

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ «ОДЕСЬКА ПОЛІТЕХНІКА»
НАВЧАЛЬНО-НАУКОВИЙ
ІНСТИТУТ КОМП'ЮТЕРНИХ СИСТЕМ
КАФЕДРА ПРИКЛАДНОЇ МАТЕМАТИКИ ТА
ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ

КОНСПЕКТ ЛЕКЦІЙ
з дисципліни «ТЕОРІЯ РОЗПІЗНАВАННЯ ОБРАЗІВ»
для здобувачів першого (бакалаврського) рівня вищої освіти
за спеціальністю
124 – Системний аналіз
113 – Прикладна математика

Одеса – 2025

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ «ОДЕСЬКА ПОЛІТЕХНІКА»
НАВЧАЛЬНО-НАУКОВИЙ
ІНСТИТУТ КОМП'ЮТЕРНИХ СИСТЕМ
КАФЕДРА ПРИКЛАДНОЇ МАТЕМАТИКИ ТА
ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ

КОНСПЕКТ ЛЕКЦІЙ
з дисципліни «ТЕОРІЯ РОЗПІЗНАВАННЯ ОБРАЗІВ»
для здобувачів першого (бакалаврського) рівня вищої освіти
за спеціальністю
124 – Системний аналіз
113 – Прикладна математика

Затверджено
на засіданні кафедри ПМІТ
Протокол № 1 від 30.08.2024 р.

Одеса – 2025

Конспект лекцій з дисципліни «Теорія розпізнавання образів» для здобувачів першого (бакалаврського) рівня вищої освіти за спеціальністю 124 – Системний аналіз (освітньо-професійна програма: «Системний аналіз та наука про дані» – 2025) та за спеціальністю 113 – Прикладна математика (освітньо-професійна програма: «Математичне забезпечення комп’ютерних систем» – 2025)/ Укл.: *М. В. Полякова*. – Одеса: Одеська політехніка, 2025. – 39 с.

Укладач: **Полякова М. В.**, докт. техн. наук, доц.

ЛЕКЦІЯ 1. Класифікація інтелектуальних завдань. Класифікація інтелектуальних систем.	6
ЛЕКЦІЯ 2. Структурна схема систем цифрового оброблення зображень. Основні процедури систем цифрового оброблення зображень.	7
Лекція 3. Моделі перепадів інтенсивності зображень. Диференціальні методи контурної сегментації зображень.	9
Лекція 4. Кореляційно-екстремальні методи контурної сегментації зображень.	12
Лекція 5. Метод Кані контурної сегментації зображень.	13
Лекція 6. Метод активних контурів. Зовнішня та внутрішня енергія контуру. Уточнення контурів будівель методом активних контурів.	15
Лекція 7. Метод нарощування областей сегментації зображень. Метод злиття-зщеплення сегментації зображень.	16
Лекція 8. Ознаки текстурних зображень, що обчислюються за гістограмою інтенсивності.	19
Лекція 9. Матриця сполучуваності інтенсивностей зображень. Ознаки текстурних зображень, що обчислюються за матрицею сполучуваності інтенсивностей.	20
Лекція 10. Подання зображення спрямованими відрізками та деревами.	23
Лекція 11. Поняття дерева рішень та вирішального правила. Контроль складності дерева рішень.	28
Лекція 12. Показник ентропії для побудови вирішального правила дерева рішень. Індекс Джині та CART-міра для побудови вирішального правила дерева рішень.	29
Лекція 13. Оцінка якості класифікації за допомогою матриці помилок, показників ассурасу, error rate.	31
Лекція 14. Оцінка якості класифікації за допомогою precision, recall, F-міри.	33
Лекція 15. Показники якості бінарної класифікації.	35
Рекомендована література	38

Вступ

Аналіз даних є важливою складовою сучасної науки, бізнесу та інформаційних технологій, оскільки дозволяє перетворювати великі обсяги інформації на корисні знання. Умови стрімкого розвитку цифрових технологій зумовлюють зростання ролі методів збору, обробки та інтерпретації даних для прийняття обґрунтованих рішень у різних сферах діяльності людини.

Метою аналізу даних є виявлення закономірностей, тенденцій та прихованих зв'язків у наборах даних за допомогою статистичних, математичних та обчислювальних методів. Аналіз даних поєднує елементи математичної статистики, теорії ймовірностей, машинного навчання та візуалізації, що забезпечує комплексний підхід до дослідження інформації.

Даний конспект лекцій спрямований на формування теоретичних знань і практичних навичок роботи з даними, починаючи від їх попередньої обробки та аналізу, до побудови моделей і інтерпретації отриманих результатів. У матеріалах курсу розглядаються основні поняття аналізу даних, методи дослідження та приклади їх застосування у реальних задачах.

ЛЕКЦІЯ 1. КЛАСИФІКАЦІЯ ІНТЕЛЕКТУАЛЬНИХ ЗАВДАНЬ. КЛАСИФІКАЦІЯ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ.

1. Класифікація інтелектуальних завдань.
2. Класифікація інтелектуальних систем.

1. Класифікація інтелектуальних завдань.

Інтелектуальні завдання та інтелектуальні системи є важливими складовими сучасних інформаційних технологій і тісно пов'язані з розвитком штучного інтелекту та аналізу даних. Інтелектуальними називають такі завдання, для розв'язання яких необхідне використання елементів людського мислення, зокрема аналізу, узагальнення, навчання та прийняття рішень в умовах невизначеності. До основних класів інтелектуальних завдань належать задачі розпізнавання та класифікації об'єктів, прогнозування майбутніх станів або значень, діагностики технічних і біологічних систем, а також задачі прийняття рішень і оптимізації. Залежно від умов постановки такі завдання можуть бути детермінованими, стохастичними або нечіткими, а за способом розв'язання — такими, що передбачають навчання на основі даних, або такими, що ґрунтуються на заздалегідь заданих правилах і моделях.

2. Класифікація інтелектуальних систем.

Інтелектуальні системи широко застосовуються в різних галузях діяльності людини та охоплюють широкий спектр задач — від аналізу даних до керування складними технічними об'єктами. Одним із класичних прикладів інтелектуальних систем є експертні системи, які імітують процес мислення фахівця в певній предметній області. Такі системи використовуються, наприклад, у медицині для підтримки діагностики захворювань, у технічній діагностиці для виявлення несправностей обладнання, а також у фінансовій сфері для оцінювання ризиків.

Іншим поширеним прикладом є системи розпізнавання образів, що застосовуються для оброблення зображень і відео. До них належать системи розпізнавання облич, номерних знаків, рукописного тексту або медичних знімків. Такі інтелектуальні системи базуються переважно на методах машинного навчання та нейронних мереж і здатні автоматично навчатися на великих обсягах даних.

Важливе місце займають системи підтримки прийняття рішень, які аналізують дані та пропонують оптимальні або рекомендовані варіанти дій. Вони використовуються в економіці, управлінні підприємствами, логістиці та плануванні, де необхідно враховувати велику кількість факторів і обмежень. Подібні системи допомагають людині швидше та обґрунтованіше приймати рішення.

До прикладів інтелектуальних систем також належать інтелектуальні інформаційно-пошукові системи, які здійснюють пошук і аналіз інформації з урахуванням змісту запиту, а не лише ключових слів. Такі підходи використовуються у сучасних пошукових системах, рекомендаційних сервісах та системах аналізу текстів. Окрему групу становлять інтелектуальні навчальні системи, що адаптують навчальний процес до рівня знань користувача та забезпечують індивідуальний підхід до навчання.

Крім того, широко застосовуються інтелектуальні керуючі системи, які забезпечують автоматизоване керування складними об'єктами та процесами, наприклад у робототехніці, промисловій автоматизації або транспортних системах. Такі системи здатні реагувати на зміни середовища, прогнозувати наслідки своїх дій і коригувати поведінку в реальному часі. Наведені приклади демонструють різноманіття інтелектуальних систем і їх важливу роль у сучасних інформаційних технологіях.

Наведемо конкретні приклади інтелектуальних систем із зазначенням сфери застосування та їх призначення.

У медицині. Система MYCIN — одна з перших експертних систем, призначена для діагностики бактеріальних інфекцій та рекомендацій щодо лікування. Сучасні системи аналізу медичних зображень, наприклад програмні комплекси для автоматичного виявлення пухлин на МРТ або рентгенівських знімках, використовують нейронні мережі для підтримки лікарських рішень.

У транспорті. Системи автоматичного розпізнавання номерних знаків застосовуються для контролю дорожнього руху, оплати паркування та фіксації порушень. Інтелектуальні системи керування дорожнім рухом аналізують потоки автомобілів і оптимізують роботу світлофорів у реальному часі.

У фінансах. Банківські системи виявлення шахрайства аналізують транзакції клієнтів і автоматично визначають підозрілі операції. Рекомендаційні системи кредитування оцінюють платоспроможність позичальників на основі історичних даних.

У промисловості. Системи предиктивної діагностики обладнання прогнозують можливі відмови верстатів або механізмів на основі даних датчиків. Інтелектуальні системи керування технологічними процесами застосовуються на електростанціях, хімічних і металургійних підприємствах.

В інформаційних технологіях. Пошукові системи, такі як Google Search, використовують інтелектуальні алгоритми для аналізу змісту вебсторінок і формування релевантних результатів. Рекомендаційні системи YouTube або Netflix пропонують користувачам контент відповідно до їхніх уподобань.

В освіті. Адаптивні навчальні платформи, наприклад Duolingo, аналізують успішність користувача та підлаштовують складність завдань. Інтелектуальні системи тестування автоматично оцінюють відповіді та формують індивідуальні траєкторії навчання.

Ці приклади демонструють, що інтелектуальні системи не лише існують у теорії, а й активно використовуються на практиці для розв'язання реальних задач у різних галузях.

Контрольні питання

1. Наведіть приклади інтелектуальних систем в освіті і медицині.
2. Що таке інтелектуальна задача?
3. Що таке інтелектуальна система?
4. Що таке інтелектуальні завдання і які їх основні типи?
5. Які класи інтелектуальних систем існують і за якими ознаками вони класифікуються?
6. Наведіть приклади систем для різних типів інтелектуальних завдань.

ЛЕКЦІЯ 2. СТРУКТУРНА СХЕМА СИСТЕМ ЦИФРОВОГО ОБРОБЛЕННЯ ЗОБРАЖЕНЬ. ОСНОВНІ ПРОЦЕДУРИ СИСТЕМ ЦИФРОВОГО ОБРОБЛЕННЯ ЗОБРАЖЕНЬ.

1. Структурна схема систем цифрового оброблення зображень.
2. Основні процедури систем цифрового оброблення зображень.

1. Структурна схема систем цифрового оброблення зображень.

Системи цифрового оброблення зображень призначені для отримання, перетворення, аналізу та інтерпретації візуальної інформації у цифровій формі.

Узагальнена структурна схема такої системи складається з кількох взаємопов'язаних функціональних блоків. Початковим етапом є пристрій отримання зображення, до якого належать камери, сканери, сенсори або інші засоби реєстрації. На цьому етапі відбувається перетворення оптичного зображення в електричний сигнал. Далі сигнал надходить до блоку оцифрування, де здійснюється дискретизація та квантування, внаслідок чого формується цифрове зображення.

Отримане цифрове зображення передається до блоку попередньої обробки, основним призначенням якого є покращення якості даних та усунення завад. Після цього зображення може зберігатися у пам'яті системи або передаватися до блоку оброблення та аналізу, де виконуються складніші алгоритми, пов'язані з виділенням ознак, сегментацією та розпізнаванням об'єктів. Завершальним елементом структурної схеми є блок виведення результатів, який забезпечує відображення зображення на екрані, передачу даних іншим системам або формування керуючих сигналів для виконавчих пристроїв. У сучасних системах також важливу роль відіграє блок керування, що координує роботу всіх компонентів.

2. Основні процедури систем цифрового оброблення зображень.

Основні процедури цифрового оброблення зображень можна поділити на кілька послідовних груп. Першою групою є процедури попередньої обробки, які включають фільтрацію шумів, корекцію яскравості та контрасту, нормалізацію, геометричні перетворення та виправлення спотворень. Ці процедури спрямовані на підвищення якості зображення та підготовку його до подальшого аналізу.

