

УДК 004.8:004.6:339.1

МЕТОДИ ХАІ В МУЛЬТИМОДАЛЬНОМУ АНАЛІЗІ INSTAGRAM-КОНТЕНТУ ДЛЯ МАЛОГО РИТЕЙЛІУ

Постельняк С.А., Гриньова О.Є.

e-mail: sofiiia.postelniak@nure.ua, olena.hrynova@nure.ua

Харківський національний університет радіоелектроніки, каф. ШІ
м. Харків, Україна

This research proposes a hybrid Instagram data collection pipeline for small retail trend analysis, integrating official API metrics with web scraping. By utilizing a multimodal AI framework – BERT for semantic text analysis and CLIP for visual feature extraction – the system transforms unstructured social media streams into structured style embeddings and time-series proxy metrics. These outputs provide a robust mathematical foundation for intelligent decision support, enabling small enterprises to identify latent demand vectors and adapt assortment strategies to volatile market environments.

У малому ритейлі управлінські рішення часто приймаються інтуїтивно через обмежений доступ до аналітичних інструментів і нестачу даних. Через низький рівень цифровізації стає складно вчасно помічати зміни в попиті, ідентифікувати латентні вектори попиту та здійснювати предиктивне планування асортименту, що особливо відчутно в умовах воєнного стану та швидкої зміни купівельної поведінки споживачів. Більшість аналітичних підходів орієнтовані на великі компанії з великими обсягами даних, тому вони малопридатні для малого бізнесу, а нестача даних знижує можливості аналізу та прогнозування попиту. За таких умов важливу роль відіграють зовнішні джерела інформації, зокрема соціальні мережі. Instagram є однією з найпопулярніших платформ для малого бізнесу у сфері торгівлі в Україні у fashion-сегменті та використовується як вітрина товарів і канал комунікації з покупцями. Активність користувачів у цій соціальній мережі дозволяє помічати зміни споживчих інтересів і появу нових трендів раніше, ніж вони відображаються у фактичних показниках продажів, що дає змогу частково компенсувати нестачу власних даних малого бізнесу.

Об'єктом дослідження є система формування та циркуляції інформаційних сигналів у сфері малої роздрібної торгівлі України, яка відображає зміни споживчих уподобань і модних трендів у цифровому середовищі. Предметом дослідження є процеси збору, структуризації та підготовки даних з Instagram, які використовуються в математичних моделях для аналізу та прогнозування трендових явищ. Метою дослідження є розробка методів ідентифікації споживчих трендів на основі мультимодального аналізу даних із застосуванням технологій ХАІ для забезпечення інтерпретованості та прозорості управлінських рішень у малому ритейлі.

Офіційний доступ до даних Instagram через Graph API є обмеженим для наукових і прикладних задач, пов'язаних із розпізнаванням трендів, сегментацією та прогнозуванням. API надає доступ переважно до базових метрик залученості (engagement metrics) і обмеженої історії публікацій. При цьому він не охоплює контентну семантику, візуальні характеристики, поведінкові патерни та неформальні ринкові сигнали [1]. Використання офіційного API для моделей, які потребують повного інформаційного шару, є обмеженим. Для подолання обмежень офіційного доступу до даних Instagram в роботі запропоновано гібридний підхід підготовки даних (рис. 1).

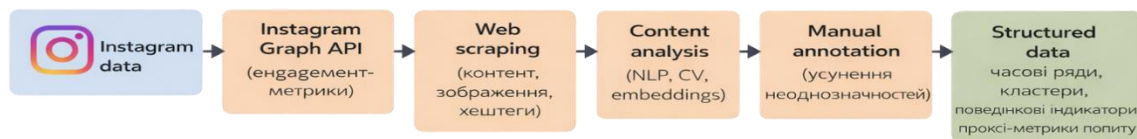


Рисунок 1 – Етапи збору та підготовки даних Instagram для аналізу трендів

На першому рівні застосовуються відкриті Instagram Graph API для формальної верифікації engagement-метрик (impressions, reach, likes, comments, saves, shares), а також поглиблені характеристики користувачів профілів (геодемографічна структура, часові патерни активності, параметри охоплення), що є структурованими, мають чітко визначену семантику та можуть безпосередньо використовуватися в математичних моделях. Водночас, офіційний API не охоплює повний контентний шар платформи (семантика візуального матеріалу, спільна поява (co-occurrence) стилів, еволюція хештегів, поведінкові зв'язки між об'єктами).

