

МОДЕЛЬ ТОРГІВ В APPLE SEARCH ADS



СТУДЕНТКА БП4 ПМ
ЖУРАВЛЬОВА АНАСТАСІЯ
КЕРІВНИК
ДРІНЬ СВІТЛАНА СЕРГІЇВНА

Мета

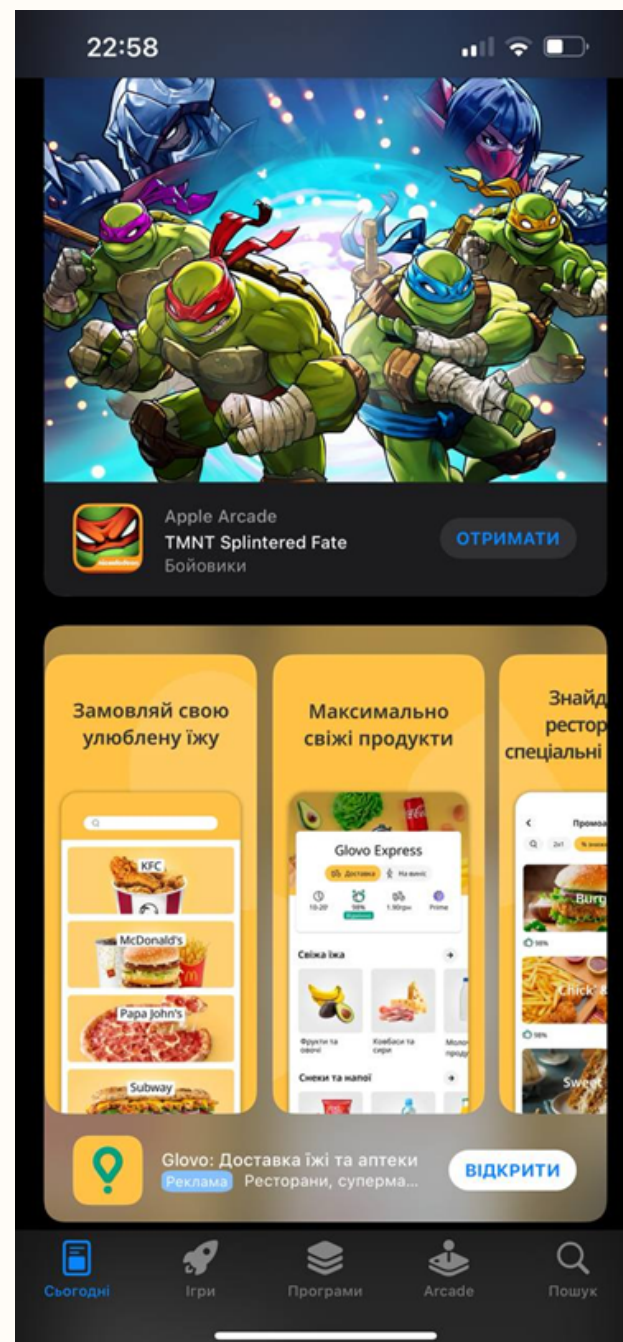
порівняти різні види торгів для
Apple Search Ads