Другою важливою групою є процедури сегментації, що полягають у розбитті зображення на окремі області або об'єкти відповідно до певних критеріїв, таких як яскравість, колір або текстура. Після сегментації виконуються процедури виділення ознак, у процесі яких обчислюються характеристики об'єктів, наприклад форма, розміри, контури чи статистичні параметри. Наступним етапом є аналіз та розпізнавання, що включає класифікацію об'єктів і прийняття рішень на основі отриманих ознак. Завершальними є процедури інтерпретації та візуалізації результатів, які забезпечують зрозуміле подання інформації користувачеві або передачу її іншим інформаційним чи керуючим системам.

Таким чином, структурна схема та основні процедури систем цифрового оброблення зображень утворюють логічно послідовний процес перетворення первинної візуальної інформації у корисні дані, що можуть бути використані для аналізу, прийняття рішень і автоматизації різноманітних прикладних задач.

3. Приклади систем цифрового оброблення зображень.

Ось конкретні приклади систем цифрового оброблення зображень із коротким поясненням їх призначення та застосування.

У медицині широко використовуються системи цифрової обробки медичних зображень, зокрема програмні комплекси для аналізу рентгенівських знімків, комп'ютерної томографії та магнітно-резонансної томографії. Наприклад, системи автоматичного виявлення пухлин або патологій на МРТ-зображеннях виконують фільтрацію шумів, сегментацію тканин і виділення ознак для подальшого аналізу лікарем.

У промисловості застосовуються системи технічного зору для контролю якості продукції. Такі системи використовуються на конвеєрних лініях для виявлення дефектів поверхні, перевірки геометричних розмірів деталей або правильності складання виробів. Камери отримують зображення, а програмне забезпечення аналізує їх у реальному часі та приймає рішення про відповідність виробу стандартам.

У транспортній сфері поширеними є системи розпізнавання номерних знаків та дорожніх ситуацій. Вони застосовуються для автоматичного контролю руху, систем оплати паркування та безпеки дорожнього руху. Такі системи виконують попередню обробку зображень, виділення символів і їх розпізнавання.

У безпекових системах використовуються системи відеоспостереження з функціями аналізу зображень. Вони здатні виявляти рух, розпізнавати обличчя або відстежувати підозрілу поведінку. Обробка зображень у таких системах дозволяє автоматизувати моніторинг великих територій.

У космічних і геоінформаційних технологіях системи цифрової обробки супутникових зображень застосовуються для аналізу стану земної поверхні, прогнозування погодних умов, моніторингу лісів і сільськогосподарських угідь. Тут використовуються методи корекції спотворень, класифікації зображень і виділення об'єктів.

У побутовій техніці та мобільних пристроях системи цифрового оброблення зображень реалізовані в камерах смартфонів. Вони виконують автоматичне покращення якості фото, стабілізацію, розпізнавання сцен і облич, а також корекцію освітлення та кольорів.

Контрольні питання

1. Які існують системи цифрової обробки зображень?
2. Що таке сегментація зображень?
3. Наведіть приклади застосування систем цифрової обробки зображень.
4. Які основні компоненти входять до структури систем цифрового оброблення зображень?
5. Що таке процедури передобробки та обробки зображень і які вони включають?
6. Наведіть приклади типових операцій цифрової обробки зображень.

Лекція 3. Моделі перепадів інтенсивності зображень. Диференціальні методи контурної сегментації зображень.

1. Моделі перепадів інтенсивності зображень.
2. Диференціальні методи контурної сегментації зображень.

1. Моделі перепадів інтенсивності зображень.

Моделі перепадів інтенсивності зображень використовуються для опису різких або поступових змін яскравості на зображенні. Такі перепади відповідають межах об'єктів, контурам, тіням або змінам матеріалу поверхні. У цифровій обробці зображень ці моделі лежать в основі методів виділення контурів і сегментації.

1. Ступінчаста (ідеальна) модель перепаду. У цій моделі інтенсивність зображення змінюється різко — майже миттєво — з одного рівня на інший. Такий перепад описується як стрибок функції яскравості.

Приклад: межа між чорним і білим прямокутником на синтетичному зображенні; чіткий контур надрукованого тексту на світлому папері. Ця модель є спрощеною, але зручною для теоретичного аналізу та розробки операторів виявлення країв.

2. Рамкова (лінійна) модель перепаду. Перепад інтенсивності відбувається не в одній точці, а на певній ділянці — між двома близькими рівнями яскравості формується “рамка”. Такий перепад має два краї: зростання і спад інтенсивності.

Приклад: світла смуга на темному фоні, наприклад дорожня розмітка або тонка тріщина на поверхні деталі. Ця модель добре описує вузькі об'єкти або лінійні структури.

3. Похила (лінійно-змінна) модель перепаду. Інтенсивність змінюється поступово, майже лінійно, на деякому інтервалі. Такий перепад виникає через розмиття, обмежену роздільну здатність камери або плавне освітлення.

Приклад: край об'єкта на фотографії, зроблений з невеликим фокусним розмиванням; перехід від світла до тіні на гладкій поверхні. Ця модель є більш реалістичною для реальних зображень.

4. Сигмоїдна (S-подібна) модель перепаду. Перепад інтенсивності має плавну S-подібну форму: спочатку повільна зміна, потім різкіша, а далі знову повільна. Така модель часто використовується для опису реальних країв з урахуванням шуму та оптичних спотворень.

Приклад: контури об'єктів на медичних зображеннях (МРТ, КТ), де яскравість тканин змінюється поступово.

5. Модель перепаду з шумом. У реальних умовах будь-який перепад інтенсивності супроводжується випадковими коливаннями яскравості — шумом. Тому практичні моделі враховують наявність шуму, який ускладнює точне визначення краю.

Приклад: нічні знімки з камери спостереження або супутникові зображення з атмосферними завадами.

Отже, моделі перепадів інтенсивності дозволяють формалізувати поняття краю на зображенні та вибрати відповідні методи його виявлення. Ідеальні моделі застосовуються переважно в теорії, тоді як похилі та сигмоїдні моделі краще описують реальні зображення, отримані цифровими сенсорами.

2. Диференціальні методи контурної сегментації зображень.

Диференціальні методи контурної сегментації зображень ґрунтуються на аналізі просторових похідних яскравості зображення. Основна ідея цих методів полягає у тому, що контури об'єктів відповідають ділянкам різкого перепаду інтенсивності, де значення похідних досягають максимумів або мають характерні зміни знака.

У межах диференціального підходу найчастіше використовуються похідні першого та другого порядку. Похідні першого порядку дозволяють оцінити швидкість зміни яскравості, тоді як похідні другого порядку характеризують кривизну цієї зміни. Для цифрових зображень похідні обчислюються за допомогою дискретних різницевих операторів.

Методи, що базуються на першій похідній (градієнті), визначають контури як точки або області, де модуль градієнта яскравості має великі значення. Напрямок градієнта вказує напрямок найшвидшого зростання інтенсивності, а сам контур розташовується перпендикулярно до цього напрямку. До найвідоміших градієнтних операторів належать оператори Робертса, Прюїтта та Собеля. Вони використовують локальні маски згортки для наближеного обчислення похідних за координатами. Градієнтні методи є відносно простими в реалізації та ефективними для зображень із чіткими межами, однак чутливі до шуму.

Методи, що використовують другу похідну, визначають контури як точки зміни знака другої похідної, так звані нульові перетини. Найпоширенішим оператором цього класу є оператор Лапласа. Оскільки друга похідна ще більш чутлива до шуму, на практиці часто застосовують оператор Лапласа від гаусівського згладжування, який поєднує попереднє фільтрування з пошуком нульових перетинів. Такі методи дозволяють отримати тонкі контури, але можуть реагувати на дрібні коливання яскравості.

Таким чином, диференціальні методи контурної сегментації базуються на математичному аналізі змін яскравості зображення та широко застосовуються в задачах

комп'ютерного зору. Вибір конкретного методу залежить від якості зображення, рівня шуму та вимог до точності й стійкості виділених контурів.

3. Застосування методу контурної сегментації зображень, що використовують другу похідну

Приклад: Виявлення дефектів на промислових металевих деталях.

Задача: Автоматично виявити тріщини та вм'ятини на поверхні металевих деталей за допомогою цифрових зображень.

1. Попередня обробка зображення.

Реальні знімки часто містять шум від освітлення або камери. Тому першим кроком є згладжування зображення. Використовується гаусівський фільтр, який розмиває високочастотні шуми, зберігаючи загальні перепади інтенсивності.

2. Обчислення другої похідної (Лапласа).

Лапласиан визначає кривизну зміни яскравості і дозволяє виявляти перепади інтенсивності, тобто контури. Для дискретного зображення використовується маска Лапласа, наприклад:

$$\begin{bmatrix} 0 & 1 & 0 \\ 1 & -4 & 1 \\ 0 & 1 & 0 \end{bmatrix}$$

або більш чутлива 8-сусідська маска:

$$\begin{bmatrix} 1 & 1 & 1 \\ 1 & -8 & 1 \\ 1 & 1 & 1 \end{bmatrix}$$

Результатом є карта другої похідної з позитивними та негативними значеннями.

3. Виявлення нульових перетинів.

Контури об'єктів знаходяться там, де друга похідна змінює знак (нульові перетини). Пікселі, де відбувається перехід від позитивного до негативного значення (або навпаки), позначаються як контурні точки. У нашому прикладі це дозволяє виділити тонкі тріщини, які майже не помітні при прямому аналізі градієнта.

4. Порогове відсіювання. Нульові перетини, що відповідають слабким перепадам від шуму, видаляються за допомогою порогового значення. Задается мінімальна амплітуда другої похідної, нижче якої точки не вважаються реальними контурами.

5. Результат. Отримана карта контурів показує чіткі межі тріщин, вм'ятин та інших дефектів на металевій поверхні. Ця інформація може використовуватися для автоматичного контролю якості, сигналізації про браковані деталі або для подальшого аналізу.

Плюси методу: дозволяє виділяти тонкі контури навіть у шумних зображеннях, добре підходить для дефектоскопії. Мінуси: дуже чутливий до шуму, тому перед застосуванням обов'язково потрібне згладжування; можуть виділятися дрібні помилкові контури.

Контрольні питання

1. Що таке перепад інтенсивності на зображенні?
2. Які етапи притаманні диференційним методам контурної сегментації?
3. Навести математичні моделі перепадів інтенсивності.
4. Як працюють диференціальні методи детектування контурів?

5. Назвіть основні оператори, що застосовуються для диференціальної контурної сегментації.

Лекція 4. Кореляційно-екстремальні методи контурної сегментації зображень.

1. Кореляційно-екстремальні методи контурної сегментації зображень
2. Приклад шаблонного порівняння для виділення контурів.

1. Кореляційно-екстремальні методи контурної сегментації зображень.

Кореляційно-екстремальні методи контурної сегментації зображень — це група методів комп'ютерного зору, які виділяють контури об'єктів, використовуючи кореляцію між зображенням і шаблонами та пошук екстремумів (максимумів/мінімумів) певних функцій як критеріїв наявності меж. Контур об'єкта відповідає різкій зміні яскравості, кольору або текстури.

Методи цього класу обчислюють кореляційну міру між локальною областю зображення та заданим шаблоном (маскою, фільтром), шукають екстремуми (піки або спади) відгуку цієї міри, екстремуми інтерпретуються як положення контурів.

Методи порівняння з шаблонами для виділення контурів — це підхід у комп'ютерному зорі, за якого контури виявляються шляхом зіставлення локальних фрагментів зображення із заздалегідь заданими шаблонами (масками), що описують характерні переходи яскравості на межах об'єктів:

1. Задається набір шаблонів (фільтрів), які відповідають різним типам контурів (горизонтальні, вертикальні, діагональні, криволінійні).
2. Шаблон послідовно накладається на зображення.
3. Обчислюється міра подібності між шаблоном і локальною ділянкою.
4. Максимальні (або мінімальні) значення відгуку вказують на наявність контуру.

Переваги: зрозумілий і керований підхід, ефективний для контурів відомої форми, простота апаратної реалізації.

Недоліки: потребує попереднього знання форми контуру, погано працюють зі складними або зашумленими зображеннями, велика кількість шаблонів і високі обчислювальні витрати.

2. Приклад шаблонного порівняння для виділення контурів.

