

DOI: 10.20535/kpissn.2025.2.320063

УДК 004.85.032.26:004.93'1:616.5-006](045)

В.Я. Данилов, О.О. Зарицький*

КПІ ім. Ігоря Сікорського, Київ, Україна

*Відповідальний автор: zaritskiy.alexey@gmail.com

МОДЕЛЬ КЛАСИФІКАЦІЇ РАКОВИХ ЗАХВОРЮВАНЬ ШКІРИ НА ОСНОВІ ІНТЕГРАЦІЇ ДИСКРИМІНАТОРА В АРХІТЕКТУРУ СХОДОВИХ НЕЙРОННИХ МЕРЕЖ

Проблематика. Напівкероване навчання (НН) є одним із перспективних напрямів глибокого навчання, особливо для задач, де отримання міток є складним і дорогим процесом, як у випадку класифікації ракових захворювань шкіри. Наявні підходи, зокрема Adversarial Autoencoders (AAE) та Ladder Networks (LN), ефективно використовують нерозмічені дані, але мають обмеження у точності реконструкції та регуляризації.

Мета дослідження. Розробка та дослідження моделі класифікації ракових захворювань шкіри на основі інтеграції дискримінатора в архітектуру сходових нейронних мереж.

Методика реалізації. Розроблена модель поєднує регуляризуючі властивості сходових нейронних мереж із використанням функції втрат на основі дискримінатора, що оцінює якість реконструкції зображень, орієнтовану на їх структуру, форму та ключові візуальні ознаки. Експерименти проводилися на датасеті HAM10000 з різними співвідношеннями розмічених і нерозмічених даних (30 %, 10 %, 5 %).

Результати дослідження. Експерименти показали, що запропонована модель підвищила F_1 -score для класу злоякісних утворень на 4 % порівняно з базовою сходовою мережею за умов 5 % розмічених даних. За 30 % маркованих зразків F_1 -score досяг 74,8 %, що всього на 1 % менше за повністю керовану модель. Відносний показник $R_{\hat{f}}$ за 5 % розмічених даних становив 0,939, перевищуючи аналогічний коефіцієнт для STFL (0,877), що підтверджує ефективність використання нерозмічених даних у запропонованій моделі.

Висновки. Запропонована модель вдосконалює наявні методи напівкерованого навчання (LN, AAE), забезпечуючи високу ефективність використання нерозмічених даних для регуляризації латентного простору енкодера у задачі класифікації ракових захворювань шкіри. Перспективи подальших досліджень включають використання дискримінатора для порівняння латентних представлень і вдосконалення функції Reconstruction Cost для розширення її застосування в інших задачах аналізу зображень.

Ключові слова: класифікація ракових захворювань шкіри; HAM10000; напівкероване навчання; сходові нейронні мережі; дискримінатор.

Вступ

Напівкероване навчання є одним із найбільш актуальних напрямів дослідження глибокого навчання. Особливо корисним воно є для задач, де отримання міток є дуже складним або дорогим процесом, як, наприклад, у сфері медичних даних. У контексті задач розпізнавання та класифікації ракових захворювань шкіри, що є однією з найпоширеніших і небезпечних форм онкології, напівкероване навчання дає можливість ефективно використовувати велику кількість нерозмічених даних для навчання моделі. Складність отримання розмічених даних

для навчання моделі у класифікації ракових захворювань шкіри також проявляється у тому, що для точного діагнозу недостатньо візуального вивчення утворення, а необхідно проводити складні дослідження із забором тканин утворення. Таким чином, дослідження методів, що можуть використовувати дуже малу кількість розмічених даних у поєднанні з великою кількістю нерозмічених, є перспективним напрямом у цій сфері.

Автоенкодери є поширеним інструментом у напівкерованому навчанні, оскільки їх структура дозволяє використовувати нерозмічені дані для регуляризації латентного простору енкоде-

Пропозиція для цитування цієї статті: В.Я. Данилов, О.О. Зарицький, “Модель класифікації ракових захворювань шкіри на основі інтеграції дискримінатора в архітектуру сходових нейронних мереж”, *Наукові вісті КПІ*, № 2, с. 30–38, 2025. doi: 10.20535/kpissn.2025.2.320063

Offer a citation for this article: V.Y. Danilov, O.O. Zarytskyi, “A skin cancer classification model incorporating a discriminator into the ladder neural network architecture”, *KPI Science News*, no. 2, pp. 30–38, 2025. doi: 10.20535/kpissn.2025.2.320063

ра, таким чином покращуючи узагальнюючу здатність моделі. Серед найбільш інноваційних підходів до використання автоенкодерів у напівкерованому навчанні є ААЕ (змагальні автоенкодери), що використовують дискримінатор для покращення регуляризації через порівняння латентного представлення з певним апіорним заданим розподілом, та LN (сходові нейронні мережі), що використовують більш складну знешумлювальну структуру у вигляді «драбини» і дозволяють зберігати важливу інформацію на різних рівнях абстракції. Хоча LN і пропонують регуляризацію за рахунок драбинної структури та додавання шуму до вхідних шарів декодера, для порівняння латентних шарів і, відповідно, обрахунку reconstruction cost здебільшого використовують прості методи типу MSE або інші аналогічні грубі методи порівняння.

