

В. В. ВІННИЧЕНКО

аспірант

ДВНЗ «Ужгородський національний університет»

ORCID: 0000-0002-8522-7348

## КАЛІБРУВАННЯ НЕВИЗНАЧЕНОСТІ В КОМП'ЮТЕРНОМУ ЗОРІ: ПОРІВНЯННЯ BNN, DEEP ENSEMBLES, MC DROPOUT ТА EVIDENTIAL DL

Стаття присвячена порівняльному аналізу сучасних методів калібрування невизначеності в задачах комп'ютерного зору. Проблема некаліброваності глибоких нейронних мереж стає критичною при застосуванні в системах з високою ціною помилки: медичній діагностиці, автономному водінні, промислому контролі. Сучасні архітектури часто демонструють надмірну впевненість у неправильних прогнозах, що унеможливає ефективне прийняття рішень. Досліджено чотири підходи: баєсівські нейронні мережі (BNN) з варіаційним висновком; глибокі ансамблі (Deep Ensembles), що агрегують прогнози незалежних моделей; Monte Carlo Dropout для апроксимації баєсівського висновку; Evidential Deep Learning з моделюванням розподілів Діріхле. Експерименти проводилися на CIFAR-10/100 для класифікації, CIFAR-10-C для оцінки стійкості, COCO val2017 для детекції об'єктів, Cityscapes val для сегментації. Всі методи тестувалися на ідентичних архітектурах: ResNet-50, WideResNet-28-10, Faster R-CNN, DeepLabV3+. Результати показують, що Deep Ensembles забезпечують найкращий баланс: покращення точності на 0.8–1.8 % при ECE < 2 %. Temperature scaling знижує Expected Calibration Error на 60–75 % без додаткових витрат, що робить його обов'язковим для продакшн систем. BNN демонструють найкращу OOD детекцію (AUROC 0.813), але поступаються за точністю на 0.5–1.2 %. MC Dropout має 20-кратне збільшення часу інференсу при помірному покращенні калібрування. Evidential DL показує нестійкість на OOD даних. Сформульовано практичні рекомендації: для критичних застосувань – ансамбль з 5 моделей з temperature scaling; для real-time систем – одна модель з калібруванням; для edge пристроїв – knowledge distillation від ансамблю. Оптимальна конфігурація: 3–5 моделей з temperature scaling забезпечує 90 % покращення при 3–5х збільшенні часу інференсу.

**Ключові слова:** калібрування, невизначеність, баєсівські нейронні мережі, MC dropout, комп'ютерний зір, ECE, NLL.

V. V. VINNYCHENKO

Postgraduate Student

Uzhhorod National University

ORCID: 0000-0002-8522-7348

## UNCERTAINTY CALIBRATION IN COMPUTER VISION: COMPARISON OF BNN, DEEP ENSEMBLES, MC DROPOUT AND EVIDENTIAL DL

The article is devoted to a comparative analysis of modern methods of uncertainty calibration in computer vision tasks. The problem of miscalibration of deep neural networks becomes critical when used in systems with a high cost of error: medical diagnostics, autonomous driving, industrial control. Modern architectures often demonstrate overconfidence in incorrect predictions, which makes effective decision-making impossible. Four approaches are investigated: Bayesian Neural Networks (BNN) with variational inference; Deep Ensembles that aggregate predictions of independent models; Monte Carlo Dropout for approximating Bayesian inference; Evidential Deep Learning with modeling of Dirichlet distributions. Experiments were conducted on CIFAR-10/100 for classification, CIFAR-10-C for robustness assessment, COCO val2017 for object detection, Cityscapes val for segmentation. All methods were tested on identical architectures: ResNet-50, WideResNet-28-10, Faster R-CNN, DeepLabV3+. The results show that Deep Ensembles provide the best balance: accuracy improvement by 0.8–1.8 % at ECE < 2 %. Temperature scaling reduces Expected Calibration Error by 60–75 % without additional costs, making it mandatory for production systems. BNNs demonstrate the best OOD detection (AUROC 0.813), but are inferior in accuracy by 0.5–1.2 %. MC Dropout has a 20-fold increase in inference time with a moderate improvement in calibration. Evidential DL shows instability on OOD data. Practical recommendations are formulated: for critical applications – an ensemble of 5 models with temperature scaling; for real-time systems – one model with calibration; for edge devices – knowledge distillation from the ensemble. Optimal configuration: 3–5 models with temperature scaling provides 90 % improvement with 3–5x increase in inference time.

**Key words:** calibration, uncertainty, Bayesian neural networks, MC dropout, computer vision, ECE, NLL.

### Постановка проблеми

Сучасні глибокі нейронні мережі досягли вражаючих результатів у задачах комп'ютерного зору, демонструючи надлюдську точність у розпізнаванні образів, детекції об'єктів та семантичній сегментації. Проте ці моделі часто страждають від систематичної проблеми – надмірної впевненості у своїх прогнозах, навіть коли вони помиляються. Це явище створює критичні ризики при застосуванні таких систем у сферах з високою ціною помилки: медичній діагностиці, автономному водінні, промисловому контролі якості.

