

УДК 004.852:519.85

DOI <https://doi.org/10.32689/maup.it.2025.2.18>**Костянтин МІНКОВ**

аспірант кафедри програмного забезпечення комп'ютерних систем,

Чернівецький національний університет імені Юрія Федьковича

ORCID: 0009-0007-3400-050X

ВИКОРИСТАННЯ ACTIVE LEARNING ЯК ІНСТРУМЕНТУ ДЛЯ УТОЧНЕННЯ МІТОК У ЗАДАЧАХ, ЗАСНОВАНИХ НА ПРИНЦИПАХ СЛАБКОГО НАВЧАННЯ

Анотація. Було запропоновано ітеративний метод, який використовує активне навчання для виявлення та корекції помилкових міток у слабо анотованому наборі даних. Вихідним є набір даних, у якому приклади мають лише зашумлені або частково надані анотації. Запропонований підхід поступово уточнює навчальну вибірку шляхом ідентифікації зразків, які з високою ймовірністю містять помилкові помічення, та передбачає залучення цілеспрямованої перевірки, як з боку людини, так і за допомогою евристичних правил із високим рівнем довіри.

Метою статті є підвищення якості міток у слабо анотованому наборі даних із мінімальними витратами на ручну анотацію. Для цього запропоновано ітеративний підхід, що базується на активному навчанні та селективній валідації з боку людини.

Методологія. Розроблено ітеративний метод, у якому дискримінативна модель (DistilBERT) навчається на поточному наборі міток, оцінює невизначеність своїх передбачень на основі варіативності виходів, а для зразків із високою невизначеністю генерує псевдомітки. Ці мітки передаються на швидке підтвердження або виправлення анотаторам. В основу слабого анотування покладено чотири евристичні λ -функції, що формують початкові мітки, а генеративна модель враховує залежності між цими функціями через регуляризовану матрицю зв'язків. Експериментально досліджено ефективність підходу на наборі IMDb Movie Reviews, що містить 50 000 текстових відгуків, з чітким розподілом на навчальну, валідаційну й тестову вибірки.

Наукова новизна. На відміну від традиційних методів активного навчання, які передбачають ручну переанотацію найневпевненіших зразків, запропоновано гібридну стратегію, що поєднує слабе навчання, псевдоанотування та вибірку людську перевірку. Це дозволяє досягати цільової точності моделі в 2–3 рази швидше, ніж класичні підходи без слабких міток, за рахунок ефективного використання обмеженого людського ресурсу та стартової інформації з евристичних правил.

Висновки. Ітеративна стратегія, яка поєднує відбір зразків на основі розбіжності передбачень, автоматичне псевдоанотування та селективну людську валідацію, дозволяє ефективно покращити якість міток у слабо анотованих даних без повної переанотації. Запропонований метод продемонстрував конкурентні результати на реальному корпусі відгуків IMDb, забезпечуючи високу точність моделі класифікації зі зниженими витратами на ручну працю.

Ключові слова: псевдомітки, матриця переходів, евристичні правила, оракул, зважування за значущістю, впевненість моделі.

Kostiantyn MINKOV. THE USE OF ACTIVE LEARNING FOR REFINING LABELS IN WEAK SUPERVISION BASED TASKS

Abstract. An iterative method has been proposed that uses active learning to detect and correct mislabeled data in a sparsely annotated dataset. The starting point is a dataset in which examples have only noisy or partially provided annotations. The proposed approach gradually refines the training sample by identifying examples that are highly likely to contain incorrect labels and involves targeted verification, both by humans and using heuristic rules with a high level of confidence.

The goal of this paper is to improve the quality of labels in a sparsely annotated dataset with minimal manual annotation costs. To do this, we propose an iterative approach based on active learning and selective human validation.

Methodology. An iterative method has been developed in which a discriminative model (DistilBERT) is trained on the current set of labels, evaluates the uncertainty of its predictions based on the variability of outputs, and generates pseudo-labels for samples with high uncertainty. These labels are sent to annotators for quick confirmation or correction. Weak annotation is based on four heuristic λ -functions that form the initial labels, and the generative model takes into account the dependencies between these functions through a regularized connection matrix. The effectiveness of the approach was experimentally investigated on the IMDb Movie Reviews dataset, which contains 50,000 text reviews, with a clear division into training, validation, and test samples.

