

УДК 004.056:004.85

ПРОЄКТУВАННЯ СЕРВІСУ МАШИННОГО ПЕРЕКЛАДУ ЗМІННИХ ДАНИХ

І. В. Малець

*Національний університет «Львівська політехніка»,
Інститут поліграфії та медійних технологій,
вул. Степана Бандери, 12, 79013, Львів, Україна*

Показано, що зростання обсягів багатомовної комунікації в поліграфічних та освітніх середовищах супроводжується переходом від цілісних документів до структурованих наборів сегментів із змінними даними. Виявлено, що типові веб-сервіси машинного перекладу обробляють такі сегменти як звичайний текст, що призводить до руйнування маркерів персоналізації та ускладнює повторне використання перекладів. Сформульовано потребу проєктування сервісу машинного перекладу змінних даних як окремого мовного шару, орієнтованого на роботу з множиною сегментів $S_i = \{s_{ij}\}$, де кожен сегмент описується контентним ключем, мовними атрибутами та ознаками належності до змінних даних. Побудовано структуру інтеграції сервісу, що забезпечує приймання пакетів сегментів від систем керування контентом і повернення результатів у тій самій структурі. Описано реалізацію цільових REST-інтерфейсів для обробки одиничних та пакетних запитів. Наведено результати апробації прототипу у виробничих потужностях Видавничого дому «УкрПол», які підтверджують доцільність сегментного підходу в задачах автоматизованої підготовки персоналізованого контенту.

Ключові слова: машинний переклад; змінні дані; сегментна модель документа; REST-API; персоналізований контент; сервісна архітектура.

Постановка проблеми. У сучасних поліграфічних, медійних та освітніх системах підготовки матеріалів зростає частка сценаріїв, у яких документи подаються не як суцільний текст, а як множина сегментів, пов'язаних із зовнішніми джерелами даних. Такі сегменти зберігаються у веб-системах керування контентом, платформах управління взаємовідносинами з клієнтами, базах даних навчальних матеріалів та шаблонах персоналізованих повідомлень. Змінні дані у цих сегментах: як-от ідентифікатори користувачів, атрибути освітніх програм чи параметри цільових груп, мають залишатися незмінними при перекладі, а результати перекладу повинні бути однозначно пов'язані з відповідними контентними ключами.

Поширені веб-сервіси машинного перекладу орієнтовані на обробку суцільних текстових блоків і не враховують структурні особливості змінних даних. Це призводить до ризику модифікації маркерів, втрати зв'язку між перекладеним текстом та джерельними полями й фрагментації мовних ресурсів між різними системами. Унаслідок цього ускладнюється побудова єдиного корпусу перекладів для

повторного використання в різних каналах дистрибуції, а контроль якості перекладу переноситься на рівень ручних перевірок.

Окремою проблемою є розрізнена інтеграція машинного перекладу в існуючі інформаційні системи. Кожен клас платформ опрацювання контенту зазвичай реалізує власні механізми звернення до зовнішніх служб, що унеможливорює централізоване керування мовними ресурсами і гальмує впровадження нових рушіїв машинного перекладу. Таким чином, постає завдання побудови сервісу машинного перекладу змінних даних, який працює на рівні сегментів документа, зберігає маркери змінних даних та забезпечує уніфіковану взаємодію з різномірними клієнтськими системами.

Аналіз останніх досліджень та публікацій. Розвиток систем машинного перекладу (MT) від ранніх правилкових підходів до сучасних нейронних моделей послідовно висвітлено у низці праць, де сформульовано основні етапи становлення та проблемні аспекти якості перекладу спеціалізованих текстів [1]. Особливу увагу приділено питанням адекватності автоматизованих перекладів, обмеженням вільно доступних програмних продуктів та необхідності поєднання машинного перекладу з людською постредакцією і термінологічною підтримкою [2].

