

Мінков К.О.

Чернівецький національний університет імені Юрія Федьковича

ГІБРИДНА АРХІТЕКТУРА ACTIVE LEARNING ТА SEMI-SUPERVISED LEARNING ДЛЯ ЗАДАЧ КЛАСИФІКАЦІЇ З МІНІМАЛЬНИМ МАРКУВАННЯМ

У роботі представлено узагальнену архітектуру активного напівкерованого навчання (ASSL), спрямовану на ефективну класифікацію даних за умов обмежених обсягів розмічених прикладів. Метод поєднує активне навчання, яке вибирає найбільш інформативні зразки для ручної анотації, із напівкерованим навчанням, що дозволяє використовувати значні непозначені вибірки шляхом формування псевдоміток. Початкова частина роботи описує процес ініціалізації моделі невеликою підмножиною вручну позначених даних, після чого виконується послідовний цикл відбору зразків з високою невизначеністю та уточнення параметрів моделі за рахунок псевдоанотованих прикладів. Особливу увагу приділено трьом фундаментальним проблемам комбінування AL (Active Learning – AL) і SSL (Semi-Supervised Learning – SSL): накопиченню помилок у псевдомітках, темпоральній нестабільності прогнозів та неузгодженості даних між слабкими і сильними аугментаціями. Для їх подолання запропоновано інтегрований критерій вибору, що поєднує оцінку невизначеності та показник неузгодженості. Обидві метрики згладжуються за допомогою експоненційного ковзного середнього та доповнюються оцінкою Upper Confidence Bound, що підвищує стійкість вибірки в динамічному SSL-середовищі. Додатково застосовано врахування різноманітності вибірки через кластеризацію ембеддингів методом K-means++, що зменшує ризик надмірної концентрації вибраних зразків у локальних областях простору ознак. Ефективність методології ASSL і варіанта ASSL-Pіз підтверджено на наборах даних CIFAR-10, CIFAR-100, SVHN та MiniImageNet. Порівняння із сучасними підходами активного навчання показало перевагу запропонованих моделей у точності та різке зменшення витрат часу на вибірку, оскільки оцінювання метрик виконується безпосередньо під час SSL-тренування. Метод продемонстрував кращі середні результати як у режимі випадкової ініціалізації, так і за умов успадкування ваг між раундами навчання. Підкреслено потенційну придатність ASSL для задач комп'ютерного зору з високою вартістю маркування, таких як сегментація та детекція об'єктів. Окреслено напрями подальших досліджень, що стосуються ролі зразків із високою неузгодженістю, впливу їх виключення або збереження в непозначеному пулі, а також можливостей розробки нових критеріїв вибірки в межах SSL-процесів.

Ключові слова: невизначеність, псевдомітки, нестабільність, аугментація, темпоральна нестабільність.

Постановка проблеми. Сучасні системи машинного навчання демонструють високу ефективність за умов наявності великих обсягів якісно позначених даних. Проте у більшості практичних сценаріїв отримання таких даних є дорогим, трудомістким або взагалі недоступним. Це особливо актуально для задач комп'ютерного зору, де розмітка зображень потребує залучення експертів і суттєвих часових витрат. У відповідь на цю проблему активно розвиваються підходи, що дають змогу зменшувати залежність моделей від повністю позначених вибірок, зокрема напівкероване навчання та методи активного навчання.

У напівкерованому навчанні модель поєднує невелику частину позначених зразків з великим

масивом непозначених даних, використовуючи внутрішні закономірності вибірки для покращення узагальнюючої здатності. Водночас слабкі сигнали, сформовані на основі непозначених прикладів, мають імовірнісну природу та можуть містити шум, що створює ризик накопичення похибок під час тренування. Активне навчання, у свою чергу, орієнтується на оптимальний вибір зразків для маркування, дозволяючи зменшити витрати на анотацію без погіршення якості моделі. Поєднання цих двох парадигм відкриває перспективи для суттєвого підвищення ефективності тренування, однак водночас призводить до появи додаткових труднощів.

Аналіз останніх досліджень і публікацій. У науковому просторі сьогодення активно дослі-

джуються підходи напівкерованого та активного навчання, зосереджені на зменшенні маркування й підвищенні стабільності моделей. Так, Х. Яном, З. Суном, І. Кінгом та Ц. Сюй [1] виконано систематизований огляд сучасних методів глибинного напівкерованого навчання, зосереджений на класифікації моделей та порівнянні їхніх характеристик. Проаналізовано ключові підходи, зокрема генеративні моделі, методи узгодженості, графові алгоритми, псевдоміткування та гібридні архітектури, з детальним описом їхніх втрат, структур і механізмів роботи. Узагальнено понад півсотні представницьких рішень, наведено порівняння за типами регуляризаторів, особливостями реалізації та сферами застосування. Окреслено обмеження існуючих моделей, включно з чутливістю до припущень про розподіл даних і проблемами стабільності при великій кількості непозначених зразків. Автори визначають відкриті напрями подальших досліджень і пропонують можливі евристичні стратегії для подолання означених проблем.

