

РЕКОМЕНДАЦІЙНІ СИСТЕМИ ДЛЯ СТРИМІНГОВИХ ПЛАТФОРМ ВІДЕОКОНТЕНТУ

Анотація: Робота присвячена побудові та аналізу рекомендаційних систем для стримінгових платформ. Протестовано три моделі: TF-IDF, Doc2Vec та S-BERT, кожна з яких продемонструвала різний рівень відповідності вимогам сучасної рекомендаційної системи. Практичні висновки щодо доцільності використання моделей: найменш придатною виявилася Doc2Vec, базовим стартовим рішенням залишається TF-IDF, а найбільш перспективною та точною - S-BERT. Отримані результати можуть слугувати основою для подальшого вдосконалення системи, включно з інтеграцією гібридного підходу, що об'єднає переваги кожної моделі для досягнення ще вищої релевантності рекомендацій.

Ключові слова: інтелектуальний аналіз даних, рекомендаційна система, TF-IDF, DOC2VEC, S-BERT.

Вступ

У сучасному цифровому середовищі стримінгові платформи відіграють ключову роль у споживанні медіаконтенту. Щодня мільйони користувачів переглядають фільми, серіали та інші відеоматеріали, що зумовлює необхідність у впровадженні ефективних систем рекомендацій. Якісні рекомендації допомагають користувачам швидше знаходити релевантний контент, а платформам - підвищувати залученість і час перегляду.

Матеріали та методи

Серед основних підходів до побудови рекомендаційних систем виділяють колаборативне фільтрування, гібридні методи та контентно-орієнтоване фільтрування [1-4]. Саме останній підхід є особливо релевантним у випадках, коли відсутні достатні дані про активність користувача або коли важливо рекомендувати нові, ще не оцінені фільми.

Контентно-орієнтоване фільтрування (Content-Based Filtering) ґрунтується на аналізі характеристик об'єктів (жанр, сюжет, актори, режисер тощо) без даних про вподобання користувача. Такий підхід дозволяє будувати гнучкі рекомендаційні системи, які добре масштабуються та не залежать від активності інших користувачів. Також саме такий вид рекомендацій дозволяє вирішити проблему холодного старту.

Таким чином, створення рекомендаційної системи для стримінгової платформи на базі методів контентно-орієнтованого фільтрування є актуальним завданням, що поєднує знання з баз даних, машинного навчання та обробки тексту.

Для аналізу та моделювання було обрано 7 джерел відкритих даних на сайтах <https://www.kaggle.com/> та <https://grouplens.org/datasets/movielens/>. Також додатково знадобилося створити власний датасет, для цього було використано OMDb арі <https://www.omdbapi.com/>.

Набір даних з кіно та відгуками на них movielens (4 окремих пов'язаних файлів з даними, використано 1) - <https://grouplens.org/datasets/movielens/>

Набір даних imdb з кіно/акторами, рейтингами, та даними про локалізацію різних фільмів (5 окремих пов'язаних файлів з даними, використано всі) – <https://www.kaggle.com/datasets/ashirwadsangwan/imdb-dataset>

У моделі сховища за типом сніжинка спроектовано 2 таблиці фактів та 7 таблиць вимірів.