Важливим напрямом досліджень є аналіз помилок сервісів машинного перекладу та порівняння їх результатів для різних мовних пар. У роботах, присвячених порівнянню сервісів DeepL і Google Translate, запропоновано типологію лексичних, граматичних і стилістичних відхилень та підкреслено залежність якості від галузевої спеціалізації текстів [3]. Окремі праці зосереджені на практиці використання інтернет-ресурсів машинного перекладу у філологічних і навчальних контекстах, де зазначається обмеженість засобів керування структурою документів і змінними даними на рівні сервісної архітектури [4].

З технічного погляду базові архітектури нейронного машинного перекладу (NMT) розглянуто в роботах, присвячених моделям типу sequence-to-sequence на основі рекурентних мереж [6] та трансформерним моделям із механізмом уваги [5]. Узагальнений огляд методів і застосувань NMT подано у [8], де систематизовано основні архітектури, схеми навчання та аспекти оцінювання якості. Підходи до побудови відкритих перекладацьких сервісів на основі корпусів OPUS і моделей OPUS-MT демонструють можливість розгортання багатомовних сервісів із використанням уніфікованих інтерфейсів доступу [7].

Узагальнюючи результати наведених досліджень, можна констатувати, що сучасні роботи детально висвітлюють якість нейронного машинного перекладу, типи помилок і архітектури моделей, однак значно менш опрацьованими залишаються питання сервісної інтеграції машинного перекладу як окремого шару між системами керування контентом (CMS), платформами управління взаємовідносинами з клієнтами (CRM) іншими джерелами змінних даних. Саме аспект проектування сервісу машинного перекладу змінних даних на рівні сегментів документа перебуває у фокусі цієї публікації.

Мета статті полягає в проектуванні архітектурного рішення та подальшої практичної реалізації сервісу машинного перекладу змінних даних, який функціонує як

централізований мовний шар для множини сегментів документа, підтримує роботу з контентними ключами та маркерами змінних даних і надає уніфіковані REST-інтерфейси для інтеграції з веб-CMS, CRM та системами керування навчальним і публікаційним контентом.

Виклад основного матеріалу дослідження. У формальній моделі [9], покладеній в основу сервісу, публікаційні одиниці (документи, шаблони листів, сторінки веб-сайтів тощо) описуються множиною (1):

$$U = \{u_i\}, i = 1 \dots N, \quad (1)$$

де u_i – окремий документ або шаблон.

Для кожного u_i задається множина сегментів (2):

$$S_i = \{s_{ij}\}, j = 1 \dots n_i, \quad (2)$$

де сегмент s_{ij} характеризується контентним ключем, текстовим вмістом, мовною парою (ℓ_{src} , ℓ_{tg}) та типом (структурний елемент, реквізити, основний текст, сегмент зі змінними даними тощо).

Така модель забезпечує незалежну обробку сегментів та їх прив'язку до відповідних записів у зовнішніх системах. Для сегментів зі змінними даними текст містить маркери, що відповідають полям у CRM або інших базах даних. У процесі машинного перекладу ці маркери трактуються як атомарні токени, які не підлягають зміні. Для сегмента s_{ij} розглядаються два джерела перекладу: пам'ять перекладів (translation memory, TM) та пул рушіїв машинного перекладу (machine translation engines). Відповідно до формальної схеми, переклад t_{ij} обирається за правилом $t_{ij} = t_{ij}^{TM}$, якщо в TM знайдено відповідний запис, $t_{ij} = t_{ij}^{MT}$ – в іншому випадку, що забезпечує поєднання повторного використання раніше затверджених перекладів із можливістю підключення різних рушіїв для нових сегментів.

Архітектура сервісу як мовного шару. Сервіс машинного перекладу розгортається як окремий компонент, що реалізує мовний шар між клієнтськими системами. Зовнішні платформи (веб-CMS, CRM, системи керування навчальним контентом, а у поліграфічних сценаріях – VDP-рішення) подають до сервісу множини сегментів S_i , сформовані на основі внутрішніх структур документів. Сервіс повертає результати перекладу у тій самій структурі, зберігаючи контентні ключі та службові атрибути. Узагальнену структуру інтеграції сервісу як мовного шару подано на рис. 1.