Актуальність

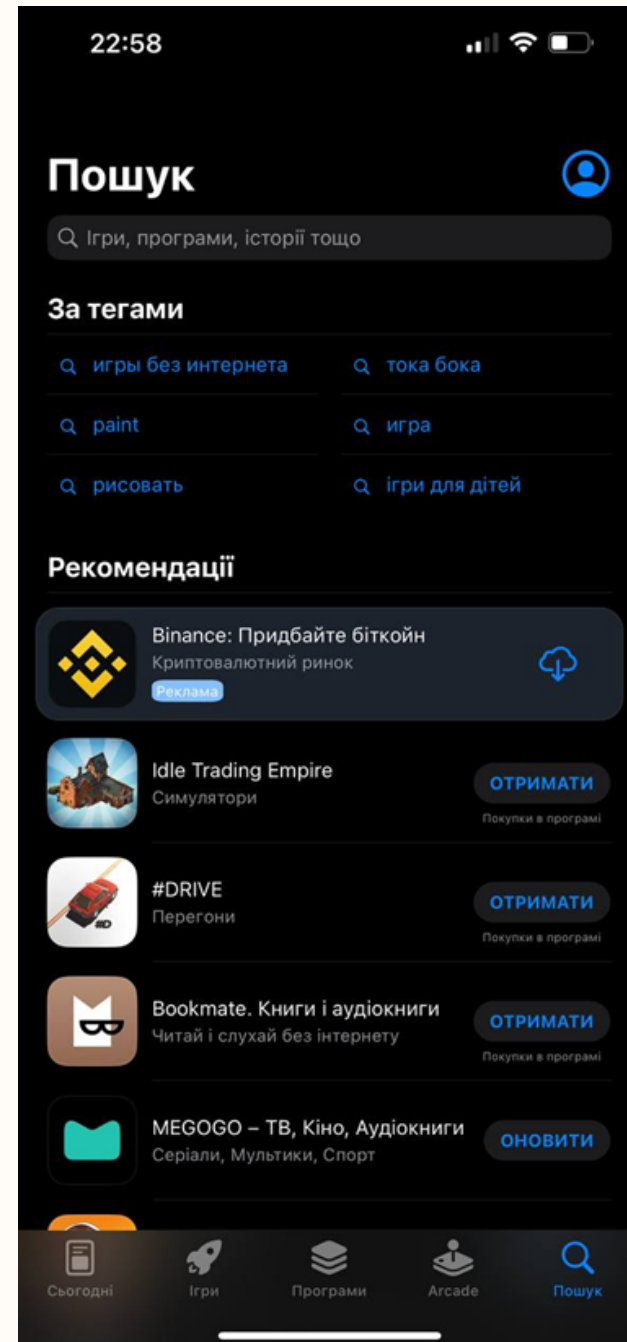
Тема замовлена одним з провідних бізнесів в Україні,
навчальний датасет - справжній кейс