Метод підсилення структурованих даних малого обсягу, що поєднує генеративний підхід WAGAN та циклічне активне напівкероване навчання SACS (Semi-supervised and Active-learning Cyclic Sampling), запропоновано Ф. Ленгом зі співавторами [2]. Модель WAGAN модифікує базовий GAN (Generative Adversarial Network) за рахунок заміни дивергенції на відстань Васерштейна, введення градієнтного штрафу та реконструкцію міток, що дозволяє уникнути колапсу режимів і підвищити різноманітність синтезованих зразків. Метод SACS вирішує проблему нестабільних та зашумлених псевдо-міток, застосовуючи два класифікатори, критерій невідповідності прогнозів і стратегію вибірки за глобальною та локальною невизначеністю. Проведено експерименти на трьох наборах даних, визначено «масштабний поріг» генерації, за яким приріст інформації стабілізується, та показано перевагу комбінованої схеми над SMOTE, ADASYN (Adaptive Synthetic), CGAN (Conditional Generative Adversarial Network) та іншими методами. Підхід забезпечив підвищення F1max на 11.5–19.6%, демонструючи ефективність у задачах із малими та незбалансованими структурованими вибірками.

У праці Х. Роди та А. Геви [3] представлено напівкерований підхід до активного навчання, що поєднує згортковий автоенкодер, глибоке кластеризування та контрастивне навчання для підвищення точності класифікації за дефіциту міток. Модель попередньо навчає автоенкодер для фор-

мування латентного простору, після чого застосовує кластеризаційний шар і запропоновану CCL (contrastive clustering loss), яка уточнює межі кластерів за рахунок обмеженої кількості позначених зразків. Активне навчання реалізоване як pool-based схема з вибіркою найневизначеніших прикладів за ентропією. Проведено експерименти на MNIST, FashionMNIST та USPS; показано, що додавання CCL суттєво підвищує якість кластерів і забезпечує приріст точності, особливо за 1–5% мічених даних. Метод перевершує VAAL, BALD і Core-Set у режимах із малими вибірками та формує чистіший латентний простір, що підтверджено візуалізаціями t-SNE. Відзначено залежність ефективності від ініціалізації кластерів і складності набору, проте підхід демонструє стійкий приріст у ресурс-обмежених сценаріях.

Цікавим також уявляється метод активного навчання Model Change (MC) для графових моделей напівкерованої класифікації, де вибір точок для маркування ґрунтується на оцінці того, наскільки зміниться класифікатор після гіпотетичного додавання мітки. Цей метод було запропоновано К. Міллером та А. Бертозі [4]. Науковцями запропоновано використання спектрального усічення: замість повного спектра лапласіана графа модель використовує лише невелику кількість найнижчих власних значень і відповідних власних векторів. Це дає змогу суттєво зменшити обчислювальні витрати під час побудови оцінок та моделювання можливих оновлень класифікатора. Розглянуто різні варіанти функцій втрат для бінарної та багатокласової класифікації. Експерименти на синтетичних даних, MNIST та двох гіперспектральних зображеннях показують, що MC стабільно перевершує методи невизначеності, випадкового вибору та популярні спектральні критерії. Метод забезпечує швидке зростання точності та кращі фінальні результати, зберігаючи помірні вимоги до ресурсів і придатність для задач із дуже малою кількістю міток.

Не дивлячись на активний розвиток методів SSL та AL, проблема впливу нестабільних псевдо-міток на надійність активного відбору станом на теперішній час є недостатньо розробленою та потребує подальших досліджень.

Постановка завдання. Метою статті є дослідження впливу слабких навчальних сигналів у задачах активного напівкерованого навчання та визначення причин зниження ефективності класичних методів вибірки за умов використання псевдо-міток. Завданням є формалізація ролі нестабільних прогнозів, аналіз темпоральної мін-

ливості псевдо-міток, оцінювання неузгодженості між слабкими та сильними аугментаціями, а також побудова критеріїв вибору зразків, стійких до цих факторів. Робота спрямована на встановлення обмежень і можливостей застосування активного навчання у середовищах зі слабким наглядом.