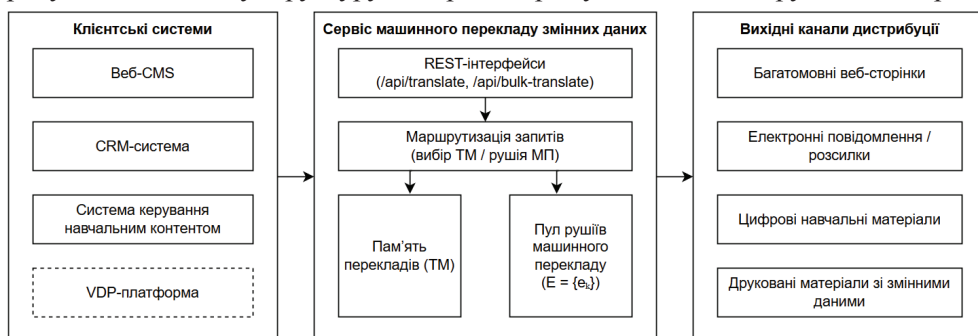


Рис. 1. Архітектура сервісу машинного перекладу змінних даних як мовного шару між клієнтськими системами

Внутрішня логіка сервісу включає модуль маршрутизації, який для кожного запиту визначає джерело перекладу (ТМ або один із рушіїв), використовуючи множину рушіїв $E = \{e_k\}$, $k = 1 \dots K$, та вагові коефіцієнти $g_k(l_{src}, l_{tgt}, d)$. Узагальнений критерій (3):

$$J(l_{src}, l_{tgt}, d) = \alpha \cdot Q - \beta \cdot C - \gamma \cdot T, \quad \alpha, \beta, \gamma > 0, \quad (3)$$

де Q – очікувана якість, C – витрати, T – час відповіді, задає основу для оптимізації вибору рушіїв за сукупністю показників.

У межах представленого прототипу головний акцент зроблено на забезпеченні коректної обробки сегментів і змінних даних; параметри критерію J можуть уточнюватися на подальших етапах розвитку системи.

Розгортання REST-інтерфейсів. Зовнішню взаємодію з сервісом найдоцільніше реалізувати у вигляді REST-інтерфейсів на основі HTTP і JSON. Для обробки одиничних сегментів використовується метод POST /api/translate, що приймає об'єкт із полями text, srcLang, tgtLang, key. Тут text містить текст сегмента s_{ij} , srcLang і tgtLang – коди мовної пари, key – контентний ключ, який однозначно пов'язує сегмент із внутрішніми структурами клієнтської системи. Відповідь сервісу включає translatedText, інформацію про джерело (ТМ або Engine), статус обробки та у разі потреби повідомлення про помилки.

Для пакетної обробки множини сегментів передбачено метод POST /api/bulk-translate, що приймає цільову мову tgtLang та масив segments. Кожен елемент масиву відповідає сегменту s_{ij} і містить key, text, sourceLang. На виході сервіс повертає масив results, у якому для кожного сегмента зазначено key, originalText, sourceLang, targetLang, translatedText, source та status. Структуру запиту та відповіді ендпоінта /api/bulk-translate для множини сегментів документа подано на рис. 2.

Сегментна модель, покладена в основу цих інтерфейсів, забезпечує незалежність сервісу від конкретних форматів зберігання документів. Оскільки взаємодія здійснюється через JSON, сервіс може бути інтегрований як із веб-додатками, так і з серверними компонентами, які формують або оновлюють багатомовний контент.