APPLE SEARCH ADS

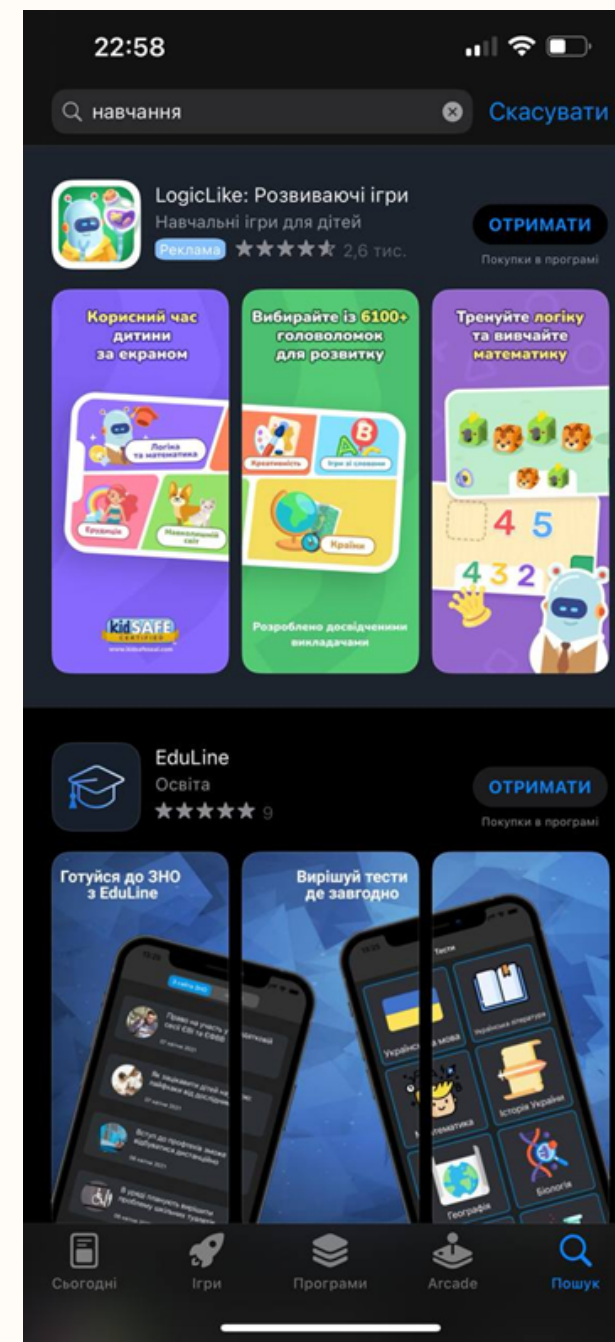
Варіанти розміщень



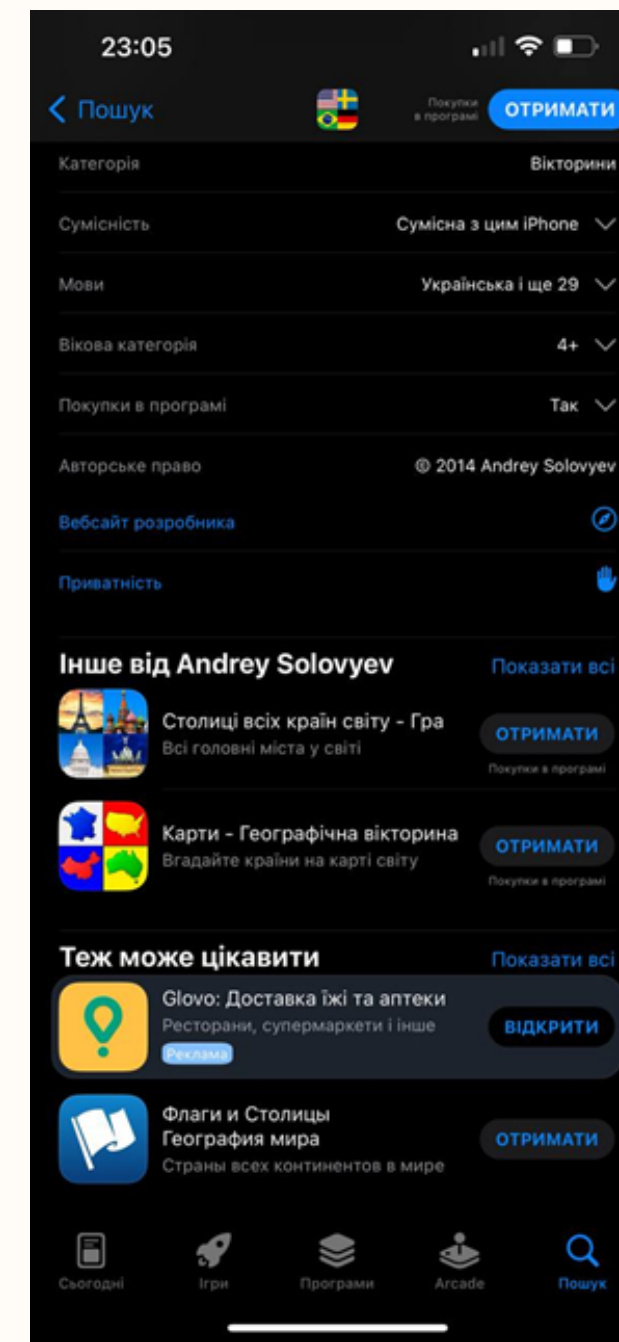
Вкладка "Сьогодні"



Вкладка "Пошук"