Виклад основного матеріалу. У цій роботі запропоновано гібридну архітектуру, що поєднує два підходи машинного навчання – активне навчання та напівкероване навчання – для ефективної класифікації даних за умов обмеженої кількості розмічених прикладів.

Основна мета полягає у мінімізації витрат на маркування при збереженні високої точності моделі. Це досягається за рахунок ітераційного процесу, у якому модель послідовно вибирає найінформативніші зразки для ручного маркування, а решту нерозмічених даних використовує для автоматичного генерування псевдоміток.

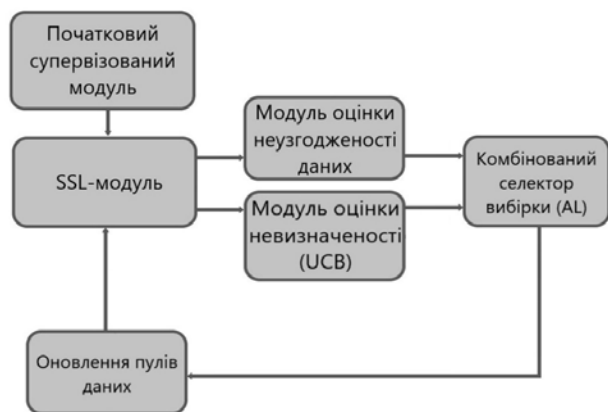


Рис. 1. Структурна схема запропонованої архітектури активного напівкерованого навчання

У задачах класифікації за дефіциту позначених даних напівкероване навчання застосовується як спосіб підвищення точності моделі за рахунок поєднання невеликої позначеної вибірки L та значного непозначеного набору U . У цьому випадку оптимізація виконується через сумарну функцію втрат, що включає supervised loss (функцію втрат під наглядом) для зразків із істинними мітками та unsupervised loss (функцію втрат без нагляду) для непозначених даних.

Псевдо-мітки формуються на основі прогнозу мережі для непозначеного зразка за правилом

$$\hat{y} = \arg \max_y p(y|x), \quad \text{якщо } \max_y p(y|x) \geq \tau \quad (1)$$

після чого використовуються в unsupervised loss як заміна справжніх класових міток.

Поєднання SSL із методами активного навчання створює низку характерних проблем. По-перше,

помилки у псевдо-мітках призводять до формування хибних керівних сигналів (supervisory signals), що підсилює ефект накопичення шуму в процесі тренування. По-друге, явище темпоральна нестабільність (temporal instability) зумовлює зміну прогнозів для одного й того самого зразка на різних кроках оптимізації, через що оцінка невизначеності стає нестійкою. По-третє, неузгодженість даних (НД – data inconsistency), що виникає через розбіжності між слабкими та сильними аугментаціями даних одного зразка, породжує суперечливі сигнали в unsupervised loss і ускладнює формування узгоджених псевдо-міток.

Перед запуском ітерацій активного напівкерованого навчання формується початкова підмножина розмічених даних L_0 , яка використовується для первинного тренування моделі. Обсяг цієї вибірки зазвичай становить 5–10% від загального пулу доступних даних, проте вона має бути достатньо репрезентативною для формування базових меж класифікації. На цьому етапі модель навчається у повністю супервізованому режимі, що дає змогу ініціалізувати ваги та задати початкові гіперпараметри, які використовуються як відправна точка для подальших циклів ASSL.

Після первинного навчання модель використовується для передбачення класів усіх нерозмічених зразків. На основі метрик визначаються зразки, щодо яких модель демонструє найменшу стабільність прогнозу.

Такі об'єкти формують «чергу на маркування» і передаються експерту для ручної анотації. У такий спосіб процес маркування концентрується на даних, що максимально збільшують інформаційну цінність навчальної вибірки.

Формально це відповідає байєсівській логіці оновлення знань, де кожна нова мітка з найвищою невизначеністю забезпечує найбільший приріст інформації для моделі.

Після кожного циклу активного навчання залишається велика кількість нерозмічених зразків. Для їх використання застосовується модуль напівкерованого навчання, який формує псевдо-мітки – тобто передбачені моделлю класи, що мають високу ймовірність (> 0.9).

Псевдомічені дані додаються до навчальної вибірки, створюючи розширений набір даних, який містить як реальні, так і псевдоанотовані приклади. Це дозволяє моделі поступово уточнювати свої параметри без додаткового ручного втручання.