Цей підхід показує класичний метод порівняння з шаблонами, коли контур описується малим еталонним шаблоном, а потім шукаються місця з максимальною подібністю. Приклад коду (Python + OpenCV):

```
import cv2
import numpy as np
import matplotlib.pyplot as plt

# 1. Зчитування зображення
img = cv2.imread("image.jpg", cv2.IMREAD_GRAYSCALE)

# 2. Згладжування
img_blur = cv2.GaussianBlur(img, (5, 5), 0)

# 3. Створення шаблону (вертикальний контур)
template = np.array([
    [-1, -1, 0, 1, 1],
    [-1, -1, 0, 1, 1],
    [-1, -1, 0, 1, 1],
```

```

        [-1, -1,  0,  1,  1],
        [-1, -1,  0,  1,  1]
    ], dtype=np.float32)

# 4. Шаблонне порівняння (нормалізована кореляція)
response = cv2.matchTemplate(
    img_blur.astype(np.float32),
    template,
    cv2.TM_CCORR_NORMED
)

# 5. Нормалізація карти відгуку
response_norm = cv2.normalize(response, None, 0, 255, cv2.NORM_MINMAX)
response_norm = response_norm.astype(np.uint8)

# 6. Порогове виділення контурів
_, edges = cv2.threshold(response_norm, 180, 255, cv2.THRESH_BINARY)

# 7. Приведення до розміру зображення
edges_full = np.zeros_like(img)
h, w = edges.shape
edges_full[:h, :w] = edges

# 8. Візуалізація
plt.figure(figsize=(12,4))
plt.subplot(1,3,1), plt.title("Оригінал"), plt.imshow(img, cmap='gray')
plt.subplot(1,3,2), plt.title("Відгук matchTemplate"),
plt.imshow(response_norm, cmap='gray')
plt.subplot(1,3,3), plt.title("Контури"), plt.imshow(edges_full, cmap='gray')
plt.show()

```

У коді `cv2.TM_CCORR_NORMED` — нормалізована кореляція, значення $\in [-1, 1]$; світлі області на карті відгуку це подібні до шаблону контури; поріг 180 (≈ 0.7) підбирається експериментально.

Контрольні питання

1. Що таке кореляційно-екстремальні методи і на чому вони базуються?
2. Як визначаються контури об'єктів за допомогою цих методів?
3. У яких випадках ці методи ефективні, а коли мають обмеження?

Лекція 5. МЕТОД КАНІ КОНТУРНОЇ СЕГМЕНТАЦІЇ ЗОБРАЖЕНЬ.

1. Етапи методу Кані контурної сегментації зображень.
2. Застосування методу Кані контурної сегментації зображень.

1. Етапи методу Кані контурної сегментації зображень.

Окреме місце серед диференціальних методів займає алгоритм Кенні, який вважається одним із найбільш ефективних методів контурної сегментації. Він поєднує гаусівське згладжування, обчислення градієнта, пригнічення немаксимумів та подвійне порогове відсіювання. Завдяки цьому контури, отримані за допомогою алгоритму Кенні, є тонкими, замкненими та менш чутливими до шуму.

Етапи методу Кані контурної сегментації зображень.

1. Згладжування зображення (Noise Reduction). На зображення застосовується гаусовий фільтр, щоб зменшити шум і дрібні нерівності, які можуть спотворити результат.

2. Обчислення градієнта інтенсивності (Gradient Calculation). Визначаються похідні зображення по горизонталі та вертикалі. Обчислюється модуль градієнта (сила краю) і напрямок градієнта (кут нахилу краю).

3. Неперервне придушення (Non-Maximum Suppression). Виконується звуження країв до одного пікселя завтовшки. Лише локальні максимуми модуля градієнта залишаються, інші пікселі обнуляються.

4. Подвійний поріг (Double Thresholding). Встановлюються два пороги: високий і низький. Пікселі з градієнтом вище високого порогу – сильні краї. Пікселі між низьким і високим порогом – слабкі краї, які можуть бути частиною контуру. Пікселі нижче низького порогу – відкидаються.

5. Відстеження країв (Edge Tracking by Hysteresis). Слабкі пікселі залишаються тільки якщо вони зв'язані з сильними краями. Так забезпечується безперервність контурів і відсікається шум.

Метод Кенні має ряд переваг. Він забезпечує високу точність визначення країв та добре підходить для шумних зображень завдяки попередньому згладжуванню. Метод створює тонкі та неперервні контури, а використання подвійного порогу дозволяє ефективно фільтрувати слабкі шуми, забезпечуючи точну сегментацію об'єктів.

Водночас метод Кенні має й недоліки. Він чутливий до вибору порогів (низького та високого), може бути обчислювально більш затратним, ніж прості оператори, такі як Собель чи Превітт. Крім того, метод може пропускати дуже слабкі краї, якщо пороги підібрані неправильно, і не завжди добре працює на текстурних або складних сценах.

2. Застосування методу Кенні контурної сегментації зображень.

Метод Кенні застосовується для різних практичних задач контурної сегментації зображень.

У картографії та аналізі супутникових знімків його використовують для виділення доріг та ліній. Метод дозволяє чітко визначати краї дорожнього полотна на фоні місцевості, що дає змогу автоматизувати створення карт і трасування дорожньої мережі.

Для сегментації будівель і споруд на аерофотознімках або знімках з дронів метод Кенні допомагає виділити контури дахів і стін, що спрощує автоматизоване створення планів забудови та аналіз урбанізації.

У медичній діагностиці метод використовується для виділення органів або патологічних утворень на МРТ і КТ. Завдяки точному визначенню контурів можна оцінити форму і розмір органів чи новоутворень для діагностики та планування лікування.

У промисловості метод Кенні допомагає контролювати якість деталей: він виділяє контури виробів, дозволяючи автоматично знаходити дефекти або відхилення від стандартної форми.

Нарешті, у робототехніці та системах машинного зору метод застосовується для орієнтації роботів у просторі та розпізнавання об'єктів. Контурна сегментація дозволяє визначати межі предметів для подальшого аналізу та навігації.

Таким чином, метод Кенні ефективний там, де необхідно виділити чіткі та неперервні контури об'єктів на зображенні для подальшого аналізу або автоматизації процесів.

Контрольні питання

1. Назвіть основні етапи методу Кенні.
2. Які переваги та недоліки методу Кенні?
3. Наведіть приклад практичного застосування методу Кенні для виділення контурів.

Лекція 6. МЕТОД АКТИВНИХ КОНТУРІВ. ЗОВНІШНЯ ТА ВНУТРІШНЯ ЕНЕРГІЯ КОНТУРУ.

Уточнення контурів будівель методом активних контурів.

1. Етапи методу активних контурів.
2. Зовнішня та внутрішня енергія контуру.
3. Уточнення контурів будівель методом активних контурів.

1. Етапи методу активних контурів.

Метод активних контурів (Active Contours, “snakes”) в обробці зображень зазвичай складається з таких етапів.

Ініціалізація контуру. Задається початковий контур (крива) поблизу об'єкта, який потрібно виділити. Це може бути коло, еліпс або довільна крива. Формується функція енергії, яку контур буде мінімізувати. Вона зазвичай містить: внутрішню енергію, яка забезпечує гладкість і неперервність контуру, зовнішню енергію, яка “тягне” контур до меж об'єкта (градієнти, краї, інтенсивність) і іноді обмежувальні або додаткові сили, наприклад, тиск, форма, знання про об'єкт.

Обчислення сил, що діють на контур. На кожную точку контуру діють сили внутрішні (згладжування, пружність) та зовнішні (від країв, текстур, регіонів).

Ітеративна деформація контуру. Контур поступово змінює форму, рухаючись у напрямку зменшення загальної енергії. Процес відбувається ітеративно.

Критерій зупинки. Алгоритм зупиняється, коли зміни контуру стають дуже малими, або досягнуто максимальну кількість ітерацій.

Отримання фінального контуру. Фінальний контур вважається межею шуканого об'єкта на зображенні.

2. Зовнішня та внутрішня енергія контуру

Внутрішня енергія відповідає за форму самого контуру. Вона прагне зробити контур гладким, без різких зламів; не дозволяє точкам контуру занадто розтягуватись або зближатись; зберігає цілісність і стабільність кривої. Тобто внутрішня енергія “стежить”, щоб контур був рівним і акуратним. Зовнішня енергія пов'язана з зображенням, на якому шукають об'єкт. Вона притягує контур до меж об'єкта; реагує на краї, контрасти, яскравість або текстуру; “штовхає” контур у напрямку інформативних ділянок зображення. Тобто зовнішня енергія “підказує” контуру, де знаходиться об'єкт. Внутрішня енергія не дає контуру стати ламаним або хаотичним. Зовнішня енергія змушує контур рухатися до реальної межі об'єкта. Разом вони дозволяють контуру знайти потрібний об'єкт та зберегти правильну і гладку форму контура.

3. Уточнення контурів будівель методом активних контурів.

Метод активних контурів застосовується для уточнення меж будівель на аерофото- та супутникових знімках після попереднього виділення об'єктів. Його основна ідея полягає в тому, що заданий початковий контур поступово деформується та прилягає до реальних меж будівлі під впливом внутрішніх і зовнішніх сил.

На початковому етапі формується контур, який приблизно окреслює будівлю. Він може бути отриманий за результатами попередньої сегментації, з бінарної маски або заданий у вигляді простого геометричного багатокутника. Цей контур зазвичай не є точним, але знаходиться поблизу реальних меж об'єкта.

Подальше уточнення відбувається за рахунок дії зовнішньої енергії, яка визначається властивостями зображення. Вона реагує на контраст між дахом будівлі та фоном, а також на лінійні краї й різкі перепади яскравості. Під її впливом контур поступово зміщується у напрямку реальних меж будівлі.

Одночасно з цим діє внутрішня енергія контуру, яка відповідає за його форму. Вона забезпечує гладкість і стабільність контуру, усуває дрібні нерівності та зменшує

вплив шумів, тіней і другорядних деталей. Завдяки цьому контур не деформується хаотично і зберігає цілісність.

Процес уточнення є ітеративним: на кожному кроці контур коригується з урахуванням балансу внутрішньої та зовнішньої енергії. Алгоритм завершується тоді, коли зміни контуру стають мінімальними. У результаті отримується точна та стабільна межа будівлі, яка добре відповідає її реальному контуру на зображенні.

Контрольні питання

1. Що таке метод активних контурів (snakes) і як він працює?
2. Чим відрізняються зовнішня і внутрішня енергія контуру?
3. Як застосовується метод активних контурів для уточнення контурів будівель?

Лекція 7. МЕТОД НАРОЩУВАННЯ ОБЛАСТЕЙ СЕГМЕНТАЦІЇ ЗОБРАЖЕНЬ. МЕТОД ЗЛИТТЯ-ЗЩЕПЛЕННЯ СЕГМЕНТАЦІЇ ЗОБРАЖЕНЬ.

1. Метод нарощування областей сегментації зображень.
2. Метод злиття-зщеплення сегментації зображень.

1. Метод нарощування областей сегментації зображень.

Суть методу полягає в тому, що спочатку обираються початкові пікселі (насіння, seed points), які належать до різних об'єктів або областей. Далі алгоритм поступово додає до кожної області сусідні пікселі, якщо їхні характеристики (наприклад, яскравість або колір) є подібними до вже включених у цю область.

На кожному кроці перевіряються сусідні пікселі (зазвичай у 4- або 8-зв'язному оточенні). Якщо різниця між значенням пікселя та середнім значенням області не перевищує заданий поріг, піксель приєднується до області. Процес триває доти, доки жоден новий піксель не може бути доданий.

У результаті зображення розбивається на однорідні області, які добре відповідають реальним об'єктам сцени.

Переваги: простота реалізації, інтуїтивно зрозумілий принцип, добре працює для однорідних ділянок. Недоліки: чутливість до шуму, залежність від вибору початкових точок, необхідність правильного підбору порога подібності.

Наведено простий приклад реалізації методу нарощування областей для сірого зображення.

1. Обирається початковий піксель (seed).
2. Створюється черга для обходу сусідніх пікселів.
3. Якщо різниця яскравості між пікселем і початковою точкою не перевищує поріг, піксель додається до області.
4. У результаті формується бінарна маска сегментованої області.

```
import numpy as np
import cv2
from collections import deque

def region_growing(image, seed, threshold=5):
    """
    image      — cipe зображення (numpy array)
    seed       — початкова точка (x, y)
    threshold  — поріг подібності
    """
    h, w = image.shape
    segmented = np.zeros((h, w), np.uint8)

    seed_value = image[seed]
```



```

queue = deque([seed])
segmented[seed] = 255

# 8-зв'язний окол
neighbors = [(-1, -1), (-1, 0), (-1, 1),
              (0, -1),          (0, 1),
              (1, -1),  (1, 0),  (1, 1)]

while queue:
    x, y = queue.popleft()

    for dx, dy in neighbors:
        nx, ny = x + dx, y + dy

        if 0 <= nx < h and 0 <= ny < w:
            if segmented[nx, ny] == 0:
                if abs(int(image[nx, ny]) - int(seed_value)) <=
threshold:
                    segmented[nx, ny] = 255
                    queue.append((nx, ny))

    return segmented

# Завантаження зображення
image = cv2.imread("image.png", cv2.IMREAD_GRAYSCALE)

# Початкова точка (seed)
seed_point = (100, 150)

# Сегментація
result = region_growing(image, seed_point, threshold=10)

# Відображення результату
cv2.imshow("Original", image)
cv2.imshow("Region Growing", result)
cv2.waitKey(0)
cv2.destroyAllWindows()

```

2. Метод злиття-зщеплення сегментації зображень.