У цій статті пропонується об'єднати регуляризуючу здатність обох цих підходів, використавши дискримінатори для визначення reconstruction cost у порівнянні латентних представлень енкодера та відповідних реконструйованих представлень декодера у сходовій нейронній мережі.

Для експериментів було вибрано відомий датасет HAM10000, що містить вибірку ракових захворювань шкіри різних класів, а також багато прикладів доброякісних утворень, таких як родимки. Вибір датасету пояснюється тим, що автори статей, які є основою цього дослідження, демонстрували результати експериментів тільки на «неприкладних» датасетах, таких як MNIST, де припущення напівкерованого навчання виконуються занадто очевидно й ефективність цих підходів для задачі класифікації складніших даних (а саме фотографій ракових захворювань шкіри) є недослідженою.

Постановка задачі

Метою роботи є дослідження і розробка моделі до напівкерованого навчання для задачі класифікації ракових захворювань шкіри за допомогою поєднання регуляризаційних властивостей сходових нейронних мереж із використанням дискримінатора для оцінювання якості реконструкції.

Огляд наявних досліджень

Нині у відкритому доступі не виявлено публікацій, у яких сходові нейронні мережі або інші denoising-autoencoder архітектури порівнювали саме на датасеті HAM10000. Єдиним винятком є робота [1], результати якої і будуть використані для прямого порівняння з новою моделлю.

Поза межами тематики дослідження є багато методів НН, що демонструють результати на HAM10000 (STFL [2], MixMatch [3], Mean-Teacher [4]), тому для зовнішнього порівняння було обрано одну з таких архітектур НН — Self-feedback Threshold Focal Learning (STFL), що показала стабільні результати за невеликої кількості розмічених даних (500 розмічених зразків) і використовує ResNet-50 як базову мережу. Головна ідея цієї моделі ґрунтується на автоматичному коригуванні порогів впевненості псевдоміток, а також використанні focal loss для боротьби з дисбалансом вибірки. STFL було вибрано для порівняння тому, що:

1. Наскільки нам відомо, ця стаття є найостаннішою публікацією з методів НН для класифікації ракових захворювань шкіри.
2. Стаття містить відкрито опубліковані метрики F_1 , асигнасу тощо саме на датасеті HAM10000, що дає змогу їх порівняти з моделлю, запропонованою у цій роботі. Подальші деталі про те, як будуть зіставлятися метрики, подано в розділі «Результати експериментів».

Таким чином, далі розглянемо теоретичні засади двох ліній досліджень — ААЕ та LN — щоб показати, як саме запропонована модель поєднує їхні регуляризуючі здібності та розширює можливості НН для задач класифікації ракових захворювань шкіри.

Adversarial Autoencoders

Adversarial Autoencoders [5] — це ймовірнісні автоенкодери, що використовують генеративно змагальну мережу [6] для накладання заданого апіорного розподілу на латентний простір автоенкодера. Нехай виходом енкодера в автоенкодерній архітектурі є деякий апостеріорний розподіл $q(z)$. Автори статті [5] визначають його таким чином:

$$q(z) = \int_x q(z|x) p_{data}(x) dx,$$

де $q(z|x)$ — це латентний розподіл, що генерує енкодер; $p_{data}(x)$ — це розподіл вхідних даних.

Формально навчання змагального автоенкодера можна описати у два етапи:

1. Фаза реконструкції, де автоенкодер мінімізує помилку реконструкції відновленого зображення відносно оригіналу.

2. Фаза регуляризації, де дискримінатор порівнює апостеріорний розподіл $q(z)$ з деяким

априорним розподілом $p_{prior}(z)$ і намагається визначити, який із них є оригіналом, а який є результатом енкодера.

Архітектуру ААЕ зображено на рис. 1. У цій взаємодії роль генератора виконує енкодер, а дискримінатор додається окремо для порівняння априорного та апостеріорного розподілу. Оцінювання дискримінатора у цій системі дає змогу енкодеру (тобто генератору) виділити ознаки таким чином, щоб апостеріорний розподіл був дуже схожим на априорний, з яким його порівнює дискримінатор. Автори показують, що такий спосіб регуляризації може суттєво покращити здатність моделей виділяти важливі ознаки, ефективно «підказуючи» їм, як приблизно повинен виглядати цільовий розподіл.

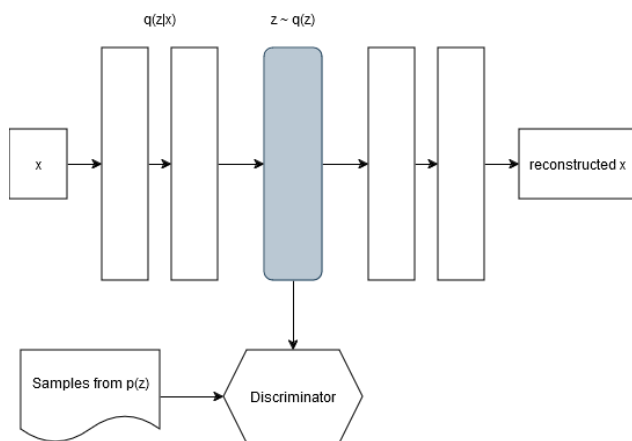


Рис. 1. Архітектура змагального автоенкодера. Схему створено за описом у [3]

Хоча заздалегідь невідомий точний розподіл даних, проте можна припустити, який це тип розподілу, і підказати за допомогою дискримінатора, як краще енкодеру виділяти ознаки, щоб апостеріорний розподіл відповідав гіпотезі.