Scientific novelty. Unlike traditional active learning methods, which involve manual re-annotation of the most uncertain samples, a hybrid strategy is proposed that combines weak learning, pseudo-annotation, and selective human verification. This allows the target model accuracy to be achieved 2–3 times faster than classical approaches without weak labels, due to the efficient use of limited human resources and initial information from heuristic rules.

Conclusions. An iterative strategy that combines sample selection based on prediction divergence, automatic pseudo-annotation, and selective human validation allows for effective improvement of label quality in weakly annotated data without complete re-annotation. The proposed method demonstrated competitive results on a real-world IMDb review corpus, providing high classification model accuracy with reduced manual labor costs.

Key words: pseudo-labels, transition matrix, heuristic rules, oracle, importance weighting, model confidence.

© К. Мінков, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

Вступ. У машинному навчанні одним із головних викликів залишається потреба у великій кількості високоякісних анотованих даних. Повноцінна ручна анотація таких наборів даних є трудомісткою, дорогою та часто нездійсненною у прикладних умовах. Саме тому дедалі більшої уваги набувають методи, що дозволяють ефективно навчатися на частково анотованих або зашумлених вибірках. Зокрема, підходи, що ґрунтуються на слабкому навчанні, дають змогу використовувати евристичні, правила, автоматично згенеровані сигнали чи агреговані джерела як сурогатні мітки.

Проте слабкі мітки можуть бути ненадійними або хибними, що без відповідної корекції призводить до значного зміщення у моделі. Щоб подолати цю проблему, останні дослідження все активніше застосовують активне навчання як механізм, що дозволяє вибірково запитувати інформацію про найбільш невизначені зразки.

Аналіз останніх досліджень і публікацій. Бігель С. з колегами [2] запропонували метод Active WeaSuL, який інтегрує активне навчання у парадигму слабого навчання задля підвищення якості маркування в умовах обмеженого доступу до вручну анотованих даних. У рамках Active WeaSuL передбачено залучення обмеженої кількості справжніх міток, що надаються експертами вибірково, тобто лише для тих зразків, які викликають найбільшу невизначеність у моделі. Модель базується на генеративному підході до об'єднання слабких міток, який модифікується за допомогою додаткової пеналізації до функції втрат. Ця пеналізація змушує модель узгоджувати ймовірнісні мітки зі справжніми мітками, отриманими в межах активного навчання. Замість локального переписування міток, запропонований метод коригує саму логіку агрегації слабких сигналів. Щоб ефективно визначати найбільш інформативні приклади для запиту, використовується стратегія вибірки на основі максимального розходження між ймовірнісним прогнозом моделі і частотним розподілом експертних міток у відповідному кластері (bucket), що оцінюється через дивергенцію Кульбака-Лейблера. Експериментальні дослідження охоплюють як синтетичні дані (двовимірні гаусіани), так і реальні набори даних, тобто виявлення візуальних відносин і класифікацію спаму на YouTube. У всіх випадках модель демонструє стабільну перевагу над базовими стратегіями [2].

В межах дослідження Ванга Ю. та ін. [5] представлено новий підхід до задачі виявлення об'єктів, який поєднує активне навчання з методами як зі слабким так і неповним підкріпленнями. ALWOD реалізує механізм warm-start активного навчання, який ґрунтується на попередньому тренуванні моделей за допомогою невеликого підмноження повністю анотованих зображень і великого слабо анотованого набору, доповненого синтетично згенерованими зображеннями. Суттєвою інновацією є використання генератора зображень, який дозволяє створити допоміжний набір повністю анотованих зображень шляхом вставки об'єктів з невеликої кількості справжніх FA-зображень у фонові сцени. Це дозволяє навчити першу модель в умовах обмеженої анотації, що забезпечує високу стартову продуктивність. На відміну від традиційних підходів, ALWOD не вимагає повного створення нових анотацій: замість цього анотатори виправляють або уточнюють вже запропоновані детектором межі об'єктів та категорії [5].