Ідея методу полягає в тому, що зображення спочатку розглядається як єдина область. Якщо ця область є неоднорідною (наприклад, має велику зміну яскравості), вона розщеплюється на менші частини. Зазвичай розщеплення виконується рекурсивно на чотири квадранти (квадродерево).

Після завершення розщеплення виконується етап злиття. На цьому етапі сусідні області аналізуються, і якщо вони є достатньо подібними за своїми характеристиками (середня яскравість, дисперсія тощо), вони об'єднуються в одну область.

Таким чином, метод забезпечує баланс між локальним аналізом (розщеплення) і глобальним узгодженням областей (злиття). У результаті отримується сегментація зображення на однорідні, логічно пов'язані області.

Метод злиття-розщеплення має низку переваг. Він дозволяє сегментувати зображення без попереднього задання початкових точок, оскільки процес починається з аналізу всього зображення як єдиної області. Метод забезпечує ієрархічний підхід до сегментації, що дає змогу аналізувати структуру зображення на різних рівнях деталізації. Він добре працює для зображень, які містять великі однорідні області, та є стійкішим до випадкового шуму порівняно з методами локального нарощування областей.

Водночас метод має й недоліки. Через використання квадродеревного розщеплення форма отриманих сегментів часто обмежується квадратами або прямокутниками, що може погано відповідати реальним контурам об'єктів. Результат

сегментації значною мірою залежить від вибору критерію однорідності та порогових значень. Крім того, алгоритм може бути обчислювально затратним для зображень великого розміру, а на зображеннях зі слабкою однорідністю призводити до надмірного розщеплення або, навпаки, недостатнього злиття областей.

Розглянемо приклад методу розщеплення з подальшим злиттям для сірого зображення. Зображення рекурсивно ділиться на чотири частини, якщо область неоднорідна. Критерієм однорідності є дисперсія (стандартне відхилення) яскравості. Якщо область однорідна, вона замінюється середнім значенням. Злиття тут реалізоване неявно через однакові значення сусідніх областей.

```
import cv2
import numpy as np

# Критерій однорідності області
def is_homogeneous(region, threshold):
    return np.std(region) < threshold

# Рекурсивне розщеплення
def split(image, x, y, size, threshold, result):
    region = image[x:x+size, y:y+size]

    if size <= 8 or is_homogeneous(region, threshold):
        mean_value = int(np.mean(region))
        result[x:x+size, y:y+size] = mean_value
    else:
        half = size // 2
        split(image, x, y, half, threshold, result)
        split(image, x+half, y, half, threshold, result)
        split(image, x, y+half, half, threshold, result)
        split(image, x+half, y+half, half, threshold, result)

# Завантаження зображення
image = cv2.imread("image.png", cv2.IMREAD_GRAYSCALE)
h, w = image.shape

# Підготовка результату
result = np.zeros_like(image)

# Запуск розщеплення
split(image, 0, 0, min(h, w), threshold=10, result=result)

# Відображення
cv2.imshow("Original", image)
cv2.imshow("Split and Merge (simplified)", result)
cv2.waitKey(0)
cv2.destroyAllWindows()
```

Контрольні питання

1. Як працює метод нарощування областей сегментації?
2. Що таке метод злиття-розщеплення і в чому його ідея?
3. Порівняйте переваги та недоліки цих двох методів сегментації.

Лекція 8. Ознаки текстурних зображень, що обчислюються за гістограмою інтенсивності.

1. Ознаки текстурних зображень, що обчислюються за гістограмою інтенсивності.
2. Застосування ознак текстурних зображень, що обчислюються за гістограмою інтенсивності.

1. Ознаки текстурних зображень, що обчислюються за гістограмою інтенсивності.

Для аналізу текстурних зображень часто використовують характеристики, обчислені на основі гістограми інтенсивності пікселів, яка показує, скільки пікселів кожного рівня яскравості присутнє на зображенні. Ці ознаки дозволяють кількісно оцінити текстуру і використовуються для класифікації, сегментації та аналізу зображень.

Середнє значення (Mean). Середнє значення інтенсивності відображає загальну яскравість зображення або його ділянки. Воно допомагає визначити, чи переважають світлі або темні пікселі.

Приклад: на зображенні сонячного неба середня яскравість буде високою, а на зображенні густого лісу — низькою.

Дисперсія (Variance). Дисперсія характеризує розсіяння інтенсивностей пікселів навколо середнього значення, що визначає контрастність текстури. Висока дисперсія свідчить про різноманіття відтінків, низька — про однорідність.

Приклад: текстура трав'яного покриву має високу дисперсію через різні відтінки зеленого, а рівна поверхня стіни — низьку.

Енергетична ознака (Energy, Uniformity). Енергія відображає однорідність розподілу інтенсивностей. Чим більш однорідна область, тим вища енергія.

Приклад: гладка стіна або папір мають високу енергію, тоді як поверхня каменю або листя — нижчу через різні відтінки.

Ентропія (Entropy). Ентропія вимірює ступінь хаотичності або непередбачуваності текстури. Висока ентропія означає наявність багатьох різних відтінків, низька — однорідність.

Приклад: щільний листяний кущ або кам'яна кладка мають високу ентропію, а гладка стіна або папір — низьку.

Коефіцієнт асиметрії (Skewness). Асиметрія показує, чи зміщена гістограма в бік темних або світлих пікселів.

Приклад: зображення сонячного неба з кількома хмарами буде мати позитивне зміщення, бо переважають світлі пікселі; зображення нічного міста може мати негативне зміщення.

Коефіцієнт сплюсненості (Kurtosis). Куртозис визначає гостроту піків гістограми і показує, наскільки значення інтенсивності зосереджені навколо середнього.

Приклад: однорідна поверхня має високий піковий показник (гострий пік гістограми), а текстура листя чи трави — низький, оскільки значень багато і вони розподілені ширше.

2. Застосування ознак текстурних зображень, що обчислюються за гістограмою інтенсивності.

Ознаки текстурних зображень, обчислені за гістограмою інтенсивності пікселів, широко використовуються у різних сферах для кількісного аналізу текстури, виділення однорідних і неоднорідних ділянок та автоматизації прийняття рішень.

У сільському господарстві ці ознаки допомагають оцінювати стан посівів на супутникових або дрон-знімках. Наприклад, середнє значення інтенсивності і дисперсія дозволяють визначити щільність та здоров'я рослинності на полі. Зони з однорідною та здоровою рослинністю матимуть високу енергію та низьку ентропію, оскільки відтінки зеленого практично не розкидані, а ділянки з нерівномірним ростом або хворобами будуть мати високі дисперсію та ентропію. Це дозволяє автоматично виділяти проблемні ділянки для подальшого догляду або обстеження.

У медичній сфері ознаки текстури використовуються для аналізу медичних знімків, таких як МРТ або КТ. Дисперсія та ентропія допомагають відокремлювати однорідні тканини від патологічних зон, де відтінки інтенсивності змінюються нерівномірно. Наприклад, нормальні тканини печінки або м'язів мають відносно низьку дисперсію та ентропію, тоді як пухлини або запальні процеси створюють ділянки з високою неоднорідністю. Це дозволяє лікарям виділити потенційно проблемні зони для більш детального аналізу та постановки діагнозу.

У промисловості текстурні ознаки використовуються для контролю якості поверхні виробів. Енергетична ознака та дисперсія допомагають визначати рівномірність покриття та однорідність матеріалу. Однорідні поверхні, такі як гладкі металеві пластини або текстильні тканини, мають високу енергію і низьку дисперсію, а дефектні або нерівні поверхні — високу дисперсію та ентропію. Це дозволяє автоматично виявляти браковані деталі на конвеєрі без участі людини.

У урбаністиці та дистанційному зондуванні текстурні ознаки застосовуються для аналізу міських територій та виділення об'єктів на аерофотознімках. Середні значення та дисперсія дозволяють відрізнити будівлі, дороги та зелені зони. Наприклад, будівлі та дороги характеризуються високим контрастом і великою дисперсією інтенсивності, а зелені зони — відносно однорідною текстурою з високою енергією. Це допомагає створювати карти землекористування, автоматично виділяти інфраструктуру та оцінювати стан міських територій.

Таким чином, ознаки текстурних зображень, що обчислюються за гістограмою інтенсивності, дозволяють кількісно описати текстуру, виділити однорідні та неоднорідні області і автоматизувати аналіз зображень у багатьох галузях, від сільського господарства та медицини до промисловості та урбаністики.

Контрольні питання

1. Які основні ознаки текстури обчислюються за гістограмою інтенсивності?
2. Як середні значення, дисперсія та ентропія характеризують текстуру зображення?
3. Наведіть приклади практичного застосування цих ознак.

Лекція 9. Матриця сполучуваності інтенсивностей зображень. Ознаки текстурних зображень, що обчислюються за матрицею сполучуваності інтенсивностей.

1. Матриця сполучуваності інтенсивностей зображень.
2. Ознаки текстурних зображень, що обчислюються за матрицею сполучуваності інтенсивностей.

1. Матриця сполучуваності інтенсивностей зображень.

Матриця сполучуваності інтенсивностей зображень, відома також як матриця спільної зустрічальності рівнів сірого (Gray-Level Co-occurrence Matrix, GLCM), є одним з класичних статистичних методів аналізу текстури цифрових зображень. Вона використовується для опису просторових взаємозв'язків між пікселями та дозволяє формалізувати поняття текстури, яке важко визначити лише за окремими значеннями яскравості.

Основна ідея матриці сполучуваності полягає в аналізі того, як часто пари значень інтенсивності зображення з'являються поруч одна з одною при заданому просторовому співвідношенні. На відміну від звичайної гістограми, яка враховує лише кількість

пікселів з певною яскравістю, матриця сполучуваності враховує просторову організацію зображення, тобто взаємне розташування пікселів.

Для побудови матриці сполучуваності зображення зазвичай переводять у градації сірого та, за необхідності, виконують квантування рівнів яскравості. Квантування означає зменшення кількості можливих значень інтенсивності, наприклад з 256 до 16 або 8 рівнів, що дозволяє зменшити розмір матриці та підвищити обчислювальну ефективність. Кількість рівнів яскравості безпосередньо визначає розмір матриці сполучуваності: якщо використовується N рівнів, то матриця має розмір $N \times N$.

Далі задається просторове співвідношення між пікселями, яке визначається відстанню та напрямком. Відстань зазвичай дорівнює одному або кільком пікселям, а напрямок задається кутом, найчастіше 0° , 45° , 90° або 135° . Напрямок 0° відповідає горизонтальному сусідству пікселів, 90° — вертикальному, а 45° і 135° — діагональним напрямкам. Вибір напрямку є важливим, оскільки текстура зображення може мати анізотропні властивості, тобто залежати від орієнтації.

Кожен елемент матриці сполучуваності $P(i, j)$ показує, скільки разів у зображенні піксель з інтенсивністю i має сусіда з інтенсивністю j на заданій відстані та в заданому напрямку. Таким чином, матриця є двовимірним статистичним розподілом пар значень яскравості. У багатьох випадках матрицю нормалізують, поділяючи всі її елементи на загальну кількість розглянутих пар пікселів, у результаті чого елементи матриці можна інтерпретувати як імовірності.

Матриця сполучуваності може бути симетричною або несиметричною залежно від способу побудови. Якщо враховуються лише пари пікселів у заданому напрямку, матриця буде несиметричною. Якщо ж додатково враховуються пари у протилежному напрямку, матриця стає симетричною, що часто використовується для отримання більш стійких статистичних характеристик текстури.

Завдяки здатності відображати просторові закономірності зображення матриця сполучуваності інтенсивностей широко застосовується в задачах комп'ютерного зору, розпізнавання образів, медичної візуалізації, аналізу супутникових знімків та автоматичної класифікації поверхонь і матеріалів. Вона дозволяє перейти від простого аналізу яскравостей до більш глибокого статистичного опису структури зображення, що робить її потужним інструментом для аналізу текстур.