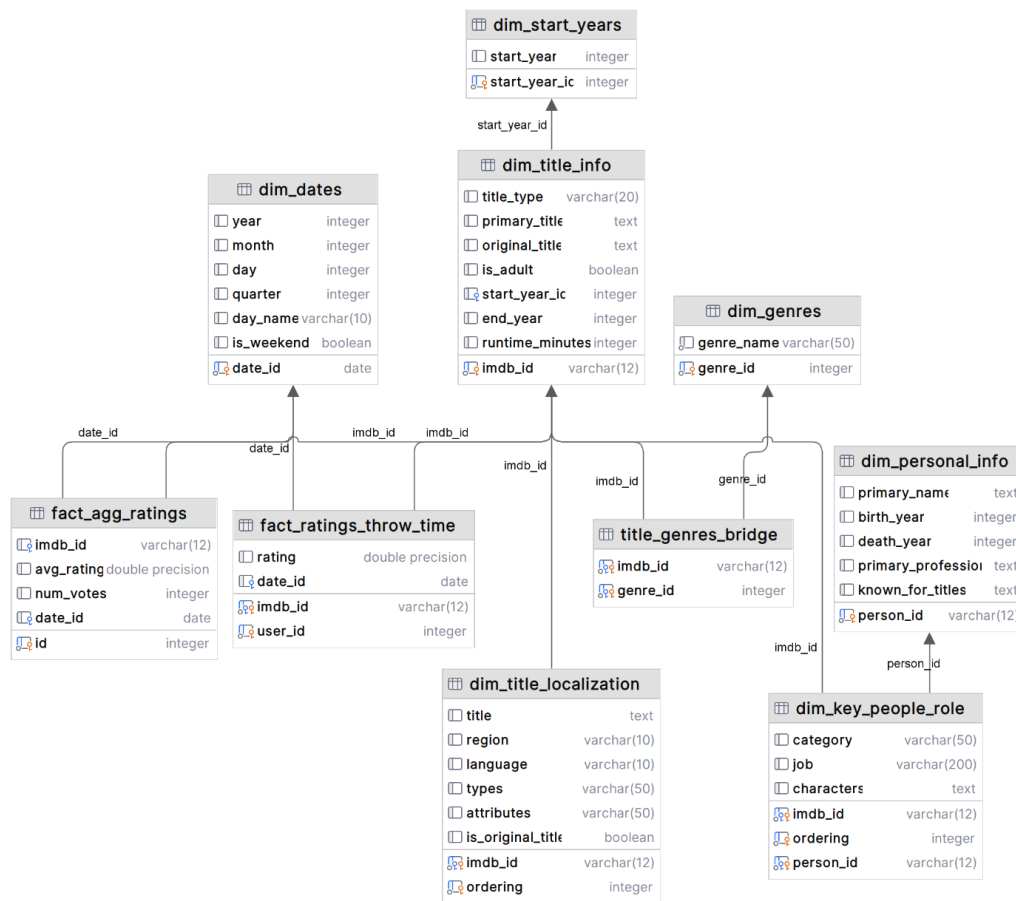


Рисунок 1. Модель бази даних, створена засобами PostgreSQL

Для реалізації рекомендаційної системи не вистачало сюжету кіно, тому довелося створити датасет з сюжетами самостійно. Для цього використано OMDb арі <https://www.omdbapi.com/>. Окрім сюжетів також вдалося отримати постери фільмів, що дозволило побудувати більш якісні візуалізації.

Для формування власного датасету використано такий алгоритм:

1. Отримати список усіх imdb_id, що були в існуючих даних.
2. Надіслати запит до OMDb арі для кожного id.
3. Записати відповідь до csv після кожного запиту.
4. Після отримання усіх відповідей отримати список усіх imdb_id, що були в існуючих даних, і отримати список imdb_id для яких вдалося отримати відповідь.
5. Знайти спільні imdb_id для цих двох сукупностей.
6. Видалити ті, які не мають опису сюжету.
7. Зберегти дані до csv файлу.

Перед виконанням дослідження було проаналізовано багато статей з відкритих джерел з питання content-based filtering. Ось деякі з них: [1-9]. В таблиці 1 можна побачити моделі, які були претендентами на випробування.

Таблиця 1

Модель	Складність	Інтерпретованість	Продуктивність
TF-IDF	низька	висока	висока
Doc2Vec	середня	середня	середня
S-BERT	висока	середня	висока
LDA/NMF	середня	висока	середня
FastText	середня	середня	висока
BM25	середня	середня	висока
USE	середня	середня	висока

Після детального аналізу кожної моделі було обрано 3 моделі:

- TF-IDF: оскільки вона є самою простою у реалізації і добре підходить як початкова точка у дослідженні, ця модель буде моделлю з якою будуть порівнюватися всі подальші результати [10,12,14];
- Doc2Vec: оскільки це компромісна модель, яка не важка і не повільна в порівнянні зі складнішими моделями і враховує контекст і порядок слів [14];
- S-BERT: оскільки на даний момент є одною із найкращих рішень для задач такого формату, але важка і повільна в порівнянні з LDA, NMF, BM25, TF-IDF [11-13].

Результати та обговорення

На жаль, не вдалося знайти загальновизнаних бенчмарків, які б напряду порівнювали ефективність TF-IDF, Doc2Vec і Sentence-BERT для задач побудови рекомендацій. Проте було знайдено дослідження, у якому ці моделі були порівняні в контексті задачі кластеризації трендових тем у Twitter [15].

