико-синтаксичних особливостей дискурсу Стармера щодо зміни клімату. Водночас, багато дискурсивних та риторичних аспектів дискурсу Стармера щодо зміни клімату ще не було проаналізовано.

Мета. У статті подано результати кількісного дослідження, яке вивчає частоту самозгадок, що проявляються через займенники першої особи в промовах Стармера про зміну клімату. Стверджується, що самозгадки можуть допомогти розкрити те, як Стармер формулює проблему зміни клімату, наприклад, чи позиціонує він себе як командного гравця, який розглядає зміну клімату як виклик, що стоїть перед усім урядом та/або всією британською нацією, або, навпаки, його використання самозгадок може свідчити про те, що він зображує себе як єдиного борця проти негативних наслідків зміни клімату. У цьому світлі стаття має на меті відповісти на таке дослідницьке питання: якою є частота самозгадок (тобто, “I”, “my”, “me”, “mine”, “myself”, “we”, “our”, “us”, “ours”, and “ourselves”) у політичних промовах Кіра Стармера про зміну клімату.

Методи. Методологічно дослідження передбачає кількісну процедуру обчислення частоти самозгадувань у корпусі. Для цього використано комп'ютерну програму AntConc для обчислення загальної кількості випадків самозгадувань у корпусі. Для цього промови Стармера в корпусі оброблено окремо в AntConc, а потім обчислено середні значення та стандартні відхилення кожного типу самозгадувань.

Результати. Результати аналізу показують, що промови Стармера з питань зміни клімату містять такі типи самозгадок: “I”, “me”, “my”, “myself”, “we”,

“us”, “our”, and “ourselves”. Судячи з цих висновків, Стармер та/або його спічрайтери не використовують самозгадки “mine” та “ours”, відповідно. Натомість вони, здається, використовують “we”, “our” та “I”.

Дискусія. Самозгадування “we” як в абсолютних значеннях, так і в нормалізованих значеннях на 1000 слів є найчастішим самозгадуванням у корпусі. Цей висновок опосередковано підтверджує попередні дослідження (Abdulla & Ahmed, 2024; Junianto, 2025; Liu, 2024; Phanthaphoommee & Munday, 2024; Romadlani, 2024; Williams & Wright, 2024), в яких стверджується, що “we” є дуже частим самозгадуванням у політичних дискурсах, особливо в англomовних. Стармер досить часто використовує “we” у його ексклюзивній формі. У зв’язку з цим у літературі (Hyland, 2001, 2002, 2004; Kapranov, 2024c) розрізняють дві форми самозгадування “we”: так зване інклюзивне “we” та ексклюзивне “we” відповідно. Інклюзивне “we” передбачає такі комунікативні ситуації, в яких мовець залучає адресата до тієї ж сфери, тобто створює дискурсивну єдність з аудиторією, в якій мовець та аудиторія розглядаються як єдиний колектив людей. На відміну від інклюзивного “we”, ексклюзивне “we” стосується таких комунікативних ситуацій, в яких мовець не залучає адресата до “спільного we”, а підкреслює той факт, що мовець та група людей, яку він представляє, не є тим самим, що й адресат, тобто мовець відрізняється від аудиторії. Використання Стармером ексклюзивного «ми» підтверджено результатами аналізу корпусу, в якому майже всі випадки вживання “we” є ексклюзивними.

Ключові слова: дискурс щодо зміни клімату, Кір Стармер, політичні промови, прем’єр-міністр, самозгадки, Велика Британія.

Abstract

Kapranov Oleksandr

SELF-MENTIONS IN BRITISH PRIME MINISTER’S SPEECHES ON CLIMATE CHANGE

Background. In the United Kingdom (the UK), the issue of climate change has become a contentious problem that is regularly elucidated by the British media, the public at large, and by a number of climate change protest movements. Whilst there are competing discursive voices on climate change in the UK, it seems quite sensible to heed to the discourse on climate change by the current Labour government, which is led by Keir Starmer. Whereas his premiership is short, it seems, nevertheless, pertinent to look at Starmer’s discourse on climate change in more detail, given that Starmer determines the UK’s government policies associated with the issue of climate change. In this regard, it should be mentioned that there are several recent studies on the lexico-syntactic peculiarities of Starmer’s climate change discourse. At the same time, however, a score of other discursive and rhetorical aspects of Starmer’s discourse on climate change have not been analysed yet.

Purpose. The article involves a quantitative study that examines the frequency of self-mentions, which are manifested by the first person pronouns in Starmer’s speeches on climate change. It is argued in the study self-mentions may help to uncover the way Starmer frames the issue climate change, for instance, whether or not he presents himself as a team player, who looks at climate change as a challenge posed to the entire government and/or the entire British nation, or, alternatively, his use of self-mentions may reveal that he portrays himself as the sole fighter against

the negative consequences of climate change. In this light, the study aims at answering the following **research question**: What is the frequency of self-mentions (i.e., “I”, “my”, “me”, “mine”, “myself”, “we”, “our”, “us”, “ours”, and “ourselves”) in Keir Starmer’s political speeches on climate change?

Methods. Methodologically, the study involves a quantitative procedure of computing the frequency of self-mentions in the corpus. To that end, the computer program AntConc is used in order to calculate the total number of occurrences of self-mentions in the corpus. To do so, Starmer’s speeches in corpus are processed separately in AntConc and, thereafter, the respective means and standard deviations of each type of self-mentions are computed.

Results. The results of the corpus analysis reveal that Starmer’s speeches on the issue of climate change involve the following types of self-mentions: “I”, “me”, “my”, “myself”, “we”, “us”, “our”, and “ourselves”. Judging from these findings, Starmer and his speechwriters do not employ the self-mentions “mine” and “ours”, respectively. Instead, they seem to capitalise on the frequent use of the self-mentions “we”, “our”, and “I”.

Discussion. The self-mention “we” both in absolute values and in normalised values per 1 000 words is the most frequently occurring self-mention in the corpus. This finding lends indirect support to the prior literature (Abdulla & Ahmed, 2024; Junianto, 2025; Liu, 2024; Phanthaphoommee & Munday, 2024; Romadlani, 2024; Williams & Wright, 2024), which posits that “we” seems to be a highly frequent self-mention in political discourses, especially in the Anglophone ones. “We” is employed by Starmer in its exclusive form rather often. In this regard, the literature (Hyland, 2001, 2002, 2004; Kapranov, 2024c) distinguishes between two forms of the self-mention “we”, the so-called inclusive “we” and exclusive “we”, respectively. The inclusive “we” presupposes such communicative situations, in which the speaker involves the addressee into the same domain, i.e. creates discursive togetherness with the audience, in which the speaker and the audience are treated as one collective body of people. In contrast to the inclusive “we”, the exclusive “we” pertains to such communicative situations, in which the speaker does not involve the addressee into the “communal we”, but emphasises the fact that the speaker and the group of people that the speaker represents are not the same as the addressee, i.e. the speaker is different from the audience. Starmer’s use of the exclusive “we” is evident from the corpus, in which nearly all of the instances of “we” are exclusive.

Keywords: climate change discourse, Keir Starmer, political speeches, prime minister, self-mentions, the UK.

Відомості про автора

Капранов Олександр, доктор філософії, доцент, NLA Коледж в Осло (Норвегія), e-mail: oleksandr.kapranov@nla.no

Kapranov Oleksandr, Dr, associate professor, NLA University College (Norway), e-mail: oleksandr.kapranov@nla.no

ORCID 0000-0002-9056-331

Надійшла до редакції 02 листопада 2025 року
Прийнято до друку 16 грудня 2025 року

UDC 81'22 : 651.926

DOI: 10.24025/2707-0573.12.2025.346654

Ліана Голетіані



МЕТАФОРА В УКРАЇНСЬКІЙ ТЕРМІНОЛОГІЇ
ВТОРИННОГО ПРАВА ЄС:
ПОРІВНЯЛЬНЕ ДОСЛІДЖЕННЯ

У статті досліджено метафоричні терміни та термінологічні словосполучення, які містять метафоричні компоненти. У центрі уваги знаходяться українські відповідники англomовних термінів, використаних в офіційних актах вторинного права ЄС. Для порівняльного аналізу було обрано термінологічний матеріал англійської, італійської, німецької, польської та словацької мов. Показано, що під час перекладу метафоричних термінів можуть відбуватися такі процеси: 1) збереження метафори; 2) заміна іншою метафорою; 3) дeметафоризація; 4) додаткова метафоризація термінів або компонентів термінологічних словосполучень. У виняткових випадках під час перекладу може виникати пряме запозичення, гібридний термін та термін-пояснення.

Ключові слова: українська термінологія права ЄС, переклад права ЄС, мовний контакт, метафора у термінології, метафора у спеціальному перекладі.