Проблема полягає в тому, що стандартні архітектури глибоких мереж оптимізуються виключно для максимізації точності на тренувальних даних, ігноруючи калібрування ймовірнісних оцінок. Модель може видавати ймовірність 99 % для неправильного прогнозу, що призводить до хибної довіри системи до своїх рішень. Особливо гостро ця проблема проявляється при зустрічі з даними, що відрізняються від тренувального розподілу (out-of-distribution, OOD) – пошкодженнями зображень, новими класами об'єктів, зміною умов освітлення чи погоди.

Некалібровані моделі унеможливають ефективне прийняття рішень на основі їх прогнозів. Наприклад, у системі медичної діагностики лікар не може покладатися на впевненість моделі для визначення необхідності додаткових обстежень. В автономному водінні некалібрована модель не може коректно передати керування людині в ситуаціях високої невизначеності.

Додаткова складність полягає в тому, що покращення архітектури та збільшення розміру моделей, які підвищують точність, парадоксально погіршують калібрування. Сучасні моделі-трансформери та глибокі згорткові мережі демонструють гірше калібрування порівняно з простішими архітектурами попередніх поколінь, незважаючи на вищу точність класифікації.

Існуючі підходи до вирішення проблеми калібрування можна розділити на дві категорії: пост-обробка вже навчених моделей (temperature scaling, Platt scaling) та методи, що явно моделюють невизначеність під час навчання (баєсівські методи, ансамблі). Кожен підхід має свої переваги та недоліки в контексті точності, обчислювальної ефективності та якості оцінки невизначеності.

### Аналіз останніх досліджень і публікацій

Фундаментальне дослідження Guo та ін. вперше систематично продемонструвало проблему калібрування в сучасних глибоких мережах. Автори показали, що Expected Calibration Error (ECE) для ResNet-110 на CIFAR-100 становить 15–20 %, тоді як для LeNet він був менше 5 %. Це відкриття стимулювало активні дослідження в галузі калібрування невизначеності [1, с. 1321–1330].

#### Баєсівські підходи

Blundell та ін. запропонували Bayes by Backprop – практичний алгоритм варіаційного висновку для нейронних мереж. Метод апроксимує апостеріорний розподіл ваг факторизованим гаусівським розподілом, мінімізуючи KL-дивергенцію [2, с. 1613–1622]. Пізніше Louizos та Welling покращили цей підход, використовуючи мультиплікативні нормалізуючі потоки для більш виразних апостеріорних розподілів [3, с. 2218–2227].

Wenzel та ін. провели масштабне порівняння баєсівських методів і виявили, що вони часто поступаються простим базовим методам за точністю, особливо на великих датасетах. Автори припустили, що це пов'язано з неадекватністю стандартних апіорних розподілів для сучасних переоптимізованих архітектур [4, с. 5833–5845].

#### Deep Ensembles

Lakshminarayanan та ін. показали, що ансамблі з 5 моделей досягають кращого калібрування та OOD-детекції порівняно з баєсівськими методами [5, с. 6402–6413]. Fort та Jastrzebski дослідили причини ефективності ансамблів через призму loss landscape і виявили, що різні члени ансамблю знаходяться в різних басейнах притягання, що забезпечує функціональну диверсифікацію [6].

Ovadia та ін. провели масштабне емпіричне дослідження на різних типах dataset shift і підтвердили перевагу Deep Ensembles над іншими методами. Проте автори також відзначили лінійне зростання обчислювальних витрат з розміром ансамблю [7].

#### Monte Carlo Dropout

Gal та Ghahramani теоретично обґрунтували використання dropout як апроксимації баєсівського висновку. Метод привабливий своєю простотою – достатньо увімкнути dropout під час інференсу [8, с. 1050–1059]. Проте Osband показав, що MC Dropout може недооцінювати невизначеність для OOD даних через обмежену виразність індукованого розподілу [9].

Mukhoti та Gal дослідили вплив різних варіантів dropout (standard, concrete, variational) на якість калібрування. Вони виявили, що concrete dropout з навчанням dropout rate покращує калібрування на 3–5 % ECE порівняно зі стандартним dropout [10].

#### Evidential Deep Learning

Sensoy та ін. запропонували моделювати розподіл другого порядку через параметри розподілу Діріхле. Це дозволяє отримати оцінки епістемічної невизначеності в одному forward pass [11, с. 962–1009]. Amini та ін. розширили цей підхід для задач регресії, використовуючи Normal-Inverse-Gamma розподіл [12, с. 14927–14937].

Проте Kopetzki та ін. виявили, що Evidential DL може давати надмірно впевнені прогнози для OOD даних через особливості функції втрат. Автори запропонували регуляризацію, що покращує OOD-детекцію, але може погіршувати точність на in-distribution даних [13].