У рамках цього дослідження було застосовано алгоритм DBSCAN у поєднанні з трьома методами векторизації: TF-IDF, Doc2Vec і Sentence-BERT. Для оцінки якості кластеризації використовувався коефіцієнт силуєту. Найвищий показник було досягнуто при використанні Doc2Vec, тоді як Sentence-BERT хоч і мав нижчий силуєт, продемонстрував кращу тематичну роздільність, виявивши два змістовно різні кластери. Модель з TF-IDF показала найгірші результати як за якістю кластеризації, так і за тематичною релевантністю [15].

Враховуючи ці результати, можна очікувати, що модель Sentence-BERT забезпечить найвищу релевантність рекомендацій у завданнях семантичного пошуку, Doc2Vec може стати надійним компромісом між якістю та ресурсними витратами, а TF-IDF - швидким, але менш точним базовим варіантом. Остаточну відповідь надасть практичне тестування, результати якого будуть розглянуті далі.

```
PLOT_SIM_WEIGHT = 0.4
GENRE_SIM_WEIGHT = 0.3
ACTOR_SIM_WEIGHT = 0.125
DIRECTOR_SIM_WEIGHT = 0.125
WRITER_SIM_WEIGHT = 0.05
```

Рисунок 2. Вагові коефіцієнти для різних полів фільму, що використані при аналізі результатів передбачень.

Для сюжетів вибрано коефіцієнт важливості 0.4, оскільки схожість сюжетів найчастіше хоче отримати кінцевий користувач під час перегляду рекомендацій. Коефіцієнт 0.3 обрано для жанрів оскільки жанри теж є дуже сильним предиктором того, як будуть розгортатись події фільму. Коефіцієнти 0.125 обрані для акторів і режисерів, оскільки вони є рівно важливими в більшості випадків і є важливішими за сценаристів для користувача.

Аналіз за фільмом

Хоча було використано 82,149 фільмів на візуалізації ми можемо побачити лише 4 в верхньому куті (рисунок 3) через те, що я застосував фільтрацію на фільми. Це зроблено для того, щоб зручніше було працювати з потрібними фільмами. На рисунку 4 можна побачити усі рекомендації.

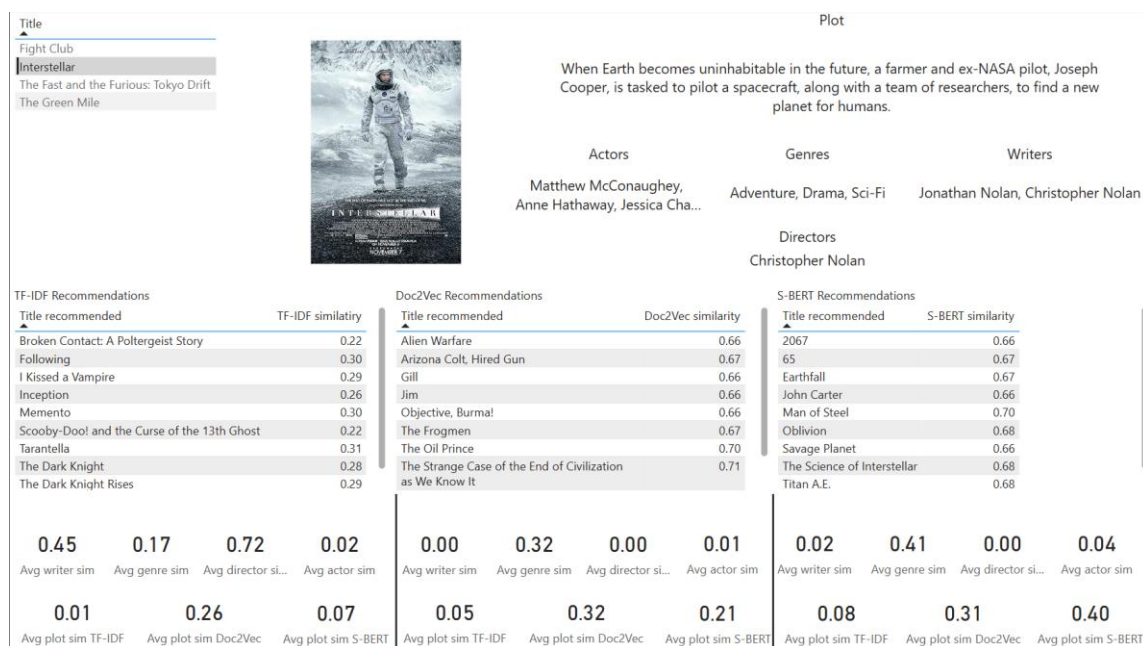


Рисунок 3. Результати рекомендацій від 3-ох моделей для фільму «Interstellar»