Розширення активного навчання, яке дозволяє вибирати найінформативніші зразки для анотації, а також контролювати точність цих анотацій, розроблено Олміном А., Дж. Ліндквістом, Свенссоном Л. та Ліндстеном Ф. [1]. Ідея науковців базується на тому, що менш точні анотації є дешевшими, тому їх можна зібрати більше в межах того самого бюджету, розширивши покриття простору ознак. Ключовим внеском є нова функція вибору (acquisition function), яка базується на взаємній інформації між слабкою анотацією та моделлю (а не між анотацією і справжньою міткою, як у попередніх підходах). Це дозволяє уникнути залежності від латентної змінної й адаптувати метод до будь-якого типу анотаційного шуму. Усі обчислення були інтегровані у фреймворк Gaussian Processes, що надає ймовірнісну оцінку та можливість ефективної побудови невизначеності. Результати експериментів на синусоїдальних функціях, наданих UCI та синтетичних задачах класифікації демонструють, що облік точності анотацій дозволяє суттєво підвищити продуктивність моделі за фіксованого бюджету. Особливо сильний ефект спостерігається у випадках, коли низькоточні анотації є значно дешевшими [1].

З огляду на проаналізовану наукову літературу, слід відзначити, що питання застосування підходу Active Learning як засобу уточнення міток у задачах, заснованих на принципах слабого навчання, досі залишається недостатньо розробленим. Недостатність ґрунтовних досліджень у цьому напрямі вказує на потребу в подальшому теоретичному осмисленні та емпіричному вивченні зазначеної проблематики з метою підвищення ефективності методів слабо контрольованого навчання.

Метою статті є підвищення якості міток в слабо поміченому наборі даних із мінімальними витратами на анотацію шляхом ітеративного відбору зразків із високою розбіжністю передбачень, генерації для них псевдоміток та залучення людини лише для підтвердження або корекції цих машинних пропозицій.

Виклад основного матеріалу дослідження. Па початку поточного дослідження у центрі уваги постає слабо помічений набір даних, який виражається як:

$$D = \left\{ \left(x_i, \hat{y}_i \right) \right\}_{i=1}^N \quad (1)$$

в якому кожна мітка $\tilde{y}_i$ може бути шумною, або вилученою за допомогою простих правил, а не мати надійне джерело походження.

Незважаючи на таку невизначеність, проводиться навчання первинної моделі f_0 , яка може бути як глибокою нейронною мережею, так і іншою параметричною моделлю. Навчання проводиться шляхом мінімізації середньої сурогатної втрати слабких міток (weak labels):

$$\hat{\theta} = \arg \min_{\theta} \frac{1}{N} \sum_{i=1}^N \ell(f_{\theta}(x_i), \tilde{y}_i) \quad (2)$$

де $\ell(\cdot, \cdot)$ може представляти як перехресну ентропію для функцій класифікації, так і середнє квадратичне похибки для функцій регресії [6].

Така стратегія «Warm-Start» забезпечує модель першим необробленим відображенням від вхідних даних до шумних цільових значень навіть за неточності певних міток [10].

Задля опису впливу шумних слабких міток та уточнення на рівні псевдоміток вводиться формальна модель шуму та отримуються межі поширення помилок для моделі порогового відсічення за впевненістю (confidence thresholding). Припускається, що кожна слабка мітка $\tilde{y}$ складається з однієї істинної мітки y в каналі шуму, залежного від класу, який описується за допомогою $C \times C$ матриці переходів T , тобто [3]:

$$T_{ij} = \Pr(\tilde{y} = j | y = i), \sum_{j=1}^C T_{ij} = 1 \forall i \quad (3)$$

Така модель охоплює як симетричний шум $T_{ij} = 1 - \eta, T_{ij} = \frac{\eta}{C-1}$ для $j \neq i$, так і асиметричний шум, тобто класово-залежний, який передбачає систематичне змішування певних класів. Згідно з цією моделлю, емпіричний ризик на основі зашумлених міток $R_{шум}(f) = E_{(x,y) \sim D} [\ell(f(x), \tilde{y})]$ залежить від «чистого» ризику $R_{чист}(f) = E[\ell(f(x), y)]$ таким чином:

$$R_{шум}(f) = \sum_{i=1}^C \sum_{j=1}^C T_{ij} E[\ell(f(x), j) | y = i] \quad (4)$$

Якщо матриця T є повноранговою, може застосовуватися T^{-1} , щоб отримати за допомогою зважування за значущістю (importance weighting) незміщену оцінку $R_{чист}(f)$:

$$\hat{R}_{чист}(f) = \frac{1}{N} \sum_{k=1}^N \sum_{i=1}^C [T^{-1}]_{k,i} \tilde{y}_{k,i} \ell(f(x_k), i) \quad (5)$$

З метою уникнення залежності від певної повнорангової матриці переходів T , для кожного класу c передбачається принаймні одна якірна множина (anchor set) [8]:

$$A_c = \{x : p(y = c | x) = 1\} \quad (6)$$

За таких умов, діагональний елемент T_{cc} може оцінюватися безпосередньо за допомогою:

$$\hat{T}_{cc} = \Pr(\tilde{y} = c | x \in A_c) \approx \frac{1}{|A_c|} \sum_{x \in A_c} 1_{\{\tilde{y}(x) = c\}} \quad (7)$$

Недіагональні елементи T_{ij} потім отримуються шляхом нормування по кожному рядку так, щоб виконувалася умова $\sum_j \hat{T}_{ij} = 1$.