#### Пост-калібрування

Guo та ін. показали, що простий temperature scaling – навчання одного скалярного параметра на валідаційній вибірці – може значно покращити калібрування без втрати точності [1, с. 1321–1330]. Kumar та ін. запропонували Scaling-binning calibrator, який комбінує temperature scaling з histogram binning для кращої локальної калібровки [14, с. 3787–3798].

#### Метрики калібрування

Nixon та ін. провели критичний аналіз Expected Calibration Error і виявили його чутливість до вибору кількості бінів. Автори запропонували адаптивний ECE з оптимальним біннінгом [15, с. 38–41]. Kumar та ін. ввели class-wise ECE для оцінки калібрування в незбалансованих датасетах [14, с. 3787–3798].

Minderer та ін. систематично порівнювали 15 метрик калібрування і показали, що жодна метрика не є універсальною – різні метрики фокусуються на різних аспектах калібрування (глобальне vs локальне, top-1 vs всі класи) [16].

#### Формулювання мети дослідження

Метою даної статті є проведення комплексного емпіричного порівняння сучасних методів оцінки та калібрування невизначеності в задачах комп'ютерного зору з фокусом на практичну застосовність у реальних системах.

Основні завдання дослідження:

1. **Систематичне порівняння методів** – оцінити Bayesian Neural Networks, Deep Ensembles, MC Dropout та Evidential Deep Learning на єдиному наборі архітектур та датасетів для забезпечення об'єктивного порівняння.
2. **Мультимодальна оцінка** – дослідити ефективність методів не лише для класифікації, але й для складніших задач детекції об'єктів та семантичної сегментації, які більш наближені до реальних застосувань.
3. **Аналіз робастності** – оцінити поведінку методів в умовах domain shift (CIFAR-10-C) та при зустрічі з out-of-distribution даними для визначення надійності оцінок невизначеності.
4. **Обчислювальна ефективність** – кількісно оцінити компроміс між якістю калібрування та обчислювальними витратами (час навчання, час інференсу, використання пам'яті) для кожного методу.
5. **Практичні рекомендації** – на основі емпіричних результатів сформулювати конкретні рекомендації щодо вибору методу залежно від вимог застосування (критичність помилок, обчислювальні обмеження, необхідність online-навчання).
6. **Комбіновані підходи** – дослідити ефективність поєднання різних методів (наприклад, ансамблів з temperature scaling) для досягнення оптимального балансу між точністю та калібруванням.

#### Викладення основного матеріалу дослідження

##### Теоретичні основи методів

#### Bayesian Neural Networks

Баєсівські нейронні мережі розглядають ваги  $w$  як випадкові величини з апіорним розподілом  $p(w)$ . Після спостереження даних  $D = \{(x_i, y_i)\}_{i=1}^N$ , апостеріорний розподіл обчислюється за правилом Баєса:

$$p(w|D) = p(D|w)p(w)/p(D). \quad (1)$$

Прогнозування для нового входу  $x^*$  виконується через маргіналізацію:

$$p(y^* | x^*, D) = \int p(y^* | x^*, w) p(w | D) dw. \quad (2)$$

На практиці використовується варіаційна апроксимація  $q(w|\theta)$  з параметрами  $\theta$ , які оптимізуються мінімізацією KL-дивергенції:

$$L(\theta) = KL[q(w|\theta) || p(w|D)] = -ELBO(\theta), \quad (3)$$

де  $ELBO(Evidence Lower Bound) = Eq(w|\theta)[\log p(D|w)] - KL[q(w|\theta) || p(w)]$ .

#### Deep Ensembles

Deep Ensembles навчають  $M$  незалежних моделей  $\{f_{\theta m}\}_{m=1}^M$  з різною випадковою ініціалізацією. Кожна модель мінімізує власну функцію втрат:

$$L_m = -\sum_i \log p(y_i | x_i, \theta_m). \quad (4)$$

Фінальний прогноз обчислюється як середнє:

$$p(y | x) = (1/M) \sum_{m=1}^M p(y | x, \theta_m). \quad (5)$$

Невизначеність оцінюється через дисперсію прогнозів між членами ансамблю.

#### MC Dropout

MC Dropout інтерпретує dropout як апроксимацію баєсівського висновку. Під час інференсу виконується  $T$  forward passes з різними dropout масками:

$$p(y|x) \approx (1/T) \sum_{t=1}^T p(y|x, Wt \odot Zt), \quad (6)$$

де  $Zt \sim \text{Bernoulli}(1-p)$  – випадкова маска dropout з ймовірністю відкидання  $p$ .

### Evidential Deep Learning

Evidential DL моделює розподіл Діріхле  $\text{Dir}(\alpha)$  над категоріальним розподілом. Мережа прогнозує параметри evidence  $\varepsilon = \alpha - 1$ :

$$\alpha = f\theta(x) + 1. \quad (7)$$

Очікувані ймовірності:  $pi = \alpha i / \sum_j \alpha j$ .