TF-IDF Recommendations		Doc2Vec Recommendations		S-BERT Recommendations	
Title recommended	TF-IDF similarity	Title recommended	Doc2Vec similarity	Title recommended	S-BERT similarity
The Prestige	0.38	The Strange Case of the End of Civilization as We Know It	0.71	Man of Steel	0.70
Tarantella	0.31	Twilight for the Gods	0.70	Oblivion	0.68
Memento	0.30	The Oil Prince	0.70	Titan A.E.	0.68
Following	0.30	Arizona Colt, Hired Gun	0.67	The Science of Interstellar	0.68
I Kissed a Vampire	0.29	The Frogmen	0.67	Wing Commander	0.67
The Dark Knight Rises	0.29	Gill	0.66	Earthfall	0.67
The Dark Knight	0.28	Zero Hour!	0.66	65	0.67
Inception	0.26	Jim	0.66	2067	0.66
Scooby-Doo! and the Curse of the 13th Ghost	0.22	Objective, Burma!	0.66	John Carter	0.66
Broken Contact: A Poltergeist Story	0.22	Alien Warfare	0.66	Savage Planet	0.66

Рисунок 4. Всі результати рекомендацій від 3-ох моделей для фільму «Interstellar»

На зображенні представлено результати роботи трьох різних моделей побудови рекомендацій за змістом - TF-IDF, Doc2Vec та S-BERT - для фільму «Interstellar». Кожна модель пропонує список рекомендованих фільмів із відповідними значеннями схожості, а нижче надані метрики, які характеризують середню схожість за різними ознаками: сценарист, жанр, режисер, акторський склад і сюжет. У центрі знаходиться короткий опис фільму «Interstellar» з основними метаданими: жанри, актори, сценаристи, режисер. Перейдемо до аналізу моделей.

TF-IDF модель демонструє слабку якість рекомендацій. У списку рекомендованих фільмів переважають неочевидні та маловідомі стрічки, як-от «Broken Contact: A Poltergeist Story», «I Kissed a Vampire» або «Scooby-Doo! and the Curse of the 13th Ghost». Жоден із цих фільмів не відповідає тематичній глибині або науково-

фантастичному контексту «Interstellar». Лише поодинокі випадки, як «Inception» або «The Dark Knight», «The Dark Knight Rises», частково корелюють з подібністю за режисером (усі чотири зняті Крістофером Ноланом), проте TF-IDF, імовірно, не враховує цей фактор напряму. Це підтверджується відповідними метриками: середня схожість за режисером лише 0.02, за акторами - 0.02, що дуже низько, і лише сценаристи дають 0.45, ймовірно, за рахунок збігів імен Ноланів. Середня схожість сюжету (plot sim) - 0.01 - є вкрай низькою і свідчить про дуже обмежену здатність цієї моделі знаходити дійсно подібні фільми за змістом. Таким чином, TF-IDF у цьому випадку відображає механістичний підхід, що базується на точному збігу лем і термінів, але не враховує семантичні відтінки.

Модель **Doc2Vec** демонструє зовсім інший профіль рекомендацій. Більшість рекомендованих назв - «Alien Warfare», «The Oil Prince», «Twilight for the Gods», «The Strange Case of the End of Civilization as We Know It» - не є очевидними за тематикою фільму «Interstellar» і не містять жодного з пов'язаних імен. Тим не менш, середня подібність за сюжетом становить 0.26, що значно вище, ніж у TF-IDF. Це, на перший погляд, вказує на кращу здатність Doc2Vec ловити певні структурні або смислові шаблони в текстах описів. Інші метрики, однак, дуже низькі або дорівнюють нулю: режисери - 0.00, сценаристи - 0.00, актори - 0.01. Жанрова подібність становить 0.32, що є найвищим значенням серед трьох моделей. Це вказує на те, що Doc2Vec, можливо, краще уловлює жанрову відповідність, хоча й не може працювати з іменами персонажів. Загалом результати виглядають сумнівно, фінальне рішення ми зможемо зробити після аналізу загальних візуалізацій порівняння моделей.