Проектування прототипа сервісу машинного перекладу змінних даних. Прототип сервісу реалізовано у середовищі Python із використанням веб-фреймворку Flask як окремий процес, що надає REST-інтерфейси для зовнішніх клієнтів. Внутрішні структури даних відображають сегментну модель документа, а взаємодія з ТМ та рушіями машинного перекладу організована через спеціалізовані адаптери. Апробація прототипу виконана на прикладі структурованого контексту поліграфічного завдання та набору сегментів персоналізованих повідомлень, де змінні дані представлені у вигляді маркерів, пов'язаних із зовнішніми полями. У процесі тестування підтверджено, що сервіс зберігає контентні ключі, забезпечує коректну передачу маркерів змінних даних і дозволяє формувати єдиний корпус перекладів для повторного використання в різних публікаційних сценаріях. Схему веб-інтерфейсу сегментного редактора, який забезпечує взаємодію користувача з цими даними (рис. 3) впроваджено в кінцевий термінал технологічного процесу друку змінних даних [10].

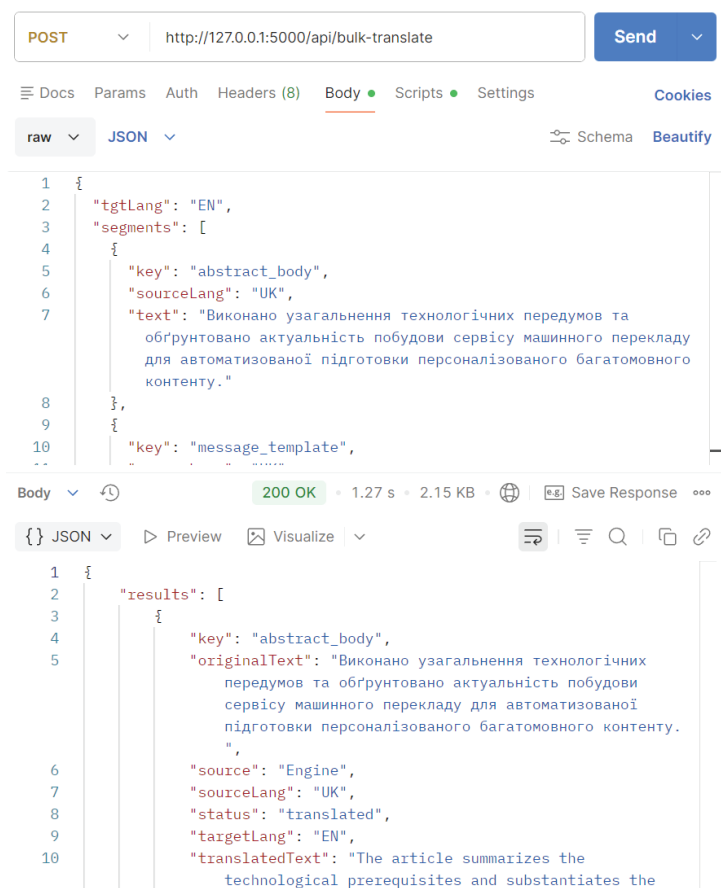


Рис. 2. Структура запиту та відповіді ендпойнта `/api/bulk-translate` для множини сегментів документа

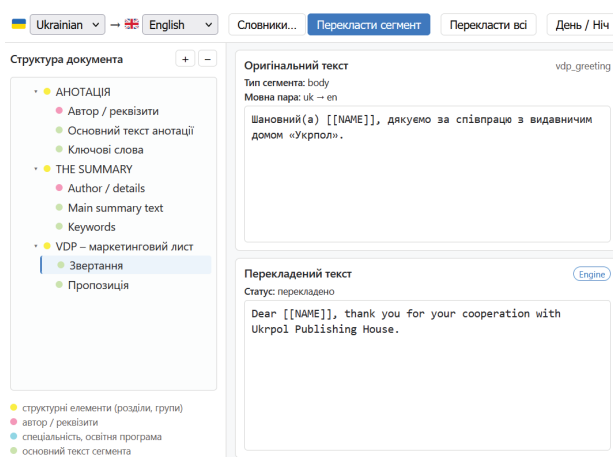


Рис. 3. Веб-інтерфейс сегментного редактора для опрацювання сегментів зі змінними даними

Отримані результати демонструють практичну реалізованість сегментного підходу та створюють основу для подальшого розширення сервісу за рахунок підключення додаткових рушіїв і впровадження формалізованих метрик якості на рівні сервісної архітектури.