Приклад матриці сполучуваності інтенсивностей зображення. Маємо зображення 3×3 з рівнями 0 і 1:

```
0 1 1
0 1 0
1 0 0
```

Для напрямку 0° (вправо) та $d = 1$:

GLCM буде 2×2 , її значення — кількість таких сусідніх пар:

```
i/j  0  1
0     1  2
1     2  1
```

2. Ознаки текстурних зображень, що обчислюються за матрицею сполучуваності інтенсивностей.

Сама по собі матриця сполучуваності рідко використовується безпосередньо для аналізу або класифікації зображень. Основну цінність становлять текстурні ознаки, які обчислюються на її основі. До найвідоміших таких ознак належать контраст, який характеризує локальні варіації яскравості; однорідність, що відображає ступінь подібності сусідніх пікселів; енергія, яка показує рівень впорядкованості текстури; кореляція, що описує лінійну залежність між інтенсивностями сусідніх пікселів; а також ентропія, яка є мірою складності та хаотичності текстури.

Ознаки текстурних зображень, що обчислюються за матрицею сполучуваності інтенсивностей, є статистичними характеристиками, які дозволяють кількісно описати структуру та просторову організацію текстури. Матриця сполучуваності інтенсивностей відображає частоти або ймовірності появи пар рівнів яскравості на фіксованій відстані та в певному напрямку, а обчислювані на її основі ознаки узагальнюють цю інформацію у вигляді числових параметрів, зручних для подальшого аналізу, порівняння або класифікації зображень.

Однією з найпоширеніших ознак є контраст. Вона характеризує ступінь локальних змін яскравості в зображенні та відображає різницю між інтенсивностями сусідніх пікселів. Високі значення контрасту відповідають текстурам із різкими переходами яскравості, наприклад, з чітко вираженими краями або грубою структурою. Низький контраст, навпаки, притаманний однорідним та гладким текстурам із поступовими змінами інтенсивності.

Іншою важливою ознакою є однорідність, яку іноді називають інверсною різницею моментів. Вона показує, наскільки близькими є значення інтенсивностей у сусідніх пікселях. Однорідність набуває великих значень у тих випадках, коли матриця сполучуваності зосереджена поблизу головної діагоналі, тобто коли пари пікселів мають однакові або близькі рівні яскравості. Така ознака добре підходить для опису рівномірних і слабоконтрастних текстур.

Енергія, або другий кутовий момент, є мірою впорядкованості текстури. Вона визначає, наскільки концентрованим є розподіл значень у матриці сполучуваності. Висока енергія відповідає регулярним і повторюваним текстурам, у яких домінують певні пари інтенсивностей. Низька енергія характерна для складних або випадкових текстур із більш рівномірним розподілом значень у матриці.

Кореляція описує ступінь лінійної залежності між інтенсивностями сусідніх пікселів. Ця ознака враховує середні значення та дисперсії рівнів яскравості й показує, наскільки передбачуваною є інтенсивність одного пікселя за значенням іншого. Висока кореляція вказує на наявність регулярних структур або напрямних елементів у текстурі, тоді як низька кореляція характерна для нерегулярних або хаотичних зображень.

Ентропія є мірою складності та невпорядкованості текстури. Вона досягає великих значень у випадках, коли матриця сполучуваності має майже рівномірний розподіл, що свідчить про велику різноманітність пар інтенсивностей. Низька ентропія характерна для простих і добре впорядкованих текстур, де кілька типів пар інтенсивностей суттєво переважають над іншими.

Окрім основних, часто використовуються й додаткові ознаки, такі як дисперсія, яка відображає розсіювання значень інтенсивностей; середнє значення та стандартне відхилення, що узагальнюють розподіл матриці; а також максимальна ймовірність, яка визначає найчастіше зустрічальну пару рівнів яскравості. У деяких задачах також аналізуються ознаки асиметрії та міри відмінності між різними напрямками, що дозволяє оцінити анізотропію текстури.

У практичних застосуваннях ознаки, обчислені за матрицею сполучуваності, зазвичай визначаються для кількох напрямків і відстаней, після чого усереднюються або комбінуються в один вектор ознак. Такий підхід підвищує стійкість опису текстури до поворотів та локальних змін зображення. Отриманий вектор текстурних ознак широко використовується в задачах класифікації, сегментації та розпізнавання текстур, забезпечуючи компактний і інформативний опис структури зображення.

Контрольні питання

1. Що таке матриця сполучуваності інтенсивностей?
2. Які ознаки текстури можна обчислити на основі матриці сполучуваності інтенсивностей?
3. Наведіть приклади задач, де застосовують ці ознаки.

Лекція 10. Подання зображення спрямованими відрізками та деревами.

1. Подання зображення спрямованими відрізками.
2. Подання зображення деревами.

1. Подання зображення спрямованими відрізками.

Подання зображення спрямованими відрізками є одним зі структурних методів аналізу зображень, у якому візуальна інформація описується не окремими пікселями, а сукупністю геометричних елементів із заданою орієнтацією. Основна ідея цього підходу полягає в тому, що форма та структура об'єктів на зображенні можуть бути ефективно представлені через лінійні сегменти, які відображають напрямні та контурні властивості зображення.

Процес подання зображення спрямованими відрізками зазвичай починається з попередньої обробки, метою якої є підвищення якості зображення та зменшення впливу шумів. На цьому етапі застосовуються методи згладжування, нормалізації яскравості та виділення країв. Виділення контурів є ключовим кроком, оскільки саме на основі контурної інформації формуються відрізки. Контури відображають місця різких змін інтенсивності, які зазвичай відповідають межах об'єктів або важливим структурним елементам сцени.

Після виділення контурів здійснюється їх апроксимація лінійними сегментами. Безперервні контурні лінії розбиваються на окремі відрізки таким чином, щоб кожен відрізок максимально точно описував локальну частину контуру. Кожен відрізок характеризується координатами початкової та кінцевої точок, довжиною та напрямком. Напрямок відрізка визначається вектором, який вказує орієнтацію сегмента у просторі зображення, або кутом відносно фіксованої осі координат. Спрямованість відрізків дозволяє зберігати інформацію про порядок обходу контуру та локальні геометричні особливості.

Отриманий набір спрямованих відрізків утворює структурний опис зображення, який значно компактніший за піксельне подання. Такий опис добре відображає форму об'єктів, їх геометричні властивості та орієнтацію елементів. Він є менш чутливим до локальних змін яскравості та незначних спотворень, оскільки базується на узагальнених геометричних характеристиках, а не на значеннях інтенсивності кожного пікселя.

Подання зображення спрямованими відрізками також дозволяє ефективно аналізувати напрямні властивості текстур і структур. Наприклад, у зображеннях із вираженою лінійною структурою, таких як дороги, межі полів або технічні креслення, спрямовані відрізки безпосередньо відображають домінуючі напрямки та взаємне розташування елементів. Це спрощує задачі порівняння, класифікації та розпізнавання форм.

Важливою особливістю такого подання є можливість подальшої організації відрізків у більш складні структури, наприклад у ланцюжки або графи, що відображають зв'язки між сусідніми сегментами. Це дозволяє відновлювати цілісні контури об'єктів і аналізувати їх форму на більш високому рівні абстракції. Крім того, на основі спрямованих відрізків можна обчислювати різноманітні ознаки, такі як розподіл орієнтацій, середня довжина сегментів або щільність відрізків у певних напрямках.

Разом із перевагами подання зображення спрямованими відрізками має й певні обмеження. Його ефективність значною мірою залежить від якості виділення контурів і коректності апроксимації. У випадку сильних шумів, розмитих меж або складних текстур можливе формування надлишкових або спотворених відрізків, що ускладнює подальший аналіз. Тому вибір алгоритмів попередньої обробки та критеріїв формування відрізків є критично важливими для практичного застосування цього підходу.

У підсумку, подання зображення спрямованими відрізками є ефективним структурним методом, який дозволяє перейти від низькорівневого піксельного опису до більш інформативного геометричного представлення. Воно широко використовується в задачах аналізу та розпізнавання зображень, де ключову роль відіграє форма, орієнтація та структура візуальних об'єктів.

Приклад подання зображення спрямованими відрізками.

Розглянемо бінарне зображення, у якому зображено просту фігуру — кут:

```
0 0 1 0 0
0 0 1 0 0
1 1 1 1 1
0 0 0 0 0
0 0 0 0 0
```

Одиниці відповідають пікселям об'єкта, нулі — фону. Це зображення містить вертикальну та горизонтальну лінії, що перетинаються.

На першому етапі зображення аналізується для виділення контурів або лінійних структур. У цьому випадку контур уже очевидний і складається з двох основних ліній: вертикальної та горизонтальної.

Далі кожна з цих ліній апроксимується спрямованими відрізками. У результаті отримаємо, наприклад, два відрізки.

Перший відрізок — вертикальний: початкова точка (2, 0), кінцева точка (2, 2), напрямок — зверху вниз, кут орієнтації 90° , довжина 3 пікселі. Другий відрізок — горизонтальний: початкова точка (0, 2), кінцева точка (4, 2), напрямок — зліва направо, кут орієнтації 0° , довжина 5 пікселів.

Таким чином, замість 25 пікселів зображення подається двома спрямованими відрізками, які зберігають ключову інформацію про форму об'єкта.

Приклад подання зображення спрямованими відрізками літери «А».

Розглянемо бінарне зображення, на якому зображена літера «А». Після попередньої обробки та бінаризації зображення можна подати у вигляді сітки пікселів, де одиниці відповідають елементам літери, а нулі — фону. Літера має дві похилі бокові лінії та одну горизонтальну перемичку.

На першому етапі виконується виділення контурів літери. Контур складається з ламаної лінії, що окреслює зовнішню форму символу та внутрішні структурні елементи. Отриманий контур розбивається на окремі ділянки з приблизно сталим напрямком.

На другому етапі кожна така ділянка апроксимується спрямованим відрізком. У результаті літера «А» може бути представлена сукупністю п'яти основних спрямованих відрізків.

Перший відрізок відповідає лівій похилій стороні літери. Він має початкову точку в нижній лівій частині символу та кінцеву точку у верхній частині. Його напрямок спрямований знизу вгору та вправо, а кут орієнтації становить приблизно 60° . Довжина цього відрізка є відносно великою, оскільки він формує основний контур літери.

Другий відрізок описує праву похилу сторону літери. Початкова точка розташована у верхній частині символу, а кінцева — внизу праворуч. Напрямок відрізка спрямований зверху вниз та вправо, з кутом орієнтації близько -60° або 120° відносно осі абсцис. За довжиною він приблизно дорівнює першому відрізку, що відображає симетрію форми.

Третій відрізок відповідає горизонтальній перемичці літери «А». Він починається на лівій похилій стороні та закінчується на правій, має напрямок зліва направо та кут орієнтації 0° . Його довжина менша, ніж у бокових відрізків, що відповідає реальній геометрії символу.

Окрім основних елементів, контур літери може містити допоміжні короткі відрізки у верхній точці з'єднання бокових сторін або у місцях переходу між лініями.

Такі відрізки мають невелику довжину та слугують для точнішого опису локальних змін напрямку контуру.

У підсумку подання літери «А» спрямованими відрізками можна записати як множину структурних елементів, кожен з яких описується координатами початкової та кінцевої точок, довжиною та напрямком. Сукупність цих відрізків утворює компактний структурний опис символу, який зберігає його форму та пропорції без необхідності зберігати інформацію про кожен окремий піксель.

Такий опис дозволяє легко порівнювати літеру «А» з іншими символами, аналізувати її геометричні властивості та виконувати розпізнавання. Наприклад, наявність двох довгих похилих відрізків із симетричними напрямками та одного коротшого горизонтального відрізка є характерною ознакою саме цього символу.

Приклад опису контуру реального об'єкта у вигляді спрямованих відрізків для дорожнього знака «Стоп».

Розглянемо зображення дорожнього знака «Стоп», який має форму правильного восьмикутника. Після попередньої обробки зображення, що включає згладжування, бінаризацію та виділення країв, отримується замкнений контур, який чітко відокремлює знак від фону. Цей контур є основним джерелом інформації про форму об'єкта.

Отриманий контур являє собою ламану лінію, що складається з послідовних прямих ділянок. Кожна така ділянка апроксимується спрямованим відрізком із практично сталим напрямком. У результаті контур знака подається у вигляді восьми спрямованих відрізків, з'єднаних послідовно та утворюючих замкнену структуру.