Після цього формулюється постановка задачі ASSL для багатокласової класифікації. На ітерації

активного навчання n використовується позначений набір даних L_n та значний непозначений набір U_n , які слугують відповідно позначеними й непозначеними вибірками для етапу тренування SSL. Після виконання навчання SSL застосовується функція вибору, що визначає K зразків $L'_n \subset U_n$, які необхідно промаркувати. Оновлений позначений набір для раунду $n+1$ задається як $L_{n+1} = L_n \cup L'_n$, тоді як непозначена вибірка оновлюється за правилом $U_{n+1} = U_n \setminus L'_n$.

Нехай множина класів визначається як $C = \{1, \dots, k\}$, а x є вхідним зображенням. Для кожного $y \in C$ вихід мережі після застосування функції softmax подається у вигляді $p(y|x; \theta) = \text{SOFTMAX}(f(x; \theta))$. Вектор ймовірностей $p(x) = [p(y=1|x; \theta), \dots, p(y=k|x; \theta)]^T$ використовується для визначення прогнозованого класу $\hat{y} = \arg \max(p(x))$, а також one-hot подання $1_{\hat{y}} \in \mathbb{R}^k$.

Для непозначених даних замість істинної мітки використовується прогнозована $\hat{y}$. Значення невизначеності $u(x)$ для зображення x вводиться як L2-розбіжність між прогнозом моделі та псевдоміткою оцінка між $p(x)$ та one-hot $\hat{y}$:

$$u(x) = \|p(x) - 1_{\hat{y}}\|_2. \quad (2)$$

Однак у контексті SSL-середовища невизначеність стає нестабільною через темпоральні коливання прогнозів, тому пряме застосування uncertainty-based AL-критеріїв призводить до погіршення якості відбору. Це обґрунтовує необхідність подальшої модифікації AL-компоненти в рамках ASSL. Для аналізу цього ефекту було відстежено зміни прогнозів на непозначених даних під час SSL-тренування, використовуючи 100 випадково позначених зразків. Невизначеність $u(x)$, визначену рівнянням (2), обчислювали для множини U кожні 5 тисяч кроків, оцінюючи стабільність її значень у часі.

Під час SSL-тренування такі коливання прогнозів зберігаються, тому подальший аналіз виконується через формальний показник темпоральної нестабільності. Для кількісної оцінки аналізувалися зміни прогнозованих міток між сусідніми моментами часу. Нехай $\hat{y}_t$ – прогноз для зразка x на кроці t . Тоді показник temporal instability задається як:

$$TI_T(x) = \sum_{t=1}^T 1(\hat{y}_t(x) \neq \hat{y}_{t-1}(x)), \quad (3)$$

де $1(\cdot)$ – індикаторна функція.

Зразки з високою невизначеністю зазвичай мають і більшу темпоральну нестабільність, що узгоджується з формальним показником $TI(x)$. Отже, на відміну від активного навчання без

участі SSL, методи відбору на основі невизначеності в ASSL втрачають ефективність. У певний момент часу невизначеність зразка не дає надійної оцінки його інформативності, що погіршує якість вибірки.

Для підвищення точності оцінки невизначеності деякі підходи, зокрема Consistency AL [5], BALD [6] та ensemble-based AL [7], усереднюють результати кількох проходів інференсу. Також використовувано EMA (exponential moving averages – експоненційне ковзне середнє) та підхід UCB (Upper Confidence Bound) [8], що походить із підкріплювального навчання, для зменшення темпоральної нестабільності та підвищення ефективності вибірки в умовах динамічного оновлення моделі.

У рамках SSL-тренування ми отримуємо прогнози непозначених зразків у різні моменти часу, тому невизначеність мінібатчів (mini-batch) обчислювалася на кожному кроці, після чого для неї визначався показник UCB на основі EMA.

Невизначеність непозначеного зразка x у момент часу t оцінювалася за слабко аугментованим зображенням $p(x_w)$. Формули EMA та UCB мають вигляд:

$$\bar{u}_t(x) = \alpha u_t(x_w) + (1 - \alpha) \bar{u}_{t-1},$$

$$\bar{v}_t''(x) = \alpha (u_t(x_w) - \bar{u}_t(x))^2 + (1 - \alpha) \bar{v}_{t-1}'', \quad (4)$$

$$u_t^{\text{UCB}}(x) = \bar{u}_t(x) + c \sqrt{\bar{v}_t''(x)}$$

де $\bar{u}_t(x)$ та $\bar{v}_t''(x)$ є EMA та експоненційною ковзною оцінкою дисперсії відповідно; початкові значення задаються як $\bar{u}_0(x) = 0$ та $\bar{v}_0''(x) = 0$. Параметри α та c визначають інтенсивність згладжування та коефіцієнт UCB.