1. Вступ

Започаткований у 90-ті роки курс України на зближення з Європейським Союзом викликав необхідність перекладу українською мовою документів європейського права. У 2004 році було прийнято загальнодержавну програму адаптації законодавства України до законодавства ЄС (ЗУ № 1629-IV). Подальші політичні рішення держави у цьому напрямі не лише зміцнювали інтеграцію з ЄС у широкому спектрі політичної, соціальної, фінансової, економічної, торговельної, наукової, освітньої та культурної співпраці, а й ставали новим стимулом у розвитку юридичної української мови та її термінології в цих секторах.

23 червня 2022 року Європейська Рада надала Україні статус кандидата на вступ до Європейського Союзу. Однією з найважливіших умов, що висуваються Європейським Союзом до держав-кандидатів, є адаптація їхнього національного законодавства до законодавства ЄС та досягнення відповідності національної правової системи *acquis communautaire*. Ключову роль у цьому складному, багатоетапному процесі відіграє складання текстів правових актів ЄС офіційною мовою країни-кандидата. Програма та порядок перекладу *acquis communautaire* українською мовою тепер щорічно затверджуються постановою

© Голетіані Л., 2025

Кабінету Міністрів України після попереднього схвалення Комітетом Верховної Ради України з питань європейської інтеграції. Переклад повинен відповідати критеріям адекватності та еквівалентності. Під час перекладу має проводитися ретельна термінологічна перевірка європейських юридичних термінів.

Як відомо, оптимальним способом перекладу є виявлення у мові перекладу (МП) еквівалента терміна мови оригіналу (МО). Застосування цього способу у галузі права можливе лише за умови, коли законодавчі системи, які користуються МП та МО, досягли одного й того самого рівня нормативного розвитку.

З огляду на сучасний стан євролектів, розрізняють зрілі та усталені терміносистеми держав-засновників Європейського Союзу (зокрема, німецьку, французьку, італійську, а також англійську як основну робочу мову) та молодші терміносистеми країн, що вступили до Європейського Союзу пізніше (наприклад, терміносистеми балтійських і слов'янських мов). Серед них наймолодшим слов'янським євролектом прийнято вважати словацький (Бель, 2024, с. 65). Особливе значення для порівняльних досліджень у галузі слов'янських мов має польський євролект.

Абсолютно природно, що переклад *acquis communautaire* є величезним викликом для України, як і для всіх країн-кандидатів (Šarčević, 2015, с. 183; Бель, 2024, с. 66). Саме на цьому етапі переклади не лише слугують основою для гармонізації національного законодавства з правом ЄС, але й відіграють важливу роль у формуванні відповідного євролекту та у створенні всієї термінології ЄС відповідною мовою. Це вимагає співпраці між перекладачами, юристами, експертами з різних технічних галузей та мовознавцями.

Мовні аспекти термінології ЄС та проблеми перекладу є об'єктом досліджень окремих українських мовознавців та фахівців з перекладу і термінології¹, разом з тим фундаментальних розвідок різної проблематики в українському перекладознавстві на сьогодні ще бракує (Касяненко, 2015, с. 125). Проблема метафори у перекладах актів європейського вторинного права українською мовою ще не була об'єктом спеціального дослідження, що визначає актуальність нашої розвідки.

2. Теоретичне підґрунтя

Оскільки українська термінологія в галузі права ЄС ще молода, у перекладах *acquis communautaire* можна очікувати значний ступінь термінологічної варіативності різного роду, яка є одною з актуальних проблем сучасної лінгвістики (Голубовська та ін., 2016) і вимагає поглиблених досліджень. До чинників, які додатково спричиняють варіативність у перекладі документів ЄС, польська дослідниця Луцзя Бель відносить метафоризацію (Biel, 2023, с. 115). Дарья Касяненко також вважає метафору джерелом поповнення терміносистеми правового євролекту (Касяненко, 2011, с. 105). Звісно,

¹Термінологічній роботі в контексті мультилінгвального права ЄС присвячено окремий розділ «Термінологія та прикладна термінологічна робота в контексті перекладу офіційних документів ЄС» (Байчич, 2024) у складі першого в Україні «Посібника з перекладу актів права ЄС українською мовою» (Байчич, Робертсон, Славова 2024). Серед критеріїв відбору найбільш вдалих відповідників термінів Мартіна Байчич називає: «прозорість / умотивованість, послідовність, відповідність (для цільових груп), мовну економію, лаконічність, дериваційний потенціал, мовну коректність та перевагу рідної мови» (там само, с. 96).

метафора існує і в питомій українській термінології. Скористаємося з цього приводу визначенням метафоричного терміна, запропонованим Максимом Вакуленком:

Термінологічне слово є результатом вторинної номінації, оскільки постає внаслідок якісних змін у семантичній структурі загальновживаної лексеми і виникає в результаті перенесення лексичного значення за подібністю (метафора) [...] Відтак слово набуває нового значення і функціонує подвійно: як загальновживана лексема і як термін (Вакуленко, 2023, с. 69).

Переклад метафоричного терміна або компонента термінологічного словосполучення може викликати серйозні труднощі, якщо відповідна загальновживана лексема в мові перекладу не має переносного значення або воно є, але не термінологізується. Адекватне відтворення метафори і дотримання принципу міжмовної термінологічної узгодженості в перекладі права ЄС може виявитись неможливим. Виникає проблема пошуку оптимального, часто компромісного, способу передавання метафори, у її вирішенні перекладач-практик має балансувати між різними критеріями. Цю проблему як стилістичну вже порушено у згаданому «Посібнику з перекладу актів права ЄС українською мовою» (2024), зокрема, у розділі «Граматичні та стилістичні особливості перекладу актів права ЄС» (Палійчук, Батіна, Кабанова, 2024). Наголошено наступне:

адекватний та якісний переклад українською мовою вимагає від перекладача проведення первинного стилістичного аналізу з урахуванням унікальних характеристик євролекту, використовуваного в офіційних мовах ЄС, а також подальшого лінгвоюридичного редагування підготовленого тексту, що є стандартною законодавчою практикою Європейського Союзу (там само, с.194).

Автори наводять кілька метафоричних термінів з україномовної версії Повідомлення Комісії 2013/С 216/01 – *здоровий фінансовий сектор, здорові банки, осередки вразливості*, які є еквівалентними відповідниками для англійських метафоричних термінів *healthy financial sector, healthy banks, pockets of vulnerability*, не заглиблюючись далі у проблему можливої відсутності метафоричного еквівалента. Ця відсутність може бути обумовлена прогалинами у концептуальних відображеннях. Згідно з підходом Золтана Кьовечеша (Kövecses, 2005), метафори, які походять з різних культурних концептуальних систем, треба перекладати, використовуючи когнітивні та прагматичні рамки. Подальшим кроком у цьому напрямку стала когнітивна гіпотеза перекладу Нілі Мандельбліта (Mandelblit, 1995), що передбачає розуміння основних концепцій культури МО та пошук прагматично еквівалентних одиниць у культурі МП.

Однією з найвідоміших праць, у якій порушено проблему перекладу метафори, стала невелика стаття Пітера Ньюмарка 1980 року «Переклад метафори» (Newmark, 1980). У ній автор спочатку пропонує категоризацію можливих типів метафор, а потім подає низку можливих стратегій перекладу, які можна застосовувати залежно від типу метафори. Дослідник розрізняє традиційні, або мертві метафори, які можна перекладати буквально, та

оригінальні, або нові метафори, які потребують більш складних стратегій, таких як переклад через пояснення, порівняння, аналогію або переклад за допомогою нової метафори, яка має подібну еквівалентну функцію. Комунікативний ефект оригінальної метафори можна відтворити за допомогою іншого стилістичного засобу, якщо прямий метафоричний переклад неможливий². Очевидно, що така стратегія, як і заміна метафори порівнянням, у перекладі спеціального тексту є неприпустимою. На наш погляд, метафоричні терміни у нормативному тексті не зовсім доцільно розглядати як стилістичну рису, вони не виконують тут риторичних функцій. На перший план виходить їхня термінологічна функція, тобто термінологічність тут важливіша за метафоричність. Тому і підходити до проблеми їхнього перекладу необхідно, перш за все, як до проблеми перекладу спеціальної лексики.

Процес прийняття конкретних рішень спеціалізованим перекладачем та обмеження, що впливають на його вибір під час перекладу спеціального тексту, стали об'єктом уваги Джузеппе Палумбо у праці 1999 року (Palumbo, 1999). Дослідник виявив такі процедури:

- аналогічний переклад;
- описовий переклад;
- пояснення;
- запозичення;
- неологізм;
- опущення.

Зазначимо, що Палумбо аналізував переклад термінології у спеціальному тексті на матеріалі навчальної літератури, тому деякі виокремлені процедури можуть бути нерелевантними для перекладу термінології нормативного тексту. Переклад метафори в офіційних актах ЄС безумовно матиме інші особливості.