Adversarial Autoencoders in semi-supervised learning

Якщо у змагальних автоенкодерах використати апостеріорний розподіл як вхідний вектор для повнозв'язного шару, то отримаємо модель напівкерованого навчання для класифікації, де розмічені дані використовуються для керованої частини, а reconstruction cost та regularization cost — для некерованої частини. Такий підхід був запропонований у статті про використання ААЕ для напівкерованого навчання [7], що розвиває ідеї, описані у [8] та [5].

Автори пропонують таку функцію втрат:

$$L = \lambda_{NLL} L_{NLL} + \lambda_{rec} L_{rec} + \lambda_{reg} L_{reg},$$

де L_{NLL} — втрати класифікації між міткою моделі та справжньою міткою; L_{rec} — вартість реконструкції між реконструйованим зображенням та оригіналом; L_{reg} — втрати регуляризації, які визначає дискримінатор.

Сходові нейронні мережі

Сходові нейронні мережі (LN) [9], [10], [11] — це сучасний архітектурний підхід до напівкерованого навчання, що ґрунтується на ідеї знешумлюючих (denoising) автоенкодерів. У загальному розумінні структура сходової нейронної мережі для використання у напівкерованому навчанні нагадує будову автоенкодера [8], проте на вході у декодер латентне представлення зашумлюється і завдання декодера полягає у відновленні оригінального зображення.

Сходова нейронна мережа складається з трьох основних компонентів:

1. Зашумлений енкодер (noisy encoder), що генерує латентні представлення за допомогою додавання гаусівського шуму.
2. Чистий енкодер створює еталонні латентні представлення без шуму.
3. Декодер, що реконструює латентні представлення на кожному кроці енкодера із зашумлених версій цих представлень.

На кожному шарі енкодера та декодера використовується батч-нормалізація (batch normalization).

Енкодер і декодер працюють у двох потоках: у чистому та зашумленому. Опис архітектури формалізується таким чином [9]:

$$\tilde{x}, \tilde{z}^1, \dots, \tilde{z}^L, \tilde{y} = \text{Encoder}_{noisy}(x);$$

$$x, z^1, \dots, z^L, y = \text{Encoder}_{clean}(x);$$

$$\hat{x}, \hat{z}^1, \dots, \hat{z}^L, \hat{y} = \text{Decoder}(\tilde{z}^1, \dots, \tilde{z}^L),$$

де x , y та $\tilde{y}$ — вхідні дані, знешумлені та зашумлені виходи; $z^1, \tilde{z}^1, \hat{z}^1$ — приховане представлення, його зашумлена версія та його реконструйоване представлення.

Кожен шар декодера реконструює зашумлене латентне представлення енкодера. Для цього він використовує реконструйоване попередніми шарами декодера латентне представлення у комбінації із зашумленим представленням енкодера із поточного рівня:

$$\hat{z}^l = g(\tilde{z}^l, u^{l+1}),$$

де $\hat{z}^l$ — реконструйоване представлення; u^{l+1} — вертикальний сигнал із наступного шару після застосування батч-нормалізації; g — функція комбінування.

Цільова функція складається із зваженої суми втрат за класифікацію та за реконструкцію:

$$L = \lambda_{sup} L_{sup} + \lambda_{rec} L_{rec}.$$

Для задачі класифікації як L_{sup} використовується Cross-Entropy Loss, а як L_{rec} використовується MSE.

Спрощену схему архітектури сходової нейронної мережі зображено на рис. 2, де CE — це Cross-Entropy Loss, RC — це функції втрат реконструкції, сума яких формує загальну функцію втрат L_{rec} . На виході незашумленого енкодера зазвичай використовують повнозв'язний шар, що використовується для класифікації.

Ця архітектура забезпечує можливість одночасного навчання моделі як на мічених, так і на немічених даних. Для немічених даних використовується шлях «*Encoder_{noisy} → Decoder*», де зворотне поширення помилки дозволяє коригувати ваги модуля *Encoder_{clean}*, покращуючи здатність моделі до виділення важливих ознак.