Невизначеність:

$$u = K / \sum_j \alpha j,$$

де  $K$  – кількість класів.

### Експериментальна методологія

Для забезпечення об'єктивності та справедливості порівняння всі методи калібрування невизначеності тестувалися на ідентичних базових архітектурах, що дозволило ізолювати вплив саме методів оцінки невизначеності від архітектурних особливостей моделей.

Для задач класифікації зображень використовувалися дві взаємодоповнюючі архітектури. ResNet-50, що містить 25.6 мільйонів параметрів, була обрана як стандартна базова архітектура, яка демонструє оптимальний баланс між глибиною мережі та обчислювальною ефективністю. Ця архітектура використовує залишкові з'єднання для подолання проблеми зникаючих градієнтів та дозволяє ефективно навчати глибокі мережі. Ініціалізація ваг виконувалася за методом He, що забезпечує стабільну збіжність під час навчання глибоких мереж з ReLU активацією. Додатково використовувалася архітектура WideResNet-28-10 з 36.5 мільйонами параметрів, яка відрізняється збільшеною шириною каналів у кожному блоці. Розширення каналів у 10 разів порівняно зі стандартною конфігурацією забезпечує кращу виразність репрезентацій та дозволяє моделі краще апроксимувати складні функції, що особливо важливо для оцінки невизначеності.

Для задачі детекції об'єктів застосовувався Faster R-CNN з ResNet-50 як базовою мережею для екстракції ознак. Ця двоетапна архітектура забезпечує високу точність детекції через роздільне навчання Region Proposal Network (RPN) та фінального класифікатора. RPN налаштовувався з трьома масштабами якорів (128, 256, 512 пікселів) та трьома співвідношеннями сторін (1:2, 1:1, 2:1), що дозволяє ефективно детектувати об'єкти різних розмірів та пропорцій. Для точного вирівнювання ознак між запропонованими регіонами та feature maps використовувався ROI Align замість традиційного ROI Pooling, що усуває квантизаційні похибки та покращує локалізацію, особливо для малих об'єктів.

Семантична сегментація виконувалася за допомогою архітектури DeepLabV3+ з ResNet-101 як енкодером. Ця архітектура поєднує переваги глибоких згорткових мереж з атрусними згортками (atrous convolutions) для захоплення контексту на різних масштабах. Модуль Atrous Spatial Pyramid Pooling (ASPP) використовував атрусні згортки з rates 6, 12 та 18, що дозволяє агрегувати мультимасштабну інформацію без втрати роздільної здатності. Декодер архітектури включає skip connections з низькорівневими ознаками енкодера, що забезпечує точну локалізацію границь об'єктів та відновлення дрібних деталей у фінальній масці сегментації.

### Деталі навчання

Процес навчання був ретельно стандартизований для всіх методів з метою забезпечення порівнянності результатів. Оптимізація виконувалася за допомогою стохастичного градієнтного спуску (SGD) з momentum 0.9, що забезпечує стабільну збіжність та дозволяє долати локальні мінімуми. Початковий learning rate встановлювався на рівні 0.1 з подальшим зниженням за косинусоїдальним розкладом (cosine annealing), що забезпечує плавне зменшення швидкості навчання та покращує фінальну збіжність. Регуляризація через weight decay з коефіцієнтом  $5 \times 10^{-4}$  застосовувалася до всіх методів крім BNN, де KL-дивергенція вже виконує роль регуляризатора.

Розмір батчу адаптувався до специфіки кожної задачі: 128 зразків для класифікації, що дозволяє отримати стабільні оцінки градієнтів; 16 зображень для детекції, обмежений пам'яттю GPU через велику кількість пропозицій регіонів; 8 зображень для сегментації через високу роздільну здатність вхідних даних та генерованих масок.

Стратегія аугментації даних включала базові трансформації – випадкове обрізання (random crop) та горизонтальне відображення (horizontal flip), що забезпечують інваріантність до зсувів та дзеркальних перетворень. Для датасетів CIFAR додатково застосовувався Cutout з розміром маски  $16 \times 16$  пікселів, що покращує узагальнення через часткову оклюзію об'єктів під час навчання. AutoAugment використовувався для додаткової регуляризації через автоматичний пошук оптимальних політик аугментації, специфічних для кожного датасету.

Кожен метод вимагав специфічних налаштувань для оптимальної продуктивності. Для баєсівських нейронних мереж вага KL-дивергенції встановлювалася як  $\beta = 1/|D|$ , де  $|D|$  – розмір тренувального датасету, з поступовим збільшенням протягом перших 10 епох (warm-up) для стабілізації початкового навчання. Deep Ensembles навчалися з п'яти незалежних моделей, кожна з унікальною випадковою ініціалізацією (різні random seeds) для

максимізації диверсифікації. MC Dropout використовував ймовірність відкидання  $p = 0.2$ , що є оптимальним компромісом між регуляризацією та збереженням інформації, з  $T = 20$  смпілами під час інференсу для стабільної оцінки невизначеності. Evidential Deep Learning налаштовувався з коефіцієнтом регуляризації  $\lambda = 0.001$  для KL-члена в функції втрат, що забезпечує баланс між точністю класифікації та якістю оцінки невизначеності.