Оскільки дійсні якірні множини можуть бути невідомими, вони можуть бути наближено обрані як сукупність зразків з високою впевненістю моделі, наприклад:

$$\{x : p_{\theta}(c | x) \geq \tau_{якір}\} \quad (8)$$

де $\tau_{якір}$ є близьким 1. Таким чином створюється послідовний оцінювач T , чия похибка зникає з підвищенням порогового значення.

Навіть якщо T є оціненою, вона може бути погано обумовленою (ill-conditioned). Внаслідок цього, прямо інверсія замінюється регуляризованою за Тихоновим псевдоінверсією [9]:

$$\hat{T}_{\lambda}^{-1} = (\hat{T}^{\top} \hat{T} + \lambda I)^{-1} \hat{T}^{\top} \quad (9)$$

при чому $\lambda > 0$ визначається за допомогою перехресного затвердження (cross validation). Таким чином відбувається зменшення посилення оціночного шуму в напрямку малих сингулярних значень та отримується чисельно стабільна та незміщена оцінка ризику з похибкою порядку $O(\lambda)$ [7].

В наступній фазі базова модель f_θ вбудовується у цикл активного навчання задля точного знаходження точок даних, за яких мітки потребують втручання людини. Спочатку, для кожного зразка x_i розраховується міра неточності в межах моделі. У завданнях класифікації для цього застосовується згаданий розподіл $p_\theta(y=c|x_i)$ для всіх класів c . Враховуючи цей розподіл обчислюється ентропія:

$$H(x_i) = -\sum_c p_\theta(y=c|x_i) \log p_\theta(y=c|x_i) \quad (10)$$

яка показує, наскільки широко модель розподіляє свої ймовірності між класами.

До того ж, визначається показник найменшої впевненості (least-confident Score):

$$LC(x_i) = 1 - \max_c p_\theta(y=c|x_i) \quad (11)$$

що визначає, наскільки низькою є впевненість класів ймовірності [4].

Після цього здійснюється сортування слабо помічених зразків за спаданням відповідно до цих обох показників. Найвищу позицію у цьому сортуванні займають дані, за яких модель є найбільш невпевненою.

Якщо призначити псевдомітку $\hat{y} = \operatorname{argmax}_c p_\theta(c|x)$ тільки тоді, коли впевненість моделі дорівнює $C(x) = \max_c p_\theta(c|x) \geq \tau$, тоді підпадають зразки з низьким оціненим рівнем шуму. Таким чином:

$$S_\tau = \{x : C(x) \geq \tau\} \quad (12)$$

що є підмножиною зразків x , для яких впевненість моделі $C(x)$ є не меншою за заданий поріг τ .

У межах класово залежної шумової моделі для кожного зразку $x \in S_\tau$, тобто такого, для якого впевненість моделі $C(x) \geq \tau$, ймовірність помилки псевдомітки (тобто того, що $\hat{y} \neq y$) обмежується зверху наступною оцінкою:

$$\Pr(\hat{y} \neq y | x \in S_\tau) \leq \frac{1-\tau}{\tau} \max_{i \neq j} T_{ij} \quad (13)$$

Ця нерівність показує, що ймовірність неправильного класифікаційного передбачення зменшується лінійно зі зростанням впевненості моделі τ , навіть у присутності шуму, описаного матрицею переходів T .

Внаслідок цього, додаткова частка ризику, яка вводиться за посередництва передбачення тільки найбільш впевнених псевдоміток, обмежується за допомогою

$$R(f_{S_\tau}) - R_{\text{ист}}(f_{S_\tau}) \leq (1-\tau) \max_{i \neq j} T_{ij} \quad (14)$$

де якщо $\tau \rightarrow 1$, шум у псевдомітках зникає.