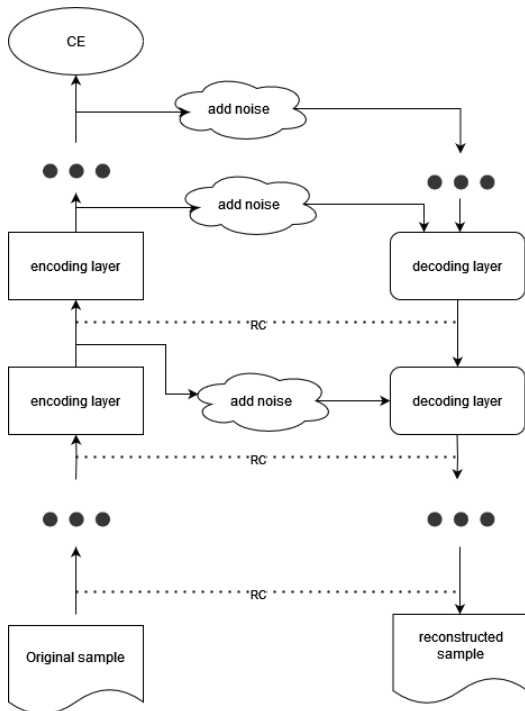


Рис. 2. Спрощена схема архітектури сходової нейронної мережі. Схему створено на основі опису з [9]

Для мічених даних додається стандартний енкодер із повнозв'язним шаром на виході, який навчається паралельно. У результаті модель має два окремі потоки: «*Encoder_{noisy} → Decoder*» для обробки немічених даних і «*Encoder_{clean} → повнозв'язний шар*» для роботи з міченими зразками.

Запропонована модель

Після огляду наявних підходів було встановлено ключовий недолік: більшість сучасних методів НН, розроблених для класифікації ракових захворювань шкіри, покладаються на евристичні методи псевдомаркування, при цьому структура латентного простору залишається поза увагою. Ці спостереження підказують, що перспективними є дослідження моделей з регуляризациєю, таких як сходові нейронні мережі та змагальні автоенкодері.

У [7] автори розглядають використання дискримінатора виключно для формування L_{reg} , тобто втрат за регуляризацию латентного простору на виході енкодера. Якщо використати дискримінатор для оцінювання втрат звичайної реконструкції L_{rec} , то це не буде мати сенсу, оскільки оригінальний автоенкодер має повну інформацію про оригінальне зображення, а отже, цінність дискримінатора в цій архітектурі невелика. Утім сходові нейронні мережі використовують зашумлення на етапі передачі латентного представлення до декодера, тому суттєва частина інформації втрачається. Як наслідок, використання дискримінатора в цій моделі набуває сенсу.

Модель класифікації ракових захворювань шкіри на основі інтеграції дискримінатора в архітектуру сходових нейронних мереж

У нашому дослідженні пропонується використовувати дискримінатор як більш м'який (у сенсі soft computing) спосіб оцінювання вартості реконструкції сходової нейронної мережі.

Запропонована функція реконструкції у сходовій нейронній мережі може бути виражена таким чином:

$$L_{rec} = E_{x \sim p_{data}}(x) [\log D(\tilde{x})], \quad (1)$$

де $\tilde{x}$ — частково реконструйоване зображення, яке декодер намагається відновити із зашумленого латентного представлення.

В (1) D — це дискримінатор, що навчається розрізняти відновлені зображення від справжніх,

що дозволяє йому «м'яко» контролювати якість реконструкції, орієнтуючись на глобальні ознаки, такі як текстура і форма, замість грубих порівнянь, які використовуються в оригінальній архітектурі схожих нейронних мереж.

Використання зашумлених даних у схожих нейронних мережах створює ефективну навчальну задачу для дискримінатора, оскільки він оцінює відновлення зображення у складніших умовах, ніж у звичайному автоенкодері. Водночас дискримінатор своїм оцінюванням спонукає генератор (тобто декодер) відновити важливі ознаки замість того, щоб грубо намагатися мінімізувати загальну похибку між відновленим зображенням та оригіналом.

Пропонується така функція втрат для цієї моделі:

$$L = \lambda_{NLL} L_{NLL} + \lambda_{rec} L_{rec}, \quad (2)$$

де L_{rec} можна використовувати як подання (1), так і комбінований варіант, що також частково враховує грубі методи оцінювання реконструкції:

$$L_{rec} = E_{x \sim p_{data}}(x) [\log D(\tilde{x})] + \lambda_s E_{x \sim p_{data}}(x) R(x, \hat{x}), \quad (3)$$

де $R(x, \hat{x})$ – це функція грубого оцінювання реконструкції (наприклад, MSE); λ_s – коефіцієнт, що набуває значень $[0,1]$.

Фрагмент моделі запропонованої мережі схематично зображено на рис. 3.