3. Мета та завдання статті

Головною метою пропонованої статті є виявити закономірності функціонування метафори в багатомовній термінології ЄС і вплив перекладу на утворення цієї термінології. Для досягнення поставленої мети потрібно виконати наступні завдання:

- 1) на емпіричному матеріалі української, англійської, італійської, німецької, польської та інколи інших мов встановити перекладацькі відповідники термінів вторинного права ЄС, що містять метафоричні значення;
- 2) проаналізувати встановлені відповідники з погляду конвергентного / дивергентного передавання метафори й запропонувати типологію способів перекладу;
- 3) з'ясувати вплив перекладу метафоричних термінів на виникнення варіативності в українській термінології;
- 4) скласти уявлення про реально наявну міжмовну узгодженість термінології права ЄС щодо метафори,
- 5) ввести проблему метафори під час перекладу термінології офіційних актів ЄС до кола проблем теоретичного обговорення, зокрема, з метою

²Загальний огляд інших наявних публікацій, які розглядають методи перекладу метафори, подано в роботі (Oliynyuk, 2014).

подальшого вироблення науково обґрунтованих настанов для перекладачів-практиків та пропозицій щодо впорядкування, стандартизації і гармонізації української термінології в галузі права ЄС.

4. Методи дослідження

Дослідження проведено за допомогою порівняльного методу, який є одним із найактуальніших як у вивченні термінології взагалі (Артикуца, 2019; Пчелінцева, 2019), так і у структурно-семантичному аналізі метафори в термінології (Кришталь, 2003). Для вивчення термінологічної варіативності в контексті перекладу права ЄС значення порівняльного методу зростає (Biel, 2023; Kerremans, 2010; Mori, 2018). Фахівці з перекладу актів ЄС рекомендують «використовувати багатомовний підхід, беручи до уваги терміни у споріднених мовах, і не обмежуватися лише англійськими відповідниками» (Байчич, 2024, с. 95). Саме тому в нашому дослідженні українські терміни порівнюватимемо з їх відповідниками не лише в англійській мові, але також в італійській, німецькій, польській, словацькій та деяких інших мовах. З іншого боку, в окремих випадках українські відповідники потребують зіставлення з термінами внутрішнього національного законодавства України, зокрема, у мікродіахронічному аспекті. Для встановлення тотожності понять, які виражаються різними варіантами термінів, будуть залучені їхні дефініції, надані законодавцем. Крім того, використано процедури перекладознавчого аналізу та аналізу семантичної структури терміна.

Термінологічний матеріал зібрано з трьох типів документів європейського вторинного права, зокрема – з регламентів, директив та рішень, на початку яких законодавець визначає використання в документі термінологію.

Іншомовні терміни були зібрані з відповідних версій офіційних актів. Усі мовні версії доступні на вебсайті Європейського Союзу (<https://eur-lex.europa.eu>), для кожного прикладу подано посилання на англomовну назву документа. Українські терміни були зібрані з офіційних перекладів документів ЄС, які опубліковано на вебсайті Верховної Ради України (<https://www.rada.gov.ua>).

5. Результати дослідження

Першим важливим проміжним результатом дослідження є відсутність в аналізованих українських перекладах помилок, зумовлених буквальним трактуванням переносного значення або його неправильним виведенням.

Порівняльний аналіз встановлених відповідників з різномовних версій документів права ЄС виявив більшу регулярність конвергентного і значно меншу регулярність дивергентного передавання метафори.

5.1. Типологія за способом перекладу метафори

До конвергентних способів ми відноситимемо переклад зі збереженням тієї самої метафори, а також переклад, за якого метафору, що лежить в основі англійського терміна, замінюють на функціонально аналогічну метафору у відповідному українському, італійському, німецькому, польському або словацькому терміні. Дивергентними ми вважатимемо спосіб, яким метафору вилучено, та спосіб, при якому термін з буквальним значенням у МО передається метафоричним терміном у МП. Взагалі, беручи за вихідну

англомовну термінологію права ЄС, можна говорити про такі регулярні типи передавання метафор:

- 1) збереження метафори;
- 2) заміна іншою метафорою;
- 3) деме́тафоризація;
- 4) додаткова метафоризація.

Розглянемо кожен із виокремлених типів на прикладах із корпусу. Віднесення прикладів до одного з типів відбуватиметься, насамперед, з огляду на співвідношення між українським та англійським терміном.

5.2. Збереження метафори

Збереження метафори терміна МО можливе лише тоді, якщо наявне переносне значення, і воно може термінізуватися у відповідній лексемі МП. Це спостерігається насамперед тоді, коли йдеться про універсальну або загальнолюдську концептуальну метафору, поняття якої було запропоновано в когнітивній теорії метафори Джорджа Лакоффа і Марка Джонсона (Lakoff & Johnson, 1980). Однією із таких загальнолюдських метафор є метафора сім'ї, тому не дивно, що для англійського терміна *engine family* (1), в основі якого лежить сімейна метафора, є еквівалентні відповідники у терміносистемах усіх долучених до дослідження мов. Це дало змогу зберегти метафору у відповідних терміносистемах, як показано у прикладах (2)–(6)

- (1) англ. engine family
(DIRECTIVE 2013/53/EU OF THE EUROPEAN PARLIAMENT
AND OF THE COUNCIL of 20 November 2013)
- (2) укр. сімейство двигунів
- (3) італ. famiglia di motori (букв. сімейство двигунів)
- (4) нім. Motorenfamilie (букв. сімейство двигунів)
- (5) пол. rodzina silników (букв. сімейство двигунів)
- (6) слов. rodina motorov (букв. сімейство двигунів)

Ще один випадок збереження сімейної метафори терміна *gas family* (7) під час перекладу українською та іншими мовами ілюструють приклади (8)–(11):

- (7) англ. gas family
(REGULATION (EU) 2016/426 OF THE EUROPEAN PARLIAMENT
AND OF THE COUNCIL of 9 March 2016)
- (8) укр. сімейство газів
- (9) італ. famiglia di gas (букв. сімейство газів)
- (10) нім. Gasfamilie (букв. сімейство газів)
- (11) пол. rodzina gazów (букв. сімейство газів)
- (12) слов. trieda plynu (букв. клас газу)

Підкреслимо, що не в усіх мовних версіях було використано метафоричний термін. Так, у словацькій мові, як свідчить функціональний відповідник *trieda plynu* (букв. клас газу) у прикладі (12), відбувається деме́тафоризація. Щодо українського терміна *сімейство газів*, то варто зауважити, що його послідовно застосовують для вираження цього поняття і у внутрішньому законодавстві, про що свідчать визначення в прикладах (13) та (14):

- (13) Сімейство газів означає групу газоподібних палив із подібною поведінкою під час горіння, пов'язаних між собою діапазоном чисел Воббе (РЕГЛАМЕНТ ЄВРОПЕЙСЬКОГО ПАРЛАМЕНТУ І РАДИ (ЄС) № 2016/426 від 9 березня 2016 року)
- (14) Сімейство газів – група газоподібного палива з подібною поведінкою під час горіння, пов'язана між собою діапазоном чисел Воббе (ПОСТАНОВА КАБІНЕТУ МІНІСТРІВ УКРАЇНИ від 4 липня 2018 р. № 814)

Крім лексеми *сімейство*, в утворенні терміна, що базується на переносному значенні, використовується також його синонім *сім'я* (пор., напр., *сім'я посад, бджолина сім'я* у національній термінології). Саме цей варіант використано як відповідник для терміна *family of benchmarks* (15), що ілюструє приклад (16):

- (15) англ. family of benchmarks
(REGULATION (EU) 2016/1011 OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL of 8 June 2016)
- (16) укр. сім'я бенчмарків
- (17) італ. famiglia di indici di riferimento (букв. *сім'я бенчмарків*)
- (18) нім. Referenzwert-Familie (букв. *сім'я бенчмарків*)
- (19) пол. rodzina wskaźników referencyjnych (букв. *сім'я бенчмарків*)
- (20) слов. skupina referenčných hodnôt (букв. *група бенчмарків*)

Приклади (17)–(19) показують, що і в цьому випадку в більшості мов – подібно до української – збережено метафору, і тільки відповідник у словацькій мові *skupina referenčných hodnôt* (букв. *група бенчмарків*) у прикладі (20) є деме́тафоризацією.