Перший відрізок відповідає верхній горизонтальній стороні восьмикутника. Його початкова точка розташована у верхньому лівому куті, а кінцева — у верхньому правому. Напрямок відрізка спрямований зліва направо, а кут орієнтації становить 0° . Довжина цього відрізка приблизно дорівнює довжині інших сторін, що відображає регулярність форми знака.

Другий відрізок описує праву верхню похилу сторону. Він має напрямок зверху вниз і вправо з кутом орієнтації близько -45° . Цей відрізок з'єднує верхню горизонтальну сторону з правою вертикальною частиною контуру.

Третій відрізок є майже вертикальним і відповідає правій стороні знака. Його напрямок спрямований зверху вниз, а кут орієнтації близький до -90° . Четвертий відрізок описує праву нижню похилу сторону з напрямком зправа наліво і вниз, маючи кут орієнтації приблизно -135° .

П'ятий відрізок відповідає нижній горизонтальній стороні восьмикутника. Він спрямований справа наліво, а його кут орієнтації становить 180° . Шостий відрізок описує ліву нижню похилу сторону з напрямком знизу вгору і вліво та кутом орієнтації близько 135° .

Сьомий відрізок є майже вертикальним і відповідає лівій стороні знака, маючи напрямок знизу вгору та кут орієнтації близький до 90° . Восьмий, останній, відрізок описує ліву верхню похилу сторону з напрямком зліва направо і вгору та кутом орієнтації близько 45° , замикаючи контур.

Таким чином, контур дорожнього знака «Стоп» подається замкненою послідовністю спрямованих відрізків, кожен з яких характеризується напрямком, довжиною та взаємним положенням. Сукупність цих відрізків утворює компактний геометричний опис форми об'єкта, який легко відрізняється від контурів інших знаків, наприклад трикутних або круглих.

Таке подання дозволяє ефективно розпізнавати дорожні знаки на зображеннях, оскільки форма восьмикутника є унікальною та може бути ідентифікована за кількістю відрізків, їх орієнтаціями та послідовністю з'єднання.

2. Подання зображення деревами.

Подання зображення деревами є структурним підходом до аналізу та опису зображень, у якому візуальна інформація організується у вигляді ієрархічної структури типу дерева. На відміну від піксельного або суто статистичного подання, деревоподібне представлення дозволяє відобразити не лише окремі елементи зображення, а й взаємозв'язки між ними, їх підпорядкованість та рівні деталізації.

Основою для побудови деревоподібного подання зображення зазвичай є результати попередньої обробки, зокрема сегментації або виділення контурів. Зображення розбивається на окремі області, об'єкти або структурні елементи, які надалі розглядаються як вузли дерева. Кожен вузол відповідає певному елементу зображення, а ребра дерева відображають відношення між елементами, наприклад належність частини до цілого, просторову вкладеність або послідовність з'єднання.

Коренем дерева зазвичай є весь об'єкт або сцена загалом. На наступному рівні ієрархії розташовуються основні структурні компоненти, наприклад контури великих об'єктів або великі сегментовані області. Кожен із цих компонентів може, у свою чергу, розгалужуватися на дрібніші елементи, такі як окремі контурні ділянки, відрізки, кути або локальні особливості. Таким чином формується багаторівнева структура, яка відображає організацію зображення від грубого до детального рівня.

Одним із поширених варіантів деревоподібного подання є контурне дерево. У такому дереві кореневий вузол відповідає зовнішньому контуру об'єкта, а дочірні вузли — внутрішнім контурам або отворам, наприклад вікнам у будівлі чи символам усередині знака. Така структура дозволяє природно описувати вкладеність форм і зберігати інформацію про топологію зображення.

Іншим прикладом є дерево спрямованих відрізків, у якому вузлами є лінійні сегменти, а зв'язки між ними відображають їх з'єднання або послідовність у межах контуру. У цьому випадку дерево дозволяє описати форму об'єкта як ієрархію лінійних елементів, що особливо зручно для аналізу складних геометричних структур, таких як рукописні символи або технічні креслення.

Перевагою подання зображення деревами є можливість багаторівневого аналізу. На верхніх рівнях дерева можна оцінювати загальну форму та структуру об'єкта, тоді як на нижчих рівнях — аналізувати дрібні деталі. Це робить деревоподібне подання зручним для задач розпізнавання, оскільки дозволяє виконувати порівняння об'єктів на різних рівнях узагальнення.

Крім того, деревоподібні структури добре підходять для застосування алгоритмів пошуку та зіставлення, таких як порівняння дерев або піддерев. Це дозволяє знаходити подібні об'єкти навіть у разі часткових спотворень або втрати окремих елементів зображення. За відповідного нормування параметрів вузлів дерева можна досягти інваріантності до зсувів, масштабування та поворотів.

Разом із тим, побудова коректного деревоподібного подання зображення є нетривіальним завданням. Вона значною мірою залежить від якості сегментації та правильного визначення структурних зв'язків між елементами. Помилки на ранніх етапах обробки можуть призводити до спотворення всієї структури дерева, що негативно впливає на подальший аналіз.

У підсумку, подання зображення деревами є потужним структурним методом, який дозволяє представити складні зображення у вигляді впорядкованої ієрархії елементів. Воно забезпечує наочний і компактний опис структури об'єктів та широко використовується в задачах комп'ютерного зору, аналізу сцен і розпізнавання образів, де важливу роль відіграють відношення між частинами зображення, а не лише їх індивідуальні властивості.

Приклад подання реального об'єкта у вигляді дерева на прикладі будівлі з вікнами та дверима.

Розглянемо зображення фасаду будівлі. Після попередньої обробки та сегментації зображення виділяється основний об'єкт — будівля, а також її складові елементи: контур

фасаду, вікна та двері. На основі цих елементів формується деревоподібне подання, яке відображає ієрархічну структуру об'єкта.

Кореневим вузлом дерева є весь об'єкт «Будівля». Цей вузол відповідає загальному контуру фасаду та містить узагальнену інформацію про розміри та положення будівлі на зображенні.

На другому рівні ієрархії розташовуються основні структурні компоненти будівлі. До них належить зовнішній контур фасаду, а також групи повторюваних елементів — вікна та двері. Таким чином, кореневий вузол має три дочірні вузли: «Зовнішній контур», «Вікна» та «Двері».

Вузол «Зовнішній контур» описує форму будівлі як цілого. Він може бути представлений замкненою ламаною лінією, що складається з кількох спрямованих відрізків, наприклад вертикальних і горизонтальних сторін фасаду та даху. Цей вузол не має підлеглих елементів або містить лише допоміжні вузли, що уточнюють форму контуру.

Вузол «Вікна» є узагальнюючим вузлом, який об'єднує всі віконні прорізи будівлі. Він має кілька дочірніх вузлів, кожен з яких відповідає окремому вікну. Кожне вікно, у свою чергу, може бути описане як прямокутний контур і мати власні підвузли, наприклад «Рама» та «Скло». Контур вікна зазвичай складається з чотирьох спрямованих відрізків, орієнтованих горизонтально та вертикально.

Аналогічно, вузол «Двері» об'єднує всі дверні прорізи. Для кожних дверей формується окремий вузол, який містить інформацію про форму дверного отвору. Цей вузол може мати підвузли, що описують дверну коробку, полотно дверей або декоративні елементи.

У підсумку деревоподібне подання фасаду будівлі можна описати у вигляді ієрархії:

- Корінь — «Будівля»
- «Зовнішній контур»
- «Вікна»
 - «Вікно 1»
 - «Вікно 2»
 - ...
- «Двері»
 - «Двері 1»

Кожен вузол дерева містить параметри, що характеризують відповідний елемент зображення, такі як геометричні розміри, положення, орієнтація та форма. Така структура дозволяє легко аналізувати як загальні властивості будівлі, так і окремі деталі, а також виконувати порівняння різних будівель за їх структурою.

Таким чином, подання зображення деревами забезпечує компактний і логічно впорядкований опис реального об'єкта, що відображає відношення «ціле—частина» та спрощує задачі розпізнавання, аналізу та класифікації зображень.

Контрольні питання

1. Що означає подання зображення спрямованими відрізками?
2. Як використовуються дерева для структурного подання зображень?
3. Наведіть приклади практичного застосування такого подання.

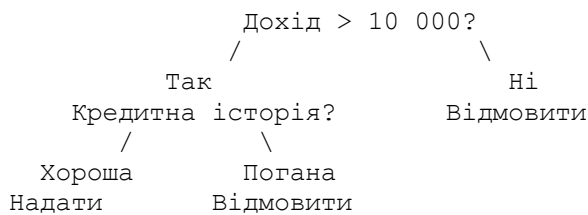
Лекція 11. Поняття дерева рішень та вирішального правила. Контроль складності дерева рішень.

1. Поняття дерева рішень.
2. Поняття вирішального правила.
3. Контроль складності дерева рішень.

1. Поняття дерева рішень.

Дерево рішень — це метод машинного навчання та аналізу даних, який застосовується для задач класифікації та регресії. Воно має ієрархічну структуру у вигляді дерева, в якій процес прийняття рішення відбувається шляхом послідовної перевірки значень ознак об'єкта. У структурі дерева виділяють кореневий вузол, що відповідає першій ознаці розбиття всієї вибірки, внутрішні вузли, у яких перевіряються умови над ознаками, гілки, що відповідають результатам цих перевірок, та листки, які містять кінцеве рішення — клас об'єкта або числове значення у випадку регресії.

Приклад дерева рішень для задачі визначення, чи видавати кредит клієнту. Ознаки: вік, рівень доходу, кредитна історія.



2. Поняття вирішального правила.

Кожному листку дерева відповідає вирішальне правило. Вирішальне правило — це правило прийняття рішення для конкретного об'єкта на основі його ознак, яке формується як шлях від кореня дерева до листка. Таким чином, дерево рішень можна розглядати як сукупність умовних правил типу «якщо — то». Наприклад, у задачі кредитного скорингу вирішальне правило може мати вигляд: якщо дохід клієнта перевищує певний поріг і його кредитна історія є хорошою, то кредит надається; в іншому випадку приймається рішення про відмову. Формально вирішальне правило задає відповідність між вектором ознак об'єкта та класом або числовим значенням, яке є результатом роботи моделі.

Побудова дерева рішень здійснюється рекурсивно. На кожному кроці алгоритм обирає таку ознаку та відповідний поріг або категорію, які найкраще розділяють навчальну вибірку на підмножини. Для оцінювання якості розбиття використовуються спеціальні критерії. Найпоширенішими з них є ентропія та приріст інформації, а також індекс Джині. Ентропія характеризує ступінь невизначеності у вузлі: чим вона більша, тим більш змішаними є класи. Приріст інформації показує, наскільки зменшується невизначеність після розбиття за певною ознакою. Індекс Джині також оцінює «нечистоту» вузла і широко застосовується в алгоритмах типу CART. У задачах регресії замість цих критеріїв зазвичай використовуються показники похибки, такі як середньоквадратична помилка.

3. Контроль складності дерева рішень.

Однією з основних проблем дерев рішень є їх схильність до перенавчання. Якщо не накладати обмежень, дерево може стати дуже глибоким і детальним, практично повністю відтворюючи навчальні дані, але погано узагальнюючи нові приклади. Тому важливим етапом є контроль складності дерева рішень.

Контроль складності може здійснюватися двома основними способами: попереднім та постфактум. Попередній контроль, або pre-pruning, передбачає введення обмежень уже на етапі побудови дерева. До таких обмежень належать максимальна глибина дерева, мінімальна кількість об'єктів у вузлі або в листку, а також мінімальний

приріст інформації, за якого дозволяється виконувати поділ. Ці заходи зменшують складність моделі та прискорюють навчання, однак можуть призвести до недонавчання, якщо обмеження є надто жорсткими.

Постфактум контроль складності, або *post-pruning*, застосовується після побудови повного дерева. У цьому випадку з дерева видаляються або спрощуються ті гілки, які не роблять суттєвого внеску в точність на валідаційних або тестових даних. Одним із поширених підходів є обрізання за критерієм «витрати–складність», де якість дерева оцінюється як сума похибки та штрафу за кількість листків. Зі зростанням штрафного коефіцієнта дерево стає простішим, що зазвичай покращує його узагальнювальну здатність.

Ознаками перенавчання дерева рішень є надмірно велика глибина, наявність листків з дуже малою кількістю об'єктів, а також значна різниця між точністю на навчальній і тестовій вибірках. Незважаючи на ці недоліки, дерева рішень мають важливі переваги: вони є наочними, легко інтерпретуються, можуть працювати з числовими та категоріальними ознаками і не потребують попередньої нормалізації даних.