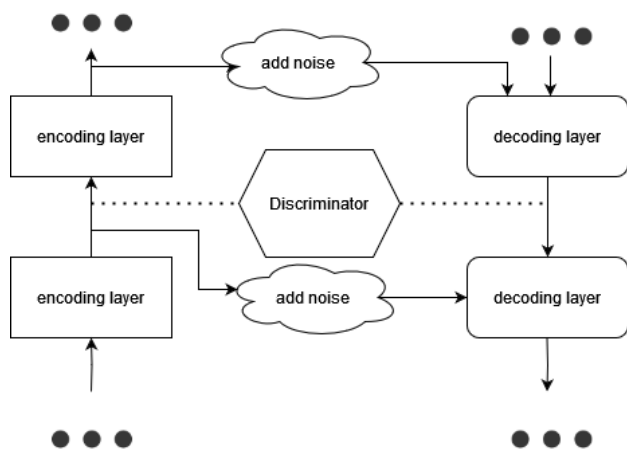


Рис. 3. Архітектура фрагмента схожої мережі із дискримінатором для оцінювання реконструкції

Тут RC замінена дискримінатором, який виконує функцію оцінювання якості реконструкції. Заміна функції втрат реконструкції на дискримінатор можна реалізувати як на кожному

рівні схожої нейронної мережі, так і вибірково. Наприклад, можна використати класичний спосіб оцінювання якості реконструкції для латентних представлень, а спосіб із дискримінатором використати для порівняння реконструйованого зображення з оригіналом.

Результати експериментів

Експерименти проводилися на датасеті HAM10000, що містить зображення ракових захворювань шкіри (6 типів), а також здорових утворень (родимок) і відповідних міток цих зображень. Усі експерименти проводилися для різних співвідношень розмічених і нерозмічених даних у вибірці: 30 %, 10 % та 5 % розмічених даних. Оскільки для дослідження моделей напівкерованого навчання необхідно припустити, що більшість зображень є нерозміченими, дані було поділено на тип «злаякісні» та «доброякісні», тобто всі шість типів злаякісних утворень було згруповано в один клас (рис. 4).

Для порівняння розглянутих моделей НН створена спрощена згорткова нейронна мережа із чотирьох згорткових шарів, а також одного лінійного шару. Було використано всього один лінійний шар для наочності демонстрації регуляризуючих здібностей автоенкодера. Ця базова модель використовувалася як еталонна для порівняння якості запропонованих моделей НН. Еталонна «погана» модель – це модель, навчена тільки на розміченій частині даних. Еталонна «хороша» модель – це модель, навчена так, ніби всі 100 % даних є розміченими. Якщо метричні показники автоенкодера будуть наближатися до показників еталонної «поганої» моделі, то це означатиме, що модель НН є неефективною, оскільки використання нерозмічених даних не покращує результат. Аналогічно, якщо якісні показники автоенкодера близькі до еталонної «хорошої» моделі, то можна зробити висновок, що запропонована модель НН є ефективною.

Базова модель використовувалася як основа для енкодера у всіх моделях автоенкодера із мінімально можливими модифікаціями, які вимагаються для реалізації. У побудові автоенкодерів для розрахунку Reconstruction Cost використано звичайний MSE для всіх прихованих шарів і дискримінатор для порівняння зображення з оригіналом. На рис. 5 зображено схематичний опис моделі згорткового автоенкодера та схожої нейронної мережі з використанням дискримінатора.