Загалом ці порогові значення показують, що щонайбільше $(1-\tau) \max_{i \neq j} T_{ij}$ псевдоміток є хибними; та що в результаті зважуваного за значущістю донавчання отримується незміщена оцінка дійсного ризику.

Хоча порогове відсічення за рівнем впевненості обмежує похибку у випадку з псевдомітками, зразки з найменшим рівнем впевненості також мають надсилатися до зовнішнього анотатора, тобто «оракула». Під час кожної анотації циклу активного навчання обираються ті вхідні дані x_i , чий показники неточності знаходяться за межами встановленого порогу яким потім оракул призначає нові мітки.

Оракул моделюється як анотатора з шумом, що характеризується клас-залежною матрицею помилок:

$$O_{ij} = \Pr(y_{\text{оракул}} = j | y_{\text{дійсн}} = i), \sum_j O_{ij} = 1 \quad (15)$$

Така матриця оцінюється за допомогою малого відкладеного калібрувального набору даних, тобто оракулу надаються зразки з відомою міткою а його помилки – фіксуються. Під час уточнення, кожна підтверджена псевдомітка додаткового зважується відповідно до надійності оракула для цього класу. Таким чином міткам з класів, у випадку з якими оракул часто допускає помилок, призначаються менші вагові коефіцієнти.

Оскільки кожен запит до оракула спричиняє витрати c_i та затримку ℓ_i , визначаються загальний бюджет B і допустима максимальна затримка L . Тому вибір множини запитів формулюється як задача оптимізації:

$$\max_{S \subseteq \{x_i\}} \sum_{x_i \in S} \Delta_i \text{ за умов } \sum_{x_i \in S} c_i \leq B, \max_{S \subseteq \{x_i\}} \ell_i \leq L \quad (16)$$

де Δ_i є очікуваним зменшенням узагальненої помилки внаслідок запиту зразку x_i , що оцінюється за допомогою функцій впливу, або невизначеності урахуванням різноманіття.

Якщо індекс i обирається на ітерації t , його вихідна слабка мітка $\tilde{y}_i$ замінюється новою отриманою міткою та множина навчання оновлюється відповідним чином:

$$D_{t+1} = \{(x_j, \tilde{y}_i)\}_{j \in S_t} \cup \{(xi, y_i^{нов})\}_{i \in S_t} \quad (17)$$

де S_t відповідає множині індексів, запитуваних на часовому кроці t .

Після повторного призначення міток, знову проводиться навчання f_θ на D_{t+1} шляхом мінімізації:

$$\hat{\theta}_{t+1} = \operatorname{argmin}_\theta \frac{1}{|D_{t+1}|} \sum_{(x,y) \in D_{t+1}} \ell(f_\theta(x), y) \quad (18)$$

що поширює експертні або відповідні правилам уточнення по загальному середовищу прийняття рішень моделі.

В межах циклу уточнення з участю людини, кожна слабка мітка не піддається переписуванню новою отриманою від оракула міткою. На практиці, цей процес переписування міток уточнюється шляхом застосування механізму селективного виправлення в ході якого вирішується, чи довіряти мітці, виправляти її чи поєднувати з альтернативним варіантом в залежності від впевненості моделі та схвалення оракула. Після запитування оракула на предмет зразка x_i , оцінений рівень впевненості моделі $C_i = \max_c p_\theta(y = \# x_i)$ порівнюється відносно порогового значення $\tau \in (0,1)$. Якщо $C_i \geq \tau$, а мітка оракула $y_i^{нов}$ співпадає з вихідною слабкою міткою, остання приймається та розглядається в якості достатньо надійної, згодом не піддаючись змінню. У випадках коли $C_i < \tau$, або якщо судження оракула протирічить слабкому анотуванню, похібний сигнал був достатньо шумним щоб потребувати виправлення.

Крім то, коли як слабка мітка, так і мітка оракула містять взаємодопнювальну інформацію, тобто, якщо слабкий анотатор є точним для широкої підожини зразків, але систематично помиляється на окремих граничних випадках, ці два сигнали об'єднуються з врахуванням зважування за впевненістю. Позначаючи через $\alpha \in [0,1]$ ваговий коефіцієнт, призначений початковій слабкій мітці, обчислюється фінальна цільова мітка для навчання:

$$y_i^{fin} = \alpha \tilde{y}_i + (1 - \alpha) y_i^{нов} \quad (19)$$

де α може бути сама функцією C_i . Таким чином, якщо $\alpha = \frac{C_i}{C_i + (1 - C_i)}$, тоді до слабкої мітки при-

значається більше довіри за умови впевненості моделі, а вага змінюється до $y_i^{нов}$ при зменшенні рівня впевненості. Таким чином, механізм уточнення плавно переходить між повністю правилами на базі міток та виправленнями оракула. Після завершення кожної ітерації модель перенавчається на вибірково скоригованих цільових мітках, і таким чином лише справді сумнівні мітки змінюються і достовірність тих зразків зберігається, які вже вважаються надійними.