5.3. Заміна іншою метафорою

Другим конвергентним способом можна вважати такий спосіб перекладу, коли метафору, що є в основі терміна МО, замінюють на іншу метафору у відповідному терміні МП. Зазвичай між двома метафорами наявний тісний концептуальний зв'язок. Наприклад, поняття, що виражені двома метафорами, можуть бути суміжними. Це спостерігається, якщо ці поняття перебувають у родо-видових відношеннях. Такий випадок ми маємо у передачі англійського терміна *parent undertaking* (21) за допомогою відповідників, що використовують метафору материнства майже у всіх мовах, зокрема в українській, як показано у прикладах (22), (23), (24) та (26):

- (21) англ. parentundertaking (букв. *батьківська компанія*)
(REGULATION (EU) No 1227/2011 OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL of 25 October 2011)
- (22) укр. материнська компанія
- (23) італ. impresa madre (букв. *материнська компанія*)
- (24) нім. Mutterunternehmen (букв. *материнська компанія*)
- (25) пол. jednostka dominująca (букв. *головна/домінуюча компанія*)
- (26) слов. materská spoločnosť (букв. *материнська компанія*)

Як бачимо, тільки в польській версії документа метафору не збережено, тут відбулася деме́тафоризація терміна (25). Наступний приклад ілюструє,

навпаки, випадок, коли для передавання видового поняття *міст* (англ. термін *bridge istitution* в (27)) використано родове поняття *перехід* (див. термін *перехідна установа* в (28)). Заміна здається обґрунтованою: попри те, що, згідно з СУМ-11 і СУМ-20, переносне значення у загальноживаній лексемі *міст* є, вторинної номінації в інституціональній галузі не відбулося, принаймні, поки що. Проте, як показують приклади (29)–(32), у багатьох інших мовах відповідний термін зберігає метафору:

- (27) англ. bridge institution (букв. *установа-міст*)
(DIRECTIVE 2014/59/EU OF THE EUROPEAN PARLIAMENT
AND OF THE COUNCIL of 15 May 2014)
- (28) укр. перехідна установа
- (29) італ. ente-ponte (букв. *установа-міст*)
- (30) нім. Brückeninstitut (букв. *установа-міст*)
- (31) пол. instytucja pomostowa (букв. *мостова установа*)
- (32) слов. preklenovacia inštitúcia (букв. *мостова установа*)

Наголосимо, що україномовна версія – не єдина, яка не зберегла цю метафору. Розширення кола мов показало, що в латиській і литовській версіях було використано функціональні відповідники, і в цьому термінологічному словосполученні відбулась заміна метафоричного компонента *мостовий* на неметафоричний *тимчасовий*, див. приклади (33) і (34):

- (33) лат. paga iduiestāde (букв. *тимчасова установа*)
- (34) лит. laikina įstaiga (букв. *тимчасова установа*)

5.4. Деметафоризація

За деметафоризації відбувається заміна терміна чи компонента термінологічного словосполучення таким терміном чи компонентом, у якому лексему використано у прямому значенні. На прикладах з різних мов (12), (20), (25), (33) та (34) ми вже могли спостерігати спосіб деметафоризації. Цікаво простежити цей механізм передавання метафори на ширшому колі прикладів. Звичайно йдеться про функціональні відповідники або описовий переклад, що застосовують через неможливість використання в ПМ переносного значення вихідної одиниці МО. Наведемо приклади перекладацьких відповідників метафоричного терміна *shell bank* (35):

- (35) англ. shell bank
(Direttiva (UE) 2015/849 del Parlamento europeo e del Consiglio,
del 20 maggio 2015)
- (36) укр. номінальний банк (Редакція від 09.07.2018)
- (37) укр. банк-оболонка (Редакція від 30.06.2021)
- (38) італ. banca di comodo
- (39) нім. Bank-Mantelgesellschaft (shell bank)
- (40) пол. bank fikcyjny
- (41) слов. shell banka

Як бачимо, спосіб деметафоризації використано в першій українській редакції від 9 липня 2018 року (36). Проте в поточній редакції, яку затверджено рішенням Урядового комітету з питань європейської та євроатлантичної

інтеграції, міжнародного співробітництва, правової політики та правоохоронної діяльності від 20 березня 2025 року, в українській версії використано конвергентний переклад за допомогою еквівалентної метафори *банк-оболонка* (37). Доцільно порівняти їхні дефініції (42) і (43), щоб переконатися в тотожності поняття, яке позначене за допомогою двох термінів:

- (42) *Номінальний банк* означає кредитну чи фінансову установу, або установу, яка здійснює діяльність, еквівалентну діяльності, яку здійснюють кредитні і фінансові установи, створені у юрисдикції, де вони не мають фізичної присутності, з огляду на розумність та управління, і яка не є афілійованою з регульованою фінансовою групою (ДИРЕКТИВА ЄВРОПЕЙСЬКОГО ПАРЛАМЕНТУ І РАДИ (ЄС) № 2015/849 від 20.05.2015; укр. ред. від 9.07. 2018)
- (43) *Банк-оболонка* означає кредитну чи фінансову установу, або установу, яка здійснює діяльність, еквівалентну діяльності, що здійснюється кредитними та фінансовими установами, яка створена в юрисдикції, де вона не має фізичної присутності, у тому числі реального центру прийняття рішень та управління, і яка не є афілійованою з регульованою фінансовою групою (ДИРЕКТИВА ЄВРОПЕЙСЬКОГО ПАРЛАМЕНТУ І РАДИ (ЄС) № 2015/849 від 20.05.2015; укр. ред. від 20.03.2025)

На нашу думку, варіант *банк-оболонка* є більш вдалим, оскільки він задовольняє вимогу горизонтальної міжмовної узгодженості з англomовною термінологією ЄС. Водночас він враховує семантичну структуру української лексики *оболонка*, яка може мати переносне значення. Згідно з СУМ-11, воно визначається як «зовнішній вигляд, зовнішня форма чого-небудь, за якою криється певний внутрішній зміст» (<https://sum.in.ua/s/obolonka>). Крім того, в українській мові ця лексема вже функціонує як анатомічний і технічний терміни (там само).

Щодо відповідних термінів з інших мов, то в них збереження англійської метафори *shell* (букв. *раковина; оболонка; каркас*) не відбулося. У двох мовах було прямо запозичено англomовний термін: у німецькій версії документа повне запозичення *shell bank* (41) вказано в дужках як варіант німецького, також метафоричного терміна (заміна метафори) *Bank-Mantelgesellschaft* (Mantel = букв. *пальто*), а у словацькій версії англійський метафоричний компонент *shell* став частиною гібридного термінологічного сполучення *shell banka* (43).

Для ширшої картини доповнимо список відповідників термінами з двох балтійських мов – латиської (44) і литовської (45), а також чотирьох слов'янських мов – чеської (46), словенської (47), болгарської (48) і хорватської (49):

- (44) лат. fiktīva banka
 (45) лит. fiktyvus bankas
 (46) чes. banka nevývjející žádnou činnost (tzv. „shellbank“)
 (47) словен. navidezna banka
 (48) болг. банка фантом
 (49) хорв. fiktivna banka

Як бачимо, і в цих мовах збереження метафори *shell* з МО було неможливе. У більшості мовних версій цей метафоричний компонент передано за допомогою функціонального відповідника: у латиській (44), литовській (45) і хорватській (49) – з буквальним значенням *фіктивний*, у словенській (47) – з буквальним значенням *віртуальний*, у болгарській – з буквальним значенням *фантомний*. У чеській мові метафоричний компонент передано поясненням (букв. *банк, який не здійснює жодної діяльності*), що є досить рідкісним способом у процесі перекладу термінів. Крім того, тут так само, як і в німецькій мові, в дужках як варіант вказано пряме запозичення з англійської з додаванням металінгвистичної позначки *так званий* (46). На тлі розмаїття перекладацьких рішень з деме́тафоризації стає закономірною відсутність єдиної практики серед українських перекладачів. Наголосимо, що попередній український варіант з деме́тафоризацією *номінальний банк* зараз не використовують, тоді як метафоричний компонент *оболонка* входить не тільки до складу терміна *банк-оболонка*, але й до складу терміна *компанія-оболонка*. Обидва терміни широко використовують у внутрішньому законодавстві, наприклад, у таких документах: Закон України від 07.12.2000 № 2121-III; Закон України від 06.12.2019 № 361-IX, Постанова Національного банку України від 26.06.2015 № 417 та інших.

5.5. Додаткова метафоризація

Останній перекладацький спосіб, який доцільно розглянути у зв'язку з наведеними раніше, стосується передавання термінів, які в МО не виникають у результаті вторинної номінації переносного значення. Йдеться про поняття, що виражають за допомогою буквального значення в МО, і їхнє вираження у МП також могло бути буквальним. Попри це, перекладач вдається до метафоричного відповідника, що можна пояснити високим ступенем терміноутворювального потенціалу метафори, усталеністю вторинної номінації переносного значення. Такий випадок в українській мові наочно ілюструє термін *вузол*, що входить до складу кількох десятків термінологічних словосполучень, наприклад: *вузол обліку*; *вузол е-взаємодії*; *газетний вузол*; *транзитний вузол*; *сортувальний вузол кореспонденції* тощо. Саме цю метафору використано в утворенні відповідника (51) для неметафоричного англійського терміна *subsystem* (50).