Висновки. У представленому дослідженні сформульовано та розв'язано задачу проектування сервісу машинного перекладу змінних даних, орієнтованого на сегментну модель документа. Публікаційні одиниці подано у вигляді множин сегментів з контентними ключами, мовними атрибутами та ознаками належності до змінних даних, що дало змогу забезпечити їх незалежну обробку й однозначну прив'язку до записів у клієнтських системах. Обґрунтовано доцільність розгортання сервісу машинного перекладу як окремого мовного шару між системами керування контентом, що уможливило збереження структури документа, уніфікацію використання контентних ключів і централізоване керування мовними ресурсами. На цій основі описано реалізацію цільових REST-інтерфейсів, призначених для обробки одиничних і пакетних запитів відповідно з передаванням структурованих сегментів та поверненням результатів у тій самій структурі.

Реалізований прототип сервісу на основі веб-фреймворку Flask забезпечує коректну обробку сегментів зі змінними даними, збереження маркерів персоналізації та формування єдиного корпусу перекладів для різних каналів дистрибуції персоналізованого контенту. Отримані результати засвідчують практичну реалізованість сегментного підходу до машинного перекладу змінних даних і створюють підґрунтя для подальших досліджень, пов'язаних з уточненням критерію вибору рушіїв машинного перекладу, розширенням набору підключених моделей нейронного машинного перекладу та інтеграцією глосаріїв і формалізованих метрик якості на рівні сервісної архітектури.

Подяка. Автор висловлює подяку Видавничому дому «Укрпол» за надані виробничі потужності друку змінних даних, що стали базою для апробації архітектури мовного шару між клієнтськими системами та проектного веб-інтерфейсу сегментного редактора.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Хорошун О. О. Машинний переклад: історичний огляд. Нова філологія. 2021. № 82. С. 333–337. DOI: 10.26661/2414-1135-2021-82-53.
2. Любимова С. А. Проблеми автоматизованого перекладу: конспект лекцій для здобувачів вищої освіти за другим (магістерським) рівнем. Одеса: ДУ «Університет Ушинського», 2024. 49 с.
3. Моїсєєва Н., Дзикович О., Штанько А. Машинний переклад: порівняння результатів та аналіз помилок DeepL та Google Translate. Advanced Linguistics. 2023. № 11. С. 78–82. DOI: 10.20535/2617-5339.2023.11.277593.
4. Подвойська О., Козоріз І., Машинний переклад за допомогою інтернет ресурсів. Молодий вчений. 2024. № 4 (128). С. 68–71. DOI: 10.32839/2304-5809/2024-4-128-10.
5. Vaswani A., Shazeer N., Parmar N. et al. Attention Is All You Need. In: Advances in Neural Information Processing Systems. 2017. P. 5998–6008. DOI: 10.48550/arXiv.1706.03762.

6. Sutskever I., Vinyals O., Le Q. V. Sequence to Sequence Learning with Neural Networks. In: *Advances in Neural Information Processing Systems*. 2014. P. 3104–3112. DOI: 10.48550/arXiv.1409.3215.
7. Tiedemann J., Thottingal S. OPUS-MT — Building Open Translation Services for Low-Resource Languages. In: *Proc. of the 22nd Annual Conf. of the European Association for Machine Translation*. Lisbon, 2020. P. 479–480.
8. Koehn P. *Neural Machine Translation*. Cambridge: Cambridge University Press, 2020. 406 p. DOI: 10.1017/9781108608480.
9. Malets I. Review of Machine Translation Quality Assessment Methods for Personalized Content Preparation. *Youth Science: Innovation and Global Challenges*, № 2, 2025.
10. Поліграфічне обладнання УкрПол. URL: ukrpol.ua/obladnannia.