Після включення скоригованих цільових міток до навчального набору даних модель одразу перенавчається, що таким чином замикає цикл «навчання → запит → корекція → перенавчання». Оновлена оцінка параметрів формується у вигляді:

$$\theta_{t+1} = \operatorname{argmin}_\theta \frac{1}{|D_t|} \sum_{(x_i, y_i) \in D_t} \ell(f_\theta(x_i), y_i) \quad (20)$$

що мінімізує ту саму функцію втрат ℓ , що й на початку. Оскільки кожна щойно виправлена мітка стосується зразку з максимальною невизначеністю, її включення непропорційно сильно впливає на здатність моделі вирішувати подібні неточності в інших зразках. Практика показує, що вже після кількох ітерацій межа прийняття рішень моделі істотно стабілізується, а кількість зразків з високою ентропією в пулі невідомих даних зменшується.

Цикл продовжується доти, доки він не навмисно не буде зупинений. Використовуються два критерії: по-перше, значенням T_{max} обмежується загальна кількість запитів до оракула, щоб уникнути необмежених витрат на анотування; по-друге, спостерігається або функція втрат на валідаційній множині $L_{val}(\theta_t)$ або внутрішня міра якості міток $Q(D_t)$. Таким чином, цикл зупиняється, коли виконується хоча б одна з умов:

$$L_{val}(\theta_{t+1}) - L_{val}(\theta_t) < \varepsilon \text{ або } Q(D_{t+1}) - Q(D_t) < \delta \quad (21)$$

де ε та δ є малими граничними значеннями. У цей момент подальші виправлення міток дають усе менший ефект.

Щоб збалансувати швидкість анотування та якість міток, застосовується політика з подвійним порогом:

– високий поріг впливу τ_H : приклади з невизначеністю вище цього порогу одразу надсилаються до оракула;

– низький поріг впливу τL : приклади з невизначеністю між τL та τH накопичуються для періодичних пакетних запитів.

Цей підхід дозволяє обробляти найважливіші приклади в реальному часі, не перевантажуючи при цьому анотатора менш критичними випадками.

Розглядаючи виправлені мітки як часткову процедуру знешумлення у межах PAC-підходу до навчання (Probably Approximately Correct learning – ймовірісно приблизно коректне навчання), можна показати, що кожне виправлення, підтверджене оракулом, зменшує ефективний рівень шуму η у навчальному наборі, тим самим звужуючи стандартну межу узагальнення. Нехай H позначає гіпотетичний клас моделей із VC-розмірністю (Валника – Червоненкіса) d . Якщо початковий рівень шуму слабких міток дорівнює η_0 , тоді після m селективних виправлень залишковий рівень шуму η_m задовольняє умову:

$$\eta_m \leq \eta_0 \left(1 - \frac{N_m}{N} \right) \quad (22)$$

а з імовірністю не менше $1 - \delta$ узагальнена похибка вдосконаленої моделі f_{θ_m} обмежується так:

$$R(f_{\theta_m}) \leq \hat{R}^{(m)}(f_{\theta_m}) + O(N d \log \left(\frac{\sqrt{\eta_m} + \log(1/\delta)}{N} \right)) \quad (23)$$

де $R^{(m)}$ є емпіричним ризиком на частково виправленій вибірці, а N – її розміром. Таким чином, кожне виправлення звужує границю, зменшуючи η_m , якщо N є досить великим, щоб контролювати дисперсію.

Спираючись на результати з теорії складності вибірки для активного навчання, можна обмежити кількість запитів до оракула, необхідну для досягнення допустимої надлишкової помилки ϵ . Якщо гіпотетичний клас H має коефіцієнт розбіжності θ щодо розподілу зі слабким поміченням, тоді достатньо $O(\theta d \log(1/\epsilon))$ запитів, щоб зменшити загальну похибку до рівня, близького до ідеального безшумного випадку. На практиці θ відображає, як швидко невпевненість моделі концентрується навколо помилково помічених прикладів; менше θ передбачає меншу необхідну кількість запитів.