Окрім EMA та UCB, оцінка невизначеності не обмежується одним часовим вимірюванням. Підхід розширюється так, щоб забезпечити стабільнішу вибірку протягом усього SSL-тренування та зменшити вплив коливань між окремими кроками.

У ASSL вибірка зразків із високим значенням UCB-невизначеності може призводити до підвищеної supervised loss у наступному раунді. У FixMatch відбір, що ґрунтується лише на supervised loss, є обмеженим через наявність unsupervised loss, який формується consistency regularization (регуляризація узгодженості). Це зумовлює потребу врахування впливу непозначених даних саме на unsupervised loss.

Для оцінювання НД використовується дивергенція Кульбака-Лейблера (KL) між слабкою x_w та сильною x_s аугментаціями одного й того самого непозначеного зразка:

$$i(x) = \frac{\text{KL}(p(x_w)p(x_s)) + \text{KL}(p(x_s)p(x_w))}{2}. \quad (5)$$

Зміни цього показника згладжуються за допомогою EMA:

$$\bar{i}_t(x) = \alpha i_t(x) + (1 - \alpha) \bar{i}_{t-1}(x). \quad (6)$$

Зразки, для яких значення $\max(p(x)) > \tau (= 0.95)$, розглядаються як pseudo-labeled (псевдо-розмічені дані), оскільки вони мають високу модельну впевненість і беруть участь в unsupervised loss. Підрахунок кількості таких випадків на кожному 5,000 кроці дає змогу оцінити частоту появи висококонфідентних зразків. Порівняння часток pseudo-labeled з високою НД та з високою невизначеністю показує, що частка першої групи є значно більшою, наведено у Таблиця 1.

Для визначення остаточного показника НД застосовується UCB:

$$\bar{v}_t^i(x) = \alpha (i_t(x) - \bar{i}_t)^2 + (1 - \alpha) \bar{v}_{t-1}^i, \quad (7)$$

$$i_t^{\text{UCB}}(x) = \bar{i}_t(x) + c \sqrt{\bar{v}_t^i(x)}. \quad (8)$$

Фінальний AL-показник формує комбіновану оцінку, яка враховує як невизначеність, так і неповноту даних:

$$\text{Score}(x) = u^{\text{UCB}}(x) \times i^{\text{UCB}}(x). \quad (9)$$

Для оцінювання методу ASSL використовувалися набори даних CIFAR – 10, CIFAR – 100 [9], SVHN [10] та MiniImageNet [11]. Для CIFAR – 10 і SVHN початковий позначений набір складався зі 100 випадково вибраних зображень, після чого на кожному раунді додатково маркувалися ще 100 зразків. Для CIFAR – 100 та MiniImageNet початкові позначені вибірки містили 1 000 зображень, і в кожному раунді додатково маркувалося 1 000 нових зразків.

Експерименти виконувалися протягом 10 раундів активного навчання. Розглядалися два варіанти ініціалізації ваг: випадкові початкові ваги (RandInit) та ваги, успадковані з попереднього раунду (ConInit). У випадку RandInit на всіх раундах використовувалася однакова початкова модель із фіксованою випадковою ініціалізацією.

Як базові моделі застосовувалися ResNet-18 та ViT-small з конфігурацією шарів, аналогічною тій,

що використовується в ALFA-Mix. На кожному раунді виконувалося 1 024 кроки тренування протягом 5 епох, що приблизно становить 1% від повної кількості кроків, зазвичай застосовуваних у масштабному SSL-тренуванні. Швидкість навчання встановлювалася $lr = 0.03$ для ResNet – 18 та $lr = 0.01$ для ViT-small. Планувальник швидкості навчання не використовувався через невелику кількість тренувальних кроків. Інші гіперпараметри SSL відповідали конфігурації FixMatch, а для аугментації непозначених даних застосовувався RandAugment.

ASSL-середовище було реалізовано на основі ALFA-Mix та доступної імплементації FixMatch у PyTorch. Для функції вибірки гіперпараметри встановлювалися таким чином: коефіцієнт EMA $\alpha = 0.8$, порогове значення для невизначеності $c_u = 0.5$, порогове значення для НД $c_i = 2.0$. Також використовувалася проста міра різноманіття: підсумковий показник кожного непозначеного зразка додатково множився на ознаку різноманітності, отриману за допомогою виділення ембеддинг-ознак (feature embedding) і кластеризації методом K-means++. Базовий варіант методу позначено як ASSL, а варіант із різноманіттям – як ASSL-Різ.