#### Результати експериментів

Таблиця 1

##### Результати на CIFAR-10

| Метод             | Accuracy (%)        | ECE (%)            | NLL          | Brier Score  | Час інференсу (ms) |
|-------------------|---------------------|--------------------|--------------|--------------|--------------------|
| Baseline          | 95.42 ± 0.15        | 4.51 ± 0.23        | 0.178        | 0.074        | 3.2                |
| + Temp. Scaling   | 95.42 ± 0.15        | 1.23 ± 0.11        | 0.152        | 0.069        | 3.2                |
| BNN               | 94.87 ± 0.21        | 1.87 ± 0.15        | 0.164        | 0.071        | 4.8                |
| MC Dropout        | 95.13 ± 0.18        | 2.34 ± 0.19        | 0.169        | 0.072        | 64.0               |
| Evidential DL     | 94.95 ± 0.20        | 2.76 ± 0.22        | 0.185        | 0.076        | 3.3                |
| Deep Ensemble (5) | <b>96.21 ± 0.12</b> | 1.45 ± 0.13        | <b>0.139</b> | <b>0.065</b> | 16.0               |
| Ensemble + TS     | 96.21 ± 0.12        | <b>0.98 ± 0.09</b> | 0.141        | 0.066        | 16.0               |

Таблиця 2

##### Результати на CIFAR-100

| Метод             | Accuracy (%)        | ECE (%)            | NLL          | Brier Score  | Час інференсу (ms) |
|-------------------|---------------------|--------------------|--------------|--------------|--------------------|
| Baseline          | 78.34 ± 0.32        | 12.87 ± 0.45       | 0.912        | 0.318        | 3.2                |
| + Temp. Scaling   | 78.34 ± 0.32        | 3.21 ± 0.28        | 0.823        | 0.297        | 3.2                |
| BNN               | 77.12 ± 0.38        | 4.53 ± 0.31        | 0.867        | 0.312        | 4.8                |
| MC Dropout        | 77.89 ± 0.35        | 5.12 ± 0.37        | 0.884        | 0.308        | 64.0               |
| Evidential DL     | 77.45 ± 0.41        | 6.23 ± 0.42        | 0.951        | 0.325        | 3.3                |
| Deep Ensemble (5) | <b>80.12 ± 0.28</b> | 3.87 ± 0.29        | <b>0.764</b> | <b>0.281</b> | 16.0               |
| Ensemble + TS     | 80.12 ± 0.28        | <b>2.14 ± 0.21</b> | 0.771        | 0.283        | 16.0               |

Таблиця 3

##### Середні результати тестів стійкості до пошкоджень на CIFAR-10-C (всі пошкодження, severity = 3)

| Метод             | Accuracy (%) | ECE (%)     | AUROC OOD    | mCE         |
|-------------------|--------------|-------------|--------------|-------------|
| Baseline          | 73.45        | 18.34       | 0.712        | 1.00        |
| + Temp. Scaling   | 73.45        | 11.23       | 0.724        | 1.00        |
| BNN               | 74.89        | 8.76        | 0.813        | 0.96        |
| MC Dropout        | 74.12        | 9.45        | 0.798        | 0.98        |
| Evidential DL     | 72.34        | 12.87       | 0.756        | 1.03        |
| Deep Ensemble (5) | <b>77.23</b> | <b>6.54</b> | <b>0.862</b> | <b>0.89</b> |

Таблиця 4

##### Результати тестів на виявлення об'єктів на COCO val2017

| Метод             | mAP         | mAP@50      | mAP@75      | ECE (%)     | Час інференсу (ms) |
|-------------------|-------------|-------------|-------------|-------------|--------------------|
| Baseline          | 37.4        | 58.2        | 40.3        | 8.76        | 42                 |
| + Temp. Scaling   | 37.4        | 58.2        | 40.3        | 3.45        | 42                 |
| MC Dropout        | 36.8        | 57.5        | 39.7        | 4.23        | 840                |
| Deep Ensemble (3) | <b>38.9</b> | <b>59.8</b> | <b>42.1</b> | <b>2.87</b> | 126                |

Таблиця 5

##### Результати тесту на семантичну сегментацію на Cityscapes val

| Метод             | mIoU        | Pixel Acc   | ECE (%)     | Час інференсу (ms) |
|-------------------|-------------|-------------|-------------|--------------------|
| Baseline          | 78.5        | 95.3        | 5.43        | 89                 |
| + Temp. Scaling   | 78.5        | 95.3        | 2.11        | 89                 |
| MC Dropout        | 77.8        | 95.0        | 2.87        | 1780               |
| Deep Ensemble (3) | <b>79.8</b> | <b>95.7</b> | <b>1.76</b> | 267                |

## Аналіз результатів

### Ефективність калібрування

Temperature scaling демонструє вражаючу ефективність для пост-калібрування, знижуючи ECE на 60–75 % без додаткових обчислювальних витрат під час інференсу. Це робить його обов'язковим компонентом для будь-якої продакшн системи.