S-BERT демонструє найсильніший загальний результат серед усіх трьох моделей. У списку рекомендованих фільмів зустрічаються «2067», «Earthfall», «John Carter», «Man of Steel», «The Science of Interstellar». Останній є документальним фільмом про сам фільм «Interstellar», що вказує на глибоку семантичну відповідність. Рекомендовані стрічки мають науково-фантастичну тематику, пов'язану з космосом, майбутнім або катастрофами - все це відповідає сюжетній лінії «Interstellar». Середня схожість за сюжетом тут - 0.31, майже така сама як у Doc2Vec, але результати кращі за змістом. Жанрова схожість - 0.41, значно вище за попередні моделі, і хоча схожість за сценаристами (0.02), режисерами (0.00) та акторами (0.04) також низька, ці значення все ж вищі, ніж у Doc2Vec. Загальна відповідність рекомендацій вказує на здатність моделі S-BERT оперувати семантичними відношеннями між описами, не зважаючи на поверхневі збіги слів. Ця модель здатна генерувати рекомендації, що мають не тільки схожий сюжет, а й атмосферу, настрій і структуру.

Оцінюючи якість моделей, можна сказати, що для цієї вибірки TF-IDF та Doc2Vec – показали однаково погані результати. Doc2Vec, хоча і показує великі значення подібності, не дає задовільних результатів - вона ігнорує персональні

метадані та пропонує фільми, що мають не схожі описи. S-BERT - єдина модель, що генерує якісні рекомендації, як тематично, так і стилістично подібні до «Interstellar», попри низькі метрики за персоніфікованими атрибутами. Семантична глибина моделі дозволяє їй відтворити суть фільму, а не лише часткові збіги слів. Саме S-BERT виглядає найбільш перспективним варіантом для використання в реальній системі рекомендацій, якщо враховувати лише цей аналіз.

Загальний аналіз

Для загального аналізу моделей було побудовано чотири візуалізації (рис. 5-8).

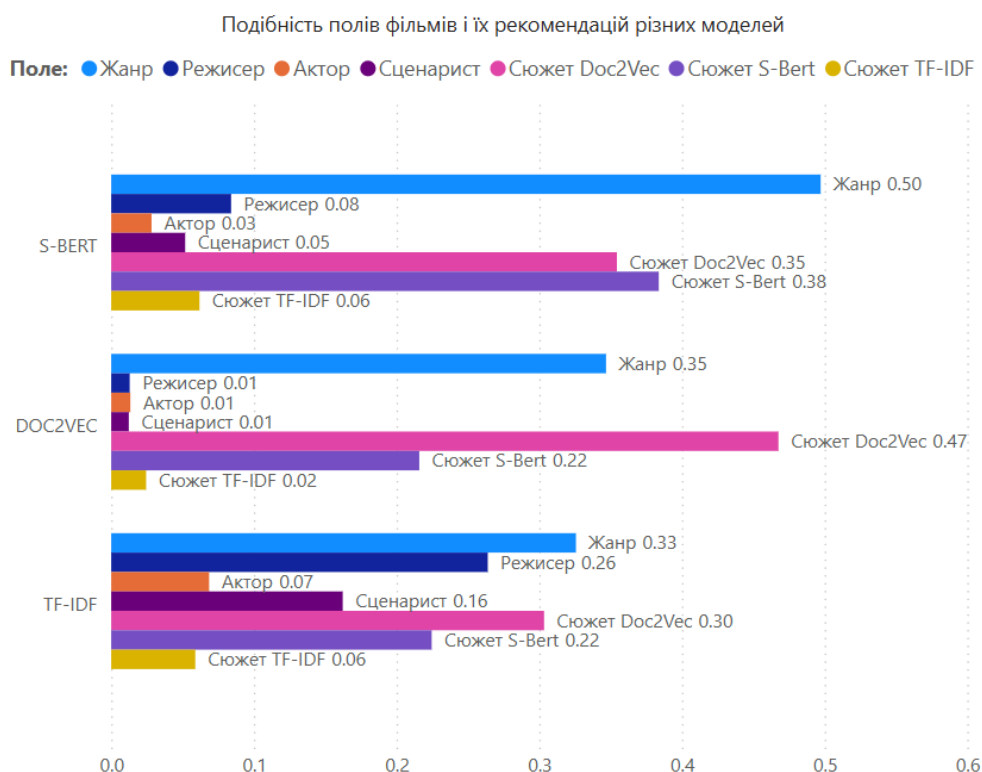


Рисунок 5. Візуалізація подібності полів фільмів та їх рекомендацій

Графік подібності полів фільмів і їх рекомендацій демонструє, що S-BERT має найвищу жанрову схожість - 0.50, що наочно свідчить про її здатність гарно класифікувати фільми за настроєм. Окрім жанру, модель S-BERT також має найвищі значення за семантичною подібністю сюжету за своєю методикою - 0.38, високу подібність за методикою Doc2Vec - 0.35 і таке саме значення подібності, як у рекомендацій запропонованих моделлю TF-IDF – 0.06.