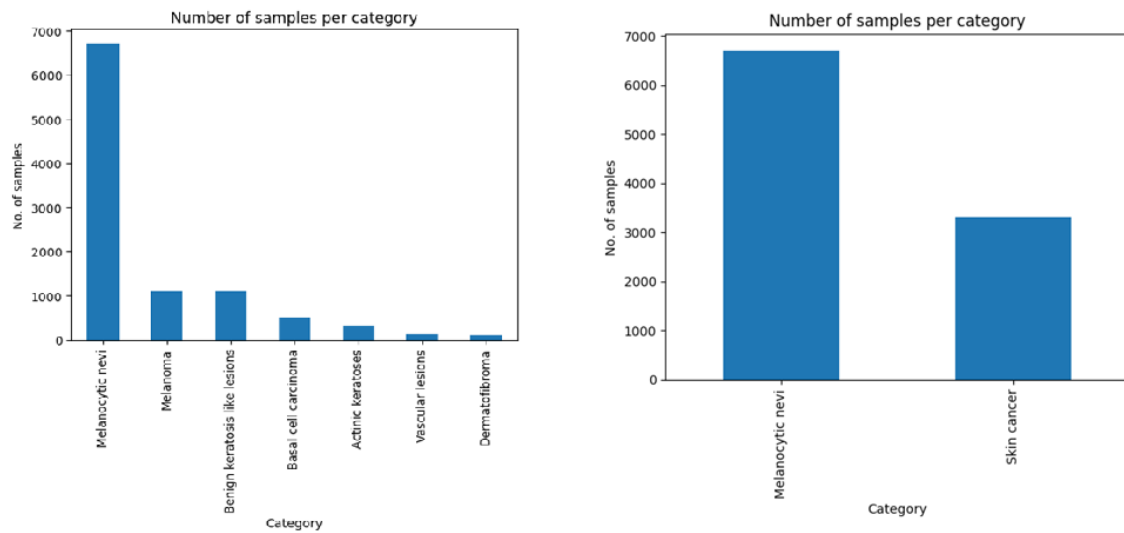


Рис. 4. Розподіл даних у датасеті до та після ребалансування вибірки

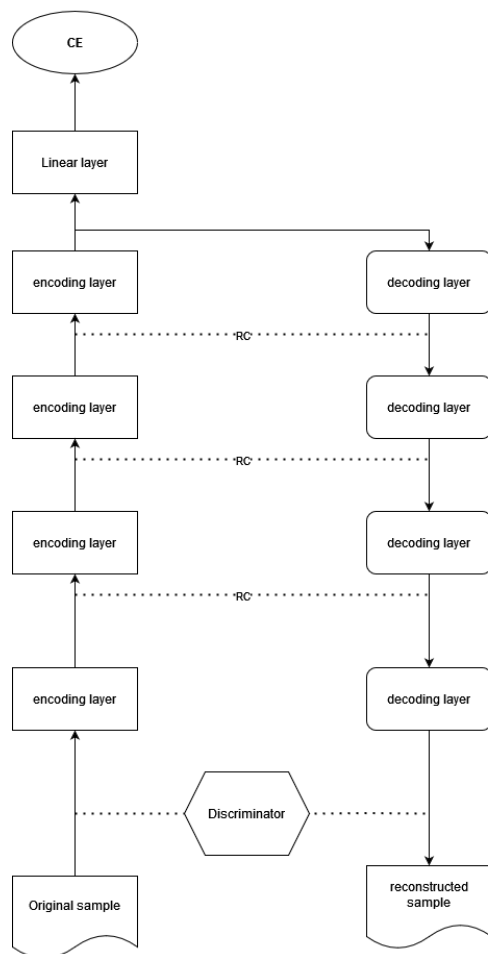
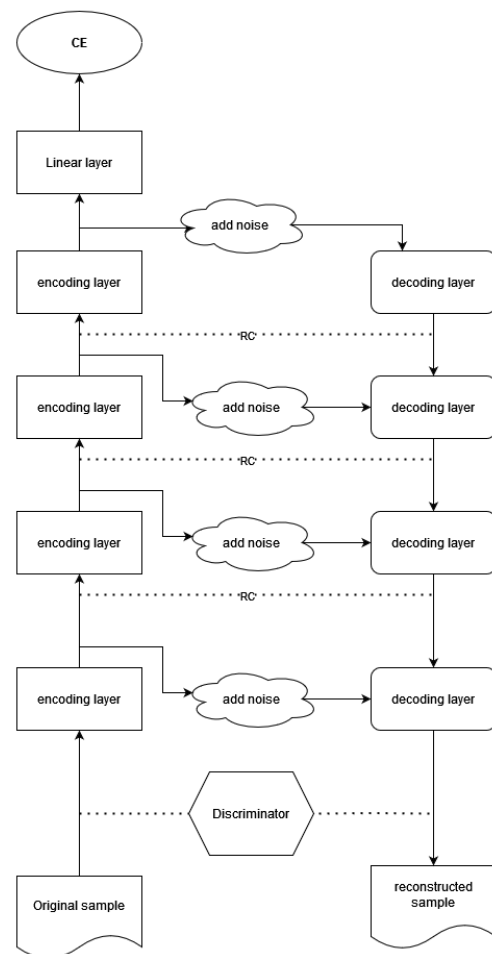
Autoencoder with discriminator for RC**Ladder Net with discriminator for RC**

Рис. 5. Автоенкодер та сходова нейронна мережа з дискримінатором

У табл. 1 відображено метричні показники моделей для різних співвідношень розмічених даних у вибірці. Тут CA – це звичайний згортковий автоенкодер, LN – це сходова нейронна мережа, simple означає, що використовується оригінальна архітектура без дискримінатора, discriminator RC – це модель із дискримінатором для оцінювання reconstruction cost (як зображено на рис. 5), combined RC – це така сама модель, але функція втрат визначається як зважена сума оцінювання дискримінатора та MSE (тобто за формулою (3)).

З результатів експериментів випливає, що використання дискримінатора виявилось ефективним і збільшило F_1 -score для класу злоякісних захворювань як для звичайного автоенкодера, так і для сходової нейронної мережі. Також із результатів експерименту видно, що використання комбінації сходових нейронних мереж із дискримінатором у більшості випадків є ефективнішим порівняно з використанням того

самого підходу для звичайного автоенкодера (F_1 -score становив 74,8 % проти 73,5 % для 30 % розмічених даних, 74,3 % проти 74,6 % для 10 %, 71 % проти 68,6 % для 5 %).

Особливо відчутна перевага такого підходу на 5 % розмічених даних, де приріст F_1 -score метрики склав 4 % порівняно зі звичайною LN та 2,4 % порівняно з аналогічною моделлю звичайного згорткового автоенкодера. Також помітно, що використання комбінованого підходу обчислення функції втрат за формулою (3) демонструє кращі результати на вибірках з малою кількістю розмічених даних. Для 30 % розмічених даних використання виключно дискримінатора без MSE (формула (2)) показало себе краще як для сходової нейронної мережі, так і для згорткового автоенкодера (73,5 % проти 71,8 % для CA, а також 74,8 % проти 72,5 % для LN), при цьому показник F_1 -score досяг 74,8 %, що всього на 1 % менше за повністю керовану модель.