Для оцінки методу в роботі у реальних умовах був проведений експеримент на наборі даних IMDb Movie Reviews, що містить 50 000 текстових відгуків про фільми, розподілених порівну між двома класами: позитивний і негативний. Для тренування було використано 25 000 прикладів, ще 5 000 – для валідації, і 20 000 – для тестування.

Були побудовані чотири слабкі функції (λ_1 – λ_4), що базуються на евристичних припущеннях, а саме:

- λ_1 – присутність позитивних слів з фіксованого словника;
- λ_2 – частота негативної лексики;
- λ_3 – емоційна забарвленість заголовку відгуку;
- λ_4 – оцінка на IMDb (у разі наявності).

Генеративна модель навчалася з урахуванням залежностей між λ -функціями через регуляризовану матрицю зв'язків. Значення гіперпараметра α було встановлено на 100. Дискримінантну модель реалізовано на базі DistilBERT, тобто попередньо натренованої трансформерної архітектури, адаптованої до завдання двокласової класифікації.

На (рис. 1) показано зміну F1-значення протягом 100 ітерацій активного уточнення міток. Початкова продуктивність моделі після лише слабого маркування становила 0.74. Після 10 активних ітерацій точність зросла до 0.87, а після 25 – до 0.91.

Для порівняння були реалізовані чисте активне навчання на базі логістичної регресії без слабких міток; та метод, у якому слабкі мітки замінюються на істинні після уточнення. Метод активного навчання без слабких міток досяг F1 = 0.91 після 60 ітерацій, а другий підхід – після 80. У порівнянні, запропонований у цьому дослідженні підхід досягає аналогічного рівня точності в 2–3 рази швидше завдяки використанню слабких правил як стартової точки.

Висновки. Було запропоновано ітеративний метод, який використовує активне навчання для виявлення та корекції помилкових міток у слабо анотованому наборі даних. Вихідним є набір даних, у якому приклади мають лише зашумлені або частково надані анотації. Запропонований підхід поступово уточнює навчальну вибірку шляхом ідентифікації зразків, які з високою ймовірністю містять помилкові помічення, та передбачає залучення цілеспрямованої перевірки, або з боку людини, або за допомогою евристичних правил із високим рівнем довіри. На кожній ітерації модель навчається на поточному наборі міток і перевіряється на узгодженість своїх передбачень: оцінюється зміна вихідних прогнозів при варіаціях вхідних даних або протягом епох навчання, при цьому висока варіативність трактується як ознака невизначеності мітки. Замість повної переанотації створюються машинно згенеровані

псевдомітки для сумнівних прикладів і подаються на швидке підтвердження чи виправлення анотаторам. Після інтеграції уточнених міток модель перенавчається, що поступово вдосконалює функцію маркування та дозволяє поетапно усувати найбільш критичні помилки. Протягом кількох циклів ця гібридна стратегія, що поєднує вибір на основі розбіжностей та вибірку людську валідацію, поступово підвищує якість анотацій без необхідності повного перепомічення всього набору даних.