Як ми вже побачили, модель Doc2Vec дуже сильно завищує значення подібності між фільмами, що майже не схожі. Це відбувається через те, що Doc2Vec - це розширення Word2Vec, яка навчає вектори не тільки для слів, а й для цілих

документів. На відміну від моделей, які оцінюють подібність поверхнево (як TF-IDF), Doc2Vec створює щільне (dense) представлення документів у вигляді векторів сталої розмірності, навіть якщо документи мають різну довжину, стиль чи лексику. Ці вектори «змушені» займати деякий компактний простір, що створює ефект згладжування. З підвищенням кількості фільмів, що аналізує модель, ми отримуємо більше згладжування, тому в деякий момент подібність починає з'являтися навіть у фільмів, що зовсім не схожі. Таким чином на цьому етапі ми можемо поставити модель Doc2Vec на останнє місце у питанні використання її для побудови рекомендацій в нашому дослідженні. Для підвищення якості роботи моделі потрібно буде постійно її перенавчати зі збільшенням розмірності векторів, що є не ефективним з точки зору використання пам'яті, часу, ресурсів системи.

TF-IDF демонструє класичну поведінку моделей на основі частоти термінів. Модель має найвищу подібність за персональними атрибутами - режисер 0.26, сценарист 0.16, актор 0.07 - що логічно, оскільки ці значення часто представлені як ключові слова у вхідному тексті. Жанрова подібність 0.33 теж досить висока, хоча й поступається S-BERT. Подібність за сюжетами за оцінками Doc2Vec і S-BERT знаходиться на рівні 0.30 і 0.22 відповідно, що можна вважати гідним показником, особливо враховуючи спрощений характер самої моделі TF-IDF. Однак TF-IDF значно слабше передає семантичну глибину сюжету, що підтверджується низьким значенням її внутрішньої сюжетної подібності - лише 0.06. Це означає, що модель здатна зловити прямі збіги, але не виявити глибинні теми чи емоційне наповнення.

	Doc2Vec	S-BERT	TF-IDF
DOC2VEC	100%	1%	1%
S-BERT	1%	100%	10%
TF-IDF	1%	10%	100%

Рисунок 6. Візуалізація схожості рекомендацій моделей

Теплокарта схожості між рекомендаціями моделей - демонструє дуже низький ступінь перетину між результатами. Doc2Vec, S-BERT і TF-IDF мають лише 1% або 10% перетину рекомендацій між собою. Це ще раз підтверджує, що кожна модель працює у власному векторному просторі з різним фокусом: TF-IDF - на лексичному збігу, Doc2Vec - на текстових структурах, S-BERT - на семантичному сенсі. Незважаючи на низьку міжмодельну схожість, саме це і є їхньою сильною стороною в контексті побудови гібридних систем, де можна адаптивно об'єднувати логіку кількох моделей, щоб створити рекомендації як зі швидким реагуванням на ключові імена, так і з глибоким розумінням смислу історії.

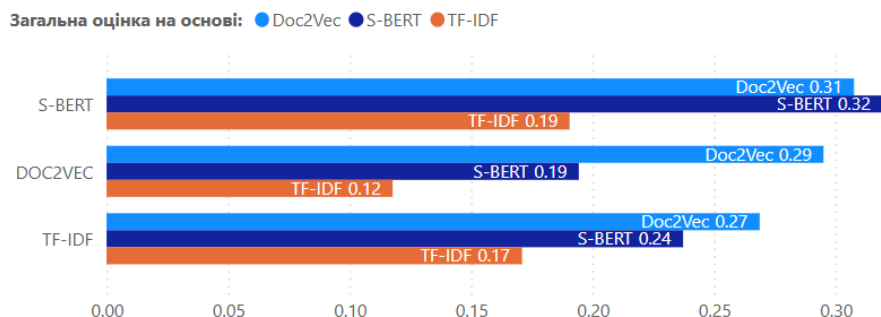


Рисунок 7. Загальні оцінки моделей

Графік, що демонструє загальні оцінки моделей на основі кожного з трьох підходів - Doc2Vec, S-BERT та TF-IDF, - чітко підкріплює перевагу моделі S-BERT. Вона отримує 0.32 за власною методикою, 0.31 за оцінкою Doc2Vec, і 0.19 за TF-IDF. Таке поєднання означає, що навіть з точки зору моделей-конкурентів, рекомендації S-BERT визнаються більш відповідними.