Таблиця 1. Метричні показники моделей, навчених різними підходами

Модель НН	Accuracy	Precision (melanocytic nevi)	Precision (skin cancer)	Recall (melanocytic nevi)	Recall (skin cancer)	F1 (melanocytic nevi)	F1 (skin cancer)
30 % розмічених даних							
CA (simple)	0,801	0,828	0,739	0,876	0,660	0,851	0,698
CA (discriminator RC)	0,827	0,844	0,789	0,902	0,687	0,872	0,735
CA (combined RC)	0,818	0,834	0,780	0,900	0,664	0,866	0,718
LN (simple)	0,807	0,824	0,767	0,896	0,641	0,858	0,699
LN (discriminator RC)	0,818	0,875	0,722	0,841	0,775	0,857	0,748
LN (combined RC)	0,807	0,854	0,721	0,849	0,729	0,852	0,725
Etalon bad model	0,773	0,779	0,751	0,908	0,519	0,839	0,614
Etalon good model	0,847	0,852	0,796	0,903	0,712	0,883	0,756
10 % розмічених даних							
CA (simple)	0,797	0,853	0,698	0,831	0,733	0,842	0,715
CA (discriminator RC)	0,816	0,874	0,720	0,839	0,775	0,856	0,746
CA (combined RC)	0,820	0,839	0,777	0,896	0,679	0,867	0,725
LN (simple)	0,794	0,844	0,702	0,839	0,710	0,841	0,706
LN (discriminator RC)	0,809	0,855	0,723	0,851	0,729	0,853	0,726
LN (combined RC)	0,805	0,889	0,685	0,800	0,813	0,842	0,743
Etalon bad model	0,743	0,786	0,654	0,844	0,552	0,819	0,591
Etalon good model	0,847	0,852	0,796	0,903	0,712	0,883	0,756
5 % розмічених даних							
CA (simple)	0,778	0,797	0,727	0,884	0,580	0,838	0,645
CA (discriminator RC)	0,786	0,813	0,722	0,871	0,626	0,841	0,671
CA (combined RC)	0,795	0,821	0,737	0,878	0,641	0,848	0,686
LN (simple)	0,802	0,803	0,799	0,922	0,576	0,858	0,670
LN (discriminator RC)	0,803	0,835	0,735	0,869	0,769	0,852	0,706
LN (combined RC)	0,794	0,850	0,696	0,831	0,725	0,840	0,710
Etalon bad model	0,691	0,913	0,534	0,582	0,897	0,711	0,670
Etalon good model	0,847	0,852	0,796	0,903	0,712	0,883	0,756

Для зовнішнього порівняння із запропонованою моделлю було обрано модель STFL [2]. Оскільки у статті як базову модель було використано ResNet-50, що має суттєво більшу кількість параметрів, ніж базова модель у нашому дослідженні, результати STFL коректно порівнювати лише відносно. У [2] наведено метричні показники моделі STFL для двох режимів: 5 % маркованих зразків (0,7462) та 100 % (0,8507). Пропонується оцінити відносну якість моделі як співвідношення показника F_1 для напівкерованої моделі ($F_{1_{SSL}}$) до відповідного показника для повністю керованої моделі ($F_{1_{SUP}}$):

$$R_{F_1} = \frac{F_{1_{SSL}}}{F_{1_{SUP}}} = \frac{0.7462}{0.8507} = 0.877.$$

Запропонована модель має відчутно вищий показник R_{F_1} для 5 % розмічених даних (0,939), що може свідчити про високу ефективність використання нерозмічених даних для покращення якості моделі.

Висновки

Досліджено та встановлено відсутність у відкритих джерелах експериментів, де denoising-

autoencoder архітектури, зокрема LN, оцінювалися на HAM10000. Виявлено, що наявні SSL-методи для класифікації ракових захворювань шкіри (FixMatch, MixMatch, Mean-Teacher, STFL) не враховують регуляризації латентного простору.

Розроблено напівкеровану модель, яка інтегрує дискримінатор ААЕ у декодер сходової нейронної мережі для комбінованого оцінювання Reconstruction Cost (MSE + дискримінатор).

Модель забезпечує суттєве покращення якості. Наприклад, за 5 % розмічених даних F_1 -score для злоякісних утворень зріс на 4 % відносно класичної LN і на 2,4 % відносно аналогічної моделі звичайного автоенкодера з дискримінатором. Відносний коефіцієнт $R_{F_1} = 0,929$ перевищив аналогічний показник STFL (0,877), що підтверджує ефективність використання нерозмічених даних у процесі напівкерованого навчання моделі.

Перспективи подальших досліджень включають вивчення використання дискримінатора для порівняння латентних представлень та аналіз впливу різних параметрів функції Reconstruction Cost.