Таким чином, дерево рішень являє собою набір вирішальних правил, організованих у вигляді ієрархічної структури. Ефективний контроль складності є ключовою умовою побудови якісної моделі, яка добре узагальнює дані та забезпечує баланс між точністю і простотою.

Контрольні питання

1. Що таке дерево рішень і для чого воно використовується?
2. Що таке вирішальне правило і як його визначають у вузлах дерева?
3. Як контролюють складність дерева рішень?

Лекція 12. Показник ентропії для побудови вирішального правила дерева рішень.

Індекс Джині та CART-міра для побудови вирішального правила дерева рішень.

1. Показник ентропії для побудови вирішального правила дерева рішень.
2. Індекс Джині та CART-міра для побудови вирішального правила дерева рішень.

1. Показник ентропії для побудови вирішального правила дерева рішень.

Ентропія є одним з ключових показників, що використовується під час побудови дерев рішень для формування вирішального правила, тобто для вибору найкращого атрибута, за яким відбувається розбиття даних у вузлі дерева.

У теорії інформації ентропія характеризує міру невизначеності або «хаотичності» системи. У задачах класифікації вона відображає, наскільки неоднорідними є класи об'єктів у певній вибірці. Якщо всі об'єкти належать до одного класу, невизначеність відсутня і ентропія дорівнює нулю. Якщо ж об'єкти рівномірно розподілені між кількома класами, ентропія досягає максимального значення.

Для вибірки, що містить k класів, ентропія визначається як сума по всіх класах добутків імовірностей належності об'єктів до відповідного класу на логарифм цих імовірностей зі знаком «мінус». Імовірність кожного класу обчислюється як частка об'єктів цього класу в загальній кількості об'єктів вибірки. Таким чином, ентропія зростає зі збільшенням невизначеності щодо того, до якого класу належить випадково обраний об'єкт.

Під час побудови дерева рішень ентропія використовується для оцінювання якості можливих розбиттів даних за різними ознаками. На кожному вузлі дерева розглядається поточна підвибірка даних і для неї обчислюється початкова ентропія. Далі для кожної кандидатної ознаки моделюється розбиття вибірки на кілька підвбірок відповідно до значень цієї ознаки. Для кожної з отриманих підвбірок обчислюється власна ентропія.

Щоб оцінити, наскільки корисним є певне розбиття, використовується поняття інформаційного виграшу. Інформаційний виграш показує, на скільки зменшується середня ентропія після розбиття даних за вибраною ознакою. Він обчислюється як різниця між ентропією до розбиття та зваженою сумою ентропій усіх підвибірок після розбиття, де вагами є частки об'єктів у відповідних підвибірках.

Ознака, для якої інформаційний виграш є максимальним, вважається найкращою для формування вирішального правила в даному вузлі дерева. Саме за цією ознакою виконується розбиття, а процес повторюється рекурсивно для кожного дочірнього вузла. Таким чином, дерево рішень будується так, щоб на кожному кроці максимально зменшувати невизначеність щодо класів об'єктів.

Отже, показник ентропії відіграє фундаментальну роль у побудові дерев рішень. Він дозволяє кількісно оцінити неоднорідність класів у даних і є основою для вибору оптимальних розбиттів, що формують ефективні вирішальні правила класифікації.

2. Індекс Джині та CART-міра для побудови вирішального правила дерева рішень.

Індекс Джині та CART-міра є альтернативними до ентропії критеріями, які широко застосовуються для побудови вирішального правила в деревах рішень, зокрема в алгоритмі CART (Classification and Regression Trees). Вони використовуються для оцінювання якості розбиття даних у вузлах дерева та вибору ознаки, що найкраще зменшує неоднорідність класів.

Індекс Джині характеризує ступінь «змішаності» класів у вибірці. Інтуїтивно він показує ймовірність того, що два випадково вибрані об'єкти з вибірки належатимуть до різних класів. Якщо всі об'єкти у вузлі належать до одного класу, індекс Джині дорівнює нулю, що означає повну однорідність. Чим більш рівномірно представлені класи, тим більшим є значення індексу Джині, а отже, тим вищою є невизначеність.

Для вибірки з k класами індекс Джині визначається через суму квадратів імовірностей належності об'єктів до кожного класу. Імовірність кожного класу обчислюється як частка об'єктів цього класу в загальній кількості об'єктів у вузлі. Таким чином, індекс Джині зростає зі збільшенням неоднорідності класів і зменшується, коли один із класів починає домінувати.

У процесі побудови дерева рішень індекс Джині використовується для вибору найкращого розбиття. Спочатку обчислюється індекс Джині для поточного вузла. Далі для кожної кандидатної ознаки розглядаються можливі варіанти розбиття вибірки, і для кожної отриманої підвибірки обчислюється власний індекс Джині. Якість розбиття оцінюється за допомогою зваженого середнього індексу Джині дочірніх вузлів, де вагами є частки об'єктів у відповідних підвибірках.

CART-міра (або зменшення індексу Джині) безпосередньо пов'язана з індексом Джині й використовується як критерій оптимальності розбиття в алгоритмі CART. Вона показує, наскільки зменшується неоднорідність класів у результаті розбиття. CART-міра обчислюється як різниця між індексом Джині батьківського вузла та зваженим індексом Джині дочірніх вузлів після розбиття. Чим більше це зменшення, тим ефективнішим вважається відповідне розбиття.

Під час побудови дерева CART на кожному вузлі перебираються всі можливі ознаки та всі допустимі варіанти їх розбиття (для числових ознак — порогові значення, для категоріальних — поділи на групи). Для кожного варіанта обчислюється CART-міра, і як вирішальне правило обирається те розбиття, яке забезпечує максимальне зменшення індексу Джині.

Отже, індекс Джині є мірою неоднорідності класів у вузлі дерева, а CART-міра — показником якості розбиття, що відображає, наскільки ефективно вибрана ознака

зменшує цю неоднорідність. Використання цих критеріїв дозволяє будувати компактні та ефективні дерева рішень, орієнтовані на максимальну чистоту класів у кожному вузлі.

Контрольні питання

1. Як використовується показник ентропії для побудови вирішального правила?
2. Що таке індекс Джині і як його застосовують у дереві рішень?
3. Як CART-міра допомагає вибрати оптимальне розщеплення?

Лекція 13. Оцінка якості класифікації за допомогою матриці помилок, показників accuracy, error rate.

1. Оцінка якості бінарної класифікації за допомогою матриці помилок, показників accuracy, error rate.
2. Оцінка якості багатокласової класифікації за допомогою матриці помилок, показників accuracy, error rate.

1. Оцінка якості бінарної класифікації за допомогою матриці помилок, показників accuracy, error rate.

Оцінка якості класифікації є важливим етапом аналізу роботи алгоритмів машинного навчання. Вона дозволяє зрозуміти, наскільки добре модель справляється з поставленим завданням і які типи помилок вона допускає. Основними інструментами такої оцінки є матриця помилок, а також показники accuracy та error rate.

Матриця помилок (confusion matrix) — це табличне подання результатів класифікації, яке показує співвідношення між істинними класами об'єктів і класами, передбаченими моделлю. У випадку бінарної класифікації матриця складається з чотирьох елементів. Істинно позитивні (True Positive, TP) — це об'єкти позитивного класу, які модель правильно віднесла до позитивного класу. Істинно негативні (True Negative, TN) — це об'єкти негативного класу, які були правильно класифіковані як негативні. Хибно позитивні (False Positive, FP) — це об'єкти негативного класу, які модель помилково віднесла до позитивного. Хибно негативні (False Negative, FN) — це об'єкти позитивного класу, які модель помилково віднесла до негативного. Саме матриця помилок дає найбільш повне уявлення про характер помилок класифікатора, оскільки дозволяє окремо аналізувати кожен тип правильних і неправильних рішень.

На основі матриці помилок обчислюється показник accuracy (точність класифікації). Accuracy визначається як відношення кількості правильно класифікованих об'єктів до загальної кількості об'єктів у вибірці. Іншими словами, це частка об'єктів, для яких модель дала правильний прогноз. Формально accuracy дорівнює сумі істинно позитивних та істинно негативних передбачень, поділених на загальну кількість усіх передбачень. Цей показник є простим для інтерпретації та часто використовується як базова міра якості моделі.

Однак показник accuracy має суттєве обмеження. Він може бути оманливим у випадках, коли класи є незбалансованими. Наприклад, якщо один клас значно переважає інший, модель може досягати високого значення accuracy, просто постійно передбачаючи домінуючий клас, але при цьому майже не розпізнавати рідкісний клас. Тому для глибшого аналізу якості класифікації accuracy зазвичай розглядають разом з матрицею помилок та іншими метриками.

Показник error rate (частота помилок) є доповненням до accuracy і показує частку неправильних класифікацій. Він обчислюється як відношення кількості хибно позитивних і хибно негативних передбачень до загальної кількості об'єктів. Фактично error rate дорівнює одиниці мінус accuracy і характеризує, наскільки часто модель помиляється. Чим менше значення error rate, тим кращою вважається модель.

Таким чином, матриця помилок є фундаментальним інструментом для аналізу результатів класифікації, оскільки вона надає детальну інформацію про всі типи правильних і неправильних рішень. Показники *assurasy* та *error rate* дозволяють узагальнити ці результати у вигляді одного числа, що спрощує порівняння різних моделей. Проте для повної та коректної оцінки якості класифікації, особливо в складних або незбалансованих задачах, їх доцільно доповнювати іншими метриками.

2. Оцінка якості багатокласової класифікації за допомогою матриці помилок, показників *assurasy*, *error rate*.

Для багатокласової класифікації оцінка якості моделі також ґрунтується на матриці помилок і узагальнювальних показниках, зокрема *assurasy* та *error rate*, однак їх інтерпретація має певні особливості через наявність більше ніж двох класів.

У багатокласовій класифікації матриця помилок має розмір $k \times k$, де k — кількість класів. Рядки цієї матриці відповідають істинним класам об'єктів, а стовпці — класам, передбаченим моделлю. Елемент матриці, що знаходиться на перетині i -го рядка та j -го стовпця, показує кількість об'єктів, які насправді належать до класу i , але були класифіковані моделлю як клас j . Значення на головній діагоналі матриці відповідають правильним класифікаціям, тобто тим об'єктам, для яких передбачений клас збігається з істинним. Усі елементи поза діагоналлю відображають різні типи помилок, зокрема плутанину між конкретними парами класів.

Матриця помилок у багатокласовому випадку дозволяє детально проаналізувати поведінку моделі. Вона дає змогу побачити, які саме класи модель найчастіше плутає між собою, які класи розпізнаються добре, а які — гірше. Це особливо важливо в задачах, де помилки між окремими класами мають різну значущість або різні наслідки.

Показник *assurasy* у багатокласовій класифікації визначається аналогічно до бінарного випадку. Він обчислюється як відношення кількості правильно класифікованих об'єктів (сума елементів головної діагоналі матриці помилок) до загальної кількості об'єктів у вибірці. Таким чином, *assurasy* показує, яка частка всіх об'єктів була віднесена моделлю до правильного класу. Цей показник є зручним для загальної оцінки якості та порівняння різних моделей.

Однак, як і у випадку бінарної класифікації, *assurasy* у багатокласових задачах може бути недостатньо інформативним, особливо коли класи представлені нерівномірно. Якщо деякі класи зустрічаються значно частіше за інші, модель може демонструвати високу *assurasy*, добре розпізнаючи домінуючі класи, але водночас погано працювати з рідкісними класами. У такій ситуації матриця помилок дозволяє виявити цю проблему, тоді як сам показник *assurasy* її не відображає.

Показник *error rate* у багатокласовій класифікації є доповненням до *assurasy* і характеризує частку неправильних передбачень. Він обчислюється як відношення кількості помилково класифікованих об'єктів до загальної кількості об'єктів, або як одиниця мінус *assurasy*. *Error rate* показує, наскільки часто модель робить помилки незалежно від того, між якими саме класами вони виникають.

Отже, у багатокласовій класифікації матриця помилок є ключовим інструментом аналізу, оскільки вона надає детальну інформацію про структуру помилок між класами. Показники *assurasy* та *error rate* дозволяють узагальнити результати класифікації в одному числі, проте для глибшого аналізу якості моделі, особливо за наявності незбалансованих класів, їх доцільно використовувати разом з матрицею помилок та іншими метриками оцінювання.

Контрольні питання

1. Що таке матриця помилок і як її інтерпретувати?
2. Як обчислюються показники *assurasy* та *error rate*?