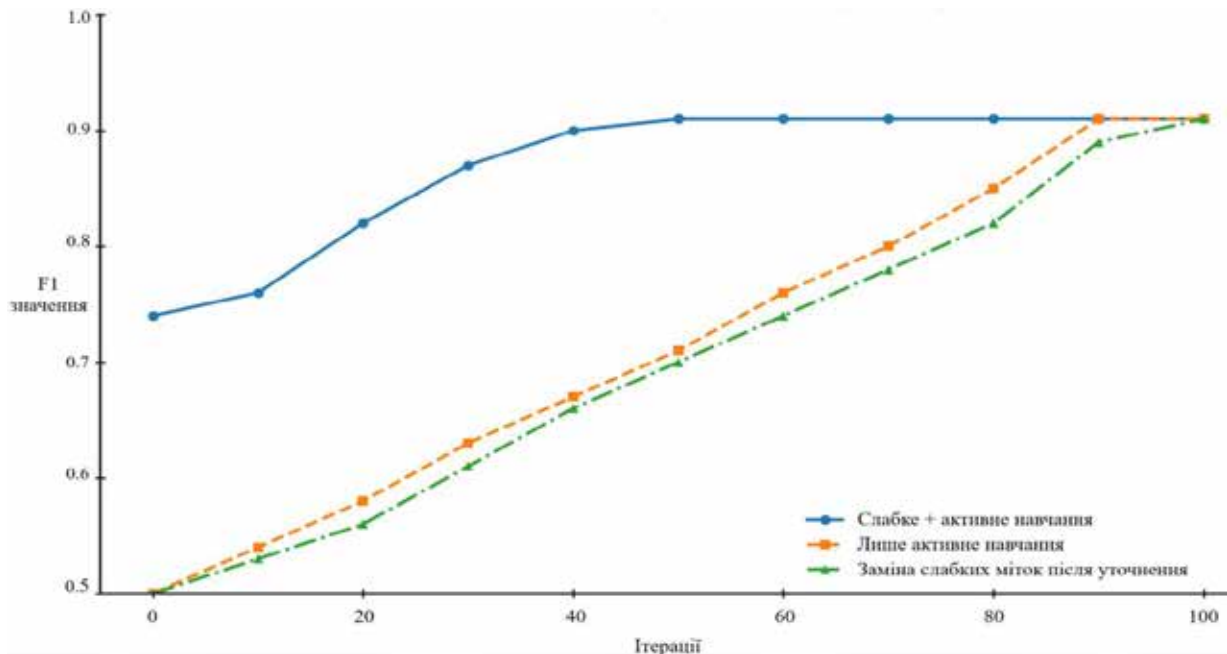


Рис. 1. Порівнянні методів з огляду на F1-значення

Список використаних джерел:

1. Active Learning with Weak Supervision for Gaussian Processes. Olmin A., Lindqvist J., Svensson L., Lindsten F. *ArXiv* : website. 2024. DOI: <https://doi.org/10.48550/arXiv.2204.08335>
2. Active WeaSuL: improving weak supervision with active learning. Biegel S., El-Khatib R., Vilas Boas L. Oliveira O., Baak M., Aben N. *ArXiv* : website. 2021. DOI: <https://doi.org/10.48550/arXiv.2104.14847>
3. Adaptive Confidence Thresholding for Monocular Depth Estimation. Choi H., Lee H., Kim S., Kim S., Kim S., Sohn K., Min D. *ArXiv* : website. 2021. DOI: <https://doi.org/10.48550/arXiv.2009.12840>
4. Agrawal A., Tripathi S., Vardhan M. Active Learning Approach Using a Modified Least Confidence Sampling Strategy for Named Entity Recognition. *Progress in Artificial Intelligence*. 2021. Vol. 10. DOI:10.1007/s13748-021-00230-w
5. ALWOD: Active Learning for Weakly-Supervised Object Detection. Wang Y., Ilic V., Li J., Kisacanin B., Pavlovic V. *ArXiv* : website. 2023. DOI: <https://doi.org/10.48550/arXiv.2309.07914>
6. Auto-generating weak labels for real & synthetic data to improve label-scarce medical image segmentation. Deshpande T., Prakash E., Ross E. G., Langlotz C., Ng A., Valanarasu J. M. J. *ArXiv* : website. 2024. DOI: <https://doi.org/10.48550/arXiv.2404.17033>
7. Cross-Validation Strategy Impacts the Performance and Interpretation of Machine Learning Models. Sweet L.-B., Müller C., Anand M., Zscheischler J. *Artificial Intelligence for the Earth Systems*. 2023. Vol. 2. P. 1–35. DOI:10.1175/AIES-D-23-0026.1
8. Lesci P., Vlachos A. AnchorAL: Computationally Efficient Active Learning for Large and Imbalanced Datasets. *ArXiv* : website. 2024. DOI: <https://doi.org/10.18653/v1/2024.naacl-long.467>
9. Vincent E. Tikhonov regularization approach to solving inverse problems for parameter learning: master's thesis. African Institute for Mathematical Sciences (AIMS), Cameroon; scientific supervisor: Dr Floriane Melo Kue. Cameroon, 2022. 30 May. 34 p. DOI:10.13140/RG.2.2.21667.43043
10. Warm Start Active Learning with Proxy Labels & Selection via Semi-Supervised Fine-Tuning / Nath V., Yang D., Roth H. R., Xu D. *ArXiv* : website. 2022. DOI: <https://doi.org/10.48550/arXiv.2209.06285>

Дата надходження статті: 10.06.2025

Дата прийняття статті: 17.06.2025

Опубліковано: 23.09.2025