Результат пошуку



Продуктова сторінка

ОСНОВНІ ПОКАЗНИКИ В ПОШУКОВИХ СИСТЕМАХ

- **CPA** – COST PER ACQUISITION, ЦІНА, ЯКУ ВИ ПЛАТИТЕ ЗА ЗАВАНТАЖЕННЯ
- **CPT** – COST PER TAP, СЕРЕДНЯ ЦІНА ЗА КЛІК НА РЕКЛАМУ
- **CR** – CONVERSION RATE, ЦЕ $\frac{\text{КІЛЬКІСТЬ ЗАВАНТАЖЕНЬ}}{\text{КІЛЬКІСТЬ КЛІКІВ}}$
- **TTR** – TAP-THROUGH-RATE, ЦЕ $\frac{\text{КІЛЬКІСТЬ НАТИСКАНЬ}}{\text{КІЛЬКІСТЬ ПОКАЗІВ}}$
- **SPEND** – КІЛЬКІСТЬ ВИТРАТ
- **IMPRESSIONS** – КІЛЬКІСТЬ ПОКАЗІВ
- **TAPS** – КІЛЬКІСТЬ НАТИСКАНЬ
- **CONVERSIONS** – КІЛЬКІСТЬ ЗАВАНТАЖЕНЬ ПІСЛЯ НАТИСКАНЬ НА РЕКЛАМНЕ ОГолошення

АУКЦІОН ДРУГОЇ ЦІНИ

Дано:

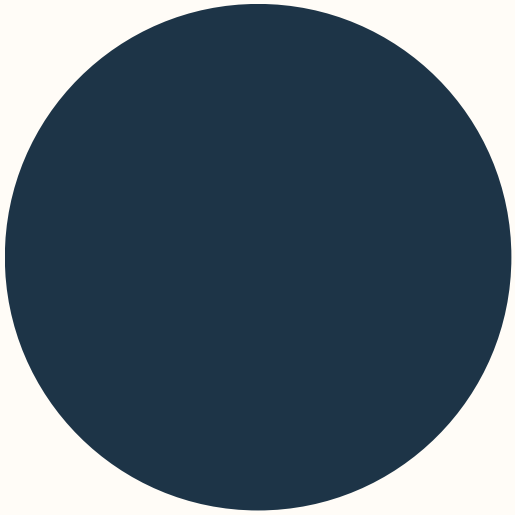
n - учасників

k - позицій

$n \geq k$

CTR_i - коефіцієнт клікабельності

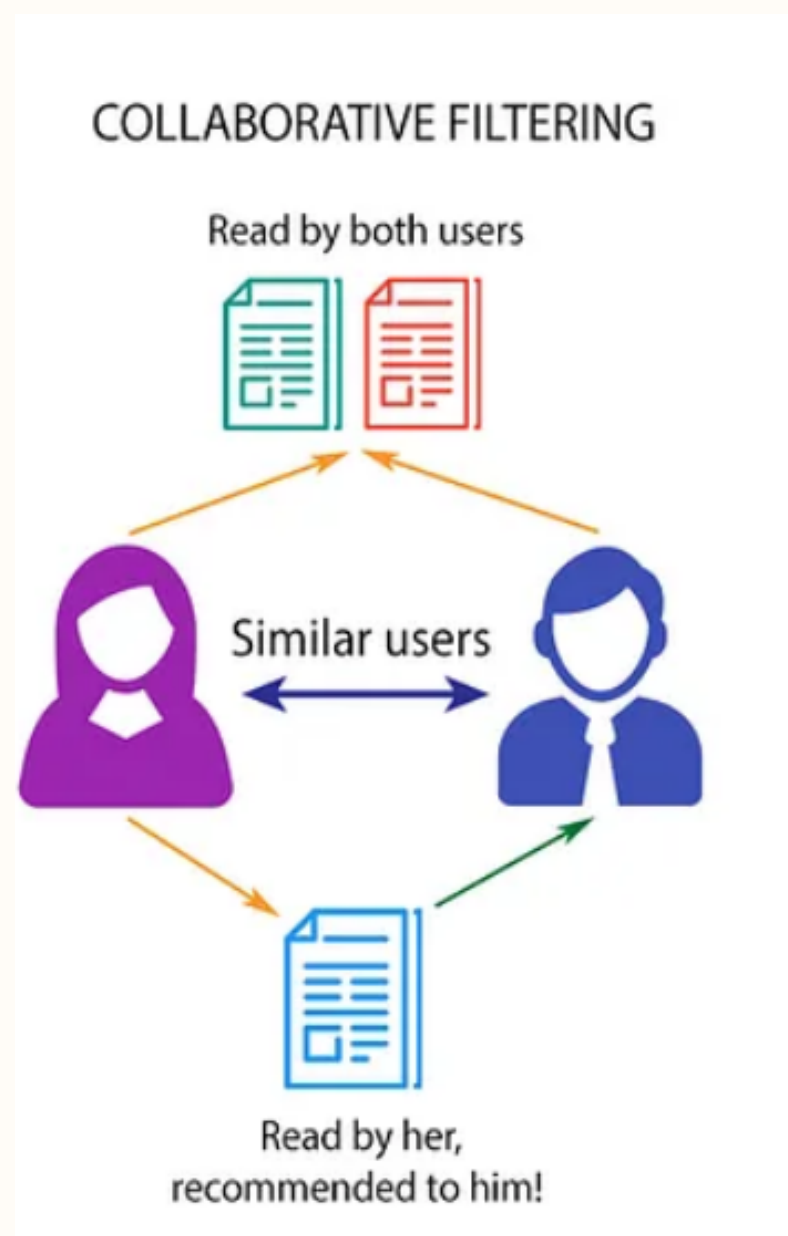




ОПИС РЕКОМЕНДАЦІЙНИХ СИСТЕМ

МЕТОД ФІЛЬТРАЦІЇ

Спільна фільтрація



01.

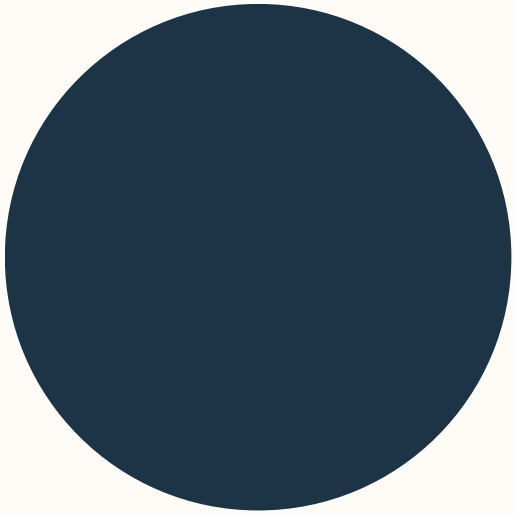
ВІДБІР КАНДИДАТІВ

02.

ПІДРАХУНОК БАЛІВ

03.

РЕЗУЛЬТАТ



МЕТОД НАЙБЛИЖЧИХ СУСІДІВ

ЯК ПРАЦЮЄ KNN?

01. ВИБРАТИ ЗНАЧЕННЯ ДЛЯ K. K ПОВИННО БУТИ НЕПАРНИМ ЧИСЛОМ

02. ЗНАЙТИ ВІДСТАНЬ ВІД НОВОЇ ТОЧКИ ДО КОЖНОГО З ТРЕНУВАЛЬНИХ ДАНИХ

03. ЗНАЙТИ K НАЙБЛИЖЧИХ СУСІДІВ ДЛЯ НОВОЇ ТОЧКИ ДАНИХ

04. ДЛЯ КЛАСИФІКАЦІЇ ПІДРАХУВАТИ КІЛЬКІСТЬ ТОЧОК ДАНИХ У КОЖНІЙ КАТЕГОРІЇ СЕРЕД K СУСІДІВ. НОВА ТОЧКА ДАНИХ НАЛЕЖАТИМЕ ДО КЛАСУ, ЯКИЙ МАЄ НАЙБІЛЬШЕ СУСІДІВ.