Результати експериментів подаються як середні значення за трьома незалежними запусками з різними початковими випадковими ініціалізаціями.

Для оцінювання запропонованого методу ASSL результати порівнювалися з випадковою вибіркою та низкою сучасних підходів активного навчання, включно з методами, такими як entropy [12], BADGE [13], CDAL [14], CoreSet [15] та ALFA-Mix [16]. Порівняння виконувалося на основі ResNet – 18 для чотирьох наборів даних, а також на основі ViT-small для CIFAR – 10 і SVHN з метою перевірки узгодженості роботи під різними архітектурами. Результати наведені у таблиці 2.

Для кількісної оцінки переваг методів використовувалася попарна матриця порівняння, де кожен елемент $e_{i,j}$ відображає кількість випадків, коли метод i перевищував за якістю метод j . Додатково обчислювався середній бал по кожному стовпцю, що дає узагальнену оцінку відносної продуктивності. Результати наведені у таблиці 3.

Таблиця 1

Порівняння частки псевдо-розмічених зразків з високою неузгодженістю даних та високою невизначеністю

Впорядковано за	Неузгодженість даних			Невизначеність		
	Перші 1%	Перші 5%	Перші 10%	Перші 1%	Перші 5%	Перші 10%
частка	54.93%	47.46%	44.20%	15.67%	24.87%	28.14%

Таблиця 2

**Порівняння точності та часу вибірки для різних методів активного
та активного напівкерованого навчання**

Методи	Точність (%)				Час (с)			
	RandInit (AL)	RandInit	ConInit	Всього	CIFAR10	SVHN	CIFAR100	MiniImageNet
Random	27.91±0.85	37.31±1.35	42.46±0.91	39.94±1.10	-	-	-	-
Entropy	27.28±1.12	30.73±1.68	38.38±1.40	34.57±1.54	25.30	30.38	25.38	51.10
CoreSet	27.28±1.07	36.55±1.73	42.62±1.27	39.56±1.50	37.16	43.23	126.13	209.54
BADGE	28.24±1.01	37.11±1.46	42.01±1.60	39.55±1.55	201.23	243.65	>3600	>3600
CDAL	28.38±0.99	37.04±1.38	42.65±1.04	39.80±1.25	27.375	33.04	64.54	134.29
ALFA-Mix	28.36±1.03	37.57±1.44	42.43±0.92	40.06±1.17	81.53	102.54	495.65	998.54
ASSL	-	37.07±1.44	42.83±0.80	39.95±1.13	0.05	0.05	0.03	0.04
ASSL-Pіз	-	37.68±1.49	42.64±1.02	40.12±1.27	37.23	42.12	124.23	205.70

Таблиця 3

Попарне порівняння методів активного навчання та ASSL за результатами експериментів

-	Random	Entropy	CoreSet	BADGE	CDAL	ALFA-Mix	ASSL	ASSL-Pіз
Random	-	11.0	3.5	2.8	2.0	2.2	3.2	2.7
Entropy	0.0	-	0.0	0.0	0.0	0.0	0.0	0.0
CoreSet	2.8	10.9	-	3.2	2.6	2.7	1.6	1.9
BADGE	3.2	11.2	3.2	-	1.9	2.1	2.6	1.7
CDAL	3.4	11.3	3.6	2.9	-	2.4	2.7	2.1
ALFA-Mix	3.8	11.4	3.4	2.8	2.2	-	2.8	1.6
ASSL	3.6	11.3	2.1	3.1	3.2	2.5	-	1.7
ASSL-Pіз	3.2	11.2	3.0	3.3	2.8	2.5	3.0	-
Середнє	2.86	11.18	2.68	2.58	2.1	2.06	2.27	1.68

Ліва частина таблиці 2 містить середню тестову точність за всіма раундами й наборами даних для двох варіантів ініціалізації ваг. Варіант ASSL-Pіз демонструє найвищі середні значення в умовах RandInit та у загальному підсумку, тоді як ASSL забезпечує найкращу середню точність у випадку ConInit. Для зіставлення якості активного навчання і ASSL застосовувався окремий AL-фреймворк із такими самими гіперпараметрами, що й у відповідних реалізаціях AL. Порівняння середніх точностей у режимі RandInit показує, що ASSL забезпечує приріст наближено 8% відносно стандартних підходів активного навчання.

Щодо часу вибірки, права частина таблиці 2 показує, що метод ASSL формує оцінку безпосередньо під час тренування, тому додаткові витрати на етапі вибірки практично відсутні. Крім того, подібно до CoreSet, варіант ASSL-Pіз забезпечує швидший процес вибірки, ніж BADGE або ALFA-Mix, для яких витрати часу залежать від кількості класів.