- (50) англ. subsystem
(Regulation (EU) 2016/424 of the European Parliament and of the Council of 9 March 2016)
- (51) укр. вузол
- (52) італ. sottosistema
- (53) нім. Teilsystem
- (54) пол. podsystem
- (55) слов. subsystém

Як бачимо на прикладах (52)–(55) в інших мовах метафоризації не відбулося і були обрані відповідники, еквівалентні англійському терміну (букв. *підсистема*). Наголосимо, що неметафоричний варіант *підсистема* вже існує в українській мові, його використовують у термінології внутрішнього українського законодавства. Те, що в цьому випадку варіанти *вузол* та *підсистема* знаходяться у відношенні повної синонімії, стає очевидним під час

порівняння їхніх визначень у перекладеному документі (56) і в національному документі (57), а також у процесі порівняння тексту першої статті двох документів (58) і (59):

- (56) *Вузол* означає систему з переліку у додатку I або комбінацію таких систем, призначену для вбудування у канатні дороги (Регламент Європейського парламенту і Ради (ЄС) 2016/424 від 9 березня 2016 року про канатні дороги та скасування Директиви 2000/9/ЄС)
- (57) *Підсистема* – складова з переліку відповідно до додатка 1 або комбінація таких складових, призначених для вбудування в канатну дорогу (Постанова Кабінету Міністрів України від 6 лютого 2019 року № 342 про затвердження Технічного регламенту канатних доріг)
- (58) Цей Регламент встановлює правила щодо надання на ринку та вільного руху пристроїв безпеки та *вузлів* для канатних доріг.
- (59) Цей Технічний регламент визначає вимоги щодо введення в обіг та надання на ринку компонентів безпеки та/або *підсистем* для канатних доріг, призначених для перевезення людей (...)

Зазначимо також, що Технічний регламент, затверджений у Постанові Кабінету Міністрів України від 6 лютого 2019 року № 342, прямо посилається на Регламент ЄС 2016/424 Європейського Парламенту та Ради від 9 березня 2016 року, на базі якого цей Технічний регламент було розроблено. Отже, обидва терміни позначають однакове поняття, тобто є дублетними. З огляду на вимогу термінологічної гармонізації було б доцільніше, якби під час перекладу Регламенту Європейського парламенту і Ради (ЄС) 2016/424 був також використаний термін *підсистема*. Відповідно застосування способу додаткової метафоризації тут навряд можна вважати виправданим.

6. Висновки

В утворенні багатомовної термінології права ЄС метафора є поширеним явищем, разом з тим її переклад викликає специфічні труднощі, пов'язані з культурними та концептуальними відмінностями. Це вимагає від перекладача використання різних перекладацьких способів.

Першим важливим (хоча й проміжним) результатом дослідження є встановлена відсутність у досліджених українських перекладах помилок, спричинених буквальною розумінням переносного значення або його неправильним виведенням.

Серед найчастіше використаних способів перекладу метафоричних термінів виявлено такі, що дають можливість зберегти метафору МО або замінити її функціонально еквівалентною іншою метафорою в МП. Рідше, але також регулярно використовують способи, наслідком яких є деме́тафоризація. В окремих випадках в різних мовах спостерігаємо додаткову метафоризацію, а також утворення гібридних термінів, терміни – прямі запозичення та терміни-пояснення, які не завжди задовольняють критерії, що їх висувають до термінів.

Особливістю української термінології, пов'язаної з впливом різних перекладацьких рішень щодо того самого метафоричного терміна МО, можна вважати високу варіативність, що динамічно змінюється. Це природно для сучасного етапу розвитку українського євролекту, проте вимагає поглибленого вивчення.

Подальше дослідження потребує перевірки цих результатів на статистично підтриманих даних, що вимагає створення репрезентативного і кількісно значущого багатомовного корпусу метафоричних термінів. Крім того, у майбутніх розвідках варто спробувати знайти кореляції між виявленими типами відповідності і лінгвістичним вираженням концептуальних метафор різними мовами. На базі таких досліджень може бути запропоновано науково обґрунтовані та стандартизовані принципи перекладу метафоричних термінів у галузі права.

Бібліографія

- Артикуца, Н.В. (2019). Юридичне термінознавство в Україні: сучасний стан, основні напрямки та перспективи розвитку. *Термінологічний вісник*, 5, 6–17.
- Байчич, М. (2024). Термінологія та прикладна термінологічна робота в контексті перекладу офіційних документів ЄС. У Байчич, М., Робертсон, К. Д., Славова, Л. (Ред.), *Посібник з перекладу актів права ЄС українською мовою: у 2 ч. Право, мова й термінологія ЄС*. (сс. 81–102). Київ: ВПЦ «Київський університет».
- Бель, Л. (2024). Переклад офіційних документів ЄС: основи теорії та практики. У Байчич, М., Робертсон, К. Д., Славова, Л. (Ред.), *Посібник з перекладу актів права ЄС українською мовою: у 2 ч. Право, мова й термінологія ЄС*. (сс. 59–80). Київ: ВПЦ «Київський університет».
- Вакуленко, М.О. (2023). *Сучасна українська термінологія: методологія, кодифікація, лексикографічна практика*. Дис. ...докт. філол. наук за спец. 10.02.01. Київ: Український мовно-інформаційний фонд НАН України.
- Голубовська, І.О., Жалай, В.Я., Линник, Т. Г., Биховец, Н.М., Пархоменко, А.Ф., Рубашова, Л. М., Кругликова, О. В., Рахманова, І. І., Бобошко, Т. М. (2016). Термінологічна варіативність: підходи до вивчення. *Лінгвістика XXI століття*, 3–22.
- Касяненко, Д.С. (2011). *Особливості перекладу та лексичної гармонізації законодавчих актів ЄС в контексті євроінтеграції України*. Дис. ... канд. філол. наук за спец. 10.02.16. Київ.
- Касяненко, Д.С. (2015). Мовні кліше та канцеляризми нормативно-правових актів ЄС крізь призму перекладу з німецької мови українською. *Наукові записки Національного університету «Острозька академія». Серія «Філологічна»*, 52, 125–127.
- Кришталь С. М. (2003). *Структурно-семантичний аналіз метафоричних термінів підмови фінансів в англійській і українській мовах*: автореф. дис. ... канд. філол. наук за спец. 10.02.04. Донецьк: Дон. нац. ун-т.
- Палайчук, Е., Батіна, І., Кабанова, С. (2024). Граматичні та стилістичні особливості перекладу права ЄС. У Байчич, М., Робертсон, К. Д., Славова, Л. (Ред.), *Посібник з перекладу актів права ЄС українською мовою: у 2 ч. Право, мова й термінологія ЄС*. (сс. 181–215). Київ: ВПЦ «Київський університет».
- Пчелінцева, О.Е. (2019). Функціонально-граматична специфіка українських віддієслівних іменників зі значенням дії на фоні інших слов'янських мов. *Термінологічний вісник*, 5, 33–39.
- СУМ-11: Словник української мови: академічний тлумачний словник в 11 томах. Київ: Наукова думка, 1970–1980. URL: <https://sum.in.ua/>.
- СУМ-20: Словник української мови у 20 томах. Київ: Наукова думка. URL: <https://slovnyk.me/dict/newsum>.
- Biel, Ł. (2023): Variation of legal terms in monolingual and multilingual contexts: Types, distribution, attitudes and causes. In Ł. Biel, & H. J. Kockaert (Eds.), *Handbook of Terminology: Volume 3. Legal Terminology* (pp. 90–123). Amsterdam – Philadelphia: John Benjamins Pub Co.
- Kerremans, K. (2010). A Comparative Study of Terminological Variation in Specialised Translation. In Carmen Heine and Jan Engberg (Eds.), *Reconceptualizing LSP. Online*

- proceedings of the XVII European LSP Symposium 2009*. Aarhus: Aarhus School of Business – Aarhus University.
- Kövecses, Z. (2005). *Metaphor in Culture: Universality and Variation*. Cambridge: Cambridge University Press.
- Mandelblit, N. (1995). The cognitive view of metaphor and its implications for translation theory. *Translation and Meaning*, 3, 483–495.
- Mori, L. (ed. By), (2018). *Observing Eurolects: Corpus Analysis of Linguistic Variation in EU Law*. Amsterdam – Philadelphia: John Benjamins Pub Co.
- Newmark, P. (1980). The Translation of Metaphor. *Babel*, 26 (2), 93–101.
- Oliynyk, T. (2014). Metaphor Translation Methods. *International Journal of Applied Science and Technology*, 4 (1), 123–126.
- Lakoff, G. & Johnson, M. (1980). *Metaphors We Live By*. Chicago – London: The University of Chicago Press.
- Palumbo, G. (1999). Aspetti della traduzione specializzata: la traduzione dall'inglese in italiano di un manuale di tecnologia dell'architettura. *Traduzione, società e cultura*, 9, 93–139.
- Šarčević, S. (2015). Basic Principles of Term Formation in the Multilingual and Multicultural Context of EU Law. In Šarčević, S. (Ed.), *Language and culture in EU law. Multidisciplinary perspectives* (pp. 183–205). Farnham: Ashgate.