References

- [1] О.О. Зарицький та В.Я. Данилов, «Порівняння ефективності методів напівкерованого навчання на основі автоенкодерів для задач класифікації фотографій ракових захворювань шкіри», *Вісник Вінницького політехнічного інституту*, 2024. № 2. с. 71–77. doi: 10.31649/1997-9266-2024-173-2-71-77.
- [2] W. Yuan, Z. Du, and S. Han, “Semi-supervised skin cancer diagnosis based on self-feedback threshold focal learning”, *Discover Oncology*, 2024, vol. 15. no. 1, 180 p. doi: 10.1007/s12672-024-01043-8.
- [3] C.B. Hansen, V. Nath, R. Gao, C. Bermudez, Y. Huo, K.L. Sandler, P.P. Massion, J.D. Blume, T.A. Lasko, and B.A. Landman, “Semi-supervised machine learning with MixMatch and equivalence classes”, *Interpretable and Annotation-Efficient Learning for Medical Image Computing* (iMIMIC, MIL3ID & LABELS Workshops, MICCAI 2020, Lima, Peru), Lecture Notes in Computer Science, 2020, vol. 12446, pp. 112–121. doi: 10.1007/978-3-030-61166-8_12.
- [4] S. Kitada and H. Iyatomi, “Skin lesion classification with ensemble of squeeze-and-excitation networks and semi-supervised learning”, arXiv preprint arXiv:1809.02568, 2018. [Online]. Retrieved from: doi: 10.48550/arXiv.1809.02568.
- [5] A. Makhzani, J. Shlens, N. Jaitly, I. Goodfellow, and B. Frey, “Adversarial autoencoders”, arXiv preprint arXiv:1511.05644, 2015. [Online]. Retrieved from: <https://arxiv.org/abs/1511.05644>. doi: 10.48550/arXiv.1511.05644.
- [6] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, “Generative adversarial nets”, *Advances in Neural Information Processing Systems*, 2014, vol. 27, pp. 2672–2680. doi: 10.1145/3422622.
- [7] R. Tachibana, T. Matsubara, and K. Uehara, “Semi-supervised learning using adversarial networks”, *2016 IEEE/ACIS 15th International Conference on Computer and Information Science (ICIS)*, Okayama, Japan, 2016, pp. 1–6. doi: 10.1109/icis.2016.7550881.
- [8] A. Gogna and A. Majumdar, “Semi supervised autoencoder”, *Neural Information Processing, 23rd International Conference, ICONIP 2016*, Kyoto, Japan, October 16–21, 2016, Proceedings, Part II 23, Springer International Publishing, 2016. doi: 10.1007/978-3-319-46672-9_10
- [9] M. Pezeshki, L. Fan, P. Brakel, A. Courville, and Y. Bengio, “Deconstructing the Ladder Network Architecture”, *Proceedings of the 33rd International Conference on Machine Learning*, New York, NY, USA, 2016, pp. 2368–2376. doi: 10.48550/arXiv.1511.06430.
- [10] A. Rasmus, H. Valpola, M. Honkala, M. Berglund and T. Raiko “Semi-Supervised Learning with Ladder Networks”, *NIPS’15: Proceedings of the 29th International Conference on Neural Information Processing Systems*, 2015, vol. 2, pp. 3546–3554. doi: 10.48550/arXiv.1507.02672.

- [11] D.P. Kingma, D.J. Rezende, S. Mohamed and M. Welling, “Semi-supervised learning with deep generative models”, *Advances in Neural Information Processing Systems*, 2014, vol. 27, pp. 3581–3589. doi: 10.48550/arXiv.1406.5298.

V.Y. Danilov, O.O. Zarytskyi

A SKIN CANCER CLASSIFICATION MODEL INCORPORATING A DISCRIMINATOR INTO THE LADDER NEURAL NETWORK ARCHITECTURE

Background. Semi-supervised learning (SSL) is one of the most promising areas of deep learning, especially for tasks where label acquisition is a complex and expensive process, such as skin cancer classification. Existing approaches, such as Adversarial Autoencoders (AAE) and Ladder Networks (LN), effectively use unlabelled data but have limitations in reconstruction and regularization accuracy.

Objective. Development and study of a model for classifying skin cancers based on the integration of a discriminator into the architecture of ladder neural networks.

Methods. The developed model combines the regularizing properties of ladder neural networks with the use of a discriminator-based loss function that evaluates the quality of image reconstruction based on their structure, shape, and key visual features. The experiments were conducted on the HAM10000 dataset with different ratios of labelled and unlabelled data (30 %, 10 %, 5 %).

Results. The experiments showed that the proposed model improved the F_1 -score for the malignancy class by 4 % compared to the baseline ladder network with 5 % labelled data. With 30 % labeled samples, the F_1 -score reached 74.8 %, which is only 1 % less than the fully supervised model. The relative indicator R_{F_1} at 5 % labelled data was 0.939, exceeding the similar coefficient for STFL (0.877), which confirms the effectiveness of using unlabelled data in the proposed model.

Conclusions. The proposed model improves on existing semi-supervised learning methods (LN, AAE) by providing high efficiency of using unlabelled data to regularize the encoder latent space in the task of skin cancer classification. Prospects for further research include using a discriminator to compare latent spaces and improving the Reconstruction Cost function to expand its application to other medical image analysis tasks.

Keywords. skin cancer classification; HAM10000; semi-supervised learning; ladder neural networks; discriminator.

Рекомендована Радою
Навчально-наукового інституту прикладного системного аналізу
КПІ ім. Ігоря Сікорського

Надійшла до редакції
18 січня 2025 року

Прийнята до публікації
30 червня 2025 року