Результати попарного порівняння, наведені на таблиці 3, показують, що ASSL-Pіз має найменшу

середню кількість програвів і загалом перевершує інші методи за сукупністю експериментів.

У наведених експериментах кількість позначених зразків відповідала обсягу, прийнятому в активному навчанні, що забезпечує коректність порівняння. Проте за умови поєднання етапів SSL можливим є отримання вищої точності за нижчої вартості маркування. Водночас збільшення кількості позначених даних може зменшити витрати тренування для досягнення тієї ж точності. Це формує компроміс між збільшенням вартості маркування та скороченням часу тренування. Оптимальний баланс між цими складовими залежить від конкретної задачі та може підвищувати ефективність методів ASSL.

ASSL може бути адаптований до більш широкого кола задач комп'ютерного зору, включно з детекцією об'єктів, сегментацією зображень та оцінюванням пози людини. Такі задачі зазвичай мають значно вищу вартість маркування, ніж класифікація. Використання підходів ASSL у цих сценаріях потенційно дає змогу прискорювати оновлення моделей при прийнятній вартості маркування, аналогічно результатам, отриманим у класифікаційних експериментах.

У межах аналізу непозначених даних у контексті SSL було визначено групу зразків із високою НД, які формують нестабільні псевдо-сигнали та ускладнюють навчання. Використання цього показника на етапі вибірки приводить до покращення результатів. Однак вплив таких зразків на SSL загалом залишається нетривіальним.

Можливими напрямками подальшого вивчення є:

- оцінка ефективності виключення зразків з високою НД під час тренування;
- вилучення таких зразків з непозначеного пулу на етапі вибірки;
- аналіз потенційної користі зразків із низькою НД.

Дослідження цих аспектів може дати змогу сформувати нові методології в ASSL та SSL.

Висновки. У роботі обґрунтовано гібридну архітектуру активного напівкеруваного навчання

(ASSL), що поєднує механізми вибірки зразків за невизначеністю та використання псевдоміток у межах SSL. Проведений аналіз показав, що темпоральна нестабільність прогнозів та неузгодженість даних суттєво знижують ефективність традиційних критеріїв активного навчання. Запропоноване використання ЕМА та UCB дає змогу стабілізувати оцінювання невизначеності та підвищити якість вибірки. Експерименти на наборах CIFAR-10, CIFAR-100, SVHN та MiniImageNet засвідчили переваги ASSL і ASSL-Піз у порівнянні з сучасними AL-методами як за точністю, так і за часовими витратами. Отримані результати підтверджують доцільність поєднання етапів SSL і AL для зменшення вартості маркування та підвищення продуктивності моделі, а також вказують на перспективність подальшого вивчення впливу зразків із високою НД на динаміку навчання.

Список літератури:

1. Yang X., Song Z., King I., Xu Z. A Survey on Deep Semi-supervised Learning. *IEEE Transactions on Knowledge and Data Engineering*. 2023. Vol. 35. No. 9. P. 8934–8954. DOI: 10.1109/TKDE.2022.3220219
2. Leng F., Li F., Lv W., Bao Y., Liu X., Zhang T., Yu G. A Semi-Supervised Active Learning Method for Structured Data Enhancement with Small Samples. *Mathematics*. 2024. Vol. 12. No. 17. P. 2634. DOI: 10.3390/math12172634
3. Roda H., Geva A. B. Semi-supervised active learning using convolutional auto-encoder and contrastive learning. *Frontiers in Artificial Intelligence*. 2024. Vol. 7. P. 1398844. DOI: 10.3389/frai.2024.1398844
4. Miller K. S., Bertozzi A. L. Model Change Active Learning in Graph-Based Semi-supervised Learning. *Communications on Applied Mathematics and Computation*. 2024. Vol. 6. P. 1270–1298. DOI: 10.1007/s42967-023-00328-z
5. Gao M., Zhang Z., Yu G., Arik S. O., Davis L. S., Pfister T. Consistency-based Semi-supervised Active Learning: Towards Minimizing Labeling Cost. *arXiv: website*. 2020. DOI: 10.48550/arXiv.1910.07153
6. Gal Y., Islam R., Ghahramani Z. Deep Bayesian Active Learning with Image Data. *arXiv: website*. 2017. DOI: 10.48550/arXiv.1703.02910
7. Beluch W. H., Genewein T., Nürnberger A., Köhler J. M., et al. The Power of Ensembles for Active Learning in Image Classification. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*. 2018. Vol. 2018. P. 9368–9377. DOI: 10.1109/CVPR.2018.00976
8. Sutton R. S., Barto A. G. Reinforcement Learning: An Introduction. 2nd ed. Cambridge (MA) – London (EN): The MIT Press, 2018. 552 p.
9. Krizhevsky A. Learning Multiple Layers of Features from Tiny Images. Technical Report. University of Toronto, 2009. 60 p. URL: <https://www.cs.toronto.edu/~kriz/learning-features-2009-TR.pdf> (last accessed: 11.11.2025).
10. Netzer Y., Wang T., Coates A., Bissacco A., Wu B., Ng A. Y. Reading Digits in Natural Images with Unsupervised Feature Learning. *NIPS Workshop on Deep Learning and Unsupervised Feature Learning*. 2011. URL: <https://www.semanticscholar.org/paper/Reading-Digits-in-Natural-Images-with-Unsupervised-Netzer-Wang/02227c94dd41fe0b439e050d377b0beb5d427cda> (last accessed: 12.11.2025).
11. Ravi S., Larochelle H. Optimization as a Model for Few-Shot Learning. *Published as a conference paper at ICLR 2017*. URL: <https://openreview.net/forum?id=rJY0-Kcll> (last accessed: 12.11.2025).
12. Wang D., Shang Y. A New Active Labeling Method for Deep Learning. *2014 International Joint Conference on Neural Networks (IJCNN)*. 2014. P. 112–119. DOI: 10.1109/IJCNN.2014.6889457
13. Ash J. T., Zhang C., Krishnamurthy A., Langford J., Agarwal A. Deep Batch Active Learning by Diverse, Uncertain Gradient Lower Bounds. *arXiv: website*. 2020. DOI: 10.48550/arXiv.1906.03671
14. Agarwal S., Arora H., Anand S., Arora C. Contextual Diversity for Active Learning. *arXiv: website*. 2020. DOI: 10.48550/arXiv.2008.05723
15. Sener O., Savarese S. Active Learning for Convolutional Neural Networks: A Core-Set Approach. *arXiv: website*. 2018. DOI: 10.48550/arXiv.1708.00489
16. Parvaneh A., Abbasnejad E., Teney D., Haffari R., van den Hengel A., Shi J. Q. Active Learning by Feature Mixing. *arXiv: website*. 2022. DOI: 10.48550/arXiv.2203.07034

Minkov K.O. HYBRID ARCHITECTURE OF ACTIVE LEARNING AND SEMI-SUPERVISED LEARNING FOR CLASSIFICATION TASKS WITH MINIMAL LABELING

The study presents a generalized architecture of Active Semi-Supervised Learning (ASSL) aimed at efficient data classification under limited amounts of labeled examples. The method combines active learning, which selects the most informative samples for manual annotation, with semi-supervised learning, which enables the use of large unlabeled datasets through pseudo-label generation. The initial stage describes the model initialization with a small subset of manually labeled data, followed by an iterative process of selecting high-uncertainty samples and refining model parameters using pseudo-annotated examples. Particular attention is given to three fundamental challenges of combining Active Learning (AL) and Semi-Supervised Learning (SSL): the accumulation of errors in pseudo-labels, temporal instability of predictions, and data inconsistency between weak and strong augmentations of the same sample. To address these issues, an integrated selection criterion is proposed, combining uncertainty estimation with a measure of inconsistency. Both metrics are smoothed using an exponential moving average and enhanced with an Upper Confidence Bound, which improves sampling stability in a dynamic SSL setting. Additionally, sample diversity is incorporated through feature-space clustering using the K-means++ algorithm, reducing the risk of over-concentration of selected samples in local regions of the representation space. The effectiveness of the ASSL methodology and its variant ASSL-Div is demonstrated on the CIFAR-10, CIFAR-100, SVHN, and MiniImageNet datasets. Comparison with modern active learning approaches shows the advantages of the proposed models in terms of accuracy and a substantial reduction in sampling time, as metric evaluation is performed directly during the SSL training process. The method achieved higher average performance both under random initialization and when reusing model weights across iterations. The study highlights the potential applicability of ASSL to computer-vision tasks with high annotation costs, such as image segmentation and object detection. Future research directions include examining the role of high-inconsistency samples, evaluating the effects of excluding or retaining such data in the unlabeled pool, and exploring new sample selection criteria within SSL procedures.

Key words: uncertainty, pseudo-labels, instability, augmentation, temporal instability.

Дата надходження статті: 15.11.2025

Дата прийняття статті: 04.12.2025

Опубліковано: 30.12.2025