References

- Artykutsa, N. V. (2019). Yurydychne terminoznavstvo v Ukrayini: suchasnyj stan, osnovni napryamky ta perspektyvy rozvytku. *Terminolohichnyj visnyk*, 5, 6–17.
- Bajčić, M. (2024). Terminology and Applied Terminology Work in the EU Context: An Overview. In Bajčić, M., Robertson, C. D., and Slavova, L. (Eds), *Manual on EU Legal Translation into Ukrainian. In 2 parts. Part A. EU Law, Language, and Terminology* (pp. 81-102). Kyiv: Publishing and Polygraphic Center «Kyiv University».
- Biel, Ł. (2024). EU Translation: Theoretical and Practical Foundations. In In Bajčić, M., Robertson, C. D., and Slavova, L. (Eds), *Manual on EU Legal Translation into Ukrainian. In 2 parts. Part A. EU Law, Language, and Terminology* (pp. 59-80). Kyiv: Publishing and Polygraphic Center «Kyiv University».
- Vakulenko, M.O. (2023). *Suchasna ukrayins'ka terminolohiya: metodolohiya, kodyfikatsiya, leksykohrafichna praktyka*. Dys. ... dokt. filol. nauk za spets. 10.02.01. Kyiv: Ukrayins'kyi movno-informatsynny fond NAN Ukrayiny.
- Holubovska, I.O., Zhalai, V. Ya., Lynnyk, T. H., Bykhovets, N. M., Parkhomenko, A. F., Rubashova, L. M., Kruhlykova, O. V., Rakhmanova, I. I., Boboshko, T. M. (2016). Terminological Variation: Perspectives of Studies. *21 century linguistics: new research and prospects*, 3–22.
- Kasyanenko, D.S. (2011). *Osoblyvosti perekladu ta leksychnoyi harmonizatsiyi zakonodavchyykh aktiv YeS v konteksti yevrointehratsiyi Ukrayiny*. Dys. ... kand. filol. nauk za spets. 10.02.16. Kyiv.
- Kasyanenko, D.S. (2015). Movni klyshe ta kantselyaryzmy normatyvno-pravovykh aktiv YeS kriz' pryzmu perekladu z nimets'koyi movy ukrayins'koyu. *Naukovi zapysky Natsional'noho universytetu «Ostroz'ka akademiya». Seriya «Filolohichna»*, 52, 125–127.
- Kryshchal', S.M. (2003). *Strukturno-semantychnyy analiz metaforychnykh terminiv pidmovy finansiv v anhliys'kiy i ukrayins'kiy movakh*: avtoref. dys. ... kand. filol. nauk za spets. 10.02.04. Donetsk: Don. nats. un-t.
- Paliichuk, E., Batina, I., Kabanova, S. (2024). Grammatical and Stylistic Issues of Translating EU Legal. In In Bajčić, M., Robertson, C. D., and Slavova, L. (Eds), *Manual on EU Legal Translation into Ukrainian. In 2 parts. Part A. EU Law, Language, and Terminology* (pp. 181–215). Kyiv: Publishing and Polygraphic Center «Kyiv University».

- Pchelintseva, O.E. (2019). Funktsional'no-hramatychna spetsyfika ukrayins'kykh viddiyeslivnykh imennykiv zi znachenniyam diyi na foni inshykh slov"yans'kykh mov. *Terminolohichnyy visnyk*, 5, 33–39.
- SUM-11: Slovník ukrajinskoyi movy: akademichnyy tlumachnyy slovník v 11 t. Kyiv: Naukova dumka, 1970–1980. URL: <https://sum.in.ua/>.
- SUM-20: Slovník ukrajinskoyi movy u 20 tomakh. Kyiv: Naukova dumka. URL: <https://slovník.me/dict/newsum>.
- Biel, Ł. (2023). Variation of legal terms in monolingual and multilingual contexts: Types, distribution, attitudes and causes. In Ł. Biel, & H. J. Kockaert (Eds.), *Handbook of Terminology: 3. Legal Terminology* (pp. 90–123). Amsterdam – Philadelphia: John Benjamins Pub Co.
- Kerremans, K. (2010). A Comparative Study of Terminological Variation in Specialised Translation. In Heine, C. and Engber, J. (Eds.), *Reconceptualizing LSP. Online proceedings of the XVII European LSP Symposium 2009*. Aarhus: Aarhus School of Business – Aarhus University.
- Kövecses, Z. (2005). *Metaphor in Culture: Universality and Variation*. Cambridge: University Press.
- Mandelblit, N. (1995). The cognitive view of metaphor and its implications for translation theory. *Translation and Meaning*, 3, 483–495.
- Mori, L. (ed. By), (2018). *Observing Eurolects: Corpus Analysis of Linguistic Variation in EU Law*. Amsterdam – Philadelphia: John Benjamins Pub Co.
- Newmark, P. (1980). The Translation of Metaphor. *Babel*, 26 (2), 93–101.
- Oliynyk, T. (2014). Metaphor Translation Methods. *International Journal of Applied Science and Technology*, 4 (1), 123–126.
- Lakoff, G., Johnson, M. (1980). *Metaphors We Live By*. Chicago – London: The University of Chicago Press.
- Palumbo, G. (1999). Aspetti della traduzione specializzata: la traduzione dall'inglese in italiano di un manuale di tecnologia dell'architettura. *Traduzione, società e cultura*, 9, 93–139.
- Šarčević, S. (2015). Basic Principles of Term Formation in the Multilingual and Multicultural Context of EU Law. In Šarčević, S. (Ed.), *Language and culture in EU law. Multidisciplinary perspectives* (pp. 183–205). Farnham: Ashgate.

Резюме

Голетіані Ліана

МЕТАФОРА В УКРАЇНСЬКІЙ ТЕРМІНОЛОГІЇ ВТОРИННОГО ПРАВА ЄС: ПОРІВНЯЛЬНЕ ДОСЛІДЖЕННЯ

Постановка проблеми. Частина англomовних термінів вторинного права ЄС має метафоричні компоненти. Переклад таких термінів викликає труднощі, якщо в мові перекладу не існує такої метафори, або вона є, але її не застосовують як термін. У галузі спеціального перекладу проблему передавання метафори не висвітлено. Дослідження пропонує порівняльний аналіз відповідників метафоричних термінів на матеріалі українських перекладів та укладених іншими мовами офіційних актів ЄС.

Мета статті. Метою статті є аналіз перекладу метафоричних термінів вторинного права ЄС українською мовою. Мета статті визначає наступні завдання: 1) встановити відповідники метафоричних термінів вторинного права ЄС в різних мовах; 2) виокремити конвергентні і дивергентні способи перекладу; 3) визначити вплив метафори на виникнення термінологічної

варіативності в українській мові; 4) підвищити інтерес до метафори в юридичній термінології з боку наукової спільноти.

Методи дослідження. Дослідження проведено за допомогою порівняльного методу. Українські терміни, що використовуються в перекладах актів ЄС, порівнюються з їх відповідниками, насамперед, в англійській мові, яка є робочою в ЄС, а також – в італійській, німецькій, польській, словацькій та інших мовах. Порівняльний метод використано також у внутрішньомовній перспективі: у разі співіснування в українській мові більше ніж одного відповідника застосовано порівняльний аналіз на рівні дефініцій та вживання, зокрема, в мікродіахронічному аспекті. Використано процедури перекладознавчого аналізу та структурно-семантичного аналізу.

Основні результати дослідження. Показано, що дотримання принципу міжмовної узгодженості термінології ЄС часто не можливе через наявність, з одного боку, метафоричних термінів в англійській версії права ЄС і відсутність, з іншого боку, їх повних еквівалентів в інших мовах. Під час пошуку вдалого відповідника українські перекладачі приймають різні перекладацькі рішення, в результаті чого відбуваються такі процеси: 1) збереження метафори; 2) заміна іншою метафорою; 3) демегафоризація; 4) додаткова метафоризація. Крім того, на матеріалі різних мов серед відповідників зафіксовано прямі запозичення, гібридні терміни, терміни-пояснення. У мікродіахронічній перспективі виявлено значну варіативність в українській термінології.