REFERENCES

1. Khoroshun, O. O. (2021). Mashynnyi pereklad: istorychnyi ohliad. *Nova filolohiia*, 82, 333–337. DOI: 10.26661/2414-1135-2021-82-53. (in Ukrainian).
2. Liubymova, S. A. (2024). Problemy avtomatyzovanoho perekladu: konspekt leksii dlia zdobuvachiv vyshchoi osvity za druhyh (mahisterskym) rivnem. Odesa: Derzhavnyi zaklad «Universytet Ushynskoho». (in Ukrainian).
3. Moisieieva, N., Dzykovych, O., & Shtanko, A. (2023). Mashynnyi pereklad: porivniannia rezultativ ta analiz pomylok DeepL ta Google Translate. *Advanced Linguistics*, 11, 78–82. DOI: 10.20535/2617-5339.2023.11.277593. (in Ukrainian).
4. Podvoiska, O., Kozoriz, I., & Honcharenko, N. (2024). Mashynnyi pereklad za dopomohoiu internet resursiv. *Molodyi vchenyi*, 4(128), 68–71. DOI: 10.32839/2304-5809/2024-4-128-10. (in Ukrainian).
5. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. In *Advances in Neural Information Processing Systems* (pp. 5998–6008). DOI: 10.48550/arXiv.1706.03762 DOI: 10.48550/arXiv.1409.3215 (in English).
6. Sutskever, I., Vinyals, O., & Le, Q. V. (2014). Sequence to sequence learning with neural networks. In *Advances in Neural Information Processing Systems* (pp. 3104–3112). DOI: 10.48550/arXiv.1409.3215. (in English).
7. Tiedemann, J., & Thottingal, S. (2020). OPUS-MT: Building open translation services for low-resource languages. In *Proceedings of the 22nd Annual Conference of the European Association for Machine Translation* (pp. 479–480). Lisbon. (in English).
8. Koehn, P. (2020). *Neural machine translation*. Cambridge: Cambridge University Press. (in English).
9. Malets I. (2025). Review of Machine Translation Quality Assessment Methods for Personalized Content Preparation. *Youth Science: Innovation and Global Challenges*, 2. (in English).
10. Printing Equipment UkrPol. Retrieved from ukrpol.ua/en/equipments.

DESIGN OF A MACHINE TRANSLATION SERVICE FOR VARIABLE DATA

I. V. Malets

Lviv Polytechnic National University,
Institute of Printing Art and Media Technologies
12 Stepana Bandery St., 79013, Lviv, Ukraine
e-mail: kimika247@gmail.com

The paper addresses the problem of designing a machine translation service for variable data that operates at the level of document segments rather than monolithic text. In contemporary publishing, marketing and educational workflows, documents are increasingly represented as sets of structured segments associated with content keys, language attributes and, in many cases, personalization markers inherited from CRM or content management systems. Conventional web-based machine translation services treat such fragments as plain text, which may lead to distortion or loss of personalization markers and break the logical link between translated text and its source fields. As a result, translation resources become fragmented across platforms, translation memory cannot be reused consistently, and additional human effort is required for manual verification and correction. The proposed approach models each publication unit u_i as a set of segments $S_i = \{s_{ij}\}$ that are processed by a centralized machine translation service acting as a language layer between web content management systems, CRM platforms and other client applications. For each segment s_{ij} the service combines translation memory and a pool of machine translation engines, selecting the actual output t_{ij} either from translation memory or from a chosen engine while preserving personalization markers as atomic tokens. The integration structure is implemented via RESTful endpoints for single-segment processing and for batch translation of segment sets. Both endpoints operate on JSON messages that encapsulate content keys, source and target language codes, translation results, information about the source (translation memory or engine) and processing status. A prototype of the service has been developed using the Flask web framework. Experimental use on structured abstracts of a master's thesis and fragments of personalized messages demonstrates that the service preserves document structure and variable markers, supports consistent mapping of translations to content keys and enables the creation of a unified translation corpus reusable across multiple distribution channels. The results indicate that segment-level integration of a machine translation service provides a feasible basis for centralized management of language resources and automated preparation of multilingual personalized content.

Keywords: machine translation; variable data; segment-based document model; REST API; personalized content; service architecture.

Стаття надійшла до редакції 29.09.2025.

Received 29.09.2025.