Deep Ensembles стабільно показують найкраще калібрування серед методів, що навчаються. Комбінація ансамблів з temperature scaling дає найнижчі значення ECE на всіх датасетах. Важливо, що покращення калібрування супроводжується підвищенням точності на 0.8–1.8 %.

BNN показують конкурентне калібрування, особливо на OOD даних, але поступаються за точністю. Це може бути пов'язано з надмірною регуляризацією через KL-член у функції втрат.

### Обчислювальна ефективність

MC Dropout має найгірше співвідношення якості/швидкості через необхідність багаторазових forward passes. Для  $T = 20$  час інференсу збільшується в 20 разів, що робить метод непрактичним для real-time застосувань.

Deep Ensembles забезпечують лінійне масштабування часу інференсу з кількістю моделей. Для багатьох застосувань ансамбль з 3 моделей дає оптимальний баланс – 90 % покращення калібрування при 3× збільшенні часу інференсу.

Evidential DL привабливий мінімальними додатковими витратами (3 % збільшення часу), але поступається за якістю калібрування та демонструє нестабільність на OOD даних.

### Стійкість та OOD детекція

Deep Ensembles демонструють найкращу стійкість до пошкоджень, знижуючи mCE на 11 % порівняно з baseline. Диверсифікація прогнозів між членами ансамблю забезпечує стійкість до різних типів шуму та спотворень.

BNN показують високий AUROC для OOD детекції (0.813), що підтверджує теоретичні переваги байєвського підходу для оцінки епістемічної невизначеності. Проте це не транслюється в кращу точність на корумпованих даних.

Evidential DL несподівано показує гірші результати на OOD, ніж простий baseline з temperature scaling. Це узгоджується з недавніми дослідженнями про схильність методу до overconfidence на незнайомих даних.

## Практичні рекомендації

На основі проведених експериментів сформульовано практичні рекомендації для різних сценаріїв застосування, які враховують специфічні вимоги та обмеження кожної області.

### Критичні застосування в медицині та системах автопілоту

Для критично важливих застосувань, де помилки можуть мати серйозні наслідки для здоров'я чи безпеки людей, рекомендується використовувати комплексний підхід на основі Deep Ensemble з п'яти моделей у поєднанні з temperature scaling та механізмом selective prediction. Така конфігурація забезпечує максимальну якість калібрування та стійкість системи, що повністю виправдовує додаткові обчислювальні витрати в контексті критичності застосування.

Оптимальна конфігурація передбачає використання п'яти моделей з різною ініціалізацією ваг та різними стратегіями аугментації даних для забезпечення різноманітності прогнозів. Temperature scaling слід застосовувати на окремій валідаційній вибірці, яка не використовувалася під час навчання моделей. Механізм selective prediction реалізується через встановлення порогу відхилення на основі ентропії прогнозу або дисперсії між моделями ансамблю. Важливо калібрувати цей поріг таким чином, щоб досягти цільового рівня recall, який визначається вимогами конкретного застосування.

### Системи реального часу

У випадку систем, що працюють у реальному часі та мають жорсткі обмеження на латентність, використання ансамблів або MC Dropout стає неможливим через їхні обчислювальні вимоги. Для таких сценаріїв рекомендується застосування єдиної моделі з temperature scaling та додатковою гілкою для прогнозування впевненості (learned confidence branch). Такий підхід додає менше п'яти відсотків до часу інференсу, що зазвичай є прийнятним для більшості real-time застосувань.

Архітектура включає базову модель з додатковою auxiliary head, яка спеціалізується на оцінці невизначеності прогнозів основної мережі. Temperature scaling залишається обов'язковим етапом пост-обробки для покращення калібрування. Для оптимізації продуктивності варто реалізувати кешування temperature parameter, що дозволяє уникнути повторної калібровки при кожному запуску системи.

### Застосування з обмеженими ресурсами на edge пристроях

Edge пристрої накладають суттєві обмеження як на обчислювальні ресурси, так і на обсяг доступної пам'яті. Для таких умов оптимальним рішенням є використання knowledge distillation від ансамблю моделей у поєднанні з temperature scaling. Дистиляція дозволяє ефективно перенести знання про невизначеність від складного ансамблю в одну компактну модель, придатну для розгортання на пристроях з обмеженими ресурсами.

Процес починається з навчання teacher ансамблю, що складається з трьох-п'яти повнорозмірних моделей. Далі відбувається дистиляція знань у student модель, яка містить приблизно половину параметрів від оригінальної

архітектури. Використання soft targets з temperature рівним трьом забезпечує кращий перенос інформації про невизначеність від teacher до student моделі. Після завершення дистиляції необхідно провести фінальне temperature scaling для student моделі на окремій валідаційній вибірці.