3. Які висновки можна зробити на основі цих показників?

Лекція 14. Оцінка якості класифікації за допомогою precision, recall, F-міри.

1. Оцінка якості багатокласової класифікації за допомогою precision, recall та F-міри.
2. Приклад оцінки якості багатокласової класифікації за допомогою precision, recall та F-міри.

1. Оцінка якості багатокласової класифікації за допомогою precision, recall та F-міри.

Оцінка якості багатокласової класифікації за допомогою precision, recall та F-міри є більш складною, ніж у бінарній класифікації, оскільки потрібно враховувати результати для кожного класу окремо та їхній внесок у загальну якість моделі.

У багатокласовій класифікації модель повинна віднести кожен об'єкт до одного з K класів. Для оцінки точності використовують ті самі поняття, що і в бінарній класифікації: істинно позитивні результати (TP), хибно позитивні (FP) та хибно негативні (FN). TP для класу i означає кількість об'єктів, які правильно віднесено до класу i , FP — кількість об'єктів інших класів, помилково віднесених до класу i , а FN — кількість об'єктів класу i , помилково віднесених до інших класів.

На основі цих величин розраховують precision, recall та F1-міру для кожного класу. Precision характеризує точність позитивних прогнозів для конкретного класу, recall показує, яку частку об'єктів класу модель змогла правильно виявити, а F1-міра дозволяє оцінити баланс між точністю та повнотою.

Щоб отримати загальні показники для всієї моделі, застосовують різні способи усереднення:

Macro-average обчислює показники для кожного класу окремо і бере просте середнє. Усі класи рівнозначні, незалежно від їхньої кількості в даних. Цей підхід корисний, якщо важливо оцінити, наскільки модель працює однаково для всіх класів.

Micro-average сумує TP, FP та FN по всіх класах і обчислює показники на основі загальних сум. Цей підхід підходить для сильно незбалансованих класів, оскільки переважають часті класи, а менші класи мають менший вплив.

Weighted-average схожий на macro, але усереднення зважується по кількості об'єктів у кожному класі. Він поєднує переваги micro і macro, враховуючи дисбаланс класів.

Наприклад, якщо у нас три класи (0, 1, 2), і модель дала певні передбачення, матриця помилок може виглядати так:

$\begin{bmatrix} 3 & 0 & 0 \end{bmatrix}$

$\begin{bmatrix} 1 & 1 & 1 \end{bmatrix}$

$\begin{bmatrix} 0 & 1 & 2 \end{bmatrix}$

Для класу 0:

TP = 3, FP = 0, FN = 0;

Для класу 1:

TP = 1, FP = 1, FN = 2;

Для класу 2: TP = 2, FP = 1, FN = 1.

Macro-average обчислює середнє значення precision, recall і F1 для трьох класів, micro-average підсумовує всі TP, FP і FN і обчислює показники на їх основі, а weighted-average враховує кількість об'єктів у кожному класі при усередненні.

Таким чином, оцінка якості багатокласової класифікації за допомогою precision, recall та F1-міри дозволяє глибше зрозуміти, як модель працює з різними класами, і обирати метод усереднення залежно від завдання: macro — коли важливо рівномірне

представлення всіх класів, `micro` — коли важливі загальні результати, `weighted` — коли класи мають різну кількість об'єктів.

2. Приклад оцінки якості багатокласової класифікації за допомогою `precision`, `recall` та `F-міри`.

```
import numpy as np
from sklearn.metrics import precision_score, recall_score, f1_score, confusion_matrix

# Істинні мітки класів (3 класи: 0, 1, 2)
y_true = np.array([0, 1, 2, 0, 1, 2, 0, 1, 2, 2])

# Передбачені мітки моделі
y_pred = np.array([0, 2, 2, 0, 0, 2, 0, 1, 1, 2])

# Матриця помилок
cm = confusion_matrix(y_true, y_pred)
print("Матриця помилок:")
print(cm)

# Precision
precision_macro = precision_score(y_true, y_pred, average='macro')
precision_micro = precision_score(y_true, y_pred, average='micro')
precision_weighted = precision_score(y_true, y_pred, average='weighted')

# Recall
recall_macro = recall_score(y_true, y_pred, average='macro')
recall_micro = recall_score(y_true, y_pred, average='micro')
recall_weighted = recall_score(y_true, y_pred, average='weighted')

# F1-score
f1_macro = f1_score(y_true, y_pred, average='macro')
f1_micro = f1_score(y_true, y_pred, average='micro')
f1_weighted = f1_score(y_true, y_pred, average='weighted')

print("\nPrecision:")
print("Macro:", precision_macro)
print("Micro:", precision_micro)
print("Weighted:", precision_weighted)

print("\nRecall:")
print("Macro:", recall_macro)
print("Micro:", recall_micro)
print("Weighted:", recall_weighted)

print("\nF1-score:")
print("Macro:", f1_macro)
print("Micro:", f1_micro)
print("Weighted:", f1_weighted)
```

Контрольні питання

1. Що показує precision і recall?
2. Як обчислюється F-мера і що вона відображає?
3. У яких випадках ці показники ефективні для оцінки класифікації?

Лекція 15. Показники якості бінарної класифікації.

1. Показники якості бінарної класифікації.
2. Приклад обчислення показників якості бінарної класифікації.

1. Показники якості бінарної класифікації.

У задачах бінарної класифікації модель повинна віднести кожен об'єкт до одного з двох можливих класів, наприклад «позитивний» або «негативний». Для оцінювання якості такої моделі використовують спеціальні показники, які базуються на порівнянні прогнозів моделі з істинними мітками класів.

Основою більшості показників якості є матриця помилок. Вона описує всі можливі результати класифікації. Якщо модель правильно визначає об'єкт як позитивний, це називається істинно позитивним результатом (True Positive). Якщо модель правильно визначає негативний об'єкт, маємо істинно негативний результат (True Negative). Якщо ж негативний об'єкт помилково віднесено до позитивного класу, це хибно позитивний результат (False Positive), а коли позитивний об'єкт помилково віднесено до негативного класу — хибно негативний результат (False Negative). Наприклад, у медичній діагностиці істинно позитивний результат означає, що хворобу виявлено правильно, а хибно негативний — що хворобу не виявлено, хоча вона насправді є.

Найпростішим показником якості є точність класифікації (accuracy). Вона показує частку правильно класифікованих об'єктів серед усіх об'єктів вибірки. Цей показник є зручним і наочним, однак він добре працює лише тоді, коли класи є приблизно збалансованими. У задачах із сильним дисбалансом класів висока accuracy може вводити в оману, оскільки модель може просто ігнорувати рідкісний клас.

Для детальнішої оцінки якості позитивних прогнозів використовують показник precision. Він характеризує, яка частка об'єктів, віднесених моделлю до позитивного класу, дійсно є позитивними. Цей показник є особливо важливим у задачах, де хибні позитивні рішення є небажаними, наприклад у фільтрації спаму або в системах рекомендацій, де помилкові спрацювання можуть призводити до негативного користувацького досвіду.

Іншим важливим показником є recall, який також називають повнотою або чутливістю. Recall показує, яку частку всіх позитивних об'єктів модель змогла правильно виявити. Він є критично важливим у тих задачах, де пропуск позитивного випадку може мати серйозні наслідки, наприклад у медицині або у виявленні шахрайських операцій. Високий recall означає, що модель знаходить більшість важливих випадків, навіть якщо при цьому зростає кількість хибних позитивних результатів.

Для оцінювання здатності моделі правильно розпізнавати негативні об'єкти використовують показник специфічності (specificity). Він показує, яку частку негативних прикладів модель класифікує правильно. Цей показник важливий у задачах, де необхідно мінімізувати помилкові звинувачення або необґрунтовані спрацювання, наприклад у фінансових чи юридичних системах.

Оскільки показники precision і recall часто перебувають у протиріччі, для їх узагальнення застосовують F1-міру. Вона є гармонійним середнім між precision та recall і дозволяє оцінити баланс між ними одним числом. F1-міра особливо корисна у випадках, коли класи є незбалансованими і важливо одночасно контролювати як кількість хибних позитивних, так і хибних негативних результатів.

Для аналізу роботи класифікатора при різних порогах прийняття рішення використовують ROC-криву. Вона відображає залежність між часткою правильно виявлених позитивних об'єктів і часткою хибно позитивних результатів. Узагальненим числовим показником у цьому випадку є площа під ROC-кривою, або AUC. Чим більше значення AUC, тим краще модель відрізняє позитивний і негативний класи. Цей показник широко застосовується для порівняння різних моделей, наприклад у кредитному скорингу або оцінюванні ризиків.

Ще одним показником якості є логарифмічна втрата (log-loss), яка враховує не лише правильність класифікації, а й впевненість моделі у своїх прогнозах. Вона використовується для оцінювання моделей, що повертають імовірності належності об'єкта до класу, таких як логістична регресія або нейронні мережі. Менше значення log-loss відповідає кращій якості моделі.

Отже, показники якості бінарної класифікації відображають різні аспекти роботи моделі. Вибір конкретного показника залежить від характеру задачі та від того, яка помилка — хибно позитивна чи хибно негативна — є більш критичною у практичному застосуванні.

2. Приклад обчислення показників якості бінарної класифікації.

```
# Приклад обчислення показників якості бінарної класифікації
# Використовується бібліотека scikit-learn

import numpy as np
from sklearn.metrics import (
    confusion_matrix,
    accuracy_score,
    precision_score,
    recall_score,
    f1_score,
    roc_auc_score,
    log_loss
)

# Істинні мітки класів (0 — негативний клас, 1 — позитивний клас)
y_true = np.array([1, 0, 1, 1, 0, 1, 0, 0, 1, 0])

# Передбачені моделью класи
y_pred = np.array([1, 0, 1, 0, 0, 1, 1, 0, 1, 0])

# Імовірності належності до позитивного класу
y_prob = np.array([0.9, 0.2, 0.8, 0.4, 0.3, 0.7, 0.6, 0.1, 0.85, 0.2])

# Обчислення матриці помилок
cm = confusion_matrix(y_true, y_pred)
TN, FP, FN, TP = cm.ravel()

print("Матриця помилок:")
print(cm)

# Accuracy — частка правильних відповідей
accuracy = accuracy_score(y_true, y_pred)
```

```

# Precision — частка правильних позитивних прогнозів
precision = precision_score(y_true, y_pred)

# Recall — частка знайдених позитивних об'єктів
recall = recall_score(y_true, y_pred)

# F1-score — баланс між precision і recall
f1 = f1_score(y_true, y_pred)

# ROC-AUC — здатність моделі відрізняти класи
roc_auc = roc_auc_score(y_true, y_prob)

# Log-loss — якість імовірнісних прогнозів
logloss = log_loss(y_true, y_prob)

print("\nПоказники якості:")
print("Accuracy:", accuracy)
print("Precision:", precision)
print("Recall:", recall)
print("F1-score:", f1)
print("ROC-AUC:", roc_auc)
print("Log-loss:", logloss)

```

Контрольні питання

1. Які основні показники якості бінарної класифікації використовуються?
2. Як оцінюють співвідношення TP, FP, TN, FN у бінарній класифікації?
3. Наведіть приклад використання цих показників у реальній задачі.

РЕКОМЕНДОВАНА ЛІТЕРАТУРА

1. Кутковецький В. Я. Розпізнавання образів : навч. посіб. Миколаїв : Вид-во ЧНУ ім. Петра Могили, 2017. 420 с.
2. Савицька Я. А., Смолій В. В., Місюра М. Д., Шкарупило В. В. Системи візуалізації та розпізнавання образів : навч. посіб. Київ : ЦП «Компринт», 2020. 214 с.
3. Шаповал Н. В. Розпізнавання образів. Комп'ютерний практикум : навч. посіб. Київ : КПІ ім. Ігоря Сікорського, 2022. 128 с.
4. Дудикевич В. Б., Козак Р. П. Інтелектуальні системи та машинне навчання : навч. посіб. Львів : Видавництво Львівської політехніки, 2021. 240 с.
5. Коваленко О. С., Романюк О. Н. Комп'ютерний зір та аналіз зображень : навч. посіб. Вінниця : ВНТУ, 2023. 196 с.

Конспект лекцій з дисципліни «Теорія розпізнавання образів» для здобувачів першого (бакалаврського) рівня вищої освіти за спеціальністю 124 – Системний аналіз та за спеціальністю 113 – Прикладна математика

Укладач: Марина Вячеславівна Полякова

Національний університет «Одеська політехніка»
65044, Одеса, пр. Шевченка, 1