Евклідова відстань

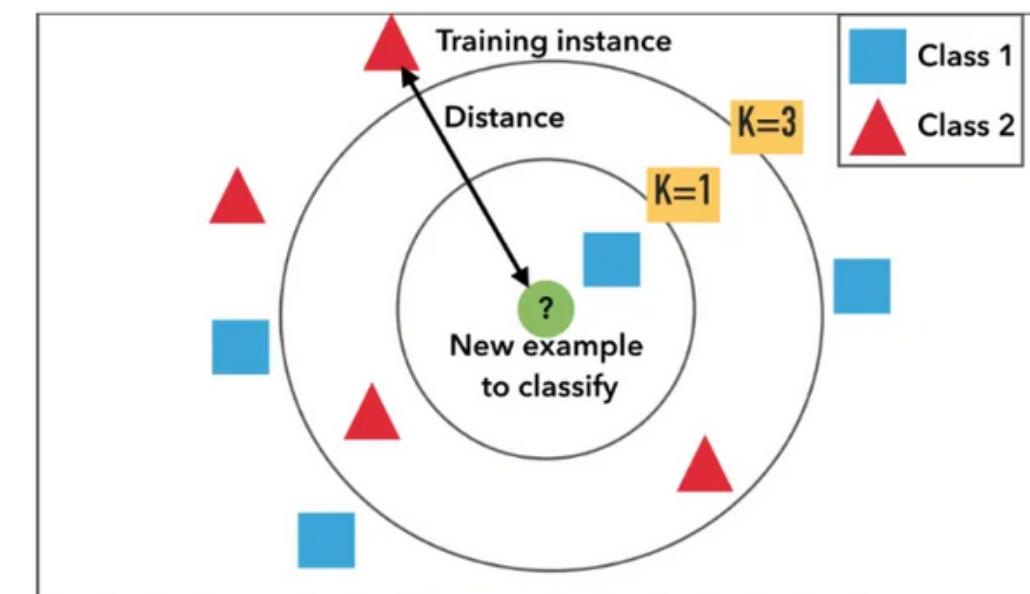
$$d(x, y) = \sqrt{\sum_{i=1}^n (y_i - x_i)^2}$$

Манхетенська метрика

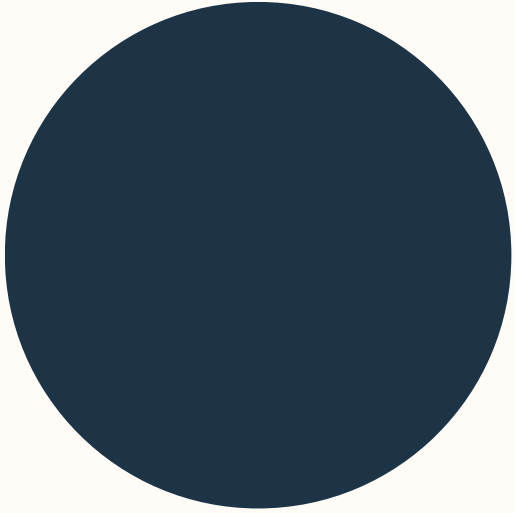
$$d(x, y) = \sum_{i=1}^m |x_i - y_i|$$

Метрика Мінковського

$$d(x, y) = \left(\sum_{i=1}^n |x_i - y_i| \right)^{1/p}$$



Для регресії значення для нової точки даних буде середнім значенням K сусідів.