#### Дослідницькі проекти

Для дослідницьких проектів, де точність оцінки невизначеності має першорядне значення, а обчислювальні обмеження є менш критичними, рекомендується гібридний підхід, що поєднує Bayesian Neural Networks з Deep Ensemble. Така комбінація дозволяє отримати найповнішу картину невизначеності, оскільки BNN ефективно моделюють епістемічну невизначеність (невизначеність моделі), тоді як ансамблі краще справляються з алеаторною невизначеністю (невизначеністю даних).

Цей підхід особливо корисний при дослідженні нових доменів або розробці інноваційних методів, де важливо мати глибоке розуміння всіх джерел невизначеності в системі. Хоча такий метод вимагає значних обчислювальних ресурсів та експертизи для правильної імплементації, він забезпечує найвищу якість оцінки невизначеності серед усіх розглянутих підходів.

#### Висновки

Проведене дослідження дозволяє зробити наступні висновки щодо калібрування невизначеності в системах комп'ютерного зору:

1. Deep Ensembles залишаються золотим стандартом для задач, де точність та калібрування є критичними. Незважаючи на обчислювальні витрати, вони забезпечують найкращий баланс між точністю прогнозування (покращення на 0.8–1.8 %), якістю калібрування ( $ECE < 2\%$ ) та робастністю до domain shift (mCE знижується на 11 %).

2. Temperature scaling є обов'язковим компонентом будь-якої продакшн системи. Простота імплементації (один скалярний параметр) та відсутність додаткових витрат під час інференсу роблять цей метод безальтернативним для базового калібрування. Зниження ECE на 60–75 % досягається без втрати точності.

3. Баєсівські методи поки не виправдовують теоретичних очікувань у практичних застосуваннях. BNN демонструють конкурентне калібрування та хорошу OOD детекцію (AUROC 0.813), але поступаються за точністю на 0.5–1.2 % та вимагають складної імплементації. Їх перевага проявляється переважно в low-data режимах та для оцінки епістемічної невизначеності.

4. MC Dropout не рекомендується для продакшн систем через неприйнятне співвідношення якості/швидкості. 20-кратне збільшення часу інференсу не виправдовується помірним покращенням калібрування. Метод може бути корисним лише для швидкого прототипування на існуючих моделях.

5. Evidential Deep Learning потребує подальших досліджень перед практичним застосуванням. Нестабільність на OOD даних та тенденція до overconfidence обмежують його використання, незважаючи на привабливу обчислювальну ефективність.

6. Оптимальна конфігурація для більшості застосувань: ансамбль з 3–5 моделей + temperature scaling + selective prediction за порогом невизначеності. Це забезпечує 90 % від максимально можливого покращення при прийнятних обчислювальних витратах (3–5х збільшення часу інференсу).

Практична значимість роботи полягає у формулюванні конкретних рекомендацій для різних сценаріїв застосування – від критичних систем до edge пристроїв. Результати демонструють, що проблема калібрування невизначеності має ефективні рішення, готові до впровадження в реальні системи комп'ютерного зору.

Подальші дослідження мають зосередитися на масштабуванні методів для transformer архітектур, розробці ефективних гібридних підходів та теоретичному обґрунтуванні емпіричної переваги ансамблів.

#### Список використаної літератури

1. Guo C., Pleiss G., Sun Y., Weinberger K. Q. On calibration of modern neural networks. 34th International Conference on Machine Learning. 2017. P. 1321–1330.
2. Blundell C., Cornebise J., Kavukcuoglu K., Wierstra D. Weight uncertainty in neural networks. 32nd International Conference on Machine Learning. 2015. P. 1613–1622.
3. Louizos C., Welling M. Multiplicative normalizing flows for variational Bayesian neural networks. 34th International Conference on Machine Learning. 2017. P. 2218–2227.
4. Wenzel F., Roth K., Veeling B. S., Światkowski J., Tran L., Mandt S. et al. How good is the Bayes posterior in deep neural networks really? Advances in Neural Information Processing Systems. 2020. Vol. 33. P. 5833–5845.
5. Lakshminarayanan B., Pritzel A., Blundell C. Simple and scalable predictive uncertainty estimation using deep ensembles. Advances in Neural Information Processing Systems. 2017. Vol. 30. P. 6402–6413.
6. Fort S., Jastrzebski S. Deep ensembles: A loss landscape perspective. arXiv preprint arXiv:1912.11370. 2019.
7. Ovadia Y., Fertig E., Ren J., Nado Z., Sculley D., Nowozin S. et al. Can you trust your model's uncertainty? Evaluating predictive uncertainty under dataset shift. Advances in Neural Information Processing Systems. 2019. Vol. 32.
8. Gal Y., Ghahramani Z. Dropout as a Bayesian approximation: Representing model uncertainty in deep learning. 33rd International Conference on Machine Learning. 2016. P. 1050–1059.