Висновки і перспективи. Переклад метафори в галузі права є багатофакторним процесом вирішення. Поглибленого вивчення вимагає когнітивно-концептуальний та культурно-визначений аспекти цього процесу. Передбачається, що закономірності, які можуть бути відкриті саме в цих галузях, можуть корелювати з виявленими у статті способами перекладу. Метафорична термінологія має знайти вагоміше місце в термінологічній перевірці під час перекладу актів ЄС і базуватись на науково обґрунтованих і стандартизованих принципах.

Ключові слова: українська термінологія права ЄС, переклад права ЄС, мовний контакт, метафора у термінології, метафора у спеціальному перекладі.

Abstract

Goletiani Liana

METAPHOR IN UKRAINIAN TERMINOLOGY OF EU SECONDARY LAW: A COMPARATIVE STUDY

Background. Some English-language terms of EU secondary law contain metaphorical components. Their translation poses difficulties when the target language lacks such a metaphor, or when the metaphor exists but is not used as a term. In this area of Ukrainian specialized translation, the problem of rendering metaphor has not yet been sufficiently addressed. The present study offers a comparative analysis of equivalents of metaphorical terms based on Ukrainian translations and official EU acts drafted in other languages.

Purpose. The purpose of the article is to analyze the translation of metaphorical terms of EU secondary law into Ukrainian. The purpose of the article defines the

following tasks: to identify equivalents of metaphorical terms of EU secondary law in different languages; to single out convergent and divergent translation strategies; to determine the impact of metaphor on the emergence of terminological variation in the Ukrainian language; to raise scholarly interest in metaphor within legal terminology.

Methods. The study employs the comparative method. Ukrainian terms used in translations of EU acts are compared with their equivalents primarily in English, as well as in Italian, German, Polish, Slovak, and other languages. The comparative method is also applied from an intralingual perspective: where more than one equivalent coexists in Ukrainian, a comparative analysis is carried out at the level of definitions and usage, including a micro-diachronic perspective. Procedures of translation analysis and structural-semantic analysis are also employed.

Results. The study demonstrates that adherence to the principle of interlingual consistency of EU terminology is often impossible due to the presence, on the one hand, of metaphorical terms in the English version of EU law and, on the other hand, the absence of their full equivalents in other languages. To render metaphorical terms, Ukrainian translators adopt different translation procedures, resulting in 1) preservation of the metaphor; 2) replacement with another metaphor; 3) metaphor neutralization; 4) additional metaphORIZATION. Moreover, across different languages, the equivalents include direct borrowings, hybrid terms, and explanatory terms. From a micro-diachronic perspective, significant variability in Ukrainian terminology has been identified.

Conclusions and prospects. The translation strategies identified in the article confirm that translating metaphor is a multifactorial decision-making process. Further in-depth research based on Ukrainian-language material is required, particularly with regard to the cognitive and cultural aspects of this process. Metaphorical terminology should take up more space in terminological verification during the translation of EU acts and should be grounded in scientifically substantiated and standardized principles.

Keywords: Ukrainian EU law terminology, translation of EU law, language contact, metaphor in terminology, metaphor in specialized translation.

Відомості про автора

Голетіані Ліана, доктор філологічних наук, асоційований професор славістики Департаменту іноземних мов, літератур і культур, Державний Університет Бергамо, Італія, e-mail: liana.goletiani@unibg.it

Goletiani Liana, Ph.D., Associate Professor of Slavic Studies, Department of Foreign Languages, Literatures and Cultures, University of Bergamo, Italy, e-mail: liana.goletiani@unibg.it

ORCID 0000-0002-1251-4258

Надійшла до редакції 03 грудня 2025 року

Прийнято до друку 25 грудня 2025 року

Олена Пчелінцева

ІНТЕНСИВНИЙ КУРС ПІДВИЩЕННЯ КВАЛІФІКАЦІЇ «ВЕЛИКІ МОВНІ МОДЕЛІ ДЛЯ ЛІНГВІСТІВ»



FRIEDRICH-SCHILLER-
UNIVERSITÄT
JENA

Інтенсивний онлайн курс з підвищення кваліфікації «Великі мовні моделі для лінгвістів» (Large Language Model) було організовано в Єнському університеті імені Фрідріха Шіллера в лютому-червні 2025 року (розробниця та викладач –

Natalia Cheilytko). Проєкт був спільно профінансований Німецькою службою академічних обмінів (нім. Deutsche Akademische Austauschdienst, DAAD), Федеральним міністерством освіти і наукових досліджень Німеччини (нім. Bundesministerium für Bildung und Forschung, BMBF), а також Єнським університетом (нім. Friedrich-Schiller-Universität Jena).

Щотижня ми зустрічались з фахівцем своєї справи, Наталією Чейлітко, яка щедро ділилась своїми знаннями про великі мовні моделі, терпляче роз'яснювала нам принципи роботи, допомагала створювати власні проєкти. Ми отримали уявлення про векторні моделі (Vector Space Model), моделі штучної нейронної мережі (Artificial Neural Network Model), методи машинного навчання (Model Transfer Learning), анотування за допомогою GPT, наявні на сьогодні українські великі мовні моделі та багато чого іншого. У доступній та зрозумілій формі ми засвоювали можливості використання цих нових для багатьох з нас інструментів у лексикографії, в семантиці, у перекладознавстві, у процесі викладання мовознавчих дисциплін. Учасники курсу з різних міст і країн змогли отримати та закріпити навички використання сучасних інформаційних технологій у власних дослідженнях та викладанні, а також уміння ставити та виконувати базові технічні завдання, користуватися сучасними програмами для прискорення та масштабування своїх наукових розвідок. До речі, з результатами двох наукових проєктів, виконаних у результаті цього навчання, можна ознайомитися у цьому номері журналу, що ви зараз читаете.

На старті у нашій команді було 14 студентів та 52 дослідники. Навчання мало практичний характер і передбачало в першу чергу власне дослідження й презентації його результатів після завершення курсу, що було обов'язковою умовою отримання сертифіката. 16 тижнів напруженої праці, перші несміливі, але наполегливі спроби роботи з інструментами штучного інтелекту, написання спеціалізованих запитів (prompt engineering), а для найбільш досвідчених – освоєння методів донавчання моделей (Fine-Tuning). До фінішу дісталися не всі☺. Але ті учасники Програми, які за результатами виконання

© Пчелінцева О., 2025

завдань і захисту власних проєктів успішно завершили курс, отримали сертифікати, що відповідають вимогам МОН України. А головне – отримали безцінний досвід і практичні навички!



[Rick Merritt. 2022. What Is a Transformer Model?](#)

Щира вдячність організаторам, спонсорам, колегам та особливо – нашій мудрій та терплячій наставниці та провіднику у світі сучасних технологій Наталії Чейлитко – за цю можливість професійного розвитку, за інтелектуальну щедрість, за цінні та надзвичайно потрібні знання!

Abstract

Pchelintseva Olena

INTENSIVE ADVANCED TRAINING COURSE «LARGE LANGUAGE MODEL FOR LINGUISTS»

The Friedrich Schiller University of Jena hosted an intensive advanced training course «Large Language Model for Linguists». The program was developed by Natalia Cheilytko, PhD (Germany – Ukraine). The project was co-funded by the DAAD, the Bundesministerium für Bildung und Forschung, and the University of Jena. The goal of the Program is to study the Large Language Model as the primary modern tool for language researchers.

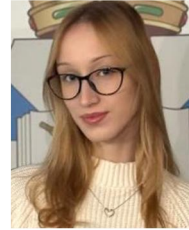
Відомості про автора

Пчелінцева Олена, доктор філологічних наук, професор, завідувач кафедри української мови та загального мовознавства Черкаського державного технологічного університету, Черкаси (Україна), e-mail: o.pchelintseva@chdtu.edu.ua

Pchelintseva Olena, Doctor of Philological Sciences, professor, Head of the Ukrainian language and general linguistics Department, Cherkasy State Technological University, Cherkasy (Ukraine), e-mail: o.pchelintseva@chdtu.edu.ua



Олена Деньга



Софія Федосєєва

**НАЦІОНАЛЬНИЙ КОД УКРАЇНСТВА
В ОСВІТНЬОМУ ПРОСТОРІ:
МИСТЕЦЬКІ ПРАКТИКИ ПОЗААУДИТОРНОЇ ДІЯЛЬНОСТІ**

У статті розглянуто проблему передавання національного коду та формування любові до української мови в освітньому просторі закладу вищої освіти. Акцентовано увагу на ролі позааудиторної творчої діяльності як ефективного засобу залучення студентської молоді до національної культурної спадщини. Проаналізовано досвід кафедри української мови та загального мовознавства Черкаського державного технологічного університету у 2025–2026 навчальному році, зокрема літературно-музичну композицію «Квіти крізь трати» та театралізований захід «Українські вечорниці». Доведено, що особистісне проживання художнього образу, участь у творчому відтворенні поетичного слова й народних традицій сприяють глибшому засвоєнню класичної та сучасної літератури, формуванню національної ідентичності та культурної пам'яті молодого покоління.