МЕТОД RANDOM FOREST

ДЕРЕВА РІШЕНЬ



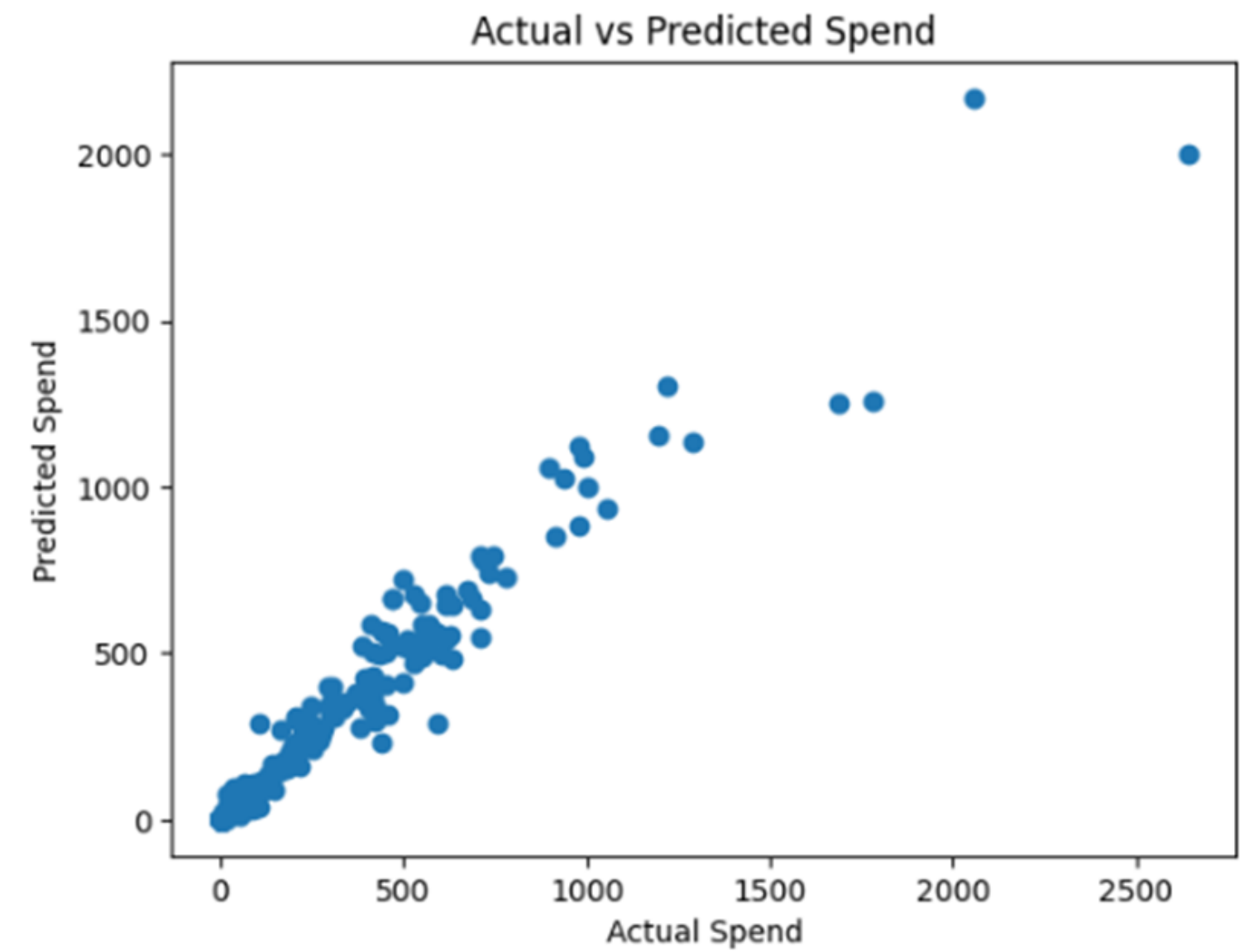
Нехай є вибірка **X** розміру **N**. Рівномірно візьмемо з вибірки **N** об'єктів з поверненням, внаслідок чого серед них можуть бути повтори. Позначимо нову вибірку через **X1**. Повторюючи процедуру **M** разів, згенеруємо **M** підвбірок **X1, ..., MX**