9. Osband I. Risk versus uncertainty in deep learning: Bayes, bootstrap and the dangers of dropout. NIPS Workshop on Bayesian Deep Learning. 2016.
10. Mukhoti J., Gal Y. Evaluating uncertainty for object detection using Monte Carlo dropout. International Conference on Computer Vision: Workshop on Uncertainty and Robustness. 2018.
11. Sensoy M., Kaplan L., Kandemir M. Evidential deep learning to quantify classification uncertainty. The Journal of Machine Learning Research. 2018. Vol. 19. No. 1. P. 962–1009.
12. Amini A., Schwarting W., Soleimany A., Rus D. Deep evidential regression. Advances in Neural Information Processing Systems. 2020. Vol. 33. P. 14927–14937.
13. Kopetzki A., Chen Y., Slesarev V. Rethinking evidential deep learning. arXiv preprint arXiv:2110.03687. 2021.
14. Kumar A., Liang P. S., Ma T. Verified uncertainty calibration. Advances in Neural Information Processing Systems. 2019. Vol. 32. P. 3787–3798.
15. Nixon J., Dusenberry M., Zhang L., Jerfel G., Tran D. Measuring calibration in deep learning. CVPR Workshops. 2019. P. 38–41.
16. Minderer M., Joslin D., Gleave A., Abeles A., Bachem O. Revisiting model calibration: Do confounding factors affect calibration metrics? arXiv preprint arXiv:2106.07998. 2021.

### References

1. Guo, C., Pleiss, G., Sun, Y., & Weinberger, K. Q. (2017). On calibration of modern neural networks. 34th International Conference on Machine Learning, 1321-1330.
2. Blundell, C., Cornebise, J., Kavukcuoglu, K., & Wierstra, D. (2015). Weight uncertainty in neural networks. 32nd International Conference on Machine Learning, 1613-1622.
3. Louizos, C., & Welling, M. (2017). Multiplicative normalizing flows for variational Bayesian neural networks. 34th International Conference on Machine Learning, 2218-2227.
4. Wenzel, F., Roth, K., Veeling, B. S., Światkowski, J., Tran, L., Mandt, S.,... & Snoek, J. (2020). How good is the Bayes posterior in deep neural networks really? Advances in Neural Information Processing Systems, 33, 5833–5845.
5. Lakshminarayanan, B., Pritzel, A., & Blundell, C. (2017). Simple and scalable predictive uncertainty estimation using deep ensembles. Advances in Neural Information Processing Systems, 30, 6402–6413.
6. Fort, S., & Jastrzebski, S. (2019). Deep ensembles: A loss landscape perspective. arXiv preprint arXiv:1912.11370.
7. Ovadia, Y., Fertig, E., Ren, J., Nado, Z., Sculley, D., Nowozin, S.,... & Wendt, M. (2019). Can you trust your model's uncertainty? Evaluating predictive uncertainty under dataset shift. Advances in Neural Information Processing Systems, 32, 13991–14002.
8. Gal, Y., & Ghahramani, Z. (2016). Dropout as a Bayesian approximation: Representing model uncertainty in deep learning. 33rd International Conference on Machine Learning, 1050-1059.
9. Osband, I. (2016). Risk versus uncertainty in deep learning: Bayes, bootstrap and the dangers of dropout. NIPS Workshop on Bayesian Deep Learning.
10. Mukhoti, J., & Gal, Y. (2018). Evaluating uncertainty for object detection using Monte Carlo dropout. International Conference on Computer Vision: Workshop on Uncertainty and Robustness.
11. Sensoy, M., Kaplan, L., & Kandemir, M. (2018). Evidential deep learning to quantify classification uncertainty. The Journal of Machine Learning Research, 19(1), 962–1009.
12. Amini, A., Schwarting, W., Soleimany, A., & Rus, D. (2020). Deep evidential regression. Advances in Neural Information Processing Systems, 33, 14927–14937.
13. Kopetzki, A., Chen, Y., & Slesarev, V. (2021). Rethinking evidential deep learning. arXiv preprint arXiv:2110.03687.
14. Kumar, A., Liang, P. S., & Ma, T. (2019). Verified uncertainty calibration. Advances in Neural Information Processing Systems, 32, 3787–3798.
15. Nixon, J., Dusenberry, M., Zhang, L., Jerfel, G., & Tran, D. (2019). Measuring calibration in deep learning. CVPR Workshops, 38–41.
16. Minderer, M., Joslin, D., Gleave, A., Abeles, A., & Bachem, O. (2021). Revisiting model calibration: Do confounding factors affect calibration metrics? arXiv preprint arXiv:2106.07998.

Дата першого надходження рукопису до видання: 24.09.2025  
 Дата прийнятого до друку рукопису після рецензування: 21.10.2025  
 Дата публікації: 28.11.2025