У сучасному українському суспільстві, що переживає складні історичні випробування, особливої актуальності набуває питання передавання новим поколінням коду нації – сукупності культурних, мовних, духовних і ціннісних орієнтирів, які забезпечують тяглість національного буття. Важливим носієм цього коду є українська мова – не лише як засіб комунікації, а як форма колективної пам'яті, духовного опору та самоідентифікації.

Практика вищої освіти засвідчує, що традиційні навчальні форми, орієнтовані передусім на відтворення програмного матеріалу, не завжди коректно формують емоційно-ціннісне ставлення студентів до мови й культури. Значно глибше запам'ятовується не те, що студент лише читає або слухає, а те, що він створює, проживає й осмислює особистісно. Саме тому важливу роль у формуванні національної свідомості відіграє позааудиторна діяльність, зорієнтована на творчу співучасть, діалог поколінь і культурну рефлексію.

Позааудиторна робота у закладі вищої освіти постає як особливий простір формування світогляду студента, де національну культуру не

© Деньга О., Федосєєва С., 2025

вивчають відсторонено, а переживають як власний досвід. Залучення молоді до мистецьких форм діяльності – літературно-музичних композицій, театралізованих дійств, реконструкції традицій – сприяє осмисленню української культури як живого, динамічного явища.

Важливим чинником ефективності таких заходів є активна роль самих студентів, які стають не лише виконавцями, а співтворцями культурного продукту. Саме в процесі творчого втілення художнього образу відбувається глибоке занурення в текст, контекст і духовний світ автора.

Яскравим прикладом реалізації зазначених принципів стала літературно-музична композиція «Квіти крізь ґрати», організована та проведена кафедрою української мови та загального мовознавства ЧДТУ у 2025–2026 навчальному році. Особливістю заходу стало те, що студенти втілили на сцені образи поетів – подружжя Ігоря та Ірини Калинців, Василя Стуса, Павла Вишебаби, проживаючи їхні долі через слово, музику та сценічну дію.

Учасники читали уривки з листів, таборової переписки, звернення до дітей, що дозволило розкрити поетів не лише як митців, а як людей – батьків, подружжя, громадян, борців. Виконання сучасних музичних інтерпретацій пісень на слова поетів створювало міст між класичним текстом і сучасним культурним контекстом, роблячи його близьким і зрозумілим молоді.

Кульмінацією композиції стало спільне виконання твору Павла Вишебаби «Мое покоління», який прозвучав як своєрідний гімн сучасної української молоді. Цей фінальний акорд символічно поєднав покоління шістдесятників і поетів-дисидентів із поколінням сучасних воїнів та студентів, утворивши духовний міст між минулим і сьогоденням, між боротьбою за слово і боротьбою за свободу. Важливу роль у створенні цілісного художнього простору відіграв візуальний супровід – картини художниці Олени Кононенко, співробітниці університету, які підсилювали образність поетичного слова й формували багатовимірний емоційний простір сприйняття.





Іншим вагомим напрямом позааудиторної діяльності кафедри стали українські вечорниці, що відтворювали народні звичаї, обряди, гумор і театральну традицію. Студенти з великим ентузіазмом і щирим гумором втілювали образи Стецька, Улянки, Карпа та Лавріна, поєднуючи фольклорну стихію з класичними текстами української літератури.

Для повноцінного входження в роль учасникам необхідно було перечитати відповідні твори, осмислити характери персонажів і поглянути на них не з позиції стороннього читача, а від першої особи, як на живих людей зі своїми почуттями, вадами й прагненнями. Такий підхід сприяв глибшому розумінню класичних текстів, розвитку емпатії та акторського мислення.

Театралізована форма вечорниць поєднала пісню, танець, гру, обрядові дійства, що перетворило захід на живу модель традиційного українського культурного простору. Особливу цінність становив міжкультурний аспект заходу – участь іноземного студента з Перу, Вела Пінто Хорхе Джеферсона, який через безпосередню взаємодію з традиціями долучився до української культурної спадщини.

Отже, досвід кафедри української мови та загального мовознавства ЧДТУ переконливо доводить, що позааудиторна творча діяльність є потужним інструментом передавання національного коду та формування любові до українського слова. Втілення образів поетів і літературних персонажів на сцені, робота з листами, поезіями, піснями й народними обрядами сприяють глибокому особистісному засвоєнню культурних смислів.

Саме через творче проживання тексту, а не лише його академічний аналіз, формується стійкий інтерес до класичної та сучасної української літератури, національних традицій і мови як духовної основи народу. Такі заходи забезпечують тяглість культурної пам'яті та створюють умови для виховання покоління, здатного не лише зберігати, а й активно розвивати українську культурну спадщину.



Abstract

Denga Olena, Fiedosieieva Sofiia

THE NATIONAL CODE OF UKRAINIAN IDENTITY IN THE EDUCATIONAL SPACE: ARTISTIC PRACTICES OF EXTRACURRICULAR ACTIVITIES

The article examines the issue of transmitting the national code and fostering love for the Ukrainian language within the educational space of a higher education institution. Particular attention is given to the role of extracurricular creative

activities as an effective means of engaging students with the national cultural heritage. The study analyzes the experience of the Department of Ukrainian Language and General Linguistics at Cherkasy State Technological University during the 2025–2026 academic year, specifically the literary and musical composition “Flowers Through the Bars” and the theatrical event “Ukrainian Vechornytsi” (Evening Gatherings). It is demonstrated that personal immersion in artistic imagery and participation in the creative interpretation of poetic texts and folk traditions contribute to a deeper understanding of classical and contemporary literature, as well as to the formation of national identity and the cultural memory of the younger generation.

Відомості про авторів

Деньга Олена, старший викладач кафедри української мови та загального мовознавства Черкаського державного технологічного університету, e-mail: helen_15_8@hotmail.com
Denga Olena, Senior Lecturer of the Ukrainian language and general linguistics Department, Cherkasy State Technological University, Cherkasy (Ukraine), e-mail: helen_15_8@hotmail.com

Федосєєва Софія, студентка, старший лаборант кафедри української мови та загального мовознавства Черкаського державного технологічного університету, e-mail: s.ye.fiedosieieva.fht24@chdtu.edu.ua
Fiedosieieva Sofiia, student, Senior Laboratory Assistant of the Ukrainian language and general linguistics Department, Cherkasy State Technological University, Cherkasy (Ukraine), e-mail: s.ye.fiedosieieva.fht24@chdtu.edu.ua

Наукове видання

LANGUAGE:
Codification • Competence • Communication
Scientific Journal
1-2(12-13)2025

МОВА:
Кодифікація • Компетенція • Комунікація
Науковий журнал
1-2(12-13)2025

(українською, англійською, німецькою, французькою, польською, чеською,
словацькою, хорватською, болгарською, словенською мовами)

Державна реєстрація:
Ідентифікатор медіа R30-04616
згідно з рішенням Національної ради України з питань телебачення
і радіомовлення від 30.05.2024 № 1916.

*За достовірність фактів, цитат, цифрових даних
та коректність посилань відповідають автори матеріалів.
Редколегія не завжди поділяє теоретико-методологічні погляди авторів.*

*Літературний редактор Людмила Усик
Переклад: Дмитро Підласий
Технічний редактор Катерина Давиденко
Дизайн обкладинки: Альона Гребенюк
Оригінал-макет виготовлено у редакційно-видавничому відділі
Черкаського державного технологічного університету*

Видавець – Черкаський державний технологічний університет,
бульвар Шевченка, 460, м. Черкаси, 18006, тел. (0472) 51-36-72, факс (0472) 71-00-94.
E-mail: chdtu@chdtu.edu.ua
Свідоцтво про внесення суб'єкта видавничої справи до Державного реєстру видавців,
виготівників і розповсюджувачів видавничої продукції Серія ДК № 896 від 16.04.2002.

Виготівник – ФОП ГОРДІЄНКО Є.І.
Свідоцтво про внесення суб'єкта видавничої справи до Державного реєстру видавців,
виготівників і розповсюджувачів видавничої продукції Серія ДК № 4518 від 04.04.2013.
Україна, 18000, м. Черкаси, тел./факс: (0472) 56-56-12, (067) 444-28-94
e-mail: book.druk@gmail.com

Формат 70x108/16. Папір офс. Друк цифровий. Гарн. Times New Roman, Arial.
Умов. друк. арк. 5,8. Обл.-вид. арк. 6,0. Наклад 50 прим. Зам. № 25-115.