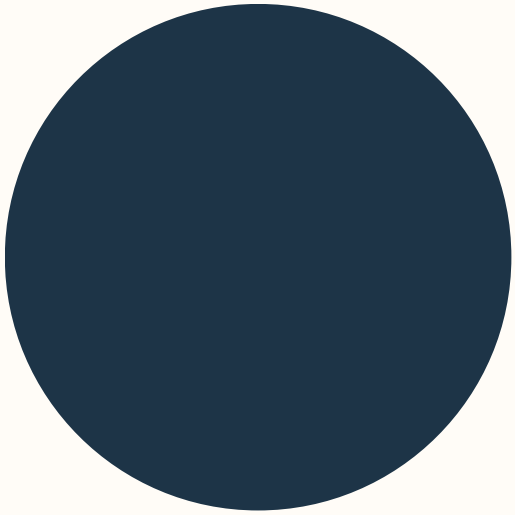
$$a(x) = \frac{1}{M} \sum_{i=1}^M a_i(x)$$

Для побудови ансамблю алгоритмів на основі бегінгу за допомогою бутстрепа генерують вибірки, на кожній з яких навчають свій класифікатор **ai(x)**

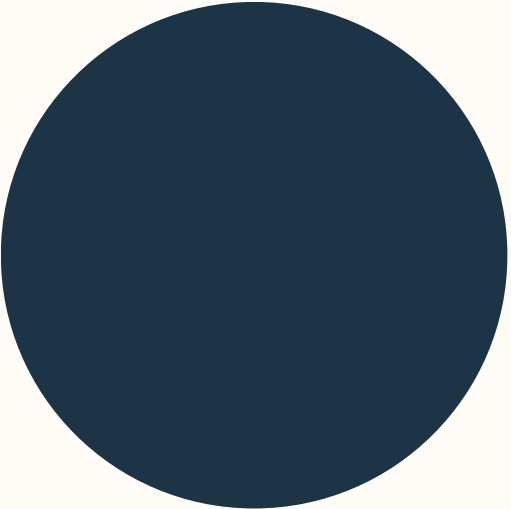
☞ Mean Squared Error (Random Forest): 0.3087189601410156
Mean Squared Error (k-Nearest Neighbors): 0.3157940295201343

☞ Mean Squared Error (Random Forest): 39.35843458258109
R-squared (Random Forest): 0.9652908125845242
Mean Squared Error (k-Nearest Neighbors): 98.848419094935
R-squared (k-Nearest Neighbors): 0.9128281309844565

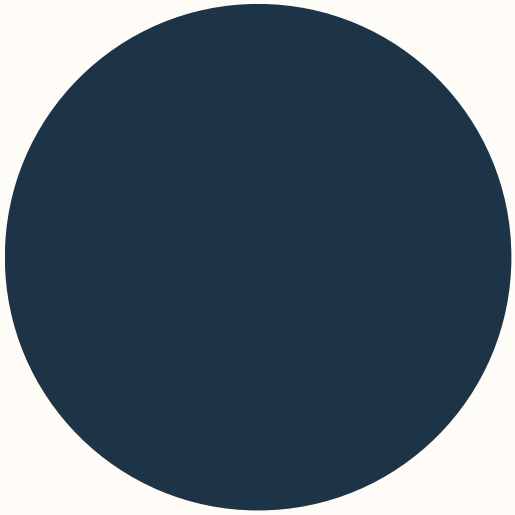




ПОДЯКА



APPENDIX



СПИСОК ЛІТЕРАТУРИ

<https://developers.google.com/machine-learning/recommendation?hl=ru>

<https://searchads.apple.com/>

<https://towardsdatascience.com/prototyping-a-recommender-system-step-by-step-part-1-knn-item-based-collaborative-filtering-637969614ea>

<https://medium.datadriveninvestor.com/k-nearest-neighbors-knn-7b4bd0128da7>

EDELMAN ET AL.: INTERNET ADVERTISING AND THE GSP AUCTION

<https://ela.kpi.ua/handle/123456789/40714>

<https://ami-ejournal.cdu.edu.ua/article/view/4158/4438>