Doc2Vec, у свою чергу, оцінює себе на рівні 0.29, але отримує 0.19 від S-BERT і лише 0.12 від TF-IDF.

TF-IDF, попри свою примітивність, отримує вищу оцінку від S-BERT - 0.24 в порівнянні з оцінкою для Doc2Vec, та оцінку 0.17 за оцінкою від самої себе, що лише на 0.02 менше від результату S-BERT.

Підсумкова загальна оцінка, що є середньою з оцінок всіх моделей, ще раз підкреслює лідерство S-BERT. S-BERT отримала найбільше значення серед усіх (0.27), і воно сформоване за рахунок сильного жанрового покриття та здатності моделі глибоко аналізувати смислові зв'язки в описах фільмів. TF-IDF посідає друге місце - 0.23, завдяки своїм показникам за структурованими полями. Doc2Vec отримує лише 0.20, причину цього вже було пояснено вище.



Рисунок 8. Загальні підсумкові оцінки моделей

Висновки

У результаті проведеного порівняльного аналізу трьох моделей побудови рекомендацій - TF-IDF, Doc2Vec та S-BERT - було виявлено суттєві відмінності в їхніх можливостях, продуктивності та здатності формувати релевантні рекомендації. Найбільш показовою виявилася неспроможність моделі Doc2Vec якісно працювати в умовах, характерних для сучасних рекомендаційних систем.

Doc2Vec, попри свою теоретичну здатність враховувати семантичні зв'язки між словами, виявилася непридатною до використання в динамічному середовищі стрімінгових платформ. Модель навчається один раз, фіксуючи розподілення векторів у заданому корпусі, і не має механізмів адаптації до змін у даних. У ситуаціях, коли додаються нові фільми, з'являються нові жанри чи змінюється мова подачі - Doc2Vec не реагує на ці зміни, залишаючись замкненою в межах застарілих патернів. Крім того, щільна структура векторного простору, у якому зберігаються подання документів, створює ефект згладжування, через що навіть абсолютно не схожі фільми отримують високу оцінку подібності. Це призводить до псевдорелевантних рекомендацій, які не мають цінності для користувача. Важливо й те, що Doc2Vec повністю ігнорує персональні атрибути, такі як актори, режисери чи сценаристи - ключові чинники у виборі фільмів.

На відміну від цього, модель TF-IDF, хоча й базується на простій статистичній концепції частотності термінів, показала значно стабільніші результати. Вона чутлива до ключових слів, тому добре ідентифікує збіги за персоналіями - режисерами, акторами, сценаристами - і здатна швидко перебудовуватись при зміні даних без потреби в перенавчанні. Її головним обмеженням є слабка здатність враховувати синонімію та глибоку семантику: TF-IDF не розуміє значення слів і не може адекватно зіставити тексти з подібним змістом, але різною лексику. Проте, завдяки високій швидкості роботи, інтерпретованості та простоті налаштування, ця модель залишається хорошою стартовою точкою або складовою гібридних систем.

S-BERT, натомість, виявилася найсильнішою та найбільш збалансованою моделлю у проведеному дослідженні. Вона створює контекстуальні векторні подання на основі моделі BERT, що дозволяє їй враховувати повноцінне значення речень. Завдяки архітектурі сіамських мереж і використанню контрастивного навчання, S-BERT ефективно розпізнає семантичні зв'язки між фільмами, навіть якщо їхні описи сформульовані по-різному. У порівнянні з TF-IDF і Doc2Vec, саме ця модель демонструє найвищу жанрову та сюжетну релевантність у рекомендаціях, хоча її оцінки подібності за персоналіями залишаються низькими. Це пояснюється тим, що модель не працює безпосередньо з метаданими, але компенсує це глибинним розумінням контексту. Попри відносну ресурсозатратність, S-BERT є найкращим вибором для задач, де ключову роль відіграє саме смислова близькість текстів.

Підсумовуючи, можна зробити висновок, що Doc2Vec - це модель, яка не відповідає вимогам сучасних рекомендаційних систем у сфері фільмів. Її негнучкість, згладжування векторного простору, низька чутливість до персональних даних і потреба у перенавчанні виключають її з числа рекомендованих інструментів для практичного застосування. Натомість TF-IDF - хороша базова модель, особливо у випадках обмежених ресурсів, а S-BERT - оптимальне рішення для точного семантичного порівняння фільмів. У результаті саме ці дві моделі доцільно розглядати як ядро ефективної системи рекомендацій на базі content-based filtering.