Другий рівень системи базується на модульному веб-скрейпінгу, який використовується для отримання інформації, недоступної через офіційний API Instagram. Такий підхід дозволяє автоматизовано збирати дані безпосередньо з веб-сторінок платформи, що є поширеною практикою у задачах аналізу цифрового контенту. До отриманих даних належать візуально-стильові ознаки товарів, семантичні маркери опису, ко-теги та ко-категорії, контекстні коментарі, структура посилань на суміжний контент, а також частотні профілі товарних ніш.

У результаті веб-скрейпінгу формується масив неструктурованих даних (зображення, текстові описи та хештеги), які в сирому вигляді не можуть бути безпосередньо використані в математичному моделюванні. Для їх подальшої обробки застосовується автоматизована контент-аналітика на основі моделей оброблення природної мови (NLP) та комп'ютерного зору (CV). Зокрема, моделі побудови семантичних векторних представлень (sentence embeddings), такі як BERT та його похідні, використовуються для трансформації лінгвістичних структур у багатовимірний векторний простір, що відображають їх зміст і контекст [2].

Для аналізу візуального контенту застосовуються моделі комп'ютерного зору, зокрема згорткові нейронні мережі (CNN), Vision Transformers та моделі типу CLIP. Останні реалізують підхід Visual Models from Natural Language Supervision, який полягає у спільному аналізі зображень і текстових описів, що дозволяє встановлювати відповідність між візуальними та семантичними характеристиками продуктів і стилів. Використання таких моделей дає змогу формувати векторні представлення візуально-семантичних ознак, придатні для подальшого аналізу та кластеризації.

Третій рівень передбачає селективну анотовану ручну експертизу, яка проводиться для частини зібраних даних з метою усунення категоріальних неоднозначностей і підвищення точності класифікації [3]. Це важливо для fashion-сегменту, оскільки стилістичні категорії часто мають нечіткі межі та залежать від контексту. Користувачі можуть використовувати різні назви для позначення одного й того самого стилю, наприклад «minimalist style», «clean aesthetic», а також «retro look» і «vintage style». Формально це різні хештеги, але за змістом вони можуть описувати схожі образи. Тому ручна анотація виконує функцію семантичної перевірки та допомагає уточнити категорії, підвищуючи точність як supervised, так і unsupervised моделей. У результаті обробки неструктуровані дані перетворюються на впорядковані структури, зокрема часові ряди, кластерні мітки, стильові embeddings, поведінкові індикатори та проксі-метрики попиту для використання в математичних моделях аналізу, прогнозування та підтримки управлінських рішень.

Етап збору та підготовки даних є ключовим, оскільки саме на ньому контент із соціальних мереж перетворюється на структуровані дані, зокрема часові ряди, категоріальні кластери, поведінкові індикатори та проксі-метрики попиту (не лише прогнозує зміну вектора попиту, наприклад, з США на Францію, а й візуалізує та аргументує причини такої зміни через виділення домінантних ознак стилю). Отримані дані можуть використовуватися в системі підтримки прийняття рішень для виявлення трендів, прогнозування попиту та формування асортиментної політики малого бізнесу.

Список використаних джерел:

1. Bekavac L., Mayer S. Auditing Meta and TikTok Research API Data Access under Article 40(12) of the Digital Services Act // arXiv. – 2026. URL: <https://arxiv.org/abs/2601.12390> (дата звернення: 25.01.2026).
2. Devlin J., Chang M.-W., Lee K., Toutanova K. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding // Proceedings of NAACL-HLT. – Minneapolis, 2019. – P. 4171–4186.
3. Filatov V. O., Yerokhin M. A. Improved multi-objective optimization in business process management using R-NSGA-II. Radio electronics, computer science, control. 2023. No. 3. P. 187. DOI: <https://doi.org/10.15588/1607-3274-2023-3-18>.