Список використаних джерел

1. Стаття Milvus «What is content-based filtering in recommender systems?». URL: <https://milvus.io/ai-quick-reference/what-is-contentbased-filtering-in-recommender-systems> (дата звернення: 20.05.2025).
2. Стаття IBM «What is content-based filtering?». URL: <https://www.ibm.com/think/topics/content-based-filtering> (дата звернення: 20.05.2025).
3. Стаття Medium «Recommendation Systems: Content-Based Filtering». URL: <https://medium.com/@zbeyza/recommendation-systems-content-based-filtering-e19e3b0a309e> (дата звернення: 20.05.2025).
4. Стаття Medium «Demystifying Latent Dirichlet Allocation: Unveiling The Power of Topic Modeling». URL: <https://ai.plainenglish.io/unveiling-the-power-of-latent-dirichlet-allocation-lda-unleashing-the-potential-of-topic-3947cacbafc2> (дата звернення: 20.05.2025).
5. Стаття Medium «https://medium.com/@readwith_emma/understanding-okapi-bm25-document-ranking-algorithm-70d81adab001». URL: https://medium.com/@readwith_emma/understanding-okapi-bm25-document-ranking-algorithm-70d81adab001 (дата звернення: 20.05.2025).
6. Стаття 360digitmg «Non-Negative Matrix Factorization : Applications & Advantages». URL: <https://360digitmg.com/blog/non-negative-matrix-factorization> (дата звернення: 20.05.2025).
7. Стаття 33rdsquare «Supercharging Text Classification with FastText: Facebook's Powerful NLP Library». URL: <https://33rdsquare.com/tech/ai/word-representations-text-classification-using-fasttext-nlp-facebook/> (дата звернення: 20.05.2025).
8. Стаття spotintelligence «What is a Universal Sentence Encoder?». URL: <https://spotintelligence.com/2024/01/10/universal-sentence-encoder-explained-how-to-tensorflow-tutorial/> (дата звернення: 20.05.2025).
9. Juni Permana A. H. J. P., Agung Toto Wibowo. Movie recommendation system based on synopsis using content-based filtering with TF-IDF and cosine similarity. International journal on information and communication technology (ijoiect). 2023. Т. 9, № 2. С. 1–14. URL: <https://doi.org/10.21108/ijoiect.v9i2.747> (дата звернення: 21.05.2025).

10. N. K., V. N. Influence of pre-processing strategies on sentiment analysis performance: leveraging bert, TF-IDF and glove features. *Journal of machine and computing*. 2025. P. 464–473. URL: <https://doi.org/10.53759/7669/jmc202505036> (дата звернення: 21.05.2025).
11. Xie S., Yang Q. A phrase disambiguation method of “quanbu V de N” based on SBERT model and syntactic rule. *Lecture notes in computer science*. Cham, 2023. C. 364–374. URL: https://doi.org/10.1007/978-3-031-28956-9_29 (дата звернення: 21.05.2025).
12. Reimers N., Gurevych I. Sentence-BERT: sentence embeddings using siamese bert-networks. *Proceedings of the 2019 conference on empirical methods in natural language processing and the 9th international joint conference on natural language processing (EMNLP-IJCNLP)*, м. Hong Kong, China. Stroudsburg, PA, USA, 2019. URL: <https://doi.org/10.18653/v1/d19-1410> (дата звернення: 21.05.2025).
13. Westin F. Time Period Categorization in Fiction: A Comparative Analysis of Machine Learning Techniques. *Cataloging & Classification Quarterly*. 2024. C. 1–30. URL: <https://doi.org/10.1080/01639374.2024.2315548> (дата звернення: 21.05.2025).
14. Multi-co-training for document classification using various document representations: TF-IDF, LDA, and Doc2Vec / D. Kim та ін. *Information Sciences*. 2019. Т. 477. С. 15–29. URL: <https://doi.org/10.1016/j.ins.2018.10.006> (дата звернення: 21.05.2025).
15. Text Vectorization Techniques for Trending Topic Clustering on Twitter: A Comparative Evaluation of TF-IDF, Doc2Vec, and Sentence-BERT / A. D. Susanto та ін. 2023 5th International Conference on Cybernetics and Intelligent System (ICORIS), м. Pangkalpinang, Indonesia, 6–7 жовт. 2023 р. 2023. URL: <https://doi.org/10.1109/icoris60118.2023.10352228> (дата звернення: 21.05.2025).